

# Specifying Concurrent Systems with $\text{TLA}^+$

Leslie Lamport

9 Feb 2000

©1999 by Leslie Lamport

**Preliminary Draft**

**Be sure to read the description of this document on  
page 3 of the Introduction.**



# Contents

|                                                                |           |
|----------------------------------------------------------------|-----------|
| Introduction . . . . .                                         | 1         |
| <b>I Getting Started</b>                                       | <b>5</b>  |
| <b>1 A Little Simple Math</b>                                  | <b>9</b>  |
| 1.1 Propositional Logic . . . . .                              | 9         |
| 1.2 Sets . . . . .                                             | 11        |
| 1.3 Predicate Logic . . . . .                                  | 12        |
| 1.4 Formulas and Language . . . . .                            | 14        |
| <b>2 Specifying a Simple Clock</b>                             | <b>15</b> |
| 2.1 Behaviors . . . . .                                        | 15        |
| 2.2 An Hour Clock . . . . .                                    | 15        |
| 2.3 A Closer Look at the Hour-Clock Specification . . . . .    | 18        |
| 2.4 The Hour-Clock Specification in TLA <sup>+</sup> . . . . . | 19        |
| 2.5 Another Way to Specify the Hour Clock . . . . .            | 21        |
| <b>3 An Asynchronous Interface</b>                             | <b>23</b> |
| 3.1 The First Specification . . . . .                          | 24        |
| 3.2 Another Specification . . . . .                            | 28        |
| 3.3 Types: A Reminder . . . . .                                | 30        |
| 3.4 Definitions . . . . .                                      | 31        |
| 3.5 Comments . . . . .                                         | 32        |
| <b>4 A FIFO</b>                                                | <b>35</b> |
| 4.1 The Inner Specification . . . . .                          | 35        |
| 4.2 Instantiation Examined . . . . .                           | 37        |
| 4.2.1 Instantiation is Substitution . . . . .                  | 37        |
| 4.2.2 Parametrized Instantiation . . . . .                     | 39        |
| 4.2.3 Implicit Substitutions . . . . .                         | 39        |
| 4.2.4 Instantiation Without Renaming . . . . .                 | 40        |
| 4.3 Hiding the Queue . . . . .                                 | 41        |

|           |                                                    |           |
|-----------|----------------------------------------------------|-----------|
| 4.4       | A Bounded FIFO . . . . .                           | 41        |
| 4.5       | What We're Specifying . . . . .                    | 43        |
| <b>5</b>  | <b>A Caching Memory</b>                            | <b>45</b> |
| 5.1       | The Memory Interface . . . . .                     | 45        |
| 5.2       | Functions . . . . .                                | 48        |
| 5.3       | A Linearizable Memory Specification . . . . .      | 50        |
| 5.4       | Tuples as Functions . . . . .                      | 51        |
| 5.5       | Recursive Function Definitions . . . . .           | 53        |
| 5.6       | A Write-Through Cache . . . . .                    | 54        |
| 5.7       | Invariance . . . . .                               | 60        |
| 5.8       | Proving Implementation . . . . .                   | 62        |
| <b>6</b>  | <b>Some More Math</b>                              | <b>65</b> |
| 6.1       | Sets . . . . .                                     | 65        |
| 6.2       | Silly Expressions . . . . .                        | 67        |
| 6.3       | Recursive Function Definitions Revisited . . . . . | 67        |
| 6.4       | Functions versus Operators . . . . .               | 69        |
| 6.5       | Using Functions . . . . .                          | 71        |
| 6.6       | Choose . . . . .                                   | 72        |
| <b>7</b>  | <b>Writing a Specification—Some Advice</b>         | <b>75</b> |
| 7.1       | Why to Specify . . . . .                           | 75        |
| 7.2       | What to Specify . . . . .                          | 76        |
| 7.3       | The Grain of Atomicity . . . . .                   | 76        |
| 7.4       | The Data Structures . . . . .                      | 78        |
| 7.5       | Writing the Specification . . . . .                | 79        |
| 7.6       | Some Further Hints . . . . .                       | 79        |
| 7.7       | When and How to Specify . . . . .                  | 82        |
| <b>II</b> | <b>More Advanced Topics</b>                        | <b>85</b> |
| <b>8</b>  | <b>Liveness and Fairness</b>                       | <b>87</b> |
| 8.1       | Temporal Formulas . . . . .                        | 88        |
| 8.2       | Temporal Tautologies . . . . .                     | 92        |
| 8.3       | Weak Fairness . . . . .                            | 95        |
| 8.4       | The Memory Specification . . . . .                 | 100       |
| 8.4.1     | The Liveness Requirement . . . . .                 | 100       |
| 8.4.2     | Another Way to Write It . . . . .                  | 101       |
| 8.4.3     | A Generalization . . . . .                         | 104       |
| 8.5       | Strong Fairness . . . . .                          | 105       |
| 8.6       | Liveness for the Write-Through Cache . . . . .     | 106       |
| 8.7       | Quantification . . . . .                           | 108       |

|           |                                               |            |
|-----------|-----------------------------------------------|------------|
| 8.8       | Temporal Logic Examined . . . . .             | 110        |
| 8.8.1     | A Review . . . . .                            | 110        |
| 8.8.2     | Machine Closure . . . . .                     | 110        |
| 8.8.3     | Machine Closure and Possibility . . . . .     | 112        |
| 8.8.4     | The Unimportance of Liveness . . . . .        | 113        |
| 8.8.5     | Temporal Logic Considered Confusing . . . . . | 113        |
| <b>9</b>  | <b>Real Time</b>                              | <b>115</b> |
| 9.1       | The Hour Clock Revisited . . . . .            | 115        |
| 9.2       | Real-Time Specifications in General . . . . . | 118        |
| 9.3       | The Real-Time Write-Through Cache . . . . .   | 122        |
| 9.4       | Zeno Specifications . . . . .                 | 127        |
| 9.5       | Hybrid System Specifications . . . . .        | 129        |
| 9.6       | Remarks on Real Time . . . . .                | 131        |
| <b>10</b> | <b>Composing Specifications</b>               | <b>133</b> |
| 10.1      | Composing Two Specifications . . . . .        | 134        |
| 10.2      | Composing Many Specifications . . . . .       | 136        |
| 10.3      | The FIFO . . . . .                            | 139        |
| 10.4      | Composition with Shared State . . . . .       | 142        |
| 10.4.1    | Explicit State Changes . . . . .              | 142        |
| 10.4.2    | Composition with Joint Actions . . . . .      | 145        |
| 10.5      | A Brief Review . . . . .                      | 149        |
| 10.5.1    | A Taxonomy of Composition . . . . .           | 149        |
| 10.5.2    | Interleaving Reconsidered . . . . .           | 149        |
| 10.5.3    | Joint Actions Reconsidered . . . . .          | 150        |
| 10.6      | Liveness and Hiding . . . . .                 | 151        |
| 10.6.1    | Liveness and Machine Closure . . . . .        | 151        |
| 10.6.2    | Hiding . . . . .                              | 152        |
| 10.7      | Open-System Specifications . . . . .          | 154        |
| 10.8      | Interface Refinement . . . . .                | 156        |
| 10.8.1    | A Binary Hour Clock . . . . .                 | 156        |
| 10.8.2    | Refining a Channel . . . . .                  | 158        |
| 10.8.3    | Interface Refinement in General . . . . .     | 161        |
| 10.8.4    | Open-System Specifications . . . . .          | 163        |
| 10.9      | Should You Compose? . . . . .                 | 165        |
| <b>11</b> | <b>Advanced Examples</b>                      | <b>167</b> |
| 11.1      | Specifying Data Structures . . . . .          | 168        |
| 11.1.1    | Local Definitions . . . . .                   | 168        |
| 11.1.2    | Graphs . . . . .                              | 170        |
| 11.1.3    | Solving Differential Equations . . . . .      | 172        |
| 11.1.4    | The Riemann Integral . . . . .                | 177        |
| 11.2      | Other Memory Specifications . . . . .         | 178        |

|            |                                                |            |
|------------|------------------------------------------------|------------|
| 11.2.1     | The Interface . . . . .                        | 178        |
| 11.2.2     | The Correctness Condition . . . . .            | 179        |
| 11.2.3     | A Serial Memory . . . . .                      | 181        |
| 11.2.4     | A Sequentially Consistent Memory . . . . .     | 188        |
| 11.2.5     | The Memory Specifications Considered . . . . . | 193        |
| <b>III</b> | <b>The Tools</b>                               | <b>197</b> |
| <b>12</b>  | <b>The Java Front End</b>                      | <b>199</b> |
| 12.1       | Finding an Error . . . . .                     | 199        |
| <b>13</b>  | <b>The TLC Model Checker</b>                   | <b>201</b> |
| 13.1       | Introduction to TLC . . . . .                  | 201        |
| 13.2       | How TLC Works . . . . .                        | 208        |
| 13.2.1     | TLC Values . . . . .                           | 208        |
| 13.2.2     | How TLC Evaluates Expressions . . . . .        | 209        |
| 13.2.3     | Assignment and Replacement . . . . .           | 211        |
| 13.2.4     | Overriding Modules . . . . .                   | 213        |
| 13.2.5     | How TLC Computes States . . . . .              | 213        |
| 13.2.6     | When TLC Computes What . . . . .               | 214        |
| 13.2.7     | Random Simulation . . . . .                    | 215        |
| 13.3       | How to Use TLC . . . . .                       | 216        |
| 13.3.1     | Running TLC . . . . .                          | 216        |
| 13.3.2     | Debugging a Specification . . . . .            | 217        |
| 13.3.3     | The <i>TLC</i> Module . . . . .                | 222        |
| 13.3.4     | Checking Action Invariance . . . . .           | 224        |
| 13.3.5     | Checking Implementation . . . . .              | 227        |
| 13.3.6     | Some Hints . . . . .                           | 228        |
| 13.4       | What TLC Doesn't Do . . . . .                  | 231        |
| 13.5       | Future Plans . . . . .                         | 233        |
| 13.6       | The Fine Print: What TLC Really Does . . . . . | 234        |
| 13.6.1     | TLC Values . . . . .                           | 234        |
| 13.6.2     | Overridden Values . . . . .                    | 235        |
| 13.6.3     | Expression Evaluation . . . . .                | 235        |
| 13.6.4     | Fingerprinting . . . . .                       | 244        |
| <b>IV</b>  | <b>The TLA<sup>+</sup> Language</b>            | <b>245</b> |
| <b>14</b>  | <b>The Syntax of TLA<sup>+</sup></b>           | <b>255</b> |
| 14.1       | The Simple Grammar . . . . .                   | 256        |
| 14.1.1     | BNF Grammars . . . . .                         | 256        |
| 14.1.2     | The BNF Grammar of TLA <sup>+</sup> . . . . .  | 260        |

|           |                                                          |            |
|-----------|----------------------------------------------------------|------------|
| 14.2      | The Complete Grammar . . . . .                           | 268        |
| 14.2.1    | Precedence and Associativity . . . . .                   | 268        |
| 14.2.2    | Alignment . . . . .                                      | 271        |
| 14.2.3    | Comments . . . . .                                       | 272        |
| 14.2.4    | Temporal Formulas . . . . .                              | 273        |
| 14.2.5    | Two Anomalies . . . . .                                  | 274        |
| 14.3      | What You Type . . . . .                                  | 274        |
| <b>15</b> | <b>The Operators of TLA<sup>+</sup></b>                  | <b>275</b> |
| 15.1      | Constant Operators . . . . .                             | 275        |
| 15.1.1    | Boolean Operators . . . . .                              | 277        |
| 15.1.2    | The Choose Operator . . . . .                            | 278        |
| 15.1.3    | The Three Interpretations of Boolean Operators . . . . . | 280        |
| 15.1.4    | Conditional Constructs . . . . .                         | 282        |
| 15.1.5    | LET . . . . .                                            | 282        |
| 15.1.6    | The Operators of Set Theory . . . . .                    | 283        |
| 15.1.7    | Functions . . . . .                                      | 285        |
| 15.1.8    | Records . . . . .                                        | 289        |
| 15.1.9    | Tuples . . . . .                                         | 290        |
| 15.1.10   | Strings . . . . .                                        | 291        |
| 15.1.11   | Numbers . . . . .                                        | 292        |
| 15.2      | Nonconstant Operators . . . . .                          | 293        |
| 15.2.1    | The Meaning of a Basic Constant Expression . . . . .     | 293        |
| 15.2.2    | The Meaning of a State Function . . . . .                | 294        |
| 15.2.3    | Action Operators . . . . .                               | 296        |
| 15.2.4    | Temporal Operators . . . . .                             | 298        |
| <b>16</b> | <b>The Meaning of a Module</b>                           | <b>301</b> |
| 16.1      | Operators and Expressions . . . . .                      | 301        |
| 16.1.1    | The Order and Arity of an Operator . . . . .             | 302        |
| 16.1.2    | $\lambda$ Expressions . . . . .                          | 303        |
| 16.1.3    | Simplifying Operator Application . . . . .               | 304        |
| 16.1.4    | Expressions . . . . .                                    | 305        |
| 16.2      | Levels . . . . .                                         | 305        |
| 16.3      | Contexts . . . . .                                       | 308        |
| 16.4      | The Meaning of a $\lambda$ Expression . . . . .          | 309        |
| 16.5      | The Meaning of a Module . . . . .                        | 311        |
| 16.5.1    | Extends . . . . .                                        | 312        |
| 16.5.2    | Declarations . . . . .                                   | 312        |
| 16.5.3    | Operator Definitions . . . . .                           | 313        |
| 16.5.4    | Function Definitions . . . . .                           | 313        |
| 16.5.5    | Instantiation . . . . .                                  | 314        |
| 16.5.6    | Theorems and Assumptions . . . . .                       | 315        |

---

|                                               |            |
|-----------------------------------------------|------------|
| 16.5.7 Submodules . . . . .                   | 315        |
| 16.6 Correctness of a Module . . . . .        | 316        |
| 16.7 Finding Modules . . . . .                | 316        |
| 16.8 The Semantics of Instantiation . . . . . | 317        |
| <b>17 The Standard Modules</b>                | <b>323</b> |
| 17.1 Module <i>Sequences</i> . . . . .        | 323        |
| 17.2 Module <i>FiniteSets</i> . . . . .       | 324        |
| 17.3 Module <i>Bags</i> . . . . .             | 324        |
| 17.4 The Numbers Modules . . . . .            | 328        |
| <b>V Appendix</b>                             | <b>333</b> |
| <b>A The ASCII Specifications</b>             | <b>335</b> |
| A.1 The Asynchronous Interface . . . . .      | 335        |
| A.2 A FIFO . . . . .                          | 336        |
| A.3 A Caching Memory . . . . .                | 337        |
| A.4 The Alternating Bit Protocol . . . . .    | 342        |
| <b>References</b>                             | <b>345</b> |

# List of Figures

|      |                                                                                     |     |
|------|-------------------------------------------------------------------------------------|-----|
| 2.1  | The hour clock specification—typeset and ASCII versions. . . . .                    | 20  |
| 3.1  | The First Specification of an Asynchronous Interface . . . . .                      | 27  |
| 3.2  | Our second specification of an asynchronous interface. . . . .                      | 30  |
| 3.3  | The hour clock specification with comments. . . . .                                 | 33  |
| 4.1  | The specification of a FIFO, with the internal variable $q$ visible. . . . .        | 38  |
| 4.2  | Specification of a FIFO buffer of length $N$ . . . . .                              | 43  |
| 5.1  | The Specification of a Memory Interface . . . . .                                   | 48  |
| 5.2  | The internal memory specification (beginning). . . . .                              | 52  |
| 5.3  | The internal memory specification (end). . . . .                                    | 53  |
| 5.4  | The memory specification. . . . .                                                   | 53  |
| 5.5  | The write-through cache specification (beginning). . . . .                          | 56  |
| 5.6  | The write-through cache specification (middle). . . . .                             | 57  |
| 5.7  | The write-through cache specification (end). . . . .                                | 58  |
| 9.1  | The real-time specification of an hour clock. . . . .                               | 119 |
| 9.2  | The <i>RealTime</i> module for writing real-time specifications. . . . .            | 122 |
| 9.3  | A real-time version of the linearizable memory specification. . . . .               | 123 |
| 9.4  | A real-time version of the write-through cache (beginning). . . . .                 | 126 |
| 9.5  | A real-time version of the write-through cache (end). . . . .                       | 127 |
| 10.1 | A noninterleaving composite specification of the FIFO. . . . .                      | 141 |
| 10.2 | A joint-action compositional specification of the linearizable mem-<br>ory. . . . . | 148 |
| 10.3 | A specification of a binary hour clock. . . . .                                     | 158 |
| 10.4 | Refining a channel. . . . .                                                         | 160 |
| 11.1 | A module for specifying operators on graphs. . . . .                                | 173 |
| 11.2 | A module for specifying the solution to a differential equation. . . . .            | 176 |
| 11.3 | A specification of the Riemann integral. . . . .                                    | 177 |
| 11.4 | A module for specifying a register interface to a memory. . . . .                   | 179 |

|       |                                                                   |     |
|-------|-------------------------------------------------------------------|-----|
| 11.5  | Module <i>InnerSerial</i> (beginning).                            | 189 |
| 11.6  | Module <i>InnerSerial</i> (middle).                               | 190 |
| 11.7  | Module <i>InnerSerial</i> (end).                                  | 191 |
| 11.8  | Module <i>InnerSequential</i> (beginning).                        | 194 |
| 11.9  | Module <i>InnerSequential</i> (end).                              | 195 |
|       |                                                                   |     |
| 13.1  | The Alternating Bit Protocol (beginning)                          | 204 |
| 13.2  | The Alternating Bit Protocol (end)                                | 205 |
| 13.3  | Module <i>MCAAlternatingBit</i> .                                 | 206 |
| 13.4  | A configuration file for module <i>MCAAlternatingBit</i> .        | 207 |
| 13.5  | The standard module <i>TLC</i>                                    | 223 |
| 13.6  | An invariant of the alternating bit protocol's next-state action. | 225 |
| 13.7  | A higher-level specification of the alternating bit protocol.     | 228 |
| 13.8  | A module for checking the alternating bit protocol.               | 229 |
| 13.9  | The constant operators.                                           | 248 |
| 13.10 | Miscellaneous constructs.                                         | 249 |
| 13.11 | Action operators.                                                 | 249 |
| 13.12 | Temporal operators.                                               | 249 |
| 13.13 | User-definable operator symbols.                                  | 250 |
| 13.14 | The precedence ranges of operators.                               | 251 |
| 13.15 | Operators defined in the standard modules.                        | 252 |
| 13.16 | The ASCII representations of typeset symbols.                     | 253 |
|       |                                                                   |     |
| 14.1  | The definition of the grammar <i>GSE</i> for the language SE.     | 260 |
| 14.2  | The module <i>BNFGrammars</i>                                     | 261 |
|       |                                                                   |     |
| 17.1  | The standard <i>Sequences</i> module.                             | 325 |
| 17.2  | The standard <i>FiniteSets</i> module.                            | 325 |
| 17.3  | The standard <i>Bags</i> module.                                  | 327 |
| 17.4  | The <i>Peano</i> module.                                          | 329 |
| 17.5  | The <i>ProtoReals</i> module (beginning).                         | 330 |
| 17.6  | The <i>ProtoReals</i> module (end).                               | 331 |
| 17.7  | The standard <i>Naturals</i> module.                              | 332 |
| 17.8  | The standard <i>Integers</i> module.                              | 332 |
| 17.9  | The standard <i>Reals</i> module.                                 | 332 |

# Introduction

*Writing is nature's way of letting you  
know how sloppy your thinking is.*

Guindon

Writing a specification for a system helps us understand it. It's a good idea to understand something before you build it, so it's a good idea to specify a system before you implement it.

Mathematics is nature's way of letting you know how sloppy your writing is. Specifications written in an imprecise language like English are usually imprecise. In engineering, imprecision is an invitation to error. Science and engineering have adopted mathematics as a language for writing precise descriptions.

Formal mathematics is nature's way of letting you know how sloppy your mathematics is. The mathematics written by most mathematicians and scientists is still imprecise. Most mathematics texts are precise in the small, but imprecise in the large. Each equation is a precise assertion, but you have to read the text to understand how the equations relate to one another, and what the theorems really mean. Logicians have developed ways of eliminating the words and formalizing mathematics.

Most mathematicians and computer scientists think that writing mathematics formally, without words, is tiresome. I've asked a number of computer scientists the following question: How long would a formal specification of the Riemann integral of elementary calculus be, assuming only arithmetic operations on real numbers? The answers I received ranged up to 50 pages. Section 11.1.4 shows how to do it in about 20 lines. Once you learn how, it's easy to express ordinary mathematics in a precise, completely formal language.

To specify systems with mathematics, we must decide what kind of mathematics to use. We can specify an ordinary sequential program by describing its output as a function of its input. So, sequential programs can be specified in terms of functions. Concurrent systems are usually described in terms of their behaviors—what they do in the course of an execution. In 1977, Amir Pnueli introduced the use of temporal logic for describing such behaviors.

Temporal logic is appealing because, in principle, it allows a concurrent system to be described by a single formula. In practice, temporal logic proved to be

cumbersome. Pnueli's temporal logic was ideal for describing some properties of systems, but awkward for others. So, it was usually combined with some more traditional way of describing systems.

In the late 1980's, I discovered TLA, the Temporal Logic of Actions. TLA is a simple variant of Pnueli's original logic that makes it practical to write a specification as a single formula. Most of a TLA specification consists of ordinary, nontemporal mathematics. Temporal logic plays a significant role only in describing those properties that it's good at describing. TLA also provides a nice way to formalize the style of reasoning about systems that has proved to be most effective in practice—a style known as *assertional* reasoning. However, the topic of this document is specification, not proof, so I will have little to say about proofs.

TLA provides a mathematical foundation for describing concurrent systems. To write specifications, we need a complete language built atop that foundation. I initially thought that this language should be some sort of abstract programming language whose semantics would be based on TLA. I didn't know what kind of programming language constructs would be best, so I decided to start writing specifications directly in TLA. I intended to introduce programming constructs as I needed them. To my surprise, I discovered that I didn't need them. What I needed was a robust language for writing mathematics.

Although mathematicians have developed the science of writing formulas, they haven't turned that science into an engineering discipline. They have developed notations for mathematics in the small, but not for mathematics in the large. The specification of a real system can be dozens or even hundreds of pages long. Mathematicians know how to write 20-line formulas, not 20-page formulas. So, I had to introduce notations for writing long formulas. What I took from programming languages were ideas for modularizing large specifications.

The language I came up with is called  $TLA^+$ . I refined  $TLA^+$  in the course of writing specifications of disparate systems. But it has changed little in the last few years. I have found  $TLA^+$  to be quite good for specifying a wide class of systems—from program interfaces (APIs) to distributed systems. It can be used to write a precise, formal description of almost any sort of discrete system. It's especially well suited to describing asynchronous systems—that is, systems with components that do not operate in strict lock-step.

One advantage of a precise specification language is that it enables us to build tools that can help us write correct specifications. There are now at least two such tools under development: a parser and a model checker, described in Part III. The parser can catch simple errors in any  $TLA^+$  specification. The model checker can catch many more errors, but it works on a restricted class of specifications—a class that seems to include most of the specifications of interest to industry today.

## The State of this Document

This document is a preliminary draft. Here is a brief description of the individual parts what and who should read them.

### Part I

These chapters are complete and shouldn't have too many errors. They are an introduction and should be read by everyone interested in using TLA<sup>+</sup>. They explain how to specify the class of properties known as *safety* properties. These properties, which can be specified with almost no temporal logic, are all that most engineers need to know about.

### Part II

Temporal logic comes to the fore in Chapter 8, where it is used to specify the additional class of properties known as liveness properties. This chapter is in pretty good shape, but probably has more errors per page than the preceding chapters. The remaining chapters in this part are rough drafts and are full of errors. Chapter 9 describes how to specify real-time properties, and Chapter 10 describes how to write specifications as compositions. Chapter 11 describes some more advanced examples.

### Part III

This part describes the parser and the TLC model checker. If you are reading this because you want to use TLA<sup>+</sup>, then you'll probably want to use these tools and should read these chapters. They are preliminary and have lots of errors. Before trying to use TLC be sure to read Section 13.4 on page 231; it describes limitations of the current version of the program.

### Part IV

This part is a reference manual for the language. It has not been read carefully and is undoubtedly full of errors. Part I should give you a good enough working knowledge of the language for most of your needs. Part IV describes the fine points of the syntax and semantics; it also contains the standard modules. Chapter 14 gives the syntax of TLA<sup>+</sup> and includes a BNF grammar (written in TLA<sup>+</sup>). Chapter 15 describes the precise meanings and the general forms of all the built-in operators of TLA<sup>+</sup>. It should answer any questions you might have about exactly a TLA<sup>+</sup> operator means. Chapter 16 describes the precise meaning of all the higher-level TLA<sup>+</sup> constructs, such as definitions and module inclusion. Chapters 15 and 16 together specify the semantics of the language.

Chapter 17 describes the standard modules—except for module *RealTime*, described in Chapter 9, and module *TLC*, described in Chapter 13. You might want to look at this chapter if you’re curious about how standard elementary mathematics can be formalized in TLA<sup>+</sup>.

You will seldom have occasion to read any of Part IV. However, it does have something you may want to refer to often: a mini-manual that compactly presents lots of useful information. Pages 248–253 list all TLA<sup>+</sup> operators, all user-definable symbols, the precedence of all operators, all operators defined in the standard modules, and the ASCII representation of symbols like  $\otimes$ .

## The Appendix

The specifications that appear in the book are typeset for easy reading by humans. To be read by a tool, a specification must be written in ASCII. The appendix includes the ASCII versions of all the specifications in Part I, as well as the specifications from Chapter 13. Comparing the ASCII and the typeset version can teach you how to write TLA<sup>+</sup> specifications in ASCII.

# Part I

## Getting Started



A system specification consists of a lot of ordinary mathematics glued together with a little bit of temporal logic. So, most of the work in writing a precise specification consists of expressing ordinary mathematics precisely. That's why most of the details of  $\text{TLA}^+$  are concerned with expressing ordinary mathematics.

Unfortunately, the computer science departments in many universities apparently believe that fluency in C++ is more important than a sound education in elementary mathematics. So, some readers may be unfamiliar with the mathematics needed to write specifications. Fortunately, this mathematics is quite simple. If overexposure to C++ hasn't destroyed your ability to think logically, you should have no trouble filling any gaps in your mathematics education. You probably learned arithmetic before being exposed to C++, so I will assume you know about numbers and arithmetic operations on them.<sup>1</sup> I will try to explain all other mathematical concepts that you need, starting in Chapter 1 with a review of some elementary math. I hope most readers will find this review completely unnecessary.

After the brief review of simple mathematics in the next section, Chapters 2 through 5 describe  $\text{TLA}^+$  with a sequence of examples. Chapter 6 explains some more about the math used in writing specifications, and Chapter 7 reviews everything and provides some advice. By the time you finish Chapter 7, you should be able to handle most of the specification problems that you are likely to encounter in ordinary engineering practice.

---

<sup>1</sup>Some readers may need reminding that numbers are not strings of bits, and  $2^{33} * 2^{33}$  equals  $2^{66}$ , not *overflow error*.



# Chapter 1

## A Little Simple Math

### 1.1 Propositional Logic

Elementary algebra is the mathematics of real numbers and the operators  $+$ ,  $-$ ,  $*$  (multiplication), and  $/$  (division). Propositional logic is the mathematics of the two Boolean values TRUE and FALSE and the five operators whose names (and common pronunciations) are:

|                            |                                         |
|----------------------------|-----------------------------------------|
| $\wedge$ conjunction (and) | $\Rightarrow$ implication (implies)     |
| $\vee$ disjunction (or)    | $\equiv$ equivalence (is equivalent to) |
| $\neg$ negation (not)      |                                         |

To learn how to compute with numbers, you had to memorize addition and multiplication tables and algorithms for calculating with multidigit numbers. Propositional logic is much simpler, since there are only two values, TRUE and FALSE. To learn how to compute with these values, all you need to know are the following definitions of the five Boolean operators:

- $\wedge$   $F \wedge G$  equals TRUE iff both  $F$  and  $G$  equal TRUE.
- $\vee$   $F \vee G$  equals TRUE iff  $F$  or  $G$  equals TRUE (or both do).
- $\neg$   $\neg F$  equals TRUE iff  $F$  equals FALSE.
- $\Rightarrow$   $F \Rightarrow G$  equals TRUE iff  $F$  equals FALSE or  $G$  equals TRUE (or both).
- $\equiv$   $F \equiv G$  equals TRUE iff  $F$  and  $G$  both equal TRUE or both equal FALSE.

*iff* stands for *if and only if*. Like most mathematicians, I use *or* to mean *and/or*.

We can also describe these operators by *truth tables*. This truth table for  $F \Rightarrow G$  gives its value for all four combinations of truth values of  $F$  and  $G$ :

| $F$   | $G$   | $F \Rightarrow G$ |
|-------|-------|-------------------|
| TRUE  | TRUE  | TRUE              |
| TRUE  | FALSE | FALSE             |
| FALSE | TRUE  | TRUE              |
| FALSE | FALSE | TRUE              |

The formula  $F \Rightarrow G$  asserts that  $F$  implies  $G$ . People often find the definition of  $\Rightarrow$  confusing. They don't understand why  $\text{FALSE} \Rightarrow \text{TRUE}$  and  $\text{FALSE} \Rightarrow \text{FALSE}$  should equal  $\text{TRUE}$ . The explanation is simple. We expect that if  $n$  is greater than 3 then it should be greater than 1, so  $n > 3$  should imply  $n > 1$ . Substituting 4, 2, and 0 for  $n$  in the formula  $(n > 3) \Rightarrow (n > 0)$  explains why  $F \Rightarrow G$  means  $F$  *implies*  $G$  or, equivalently, *if*  $F$  *then*  $G$ .

The equivalence operator  $\equiv$  is equality for Booleans. We can replace  $\equiv$  by  $=$ , but not vice versa. (We can write  $\text{FALSE} = \neg \text{TRUE}$ , but not  $2 + 2 \equiv 4$ .) It's a good idea to write  $\equiv$  instead of  $=$  to make it clear that the equal expressions are Booleans.<sup>1</sup>

Just like formulas of algebra, formulas of propositional logic are made up of values, operators, and identifiers like  $x$  that stand for values. However, propositional-logic formulas use only the two values  $\text{TRUE}$  and  $\text{FALSE}$  and the five Boolean operators  $\wedge$ ,  $\vee$ ,  $\neg$ ,  $\Rightarrow$ , and  $\equiv$ . In algebraic formulas,  $*$  has higher precedence (binds more tightly) than  $+$ , so  $x + y * z$  means  $x + (y * z)$ . Similarly,  $\neg$  has higher precedence than  $\wedge$  and  $\vee$ , which have higher precedence than  $\Rightarrow$  and  $\equiv$ , so  $\neg F \wedge G \Rightarrow H$  means  $((\neg F) \wedge G) \Rightarrow H$ . Other mathematical operators like  $+$  and  $>$  have higher precedence than the operators of propositional logic, so  $n > 0 \Rightarrow n - 1 \geq 0$  means  $(n > 0) \Rightarrow (n - 1 \geq 0)$ . Redundant parentheses can't hurt and often make a formula easier to read. If you have the slightest doubt about whether parentheses are needed, use them.

The operators  $\wedge$  and  $\vee$  are associative, just like  $+$  and  $*$ . Associativity of  $+$  means that  $x + (y + z)$  equals  $(x + y) + z$ , so we can write  $x + y + z$  without parentheses. Similarly, associativity of  $\wedge$  and  $\vee$  lets us write  $F \wedge G \wedge H$  or  $F \vee G \vee H$ . Like  $+$  and  $*$ , the operators  $\wedge$  and  $\vee$  are also commutative, so  $F \wedge G$  is equivalent to  $G \wedge F$ , and  $F \vee G$  is equivalent to  $G \vee F$ .

The truth of the formula  $(x = 2) \Rightarrow (x + 1 = 3)$  expresses a fact about numbers. To determine that it's true, we have to understand some elementary properties of arithmetic. However, we can tell that  $(x = 2) \Rightarrow (x = 2) \vee (y > 7)$  is true even if we know nothing about numbers. This formula is true because  $F \Rightarrow F \vee G$  is true, regardless of what the formulas  $F$  and  $G$  are. In other words,  $F \Rightarrow F \vee G$  is true for all possible truth values of its identifiers  $F$  and  $G$ . Such a formula is called a *tautology*.

<sup>1</sup>Section 15.1.3 explains a more subtle reason for using  $\equiv$  instead of  $=$  for equality of Boolean values.

In general, a tautology of propositional logic is a propositional-logic formula that is true for all possible truth values of its identifiers. One can prove all tautologies from a few simple axioms and rules. However, that would be like computing  $437 + 256$  from the axioms of arithmetic. It's usually easier to verify that a simple formula is a tautology by writing its truth table—that is, by directly calculating the value of the formula for all possible truth values of its components. The formula is a tautology iff it equals TRUE for all these values. To construct the truth table for a formula, we construct the truth table for all its subformulas. For example, the following truth table shows that  $(F \Rightarrow G) \equiv (\neg F \vee G)$  is a tautology.

| $F$   | $G$   | $F \Rightarrow G$ | $\neg F$ | $\neg F \vee G$ | $(F \Rightarrow G) \equiv \neg F \vee G$ |
|-------|-------|-------------------|----------|-----------------|------------------------------------------|
| TRUE  | TRUE  | TRUE              | FALSE    | TRUE            | TRUE                                     |
| TRUE  | FALSE | FALSE             | FALSE    | FALSE           | TRUE                                     |
| FALSE | TRUE  | TRUE              | TRUE     | TRUE            | TRUE                                     |
| FALSE | FALSE | TRUE              | TRUE     | TRUE            | TRUE                                     |

Propositional logic is the basis of all mathematical reasoning. It should be as familiar to you as simple algebra. Checking that  $\neg(F \vee G) \equiv \neg F \wedge \neg G$  is a tautology should be as easy as checking that  $2 * (x + 3 * y)$  equals  $2 * x + 6 * y$ .

Although propositional logic is simple, complex propositional formulas can get confusing. You may find yourself trying to simplify some formula and not being sure if the simplified version means the same thing as the original. There exist a number of programs for verifying propositional logic tautologies; some of them can be found and used on the World Wide Web.

## 1.2 Sets

Set theory is the foundation of ordinary mathematics. A set is often described as a collection of elements, but saying that a set is a collection doesn't explain very much. The concept of set is so fundamental that we don't try to define it. We take as undefined concepts the notion of a set and the relation  $\in$ , where  $x \in S$  means that  $x$  is an element of  $S$ . We often say *is in* instead of *is an element of*.

A set can have a finite or infinite number of elements. The set of all natural numbers (0, 1, 2, etc.) is an infinite set. The set of all natural numbers less than 3 is finite, and contains the three elements 0, 1, and 2. We can write this set  $\{0, 1, 2\}$ .

A set is completely determined by its elements. Two sets are equal iff they have the same elements. Thus,  $\{0, 1, 2\}$  and  $\{2, 1, 0\}$  and  $\{0, 0, 1, 2, 2\}$  are all the same set—the unique set containing the three elements 0, 1, and 2. The empty set, which we write  $\{\}$ , is the unique set that has no elements.

The most common operations on sets are:

$\cap$  intersection       $\cup$  union       $\subseteq$  subset       $\setminus$  set difference

Here are their definitions and examples of their use:

$S \cap T$  The set of elements in both  $S$  and  $T$ .  
 $\{1, -1/2, 3\} \cap \{1, 2, 3, 5, 7\} = \{1, 3\}$

$S \cup T$  The set of elements in  $S$  or  $T$  (or both).  
 $\{1, -1/2\} \cup \{1, 5, 7\} = \{1, -1/2, 5, 7\}$

$S \subseteq T$  True iff every element of  $S$  is an element of  $T$ .  
 $\{1, 3\} \subseteq \{3, 2, 1\}$

$S \setminus T$  The set of elements in  $S$  that are not in  $T$ .  
 $\{1, -1/2, 3\} \setminus \{1, 5, 7\} = \{-1/2, 3\}$

This is all you need to know about sets before we start looking at how to specify systems. We'll return to set theory in Section 6.1.

### 1.3 Predicate Logic

Once we have sets, it's natural to say that some formula is true for all the elements of a set, or for some of the elements of a set. Predicate logic extends propositional logic with the two quantifiers:

$\forall$  universal quantification (for all)  
 $\exists$  existential quantification (there exists)

The formula  $\forall x \in S : F$  asserts that formula  $F$  is true for every element  $x$  in the set  $S$ . For example,  $\forall n \in \text{Nat} : n + 1 > n$  asserts that the formula  $n + 1 > n$  is true for all elements  $n$  of the set  $\text{Nat}$  of natural numbers. This formula happens to be true.

The formula  $\exists x \in S : F$  asserts that formula  $F$  is true for at least one element  $x$  in  $S$ . For example,  $\exists n \in \text{Nat} : n^2 = 2$  asserts that there exists a natural number  $n$  whose square equals 2. This formula happens to be false.

Formula  $F$  is true for some  $x \in S$  iff  $F$  is not false for all  $x \in S$ —that is, iff it's not the case that  $\neg F$  is true for all  $x \in S$ . Hence, the formula

$$(1.1) \quad (\exists x \in S : F) \equiv \neg(\forall x \in S : \neg F)$$

is a tautology of predicate logic.<sup>2</sup>

---

<sup>2</sup>Strictly speaking,  $\in$  isn't an operator of predicate logic, so this isn't really a predicate-logic tautology.

Since there exists no element in the empty set, the formula  $\exists x \in \{\} : F$  is false for every formula  $F$ . By (1.1), this implies that  $\forall x \in \{\} : F$  must be true for every  $F$ .

The quantification in the formulas  $\forall x \in S : F$  and  $\exists x \in S : F$  is said to be *bounded*, since these formulas make an assertion only about elements in the set  $S$ . There is also unbounded quantification. The formula  $\forall x : F$  asserts that  $F$  is true for all values  $x$ , and  $\exists x : F$  asserts that  $F$  is true for at least one value of  $x$ —a value that is not constrained to be in any particular set. Bounded and unbounded quantification are related by the following tautologies:

$$\begin{aligned} (\forall x \in S : F) &\equiv (\forall x : (x \in S) \Rightarrow F) \\ (\exists x \in S : F) &\equiv (\exists x : (x \in S) \wedge F) \end{aligned}$$

The analog of (1.1) for unbounded quantifiers is also a tautology:

$$(\exists x : F) \equiv \neg(\forall x : \neg F)$$

Whenever possible, it is better to use bounded than unbounded quantification in a specification. This makes the specification easier for both people and tools to understand.

Universal quantification generalizes conjunction. If  $S$  is a finite set, then  $\forall x \in S : F$  is the conjunction of the formulas obtained by substituting the different elements of  $S$  for  $x$  in  $F$ . For example,

$$(\forall x \in \{2, 3, 7\} : x < y^x) \equiv (2 < y^2) \wedge (3 < y^3) \wedge (7 < y^7)$$

We sometimes informally talk about the conjunction of an infinite number of formulas when we formally mean a universally quantified formula. For example, the conjunction of the formulas  $x \leq y^x$  for all natural numbers  $x$  is the formula  $\forall x \in \text{Nat} : x \leq y^x$ . Similarly, existential quantification generalizes disjunction.

Logicians have rules for proving predicate-logic tautologies such as (1.1), but you shouldn't need them. You should become familiar enough with predicate logic that simple tautologies are obvious. Thinking of  $\forall$  as conjunction and  $\exists$  as disjunction can help. For example, the associativity and commutativity of conjunction and disjunction lead to the tautologies:

$$\begin{aligned} (\forall x \in S : F) \wedge (\forall x \in S : G) &\equiv (\forall x \in S : F \wedge G) \\ (\exists x \in S : F) \vee (\exists x \in S : G) &\equiv (\exists x \in S : F \vee G) \end{aligned}$$

for any set  $S$  and formulas  $F$  and  $G$ .

Mathematicians use some obvious abbreviations for nested quantifiers. For example:

$$\begin{aligned} \forall x \in S, y \in T : F &\text{ means } \forall x \in S : (\forall y \in T : F) \\ \exists w, x, y, z \in S : F &\text{ means } \exists w \in S : (\exists x \in S : (\exists y \in S : (\exists z \in S : F))) \end{aligned}$$

In the expression  $\exists x \in S : F$ , logicians say that  $x$  is a *bound variable* and that occurrences of  $x$  in  $F$  are *bound*. For example,  $n$  is a bound variable in the formula  $\exists n \in \text{Nat} : n + 1 > n$ , and the two occurrences of  $n$  in the subexpression  $n + 1 > n$  are bound. A variable  $x$  that's not bound is said to be *free*, and occurrences of  $x$  that are not bound are called *free* occurrences. This terminology is rather misleading. A bound variable doesn't really occur in a formula because replacing it by some new variable doesn't change the formula. The two formulas

$$\exists n \in \text{Nat} : n + 1 > n \quad \exists x \in \text{Nat} : x + 1 > x$$

are equivalent. Calling  $n$  a variable of the first formula is a bit like calling  $a$  a variable of that formula because it appears in the name  $\text{Nat}$ . Although misleading, this terminology is common and often convenient.

## 1.4 Formulas and Language

When you first studied mathematics, formulas were statements. The formula  $2 * x > x$  was just a compact way of writing the statement “2 times  $x$  is greater than  $x$ .” In this book, you are entering the realm of logic, where a formula is a noun. The formula  $2 * x > x$  is just a formula; it may be true or false, depending on the value of  $x$ . If we want to assert that this formula is true, meaning that  $2 * x$  really is greater than  $x$ , we should explicitly write “ $2 * x > x$  is true.”

Using a formula in place of a statement can lead to confusion. On the other hand, formulas are more compact and easier to read than prose. Reading  $2 * x > x$  is easier than reading “ $2 * x$  is greater than  $x$ ”, and “ $2 * x > x$  is true” may seem redundant. So, like most mathematicians, I will sometimes write a sentence like:

We know that  $x$  is positive, so  $2 * x > x$ .

If it's not obvious whether a formula is really a formula or is the statement that the formula is true, here's an easy way to tell. Replace the formula with a name and read the sentence. If the sentence is grammatically correct, even though nonsensical, then the formula is a formula; otherwise, it's a statement. The formula  $2 * x > x$  in the sentence above is a statement because

We know that  $x$  is positive, so Mary.

is ungrammatical. It is a formula in the sentence

To prove  $2 * x > x$ , we must prove that  $x$  is positive.

because the following silly sentence is grammatically correct:

To prove Fred, we must prove that  $x$  is positive.

## Chapter 2

# Specifying a Simple Clock

### 2.1 Behaviors

Before we try to specify a system, let's look at how scientists do it. For centuries, they have described a system with equations that determine how its state evolves with time, where the state consists of the values of variables. For example, the state of the system comprising the earth and the moon might be described by the values of the four variables  $e\_pos$ ,  $m\_pos$ ,  $e\_vel$ , and  $m\_vel$ , representing the positions and velocities of the two bodies. These values are elements in a 3-dimensional space. The earth-moon system is described by equations expressing the variables' values as functions of time and of certain constants—namely, their masses and initial positions and velocities.

A behavior of the earth-moon system consists of a function  $F$  from time to states,  $F(t)$  representing the state of the system at time  $t$ . A computer system differs from the systems traditionally studied by scientists because we can pretend that its state changes in discrete steps. So, we represent the execution of a system as a sequence of states. Formally, we define a *behavior* to be a sequence of states, where a state is an assignment of values to variables. We specify a system by specifying a set of possible behaviors—the ones representing a correct execution of the system.

### 2.2 An Hour Clock

Let's start with a very trivial system—a digital clock that displays only the hour. To make the system completely trivial, we ignore the relation between the display and the actual time. The hour clock is then just a device whose display cycles through the values 1 through 12. Let the variable  $hr$  represent the clock's

display. A typical behavior of the clock is the sequence

$$(2.1) \quad [hr = 11] \rightarrow [hr = 12] \rightarrow [hr = 1] \rightarrow [hr = 2] \rightarrow \dots$$

of states, where  $[hr = 11]$  is a state in which the variable  $hr$  has the value 11. A pair of successive states, such as  $[hr = 1] \rightarrow [hr = 2]$ , is called a *step*.

To specify the hour clock, we describe all its possible behaviors. We write an *initial predicate* that specifies the possible initial values of  $hr$ , and a *next-state relation* that specifies how the value of  $hr$  can change in any step.

We don't want to specify exactly what the display reads initially; any hour will do. So, we want the initial predicate to assert that  $hr$  can have any value from 1 through 12. Let's call the initial predicate  $HCini$ . We might informally define  $HCini$  by:

$$HCini \triangleq hr \in \{1, \dots, 12\}$$

Later, we'll see how to write this definition formally, without the “...” that stands for the informal *and so on*.

The next-state relation  $HCnxt$  is a formula expressing the relation between the values of  $hr$  in the old (first) state and new (second) state of a step. We let  $hr$  represent the value of  $hr$  in the old state and  $hr'$  represent its value in the new state. (The ' in  $hr'$  is read *prime*.) We want the next-state relation to assert that  $hr'$  equals  $hr + 1$  except if  $hr$  equals 12, in which case  $hr'$  should equal 1. Using an IF/THEN/ELSE construct with the obvious meaning, we can define  $HCnxt$  to be the next-state relation by writing:

$$HCnxt \triangleq hr' = \text{IF } hr \neq 12 \text{ THEN } hr + 1 \text{ ELSE } 1$$

$HCnxt$  is an ordinary mathematical formula, except that it contains primed as well as unprimed variables. Such a formula is called an *action*. An action is true or false of a step. A step that satisfies the action  $HCnxt$  is called an  $HCnxt$  *step*.

When an  $HCnxt$  step occurs, we sometimes say that  $HCnxt$  is *executed*. However, it would be a mistake to take this terminology seriously. An action is a formula, and formulas aren't executed.

We want our specification to be a single formula, not the pair of formulas  $HCini$  and  $HCnxt$ . This formula must assert about a behavior that (i) its initial state satisfies  $HCini$ , and (ii) each of its steps satisfies  $HCnxt$ . We express (i) as the formula  $HCini$ , which we interpret as a statement about behaviors to mean that the initial state satisfies  $HCini$ . To express (ii), we use the temporal-logic operator  $\Box$  (pronounced *box*). The temporal formula  $\Box F$  asserts that formula  $F$  is always true. In particular,  $\Box HCnxt$  is the assertion that  $HCnxt$  is true for every step in the behavior. So,  $HCini \wedge \Box HCnxt$  is true of a behavior iff the initial state satisfies  $HCini$  and every step satisfies  $HCnxt$ . This formula describes all behaviors like the one in (2.1) on this page; it seems to be the specification we're looking for.

If we considered the clock only in isolation, and never tried to relate it to another system, then this would be a fine specification. However, suppose the clock is part of a larger system—for example, the hour display of a weather station that displays the current hour and temperature. The state of the station is described by two variables:  $hr$ , representing the hour display, and  $tmp$ , representing the temperature display. Consider this behavior of the weather station:

$$\begin{aligned} \begin{bmatrix} hr &= 11 \\ tmp &= 23.5 \end{bmatrix} &\rightarrow \begin{bmatrix} hr &= 12 \\ tmp &= 23.5 \end{bmatrix} \rightarrow \begin{bmatrix} hr &= 12 \\ tmp &= 23.4 \end{bmatrix} \rightarrow \\ &\begin{bmatrix} hr &= 12 \\ tmp &= 23.3 \end{bmatrix} \rightarrow \begin{bmatrix} hr &= 1 \\ tmp &= 23.3 \end{bmatrix} \rightarrow \dots \end{aligned}$$

In the second and third steps,  $tmp$  changes but  $hr$  remains the same. These steps are not allowed by  $\Box HCnext$ , which asserts that every step must increment  $hr$ . The formula  $HCini \wedge \Box HCnext$  does not describe the hour clock in the weather station.

A formula that describes *any* hour clock must allow steps that leave  $hr$  unchanged—in other words,  $hr' = hr$  steps. These are called *stuttering steps* of the clock. A specification of the hour clock should allow both  $HCnext$  steps and stuttering steps. So, a step should be allowed iff it is either an  $HCnext$  step or a stuttering step—that is, iff it is a step satisfying  $HCnext \vee (hr' = hr)$ . This suggests that we adopt  $HCini \wedge \Box(HCnext \vee (hr' = hr))$  as our specification. In TLA, we let  $[HCnext]_{hr}$  stand for  $HCnext \vee (hr' = hr)$ , so we can write the formula more compactly as  $HCini \wedge \Box[HCnext]_{hr}$ .

The formula  $HCini \wedge \Box[HCnext]_{hr}$  does allow stuttering steps. In fact, it allows the behavior

$$[hr = 11] \rightarrow [hr = 12] \rightarrow [hr = 12] \rightarrow [hr = 12] \rightarrow \dots$$

that ends with an infinite sequence of stuttering steps. This behavior describes a clock whose display attains the value 12 and then keeps that value forever—in other words, a clock that stops at 12. In a like manner, we can represent a terminating execution of any system by an infinite behavior that ends with a sequence of nothing but stuttering steps. We have no need of finite behaviors (finite sequences of states), so we consider only infinite ones.

It's natural to require that a clock does not stop, so our specification should assert that there are infinitely many nonstuttering steps. Chapter 8 explains how to express this requirement. For now, we content ourselves with clocks that may stop, and we take as our specification of an hour clock the formula  $HC$  defined by

$$HC \triangleq HCini \wedge \Box[HCnext]_{hr}$$

I pronounce  $[HCnext]_{hr}$  as square  $HCnext$  sub  $hr$ .

## 2.3 A Closer Look at the Hour-Clock Specification

A state is an assignment of values to variables, but what variables? The answer is simple: all variables. In the behavior (2.1) on page 16,  $[hr = 1]$  represents some particular state that assigns the value 1 to  $hr$ . It might assign the value 23 to the variable  $tmp$  and the value  $\sqrt{-17}$  to the variable  $m\_pos$ . We can think of a state as representing a potential state of the entire universe. A state that assigns 1 to  $hr$  and a particular point in 3-space to  $m\_pos$  describes a state of the universe in which the hour clock reads 1 and the moon is in a particular place. A state that assigns  $\sqrt{-2}$  to  $hr$  doesn't correspond to any state of the universe that we recognize, because the hour-clock can't display the value  $\sqrt{-2}$ . It might represent the state of the universe after a bomb fell on the clock, making its display purely imaginary.

A behavior is an infinite sequence of states—for example:

$$(2.2) \quad [hr = 11] \rightarrow [hr = 77.2] \rightarrow [hr = 78.2] \rightarrow [hr = \sqrt{-2}] \rightarrow \dots$$

A behavior describes a potential history of the universe. The behavior (2.2) doesn't correspond to a history that we understand, because we don't know how the clock's display can change from 11 to 77.2. Whatever kind of history it represents is not one in which the clock is doing what it's supposed to.

Formula  $HC$  is a temporal formula. A temporal formula is an assertion about behaviors. We say that a behavior *satisfies*  $HC$  iff  $HC$  is a true assertion about the behavior. Behavior (2.1) satisfies formula  $HC$ . Behavior (2.2) does not, because  $HC$  asserts that every step satisfies  $HC_{next}$ , and the first and third steps of (2.2) don't. (The second step,  $[hr = 77.2] \rightarrow [hr = 78.2]$ , does satisfy  $HC_{next}$ .) We regard formula  $HC$  to be the specification of an hour clock because it is satisfied by exactly those behaviors that represent histories of the universe in which the clock functions properly.

If the clock is behaving properly, then its display should be an integer from 1 through 12. So,  $hr$  should be an integer from 1 through 12 in every state of any behavior satisfying the clock's specification,  $HC$ . Formula  $HC_{ini}$  asserts that  $hr$  is an integer from 1 through 12, and  $\Box HC_{ini}$  asserts that  $HC_{ini}$  is always true. So,  $\Box HC_{ini}$  should be true for any behavior satisfying  $HC$ . Another way of saying this is that  $HC$  implies  $\Box HC_{ini}$ , for any behavior. Thus, the formula  $HC \Rightarrow \Box HC_{ini}$  should be satisfied by *every* behavior. A temporal formula satisfied by every behavior is called a *theorem*, so  $HC \Rightarrow \Box HC_{ini}$  should be a theorem.<sup>1</sup> It's easy to see that it is:  $HC$  implies that  $HC_{ini}$  is true initially (in the first state of the behavior), and  $\Box[HC_{next}]_{hr}$  implies that each step either advances  $hr$  to its proper next value or else leaves  $hr$  unchanged. We can

<sup>1</sup>Logicians call a formula *valid* if it is satisfied by every behavior; they reserve the term *theorem* for provably valid formulas.

formalize this reasoning using the proof rules of TLA, but I'm not going to delve into proofs and proof rules.

## 2.4 The Hour-Clock Specification in $TLA^+$

Figure 2.1 on the next page shows how the hour clock specification can be written in  $TLA^+$ . There are two versions: the ASCII version on the bottom is the actual  $TLA^+$  specification, the way you type it; the top version is typeset the way a “pretty-printer” might display it. Before trying to understand the specification, observe the relation between the two syntaxes:

- Reserved words that appear in small upper-case letters (like `EXTENDS`) are written in ASCII with ordinary upper-case letters.
- When possible, symbols are represented pictorially in ASCII—for example,  $\square$  is typed as `[]` and  $\neq$  as `#`. (You can also type  $\neq$  as `/=`.)
- When there is no good ASCII representation,  $\TeX$  notation [1] is used—for example,  $\in$  is typed as `\in`.

A complete list of symbols and their ASCII equivalents appears in Figure 13.16 on page 253. I will usually show the typeset version of a specification; the ASCII versions of all specifications appear in the Appendix.

Now let's look at what the specification says. It starts with

```

┌────────────────── MODULE HourClock ───────────────────┐

```

which begins a module named *HourClock*.  $TLA^+$  specifications are partitioned into modules; the hour clock's specification consists of this single module.

Arithmetic operators like  $+$  are not built into  $TLA^+$ , but are themselves defined in modules. (You might want to write a specification in which  $+$  means addition of matrices rather than numbers.) The usual operators on natural numbers are defined in the *Naturals* module. Their definitions are incorporated into module *HourClock* by the statement

```
EXTENDS Naturals
```

Every symbol that appears in a formula must either be a built-in operator of  $TLA^+$ , or else it must be declared or defined. The statement

```
VARIABLE hr
```

declares *hr* to be a variable.

To define *HCini*, we need to express the set  $\{1, \dots, 12\}$  formally, without the ellipsis “...”. We can write this set out completely as

```
{1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12}
```

```

----- MODULE HourClock -----
EXTENDS Naturals
VARIABLE hr
HCini  $\triangleq$  hr  $\in$  (1 .. 12)
HCnxt  $\triangleq$  hr' = IF hr  $\neq$  12 THEN hr + 1 ELSE 1
HC  $\triangleq$  HCini  $\wedge$   $\square$ [HCnxt]hr
-----
THEOREM HC  $\Rightarrow$   $\square$ HCini
-----

----- MODULE HourClock -----
EXTENDS Naturals
VARIABLE hr
HCini == hr \in (1 .. 12)
HCnxt == hr' = IF hr # 12 THEN hr + 1 ELSE 1
HC == HCini /\ [] [HCnxt]hr
-----
THEOREM HC => []HCini
=====

```

Figure 2.1: The hour clock specification—typeset and ASCII versions.

but that’s tiresome. Instead, we use the operator “..”, defined in the *Naturals* module, to write this set as 1..12. In general  $i..j$  is the set of integers from  $i$  through  $j$ , for any integers  $i$  and  $j$ . (It equals the empty set if  $j < i$ .) It’s now obvious how to write the definition of *HCini*. The definitions of *HCnxt* and *HC* are written just as before. (The ordinary mathematical operators of logic and set theory, like  $\wedge$  and  $\in$ , are built into  $\text{TLA}^+$ .)

The line

---

can appear anywhere between statements; it’s purely cosmetic and has no meaning. Following it is the statement

THEOREM *HC*  $\Rightarrow$   $\square$ *HCini*

of the theorem that was discussed above. This statement asserts that the formula  $HC \Rightarrow \square HCini$  is true in the context of the statement. More precisely, it asserts that the formula follows logically from the definitions in this module, the definitions in the *Naturals* module, and the rules of  $\text{TLA}^+$ . If the formula were not true, then the module would be incorrect.

The module is terminated by the symbol

The specification of the hour clock is the definition of  $HC$ , including the definitions of the formulas  $HCnext$  and  $HCini$  and of the operators  $..$  and  $+$  that appear in the definition of  $HC$ . Formally, nothing in the module tells us that  $HC$  rather than  $HCini$  is the clock's specification.  $TLA^+$  is a language for writing mathematics—in particular, for writing mathematical definitions and theorems. What those definitions represent, and what significance we attach to those theorems, lies outside the scope of mathematics and therefore outside the scope of  $TLA^+$ . Engineering requires not just the ability to use mathematics, but the ability to understand what, if anything, the mathematics tells us about an actual system.

## 2.5 Another Way to Specify the Hour Clock

The *Naturals* module also defines the modulus operator, which we write  $\%$ . The formula  $i \% n$ , which mathematicians write  $i \bmod n$ , is the remainder when  $i$  is divided by  $n$ . More formally,  $i \% n$  is the natural number less than  $n$  satisfying  $i = q * n + (i \% n)$  for some natural number  $q$ . Let's express this condition mathematically. The *Naturals* module defines  $Nat$  to be the set of natural numbers, and the assertion that there exists a  $q$  in the set  $Nat$  satisfying a formula  $F$  is written  $\exists q \in Nat : F$ . Thus, if  $i$  and  $n$  are elements of  $Nat$  and  $n > 0$ , then  $i \% n$  is the unique number satisfying

$$(i \% n \in 0 .. (n - 1)) \wedge (\exists q \in Nat : i = q * n + (i \% n))$$

We can use  $\%$  to simplify our hour-clock specification a bit. Observing that  $(11 \% 12) + 1$  equals 12 and  $(12 \% 12) + 1$  equals 1, we can define a different next-state action  $HCnext2$  and a different formula  $HC2$  to be the clock specification:

$$HCnext2 \triangleq hr' = (hr \% 12) + 1 \quad HC2 \triangleq HCini \wedge \Box[HCnext2]_{hr}$$

Actions  $HCnext$  and  $HCnext2$  are not equivalent. The step  $[hr = 24] \rightarrow [hr = 25]$  satisfies  $HCnext$  but not  $HCnext2$ , while the step  $[hr = 24] \rightarrow [hr = 1]$  satisfies  $HCnext2$  but not  $HCnext$ . However, any step starting in a state with  $hr$  in  $1 .. 12$  satisfies  $HCnext$  iff it satisfies  $HCnext2$ . It's therefore not hard to deduce that any behavior starting in a state satisfying  $HCini$  satisfies  $\Box[HCnext]_{hr}$  iff it satisfies  $\Box[HCnext2]_{hr}$ . Hence, formulas  $HC$  and  $HC2$  are equivalent. It doesn't matter which of them we take to be the specification of an hour clock.

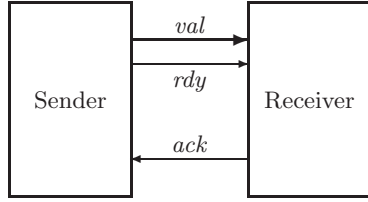
Mathematics provides infinitely many ways of expressing the same thing. The expressions  $6 + 6$ ,  $3 * 4$ , and  $141 - 129$  all have the same meaning; they are just different ways of writing the number 12. We could replace either instance of the number 12 in module *HourClock* by any of these expressions without changing the meaning of any of the module's formulas.

When writing a specification, you will often be faced with a choice of how to express something. When that happens, you should first make sure that the choices yield equivalent specifications. If they do, then you can choose the one that you feel makes the specification easiest to understand. If they don't, then you must decide which one you mean.

## Chapter 3

# An Asynchronous Interface

We now specify an interface for transmitting data between asynchronous devices. A *sender* and a *receiver* are connected as shown here:



Data is sent on *val*, and the *rdy* and *ack* lines are used for synchronization. The sender must wait for an acknowledgment (an *Ack*) for one data item before it can send the next. The interface uses the standard two-phase handshake protocol, described by the following sample behavior.

$$\begin{array}{c}
 \begin{bmatrix} val = 26 \\ rdy = 0 \\ ack = 0 \end{bmatrix} \xrightarrow{\text{Send } 37} \begin{bmatrix} val = 37 \\ rdy = 1 \\ ack = 0 \end{bmatrix} \xrightarrow{\text{Ack}} \begin{bmatrix} val = 37 \\ rdy = 1 \\ ack = 1 \end{bmatrix} \xrightarrow{\text{Send } 4} \\
 \begin{bmatrix} val = 4 \\ rdy = 0 \\ ack = 1 \end{bmatrix} \xrightarrow{\text{Ack}} \begin{bmatrix} val = 4 \\ rdy = 0 \\ ack = 0 \end{bmatrix} \xrightarrow{\text{Send } 19} \begin{bmatrix} val = 19 \\ rdy = 1 \\ ack = 0 \end{bmatrix} \xrightarrow{\text{Ack}} \dots
 \end{array}$$

(It doesn't matter what value *val* has in the initial state.)

It's easy to see from this sample behavior what the set of all possible behaviors should be—once we decide what the data values are that can be sent. But, before writing the TLA<sup>+</sup> specification that describes these behaviors, let's look at what I've just done.

In writing this behavior, I made the decision that *val* and *rdy* should change in a single step. The values of the variables *val* and *rdy* represent voltages

on some set of wires in the physical device. Voltages on different wires don't change at precisely the same instant. I decided to ignore this aspect of the physical system and pretend that the values of *val* and *rdy* represented by those voltages change instantaneously. This simplifies the specification, but at the price of ignoring what may be an important detail of the system. In an actual implementation of the protocol, the voltage on the *rdy* line shouldn't change until the voltages on the *val* lines have stabilized; but you won't learn that from my specification. Had I wanted the specification to convey this requirement, I would have written a behavior in which the value of *val* and the value of *rdy* change in separate steps.

A specification is an abstraction. It describes some aspects of the system and ignores others. We want the specification to be as simple as possible, so we want to ignore as many details as we can. But, whenever we omit some aspect of the system from the specification, we admit a potential source of error. With my specification, we can verify the correctness of a system that uses this interface, and the system could still fail because the implementor didn't know that the *val* line should stabilize before the *rdy* line is changed.

The hardest part of writing a specification is choosing the proper abstraction. I can teach you about  $\text{TLA}^+$ , so expressing an abstract view of a system as a  $\text{TLA}^+$  specification becomes a straightforward task. But I don't know how to teach you about abstraction. A good engineer knows how to abstract the essence of a system and suppress the unimportant details when specifying and designing it. The art of abstraction is learned only through experience.

When writing a specification, you must first choose the abstraction. In a  $\text{TLA}^+$  specification, this means choosing (i) the variables that represent the system's state and (ii) the granularity of the steps that change those variables' values. Should the *rdy* and *ack* lines be represented as separate variables or as a single variable? Should *val* and *rdy* change in one step, two steps, or an arbitrary number of steps? To help make these choices, I recommend that you start by writing the first few steps of one or two sample behaviors, just as I did at the beginning of this section. Chapter 7 has more to say about these choices.

## 3.1 The First Specification

Now let's specify the interface with a module *AsynchInterface*. The variables *rdy* and *ack* can assume the values 0 and 1, which are natural numbers, so our module `EXTENDS` the *Naturals* module. We next decide what the possible values of *val* should be—that is, what data values may be sent. We could write a specification that places no restriction on the data values. The specification could allow the sender to first send 37, then send  $\sqrt{-15}$ , and then send *Nat* (the entire set of natural numbers). However, any real device can send only a restricted set of values. We could pick some specific set—for example, 32-bit

numbers. However, the protocol is the same regardless of whether it's used to send 32-bit numbers or 128-bit numbers. So, we compromise between the two extremes of allowing anything to be sent and allowing only 32-bit numbers to be sent by assuming only that there is some set *Data* of data values that may be sent. The constant *Data* is a parameter of the specification. It's declared by the statement

CONSTANT *Data*

Our three variables are declared by

VARIABLES *val*, *rdy*, *ack*

The keywords VARIABLE and VARIABLES are synonymous, as are CONSTANT and CONSTANTS.

The variable *rdy* can assume any value—for example,  $-1/2$ . That is, there exist states that assign the value  $-1/2$  to *rdy*. When discussing the specification, we usually say that *rdy* can assume only the values 0 and 1. What we really mean is that the value of *rdy* equals 0 or 1 in every state of any behavior satisfying the specification. But a reader of the specification shouldn't have to understand the complete specification to figure this out. We can make the specification easier to understand by telling the reader what values the variables can assume in a behavior that satisfies the specification. We could do this with comments, but I prefer to use a definition like this one:

$$\text{TypeInvariant} \triangleq (val \in Data) \wedge (rdy \in \{0, 1\}) \wedge (ack \in \{0, 1\})$$

I call the set  $\{0, 1\}$  the *type* of *rdy*, and I call *TypeInvariant* a *type invariant*. Let's define *type* and some other terms more precisely:

- A *state function* is an ordinary expression (one with no prime or  $\Box$ ) that can contain variables and constants.
- A *state predicate* is a Boolean-valued state function.
- An *invariant* *Inv* of a specification *Spec* is a state predicate such that  $Spec \Rightarrow \Box Inv$  is a theorem.
- A variable *v* has *type* *T* in a specification *Spec* iff  $v \in T$  is an invariant of *Spec*.

We can make the definition of *TypeInvariant* easier to read by writing it as follows.

$$\begin{aligned} \text{TypeInvariant} \triangleq & \wedge val \in Data \\ & \wedge rdy \in \{0, 1\} \\ & \wedge ack \in \{0, 1\} \end{aligned}$$

Each conjunct begins with a  $\wedge$  and must lie completely to the right of that  $\wedge$ . (The conjunct may occupy multiple lines). We use a similar notation for disjunctions. When using this bulleted-list notation, the  $\wedge$ 's or  $\vee$ 's must line up precisely (even in the ASCII input). Because the indentation is significant, we can eliminate parentheses, making this notation especially useful when conjunctions and disjunctions are nested.

The formula *TypeInvariant* will not appear as part of the specification. We do not assume that *TypeInvariant* is an invariant; the specification should imply that it is. In fact, its invariance will be asserted as a theorem.

The initial predicate is straightforward. Initially, *val* can equal any element of *Data*. We can start with *rdy* and *ack* either both 0 or both 1.

$$\begin{aligned} Init &\triangleq \wedge val \in Data \\ &\wedge rdy \in \{0, 1\} \\ &\wedge ack = rdy \end{aligned}$$

Now for the next-state action *Next*. A step of the protocol either sends a value or receives a value. We define separately the two actions *Send* and *Rcv* that describe the sending and receiving of a value. A *Next* step (one satisfying action *Next*) is either a *Send* step or a *Rcv* step, so it is a  $Send \vee Rcv$  step. Therefore, *Next* is defined to equal  $Send \vee Rcv$ . Let's now define *Send* and *Rcv*.

We say that action *Send* is *enabled* in a state from which it is possible to take a *Send* step. From the sample behavior above, we see that *Send* is enabled iff *rdy* equals *ack*. Usually, the first question we ask about an action is, when is it enabled? So, the definition of an action usually begins with its enabling condition. The first conjunct in the definition of *Send* is therefore  $rdy = ack$ . The next conjuncts tell us what the new values of the variables *val*, *rdy*, and *ack* are. The new value *val'* of *val* can be any element of *Data*—that is, any value satisfying  $val' \in Data$ . The value of *rdy* changes from 0 to 1 or from 1 to 0, so  $rdy'$  equals  $1 - rdy$  (because  $1 = 1 - 0$  and  $0 = 1 - 1$ ). The value of *ack* is left unchanged.

TLA<sup>+</sup> defines UNCHANGED *v* to mean that the expression *v* has the same value in the old and new states. More precisely, UNCHANGED *v* equals  $v' = v$ , where *v'* is the expression obtained from *v* by priming all variables. So, we define *Send* by:

$$\begin{aligned} Send &\triangleq \wedge rdy = ack \\ &\wedge val' \in Data \\ &\wedge rdy' = 1 - rdy \\ &\wedge \text{UNCHANGED } ack \end{aligned}$$

(I could have written  $ack' = ack$  instead of UNCHANGED *ack*, but I prefer to use the UNCHANGED construct in specifications.)

A *Rcv* step is enabled iff *rdy* is different from *ack*; it complements the value of *ack* and leaves *val* and *rdy* unchanged. Both *val* and *rdy* are left unchanged iff

|                                                  |                                                                                                                      |
|--------------------------------------------------|----------------------------------------------------------------------------------------------------------------------|
| MODULE <i>AsynchInterface</i>                    |                                                                                                                      |
| EXTENDS                                          | <i>Naturals</i>                                                                                                      |
| CONSTANT                                         | <i>Data</i>                                                                                                          |
| VARIABLES                                        | <i>val, rdy, ack</i>                                                                                                 |
| <i>TypeInvariant</i>                             | $\triangleq \wedge val \in Data$<br>$\wedge rdy \in \{0, 1\}$<br>$\wedge ack \in \{0, 1\}$                           |
| <i>Init</i>                                      | $\triangleq \wedge val \in Data$<br>$\wedge rdy \in \{0, 1\}$<br>$\wedge ack = rdy$                                  |
| <i>Send</i>                                      | $\triangleq \wedge rdy = ack$<br>$\wedge val' \in Data$<br>$\wedge rdy' = 1 - rdy$<br>$\wedge \text{UNCHANGED } ack$ |
| <i>Rcv</i>                                       | $\triangleq \wedge rdy \neq ack$<br>$\wedge ack' = 1 - ack$<br>$\wedge \text{UNCHANGED } \langle val, rdy \rangle$   |
| <i>Next</i>                                      | $\triangleq Send \vee Rcv$                                                                                           |
| <i>Spec</i>                                      | $\triangleq Init \wedge \square [Next]_{\langle val, rdy, ack \rangle}$                                              |
| THEOREM $Spec \Rightarrow \square TypeInvariant$ |                                                                                                                      |

Figure 3.1: The First Specification of an Asynchronous Interface

the pair of values *val*, *rdy* is left unchanged. TLA<sup>+</sup> uses angle brackets  $\langle$  and  $\rangle$  to enclose ordered tuples, so *Rcv* asserts that  $\langle val, rdy \rangle$  is left unchanged. (Angle brackets are typed in ASCII as << and >>.) The definition of *Rcv* is therefore:

$$\begin{aligned}
 Rcv \triangleq & \wedge rdy \neq ack \\
 & \wedge ack' = 1 - ack \\
 & \wedge \text{UNCHANGED } \langle val, rdy \rangle
 \end{aligned}$$

As in our clock example, the complete specification *Spec* should allow stuttering steps—in this case, ones that leave all three variables unchanged. So, *Spec* allows steps that leave  $\langle val, rdy, ack \rangle$  unchanged. Its definition is

$$Spec \triangleq Init \wedge \square [Next]_{\langle val, rdy, ack \rangle}$$

Module *AsynchInterface* also asserts the invariance of *TypeInvariant*. It appears in full in Figure 3.1 on this page.

## 3.2 Another Specification

Module *AsynchInterface* is a fine description of the interface and its handshake protocol. However, it's not easy to use it to help specify a system that uses the interface. Let's rewrite the interface specification in a form that makes it more convenient to use as part of a larger specification.

The first problem with the original specification is that it uses three variables to describe a single interface. A system might use several different instances of the interface. To avoid a proliferation of variables, we replace the three variables *val*, *rdy*, *ack* with a single variable *chan* (short for *channel*). A mathematician would do this by letting the value of *chan* be an ordered triple—for example, a state  $[chan = \langle -1/2, 0, 1 \rangle]$  might replace the state with  $val = -1/2$ ,  $rdy = 0$ , and  $ack = 1$ . But programmers have learned that using tuples like this leads to mistakes; it's easy to forget if the *ack* line is represented by the second or third component. TLA<sup>+</sup> therefore provides records in addition to more conventional mathematical notation.

Let's represent the state of the channel as a record with *val*, *rdy*, and *ack* fields. If *r* is such a record, then *r.val* is its *val* field. The type invariant asserts that the value of *chan* is an element of the set of all such records *r* in which *r.val* is an element of the set *Data* and *r.rdy* and *r.ack* are elements of the set  $\{0, 1\}$ . This set of records is written:

$$[val : Data, rdy : \{0, 1\}, ack : \{0, 1\}]$$

The components of a record are not ordered, so it doesn't matter in what order we write them. This same set of records can also be written as:

$$[ack : \{0, 1\}, val : Data, rdy : \{0, 1\}]$$

Initially, *chan* can equal any element of this set whose *ack* and *rdy* fields are equal, so the initial predicate is the conjunction of the type invariant and the condition  $chan.ack = chan.rdy$ .

A system that uses the interface may perform an operation that sends some data value *d* and performs some other changes that depend on the value *d*. We'd like to represent such an operation as an action that is the conjunction of two separate actions: one that describes the sending of *d* and the other that describes the other changes. Thus, instead of defining an action *Send* that sends some unspecified data value, we define the action *Send(d)* that sends data value *d*. The next-state action is satisfied by a *Send(d)* step, for some *d* in *Data*, or a *Rcv* step. (The value received by a *Rcv* step equals *chan.val*.) Saying that a step is a *Send(d)* step for some *d* in *Data* means that there exists a *d* in *Data* such that the step satisfies *Send(d)*—in other words, that the step is an  $\exists d \in Data : Send(d)$  step. So we define

$$Next \triangleq (\exists d \in Data : Send(d)) \vee Rcv$$

The  $Send(d)$  action asserts that  $chan'$  equals the record  $r$  such that:

$$r.val = d \quad r.rdy = 1 - chan.rdy \quad r.ack = chan.ack$$

This record is written in TLA<sup>+</sup> as:

$$[val \mapsto d, \quad rdy \mapsto 1 - chan.rdy, \quad ack \mapsto chan.ack]$$

(The symbol  $\mapsto$  is typed in ASCII as  $| \rightarrow$ .) The fields of records are not ordered, so this record can just as well be written:

$$[ack \mapsto chan.ack, \quad val \mapsto d, \quad rdy \mapsto 1 - chan.rdy]$$

The enabling condition of  $Send(d)$  is that the  $rdy$  and  $ack$  lines are equal, so we can define:

$$\begin{aligned} Send(d) &\triangleq \\ &\wedge chan.rdy = chan.ack \\ &\wedge chan' = [val \mapsto d, \quad rdy \mapsto 1 - chan.rdy, \quad ack \mapsto chan.ack] \end{aligned}$$

This is a perfectly good definition of  $Send(d)$ . However, I prefer a slightly different one. We can describe the value of  $chan'$  by saying that it is the same as the value of  $chan$  except that its  $val$  component equals  $d$  and its  $rdy$  component equals  $1 - chan.rdy$ . In TLA<sup>+</sup>, we can write this value as

$$[chan \text{ EXCEPT } !.val = d, \quad !.rdy = 1 - chan.rdy]$$

Think of the  $!$  as standing for the new record that the EXCEPT expression forms by modifying  $chan$ . So, the expression can be read as the record  $!$  that is the same as  $chan$  except  $!.val$  equals  $d$  and  $!.rdy$  equals  $1 - chan.rdy$ . In the expression that  $!.rdy$  equals, the symbol  $@$  stands for  $chan.rdy$ , so we can write this EXCEPT expression as:

$$[chan \text{ EXCEPT } !.val = d, \quad !.rdy = 1 - @]$$

In general, for any record  $r$ , the expression

$$[r \text{ EXCEPT } !.c_1 = e_1, \dots, !.c_n = e_n]$$

is the record obtained from  $r$  by replacing  $r.c_i$  with  $e_i$ , for each  $i$  in  $1 \dots n$ . An  $@$  in the expression  $e_i$  stands for  $r.c_i$ . Using this notation, we define:

$$\begin{aligned} Send(d) &\triangleq \\ &\wedge chan.rdy = chan.ack \\ &\wedge chan' = [chan \text{ EXCEPT } !.val = d, \quad !.rdy = 1 - @] \end{aligned}$$

The definition of  $Rcv$  is straightforward. A value can be received when  $chan.rdy \neq chan.ack$ , and receiving the value complements  $chan.ack$ :

$$\begin{aligned} Rcv &\triangleq \\ &\wedge chan.rdy \neq chan.ack \\ &\wedge chan' = [chan \text{ EXCEPT } !.ack = 1 - @] \end{aligned}$$

|                                                                |                                                                                                                                                       |
|----------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------|
| MODULE <i>Channel</i>                                          |                                                                                                                                                       |
| EXTENDS                                                        | <i>Naturals</i>                                                                                                                                       |
| CONSTANT                                                       | <i>Data</i>                                                                                                                                           |
| VARIABLE                                                       | <i>chan</i>                                                                                                                                           |
| <i>TypeInvariant</i>                                           | $\triangleq \text{chan} \in [\text{val} : \text{Data}, \text{rdy} : \{0, 1\}, \text{ack} : \{0, 1\}]$                                                 |
| <i>Init</i>                                                    | $\triangleq \wedge \text{TypeInvariant}$<br>$\wedge \text{chan.ack} = \text{chan.rdy}$                                                                |
| <i>Send</i> ( <i>d</i> )                                       | $\triangleq \wedge \text{chan.rdy} = \text{chan.ack}$<br>$\wedge \text{chan}' = [\text{chan} \text{ EXCEPT } !.\text{val} = d, !.\text{rdy} = 1 - @]$ |
| <i>Rcv</i>                                                     | $\triangleq \wedge \text{chan.rdy} \neq \text{chan.ack}$<br>$\wedge \text{chan}' = [\text{chan} \text{ EXCEPT } !.\text{ack} = 1 - @]$                |
| <i>Next</i>                                                    | $\triangleq (\exists d \in \text{Data} : \text{Send}(d)) \vee \text{Rcv}$                                                                             |
| <i>Spec</i>                                                    | $\triangleq \text{Init} \wedge \square[\text{Next}]_{\text{chan}}$                                                                                    |
| THEOREM $\text{Spec} \Rightarrow \square \text{TypeInvariant}$ |                                                                                                                                                       |

Figure 3.2: Our second specification of an asynchronous interface.

The complete specification appears in Figure 3.2 on this page.

We have now written two different specifications of the asynchronous interface. They are two different mathematical representations of the same physical system. In module *AsynchInterface*, we represented the system with the three variables *val*, *rdy*, and *ack*. In module *Channel*, we used a single variable *chan*. Since these two representations are at the same level of abstraction, they should, in some sense, be equivalent. Section 5.8 explains one sense in which they're equivalent.

### 3.3 Types: A Reminder

As defined in Section 3.1, a variable *v* has type *T* in specification *Spec* iff *v*  $\in T$  is an invariant of *Spec*. Thus, *hr* has type 1 .. 12 in the specification *HC* of the hour clock. This assertion does *not* mean that the variable *hr* can assume only values in the set 1 .. 12. A state is an arbitrary assignment of values to variables, so there exist states in which the value of *hr* is  $\sqrt{-2}$ . The assertion does mean that, in every behavior satisfying formula *HC*, the value of *hr* is an element of 1 .. 12.

If you are used to types in programming languages, it may seem strange that  $\text{TLA}^+$  allows a variable to assume any value. Why not restrict our states to ones in which variables have the values of the right type? In other words, why

not add a formal type system to TLA<sup>+</sup>? A complete answer would take us too far afield. The question is addressed further in Section 6.2. For now, remember that TLA<sup>+</sup> is an untyped language. Type correctness is just a name for a certain invariance property. Assigning the name *TypeInvariant* to a formula gives it no special status.

## 3.4 Definitions

Let's examine what a definition means. If *Id* is a simple identifier like *Init* or *Spec*, then the definition  $Id \triangleq exp$  defines *Id* to be synonymous with the expression *exp*. Replacing *Id* by *exp*, or vice-versa, in any expression *e* does not change the meaning of *e*. This replacement must be done after the expression is parsed, not in the “raw input”. For example, the definition  $x \triangleq a + b$  makes  $x * c$  equal to  $(a + b) * c$ , not to  $a + b * c$ , which equals  $a + (b * c)$ .

The definition of *Send* has the form  $Id(p) \triangleq exp$ , where *Id* and *p* are identifiers. For any expression *e*, this defines *Id(e)* to be the expression obtained by substituting *e* for *p* in *exp*. For example, the definition of *Send* in the *Channel* module defines *Send*(−5) to equal

$$\begin{aligned} &\wedge \text{chan.rdy} = \text{chan.ack} \\ &\wedge \text{chan}' = [\text{chan} \text{ EXCEPT } !.val = -5, !.rdy = 1 - @] \end{aligned}$$

*Send(e)* is an expression, for any expression *e*. Thus, we can write the formula  $\text{Send}(-5) \wedge (\text{chan.ack} = 1)$ . The identifier *Send* by itself is not an expression, and  $\text{Send} \wedge (\text{chan.ack} = 1)$  is not a grammatically well-formed string. It's non-syntactic nonsense, like  $a + * b +$ .

We say that *Send* is an *operator* that takes a single argument. We define operators that take more than one argument in the obvious way, the general form being:

$$(3.1) \quad Id(p_1, \dots, p_n) \triangleq exp$$

where the  $p_i$  are distinct identifiers and *exp* is an expression. We can consider defined identifiers like *Init* and *Spec* to be operators that take no argument, but we generally use *operator* to mean an operator that takes one or more arguments.

I will use the term *symbol* to mean an identifier like *Send* or an operator symbol like  $+$ . Every symbol that is used in a specification must either be a built-in operator of TLA<sup>+</sup> (like  $\in$ ) or it must be declared or defined. Every symbol declaration or definition has a *scope* within which the symbol may be used. The scope of a VARIABLE or CONSTANT declaration, and of a definition, is the part of the module that follows it. Thus, we can use *Init* in any expression that follows its definition in module *Channel*. The statement `EXTENDS Naturals` extends the scope of symbols like  $+$  defined in the *Naturals* module to the *Channel* module.

The operator definition (3.1) implicitly includes a declaration of the identifiers  $p_1, \dots, p_n$  whose scope is the expression  $exp$ . An expression of the form

$$\exists v \in S : exp$$

has a declaration of  $v$  whose scope is the expression  $exp$ . Thus the identifier  $v$  has a meaning within the expression  $exp$  (but not within the expression  $S$ ).

A symbol cannot be declared or defined if it already has a meaning. The expression

$$(\exists v \in S : exp1) \wedge (\exists v \in T : exp2)$$

is all right, because neither declaration of  $v$  lies within the scope of the other. Similarly, the two declarations of the symbol  $d$  in the *Channel* module (in the definition of *Send* and in the expression  $\exists d$  in the definition of *Next*) have disjoint scopes. However, the expression

$$(\exists v \in S : (exp1 \wedge \exists v \in T : exp2))$$

is illegal because the declaration of  $v$  in the second  $\exists v$  lies inside the scope of its declaration in the first  $\exists v$ . Although conventional mathematics and programming languages allow such redeclarations, TLA<sup>+</sup> forbids them because they can lead to confusion and errors.

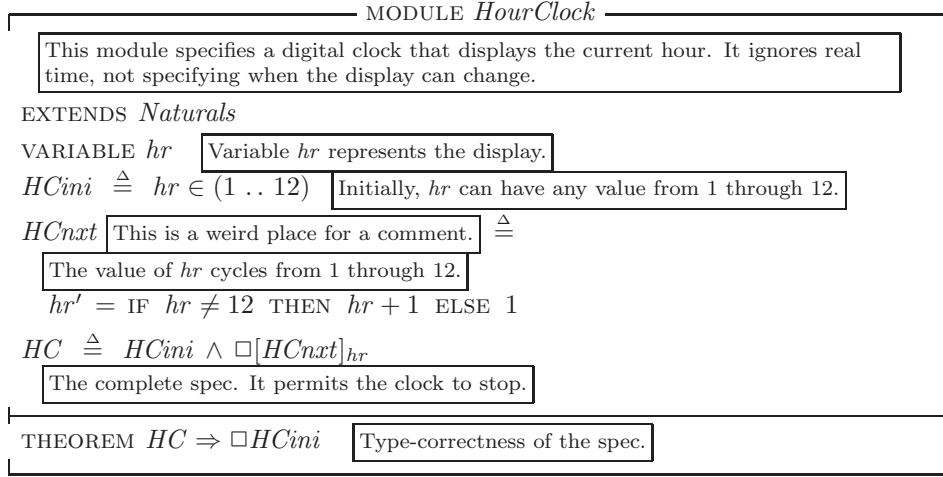
## 3.5 Comments

Even simple specifications like the ones in modules *AsynchInterface* and *Channel* can be hard to understand from the mathematics alone. That's why I began with an intuitive explanation of the interface. That explanation made it easier for you to understand formula *Spec* in the module, which is the actual specification. Every specification should be accompanied by an informal prose explanation. The explanation may be in an accompanying document, or it may be included as comments in the specification.

Figure 3.3 on the next page shows how the hour clock's specification in module *HourClock* might be explained by comments. In the typeset version, comments are distinguished from the specification itself by the use of a different font. As shown in the figure, TLA<sup>+</sup> provides two way of writing comments in the ASCII version. A comment may appear anywhere enclosed between (*\** and *\**). The text of the comment itself may not contain *\**, so these comments can't be nested. An end-of-line comment is preceded by *\\**.

A comment almost always appears on a line by itself or at the end of a line. I put a comment between *HCnext* and  $\triangleq$  just to show that it can be done.

To save space, I will write few comments in the example specifications. But specifications should have lots of comments. Even if there is an accompanying document describing the system, comments are needed to help the reader understand how the specification formalizes that description.



```

----- MODULE HourClock -----
(*****
(* This module specifies a digital clock that displays *)
(* the current hour. It ignores real time, not          *)
(* specifying when the display can change.              *)
(*****)
EXTENDS Naturals
VARIABLE hr      \* Variable hr represents the display.
HCini == hr \in (1 .. 12)  \* Initially, hr can have any
                          \* value from 1 through 12.
HCnxt (* This is a weird place for a comment. *) ==
(*****
(* The value of hr cycles from 1 through 12.          *)
(*****)
hr' = IF hr # 12 THEN hr + 1 ELSE 1
HC == HCini /\ [] [HCnxt]_hr
(* The complete spec. It permits the clock to stop. *)

-----
THEOREM HC => []HCini  \* Type-correctness of the spec.
=====

```

Figure 3.3: The hour clock specification with comments.

Comments can help solve a problem posed by the logical structure of a specification. A symbol has to be declared or defined before it can be used. In module *Channel*, the definition of *Spec* has to follow the definition of *Next*, which has to follow the definitions of *Send* and *Rcv*. But it's usually easiest to understand a top-down description of a system. We would probably first want to read the declarations of *Data* and *chan*, then the definition of *Spec*, then the definitions of *Init* and *Next*, and then the definitions of *Send* and *Rcv*. In other words, we want to read the specification more or less from bottom to top. This is easy enough to do for a module as short as *Channel*; it's inconvenient for longer specifications. We can use comments to guide the reader through a longer specification. For example, we could precede the definition of *Send* in the *Channel* module with the comment:

Actions *Send* and *Rcv* below are the disjuncts of the next-state action *Next*.

The module structure also allows us to choose the order in which a specification is read. For example, we can rewrite the hour-clock specification by splitting the *HourClock* module into three separate modules:

*HCVar*      A module that declares the variable *hr*.

*HCActions*   A module that EXTENDS modules *Naturals* and *HCVar* and defines *HCini* and *HCnext*.

*HCSpec*      A module that EXTENDS module *HCActions*, defines formula *HC*, and asserts the type-correctness theorem.

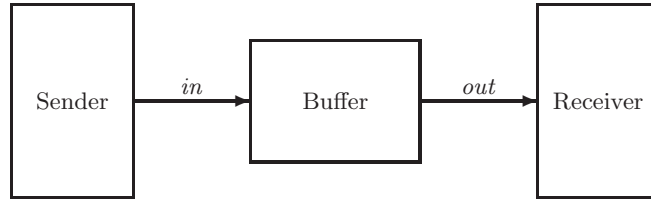
The EXTENDS relation implies a logical ordering of the modules: *HCVar* precedes *HCActions*, which precedes *HCSpec*. But the modules don't have to be read in that order. The reader can be told to read *HCVar* first, then *HCSpec*, and finally *HCActions*. The INSTANCE construct introduced below in Chapter 4 provides another tool for modularizing specifications.

Splitting a tiny specification like *HourClock* in this way would be ludicrous. But the proper splitting of modules can help make a large specification easier to read. When writing a specification, you should decide in what order it should be read. You can then design the module structure to permit reading it in that order, when each individual module is read from beginning to end. Finally, you should ensure that the comments within each module make sense when the different modules are read in the appropriate order.

## Chapter 4

# A FIFO

Our next example is a FIFO buffer, a device with which a sender process transmits a sequence of values to a receiver. The sender and receiver use two channels, *in* and *out*, to communicate with the buffer:



Values are sent over *in* and *out* using the asynchronous protocol specified by the *Channel* module of Figure 3.2 on page 30. The system's specification will allow behaviors with four kinds of nonstuttering steps: *Send* and *Rcv* actions on both the *in* channel and the *out* channel.

### 4.1 The Inner Specification

The specification of the FIFO first EXTENDS modules *Naturals* and *Sequences*. The *Sequences* module defines operations on finite sequences. We represent a finite sequence as a tuple, so the sequence of three numbers 3, 2, 1 is the triple  $\langle 3, 2, 1 \rangle$ . The sequences module defines the following operators on sequences.

*Seq*(*S*) The set of all sequences of elements of the set *S*. For example,  $\langle 3, 7 \rangle$  is an element of *Seq*(*Nat*).

*Head*(*s*) The first element of sequence *s*. For example, *Head*( $\langle 3, 7 \rangle$ ) equals 3.

- $Tail(s)$  The tail of sequence  $s$  (all but the head of  $s$ ). For example,  $Tail(\langle 3, 7 \rangle)$  equals  $\langle 7 \rangle$ .
- $Append(s, e)$  The sequence obtained by appending element  $e$  to the tail of sequence  $s$ . For example,  $Append(\langle 3, 7 \rangle, 3)$  equals  $\langle 3, 7, 3 \rangle$ .
- $s \circ t$  The sequence obtained by concatenating the sequences  $s$  and  $t$ . For example,  $\langle 3, 7 \rangle \circ \langle 3 \rangle$  equals  $\langle 3, 7, 3 \rangle$ . (We type  $\circ$  in ASCII as  $\backslash o$ .)
- $Len(s)$  The length of sequence  $s$ . For example,  $Len(\langle 3, 7 \rangle)$  equals 2.

The FIFO's specification continues by declaring the constant *Message*, which represents the set of all messages that can be sent.<sup>1</sup> It then declares the variables. There are three variables: *in* and *out*, representing the channels, and a third variable *q* that represents the queue of buffered messages. The value of *q* is the sequence of messages that have been sent by the sender but not yet received by the receiver. (Section 4.3 has more to say about this additional variable *q*.)

We want to use the definitions in the *Channel* module to specify operations on the channels *in* and *out*. This requires two instances of that module—one in which the variable *chan* of the *Channel* module is replaced with the variable *in* of our current module, and the other in which *chan* is replaced with *out*. In both instances, the constant *Data* of the *Channel* module is replaced with *Message*. We obtain the first of these instances with the statement:

$$InChan \triangleq \text{INSTANCE } Channel \text{ WITH } Data \leftarrow Message, \text{ chan} \leftarrow in$$

For every symbol  $\sigma$  defined in module *Channel*, this defines *InChan!* $\sigma$  to have the same meaning in the current module as  $\sigma$  had in module *Channel*, except with *Message* substituted for *Data* and *in* substituted for *chan*. For example, this statement defines *InChan!TypeInvariant* to equal

$$in \in [val : Message, rdy : \{0, 1\}, ack : \{0, 1\}]$$

(The statement does *not* define *InChan!Data* because *Data* is declared, not defined, in module *Channel*.) We introduce our second instance of the *Channel* module with the analogous statement:

$$OutChan \triangleq \text{INSTANCE } Channel \text{ WITH } Data \leftarrow Message, \\ \text{chan} \leftarrow out$$

The initial states of the *in* and *out* channels are specified by *InChan!Init* and *OutChan!Init*. Initially, no messages have been sent or received, so *q* should

<sup>1</sup>I like to use a singular noun like *Message* rather than a plural like *Messages* for the name of a set. That way, the  $\in$  in the expression  $m \in Message$  can be read *is a*. This is the same convention that most programmers use for naming types.

The *Channel* module appears in Figure 3.2 on page 30.

equal the empty sequence. The empty sequence is the zero-tuple (there's only one, and it's written  $\langle \rangle$ ), so we define the initial predicate to be:

$$\begin{aligned} Init &\triangleq \wedge InChan!Init \\ &\quad \wedge OutChan!Init \\ &\quad \wedge q = \langle \rangle \end{aligned}$$

We next define the type invariant. The type invariants for *in* and *out* come from the *Channel* module, and the type of *q* is the set of finite sequences of messages. The type invariant for the FIFO specification is therefore:

$$\begin{aligned} TypeInvariant &\triangleq \wedge InChan!TypeInvariant \\ &\quad \wedge OutChan!TypeInvariant \\ &\quad \wedge q \in Seq(Message) \end{aligned}$$

The four kinds of nonstuttering steps allowed by the next-state action are described by four actions:

- SSend(msg)* The sender sends message *msg* on the *in* channel.
- BufRcv* The buffer receives the message from the *in* channel and appends it to the tail of *q*.
- BufSend* The buffer removes the message from the head of *q* and sends it on channel *out*.
- RRcv* The receiver receives the message from the *out* channel.

The definitions of these actions, along with the rest of the specification, are in module *InnerFIFO* of Figure 4.1 on the next page. The reason for the adjective *Inner* is explained in Section 4.3 below.

## 4.2 Instantiation Examined

### 4.2.1 Instantiation is Substitution

Consider the definition of *Next* in module *Channel* (page 30). We can remove every defined symbol that appears in that definition by using the symbol's definition. For example, we can eliminate the expression *Send(d)* by expanding the definition of *Send*. We can repeat this process. For example the  $-$  that appears in the expression  $1 - @$  (obtained by expanding the definition of *Send*) can be eliminated by using the definition of  $-$  from the *Naturals* module. Continuing in this way, we eventually obtain a definition for *Next* in terms of only the built-in operators of  $TLA^+$  and the parameters *Data* and *chan* of the *Channel* module.

|                                                                  |                                                                                                                                                                                                                                                                                                                                                        |
|------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| MODULE <i>InnerFIFO</i>                                          |                                                                                                                                                                                                                                                                                                                                                        |
| EXTENDS <i>Naturals, Sequences</i>                               |                                                                                                                                                                                                                                                                                                                                                        |
| CONSTANT <i>Message</i>                                          |                                                                                                                                                                                                                                                                                                                                                        |
| VARIABLES <i>in, out, q</i>                                      |                                                                                                                                                                                                                                                                                                                                                        |
| <i>InChan</i>                                                    | $\triangleq$ INSTANCE <i>Channel</i> WITH <i>Data</i> $\leftarrow$ <i>Message</i> , <i>chan</i> $\leftarrow$ <i>in</i>                                                                                                                                                                                                                                 |
| <i>OutChan</i>                                                   | $\triangleq$ INSTANCE <i>Channel</i> WITH <i>Data</i> $\leftarrow$ <i>Message</i> , <i>chan</i> $\leftarrow$ <i>out</i>                                                                                                                                                                                                                                |
| <i>Init</i>                                                      | $\triangleq$ $\wedge$ <i>InChan!Init</i><br>$\wedge$ <i>OutChan!Init</i><br>$\wedge$ <i>q</i> = $\langle \rangle$                                                                                                                                                                                                                                      |
| <i>TypeInvariant</i>                                             | $\triangleq$ $\wedge$ <i>InChan!TypeInvariant</i><br>$\wedge$ <i>OutChan!TypeInvariant</i><br>$\wedge$ <i>q</i> $\in$ <i>Seq(Message)</i>                                                                                                                                                                                                              |
| <i>SSend(msg)</i>                                                | $\triangleq$ $\wedge$ <i>InChan!Send(msg)</i> <span style="border: 1px solid black; padding: 2px;">Send <i>msg</i> on channel <i>in</i>.</span><br>$\wedge$ UNCHANGED $\langle out, q \rangle$                                                                                                                                                         |
| <i>BufRcv</i>                                                    | $\triangleq$ $\wedge$ <i>InChan!Rcv</i> <span style="border: 1px solid black; padding: 2px;">Receive message from channel <i>in</i><br/>and append it to tail of <i>q</i>.</span><br>$\wedge$ <i>q'</i> = <i>Append(q, in.val)</i><br>$\wedge$ UNCHANGED <i>out</i>                                                                                    |
| <i>BufSend</i>                                                   | $\triangleq$ $\wedge$ <i>q</i> $\neq \langle \rangle$ <span style="border: 1px solid black; padding: 2px;">Enabled only if <i>q</i> is nonempty.<br/>Send <i>Head(q)</i> on channel <i>out</i><br/>and remove it from <i>q</i>.</span><br>$\wedge$ <i>OutChan!Send(Head(q))</i><br>$\wedge$ <i>q'</i> = <i>Tail(q)</i><br>$\wedge$ UNCHANGED <i>in</i> |
| <i>RRcv</i>                                                      | $\triangleq$ $\wedge$ <i>OutChan!Rcv</i> <span style="border: 1px solid black; padding: 2px;">Receive message from channel <i>out</i>.</span><br>$\wedge$ UNCHANGED $\langle in, q \rangle$                                                                                                                                                            |
| <i>Next</i>                                                      | $\triangleq$ $\vee \exists msg \in Message : SSend(msg)$<br>$\vee$ <i>BufRcv</i><br>$\vee$ <i>BufSend</i><br>$\vee$ <i>RRcv</i>                                                                                                                                                                                                                        |
| <i>Spec</i>                                                      | $\triangleq$ <i>Init</i> $\wedge$ $\square[Next]_{\langle in, out, q \rangle}$                                                                                                                                                                                                                                                                         |
| THEOREM <i>Spec</i> $\Rightarrow$ $\square$ <i>TypeInvariant</i> |                                                                                                                                                                                                                                                                                                                                                        |

Figure 4.1: The specification of a FIFO, with the internal variable *q* visible.

We consider this to be the “real” definition of *Next* in module *Channel*. The statement

$$InChan \triangleq \text{INSTANCE } Channel \text{ WITH } Data \leftarrow Message, \text{ chan} \leftarrow in$$

in module *InnerFIFO* defines *InChan!Next* to be the formula obtained from this real definition of *Next* by substituting *Message* for *Data* and *in* for *chan*. This defines *InChan!Next* in terms of only the built-in operators of  $TLA^+$  and the parameters *Message* and *in* of module *InnerFIFO*.

Let’s now consider an arbitrary INSTANCE statement

$$IM \triangleq \text{INSTANCE } M \text{ WITH } p_1 \leftarrow e_1, \dots, p_n \leftarrow e_n$$

Let  $\Sigma$  be a symbol defined in module *M* and let *d* be its “real” definition. The INSTANCE statement defines *IM!Σ* to have as its real definition the expression obtained from *d* by replacing all instances of  $p_i$  by the expression  $e_i$ , for each *i*. The definition of *IM!Σ* must contain only the parameters (declared constants and variables) of the current module, not the ones of module *M*. Hence, the  $p_i$  must consist of all the parameters of module *M*. The  $e_i$  must be expressions that are meaningful in the current module.

### 4.2.2 Parametrized Instantiation

The FIFO specification uses two instances of module *Channel*—one with *in* substituted for *chan* and the other with *out* substituted for *chan*. We could instead use a single parametrized instance by putting the following statement in module *InnerFIFO*:

$$Chan(ch) \triangleq \text{INSTANCE } Channel \text{ WITH } Data \leftarrow Message, \text{ chan} \leftarrow ch$$

For any symbol  $\Sigma$  defined in module *Channel* and any expression *exp*, this defines *Chan(exp)!Σ* to equal formula  $\Sigma$  with *Message* substituted for *Data* and *exp* substituted for *chan*. The *Rcv* action on channel *in* could then be written *Chan(in)!Rcv*, and the *Send(msg)* action on channel *out* could be written *Chan(out)!Send(msg)*.

The instantiation above defines *Chan!Send* to be an operator with two arguments. Writing *Chan(out)!Send(msg)* instead of *Chan!Send(out, msg)* is just an idiosyncrasy of the syntax. It is no stranger than the syntax for infix operators, which makes us write  $a + b$  instead of  $+(a, b)$ .

### 4.2.3 Implicit Substitutions

The use of *Message* as the name for the set of transmitted values in the FIFO specification is a bit strange, since we had just used the name *Data* for the

analogous set in the asynchronous channel specifications. Suppose we had used *Data* in place of *Message* as the constant parameter of module *InnerFIFO*. The first instantiation statement would then have been

$$InChan \triangleq \text{INSTANCE } Channel \text{ WITH } Data \leftarrow Data, chan \leftarrow in$$

The substitution  $Data \leftarrow Data$  indicates that the constant parameter *Data* of the instantiated module *Channel* is replaced with the expression *Data* of the current module. TLA<sup>+</sup> allows us to drop any substitution of the form  $\Sigma \leftarrow \Sigma$ , for a symbol  $\Sigma$ . So, the statement above can be written as

$$InChan \triangleq \text{INSTANCE } Channel \text{ WITH } chan \leftarrow in$$

We know there is an implied  $Data \leftarrow Data$  substitution because an INSTANCE statement must have a substitution for every parameter of the instantiated module. If some parameter  $p$  has no explicit substitution, then there is an implicit substitution  $p \leftarrow p$ . This means that the INSTANCE statement must lie within the scope of a declaration or definition of the symbol  $p$ .

It is quite common to instantiate a module with this kind of implicit substitution. Often, every parameter has an implicit substitution, in which case the list of explicit substitutions is empty. The WITH is then omitted.

#### 4.2.4 Instantiation Without Renaming

So far, all the instantiations we've used have been with renaming. For example, the first instantiation of module *Channel* renames the defined symbol *Send* as *InChan!Send*. This kind of renaming is necessary if we are using multiple instances of the module, or a single parametrized instance. The two instances *InChan!Init* and *OutChan!Init* of *Init* in module *InnerFIFO* are different formulas, so they need different names.

Sometimes we need only a single instance of a module. For example, suppose we are specifying a system with only a single asynchronous channel. We then need only one instance of *Channel*, so we don't have to rename the instantiated symbols. In that case, we can write something like

$$\text{INSTANCE } Channel \text{ WITH } Data \leftarrow D, chan \leftarrow x$$

This instantiates *Channel* with no renaming, but with substitution. Thus, it defines *Rcv* to be the formula of the same name from the *Channel* module, except with  $D$  substituted for *Data* and  $x$  substituted for *chan*. The expressions substituted for an instantiated module's parameters must be defined. So, this INSTANCE statement must be within the scope of the definitions or declarations of  $D$  and  $x$ .

### 4.3 Hiding the Queue

Module *InnerFIFO* of Figure 4.1 defines *Spec* to be  $Init \wedge \Box[Next] \dots$ , the sort of formula we’ve become accustomed to as a system specification. However, formula *Spec* describes the value of variable *q*, as well as of the variables *in* and *out*. The picture of the FIFO system I drew on page 35 shows only channels *in* and *out*; it doesn’t show anything inside the boxes. A specification of the FIFO should describe only the values sent and received on the channels. The variable *q*, which represents what’s going on inside the box labeled *Buffer*, is used to specify what values are sent and received. In the final specification, it should be hidden.

In TLA, we hide a variable with the existential quantifier  $\exists$  of temporal logic. The formula  $\exists x : F$  is true of a behavior iff there exists some sequence of values—one in each state of the behavior—that can be assigned to the variable *x* that will make formula *F* true. (The meaning of  $\exists$  is defined more precisely in Section 8.7.)

The obvious way to write a FIFO specification in which *q* is hidden is with the formula  $\exists q : Spec$ . However, we can’t put this definition in module *InnerFIFO* because *q* is already declared there, and a formula  $\exists q : \dots$  would redeclare it. Instead, we use a new module with a parametrized instantiation of the *InnerFIFO* module (see Section 4.2.2 above):

|                                                                                                                                                            |
|------------------------------------------------------------------------------------------------------------------------------------------------------------|
| MODULE <i>FIFO</i>                                                                                                                                         |
| CONSTANT <i>Message</i><br>VARIABLES <i>in, out</i><br><br>$Inner(q) \triangleq \text{INSTANCE } InnerFIFO$<br>$Spec \triangleq \exists q : Inner(q)!Spec$ |

Observe that the INSTANCE statement is an abbreviation for

$$Inner(q) \triangleq \text{INSTANCE } InnerFIFO \\ \text{WITH } q \leftarrow q, in \leftarrow in, out \leftarrow out, Message \leftarrow Message$$

The variable parameter *q* of module *InnerFIFO* is instantiated with the parameter *q* of the definition of *Inner*. The other parameters of the *InnerFIFO* module are instantiated with the parameters of module *FIFO*.

### 4.4 A Bounded FIFO

We have specified an unbounded FIFO—a buffer that can hold an unbounded number of messages. Any real system has a finite amount of resources, so it can

contain only a bounded number of in-transit messages. In many situations, we wish to abstract away the bound on resources and describe a system in terms of unbounded FIFOs. In other situations, we may care about that bound. We then want to strengthen our specification by placing a bound  $N$  on the number of outstanding messages.

A specification of a bounded FIFO differs from our specification of the unbounded FIFO only in that action *BufRcv* should be enabled only when there are fewer than  $N$  messages in the buffer—that is, only when  $Len(q)$  is less than  $N$ . It would be easy to write a complete new specification of a bounded FIFO by copying module *InnerFIFO* and just adding the conjunct  $Len(q) < N$  to the definition of *BufRcv*. But let's use module *InnerFIFO* as it is, rather than copying it.

The next-state action *BNext* for the bounded FIFO is the same as the FIFO's next-state action *Next* except that it allows a *BufRcv* step only if  $Len(q)$  is less than  $N$ . In other words, *BNext* should allow a step only if (i) it's a *Next* step and (ii) if it's a *BufRcv* step, then  $Len(q) < N$  is true in the first state. In other words, *BNext* should equal

$$Next \wedge (BufRcv \Rightarrow (Len(q) < N))$$

Module *BoundedFIFO* in Figure 4.2 on the next page contains the specification. It introduces the new constant parameter  $N$ . It also contains the statement

$$\text{ASSUME } (N \in Nat) \wedge (N > 0)$$

which asserts that, in this module, we are assuming that  $N$  is a positive natural number. Such an assumption has no effect on any definitions made in the module. However, it may be taken as a hypothesis when proving any theorems asserted in the module. In other words, a module asserts that its assumptions imply its theorems. It's a good idea to assert this kind of simple assumption about constants.

An ASSUME statement should only be used to assert assumptions about constants. The formula being assumed should not contain any variables. It might be tempting to assert type declarations as assumptions—for example, to add to module *InnerFIFO* the assumption  $q \in Seq(Message)$ . However, that would be wrong because it asserts that, in any state,  $q$  is a sequence of messages. As we observed in Section 3.3, a state is a completely arbitrary assignment of values to variables, so there are states in which  $q$  has the value  $\sqrt{-17}$ . Assuming that such a state doesn't exist would lead to a logical contradiction.

You may wonder why module *BoundedFIFO* asserts that  $N$  is a positive natural, but doesn't assert that *Message* is a set. Similarly, why didn't we have to specify that the constant parameter *Data* in our asynchronous interface specifications is a set? The answer is that, in  $TLA^+$ , every value is a set.<sup>2</sup> A

<sup>2</sup> $TLA^+$  is based on the mathematical formalism known as Zermelo-Fr ankel set theory, also called ZF.

---

MODULE *BoundedFIFO*

---

EXTENDS *Naturals*  
 VARIABLES *in, out*  
 CONSTANT *Message, N*  
 ASSUME  $(N \in \text{Nat}) \wedge (N > 0)$   
 $Inner(q) \triangleq \text{INSTANCE } InnerFIFO$   
 $BNext(q) \triangleq \wedge Inner(q)!Next$   
 $\quad \wedge Inner(q)!BufRcv \Rightarrow (Len(q) < N)$   
 $Spec \triangleq \exists q : Inner(q)!Init \wedge \Box[BNext(q)]_{\langle in, out, q \rangle}$

---

Figure 4.2: Specification of a FIFO buffer of length  $N$ .

value like the number 3, which we don't think of as a set, is formally a set. We just don't know what its elements are. The formula  $2 \in 3$  is a perfectly reasonable one, but  $TLA^+$  does not specify whether it's true or false. So, we don't have to assert that *Message* is a set because we know that it is one.

Although *Message* is automatically a set, it isn't necessarily a finite set. For example, *Message* could be instantiated with the set *Nat* of natural numbers. If you want to assume that a constant parameter is a finite set, then you need to state this as an assumption. (You can do this with the *IsFiniteSet* operator from the *FiniteSets* module, described in Section 6.1.) However, most specifications make perfect sense for infinite sets of messages or processors, so there is no reason to require these sets to be finite.

## 4.5 What We're Specifying

I wrote above, at the beginning of this section, that we were going to specify a FIFO buffer. Formula *Spec* of the *FIFO* module actually specifies a set of behaviors, each representing a sequence of sending and receiving operations on the channels *in* and *out*. The sending operations on *in* are performed by the sender, and the receiving operations on *out* are performed by the receiver. The sender and receiver are not part of the FIFO buffer; they form its *environment*.

Our specification describes a system consisting of the FIFO buffer and its environment. The behaviors satisfying formula *Spec* of module *FIFO* represent those histories of the universe in which both the system and its environment behave correctly. It's often helpful in understanding a specification to indicate explicitly which steps are system steps and which are environment steps. We can do this by defining the next-state action to be

$$Next \triangleq SysNext \vee EnvNext$$

where *SysNext* describes system steps and *EnvNext* describes environment steps. For the FIFO, we have

$$\begin{aligned} \text{SysNext} &\triangleq \text{BufRcv} \vee \text{BufSend} \\ \text{EnvNext} &\triangleq (\exists \text{msg} \in \text{Message} : \text{SSend}(\text{msg})) \vee \text{RRcv} \end{aligned}$$

While suggestive, this way of defining the next-state action has no formal significance. The specification *Spec* equals  $\text{Init} \wedge \Box[\text{Next}] \dots$ ; changing the way we structure the definition of *Next* doesn't change its meaning. If a behavior fails to satisfy *Spec*, nothing tells us if the system or its environment is to blame.

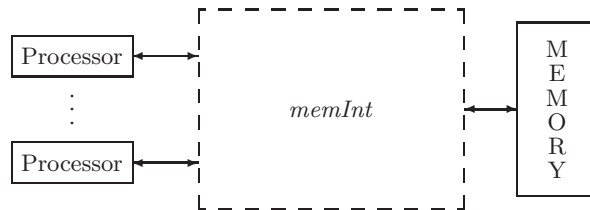
A formula like *Spec*, which describes the correct behavior of both the system and its environment, is called a *closed-system* or *complete-system* specification. An *open-system* specification is one that describes only the correct behavior of the system. A behavior satisfies an open-system specification if it represents a history in which either the system operates correctly, or it failed to operate correctly only because its environment did something wrong. Section 10.7 explains how to write open-system specifications.

Open-system specifications are philosophically more satisfying. However, closed-system specifications are a little bit easier to write, and the mathematics underlying them is simpler. So, we almost always write closed-system specifications. It's usually quite easy to turn a closed-system specification into an open-system specification. But in practice, there's little reason to do so.

## Chapter 5

# A Caching Memory

A memory system consists of a set of processors connected to a memory by some abstract interface, which we label *memInt*.



In this section we specify what the memory is supposed to do, then we specify a particular implementation of the memory using caches. We begin by specifying the memory interface, which is common to both specifications.

### 5.1 The Memory Interface

The asynchronous interface described in Chapter 3 uses a handshake protocol. Receipt of a data value must be acknowledged before the next data value can be sent. In the memory interface, we abstract away this kind of detail and represent both the sending of a data value and its receipt as a single step. We call it a *Send* step if a processor is sending the value to the memory; it's a *Reply* step if the memory is sending to a processor. Processors do not send values to one another, and the memory sends to only one processor at a time.

We represent the state of the memory interface by the value of the variable *memInt*. A *Send* step changes *memInt* in some way, but we don't want to specify exactly how. The way to leave something unspecified in a specification is to make it a parameter. For example, in the bounded FIFO of Section 4.4, we left the size of the buffer unspecified by making it a parameter *N*. We'd

therefore like to declare a parameter  $Send$  so that  $Send(p, d)$  describes how  $memInt$  is changed by a step that represents processor  $p$  sending data value  $d$  to the memory. However,  $TLA^+$  provides only `CONSTANT` and `VARIABLE` parameters, not action parameters.<sup>1</sup> So, we declare  $Send$  to be a constant operator and write  $Send(p, d, memInt, memInt')$  instead of  $Send(p, d)$ .

In  $TLA^+$ , we declare  $Send$  to be a constant operator that takes four arguments by writing

$$\text{CONSTANT } Send(-, -, -, -)$$

This means that  $Send(p, d, miOld, miNew)$  is an expression, for any expressions  $p$ ,  $d$ ,  $miOld$ , and  $miNew$ , but it says nothing about what the value of that expression is. We want it to be a Boolean value that is true iff a step in which  $memInt$  equals  $miOld$  in the first state and  $miNew$  in the second state represents the sending by  $p$  of value  $d$  to the memory.<sup>2</sup> We can assert that the value is a Boolean by the assumption:

$$\begin{aligned} \text{ASSUME } \forall p, d, miOld, miNew : \\ Send(p, d, miOld, miNew) \in \text{BOOLEAN} \end{aligned}$$

This asserts that the formula

$$Send(p, d, miOld, miNew) \in \text{BOOLEAN}$$

is true for all values of  $p$ ,  $d$ ,  $miOld$ , and  $miNew$ . The built-in symbol `BOOLEAN` denotes the set  $\{\text{TRUE}, \text{FALSE}\}$ , whose elements are the two Boolean values `TRUE` and `FALSE`.

This `ASSUME` statement asserts formally that the value of

$$Send(p, d, miOld, miNew)$$

is a Boolean. But the only way to assert formally what that value signifies would be to say what it actually equals—that is, to define  $Send$  rather than making it a parameter. We don't want to do that, so we just state informally what the value means. This statement is part of the intrinsically informal description of the relation between our mathematical abstraction and a physical memory system.

To allow the reader to understand the specification, we have to describe informally what  $Send$  means. The `ASSUME` statement asserting that  $Send(\dots)$  is a Boolean is then superfluous as an explanation. However, it could help tools understand the specification, so it's a good idea to include it anyway.

<sup>1</sup>Even if  $TLA^+$  allowed us to declare an action parameter, we would have no way to specify that a  $Send(p, d)$  action constrains only  $memInt$  and not other variables.

<sup>2</sup>We expect  $Send(p, d, miOld, miNew)$  to have this meaning only when  $p$  is a processor and  $d$  a value that  $p$  is allowed to send, but we simplify the specification a bit by requiring it to be a Boolean for all values of  $p$  and  $d$ .

A specification that uses the memory interface can use the operators *Send* and *Reply* to specify how the variable *memInt* changes. The specification must also describe *memInt*'s initial value. We therefore declare a constant parameter *InitMemInt* that is the set of possible initial values of *memInt*.

We also introduce three constant parameters that are needed to describe the interface:

*Proc* The set of processor identifiers. (We usually shorten *processor identifier* to *processor* when referring to an element of *Proc*.)

*Adr* The set of memory addresses.

*Val* The set of possible memory values that can be assigned to an address.

Finally, we define the values that the processors and memory send to one another over the interface. A processor sends a request to the memory. We represent a request as a record with an *op* field that specifies the type of request and additional fields that specify its arguments. Our simple memory allows just read and write requests. A read request has *op* field “Rd” and an *adr* field specifying the address to be read. The set of all read requests is therefore the set

$$[op : \{\text{“Rd”}\}, adr : Adr]$$

of all records whose *op* field equals “Rd” (is an element of the set {“Rd”} whose only element is the string “Rd”) and whose *adr* field is an element of *Adr*. A write request must specify the address to be written and the value to write. It is represented by a record with *op* field equal to “Wr”, and with *adr* and *val* fields specifying the address and value. We define *MReq*, the set of all requests, to equal the union of these two sets. (Set operations, including union, are described in Section 1.2.)

The memory responds to a read request with the memory value it read. We will also have it respond to a write request; and it seems nice to let the response be different from the response to any read request. We therefore require the memory to respond to a write request by returning a value *NoVal* that is different from any memory value. We could declare *NoVal* to be a constant parameter and add the assumption  $NoVal \notin Val$ . (The symbol  $\notin$  is typed in ASCII as \notin.) But it's best, when possible, to avoid introducing parameters. Instead, we define *NoVal* by:

$$NoVal \triangleq \text{CHOOSE } v : v \notin Val$$

The expression  $\text{CHOOSE } x : F$  equals an arbitrarily chosen value *x* that satisfies the formula *F*. (If no such *x* exists, the expression has a completely arbitrary value.) This statement defines *NoVal* to be some value that is not an element of

| MODULE <i>MemoryInterface</i>                                                                                                                     |                                                                                                            |
|---------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------|
| VARIABLE <i>memInt</i>                                                                                                                            |                                                                                                            |
| CONSTANTS <i>Send</i> ( $-, -, -, -$ ),                                                                                                           | A <i>Send</i> ( $p, d, memInt, memInt'$ ) step represents processor $p$ sending value $d$ to the memory.   |
| <i>Reply</i> ( $-, -, -, -$ ),                                                                                                                    | A <i>Reply</i> ( $p, d, memInt, memInt'$ ) step represents the memory sending value $d$ to processor $p$ . |
| <i>InitMemInt</i> ,                                                                                                                               | The set of possible initial values of <i>memInt</i> .                                                      |
| <i>Proc</i> ,                                                                                                                                     | The set of processor identifiers.                                                                          |
| <i>Adr</i> ,                                                                                                                                      | The set of memory addresses.                                                                               |
| <i>Val</i>                                                                                                                                        | The set of memory values.                                                                                  |
| ASSUME $\forall p, d, miOld, miNew : \wedge Send(p, d, miOld, miNew) \in \text{BOOLEAN}$<br>$\wedge Reply(p, d, miOld, miNew) \in \text{BOOLEAN}$ |                                                                                                            |
| $MReq \triangleq [op : \{\text{"Rd"}\}, adr : Adr] \cup [op : \{\text{"Wr"}\}, adr : Adr, val : Val]$                                             |                                                                                                            |
| The set of all requests; a read specifies an address, a write specifies an address and a value.                                                   |                                                                                                            |
| $NoVal \triangleq \text{CHOOSE } v : v \notin Val$                                                                                                | An arbitrary value not in <i>Val</i> .                                                                     |

Figure 5.1: The Specification of a Memory Interface

*Val*. We have no idea what the value of *NoVal* is; we just know what it isn't—namely, that it isn't an element of *Val*. The CHOOSE operator is discussed in Section 6.6.

The complete memory interface specification is module *MemoryInterface* in Figure 5.1 on this page.

## 5.2 Functions

A memory assigns values to addresses. The state of the memory is therefore an assignment of elements of *Val* (memory values) to elements of *Adr* (memory addresses). In a programming language, such an assignment is called an array of type *Val* indexed by *Adr*. In mathematics, it's called a function from *Adr* to *Val*. Before writing the memory specification, let's look at the mathematics of functions, and how it is described in TLA<sup>+</sup>.

A function  $f$  has a domain, written  $\text{DOMAIN } f$ , and it assigns to each element  $x$  of its domain the value  $f[x]$ . (Mathematicians write this as  $f(x)$ , but TLA<sup>+</sup> uses the array notation of programming languages, with square brackets.) Two functions  $f$  and  $g$  are equal iff they have the same domain and  $f[x] = g[x]$  for all  $x$  in their domain.

The *range* of a function  $f$  is the set of all values of the form  $f[x]$  with  $x$  in  $\text{DOMAIN } f$ . For any sets  $S$  and  $T$ , the set of all functions whose domain equals

$S$  and whose range is any subset of  $T$  is written  $[S \rightarrow T]$ .

Ordinary mathematics does not have a convenient notation for writing an expression whose value is a function. TLA<sup>+</sup> defines  $[x \in S \mapsto e]$  to be the function  $f$  with domain  $S$  such that  $f[x] = e$  for every  $x \in S$ .<sup>3</sup> For example,

$$succ \triangleq [n \in Nat \mapsto n + 1]$$

defines *succ* to be the successor function on the natural numbers—the function with domain *Nat* such that  $succ[n] = n + 1$  for all  $n \in Nat$ .

A record is a function whose domain is a finite set of strings. For example, a record with *val*, *ack*, and *rdy* fields is a function whose domain is the set  $\{\text{"val"}, \text{"ack"}, \text{"rdy"}\}$  consisting of the three strings “val”, “ack”, and “rdy”. The expression  $r.ack$ , the *ack* field of a record  $r$ , is an abbreviation for  $r[\text{"ack"}]$ . The record

$$[val \mapsto 42, ack \mapsto 1, rdy \mapsto 0]$$

can be written

$$[i \in \{\text{"val"}, \text{"ack"}, \text{"rdy"}\} \mapsto \\ \text{IF } i = \text{"val"} \text{ THEN } 42 \text{ ELSE IF } i = \text{"ack"} \text{ THEN } 1 \text{ ELSE } 0]$$

The EXCEPT construct for records, explained in Section 3.2, is a special case of a general EXCEPT construct for functions, where  $!c$  is an abbreviation for  $![\text{"c"}]$ . For any function  $f$ , the expression  $[f \text{ EXCEPT } ![c] = e]$  is the function  $\hat{f}$  that is the same as  $f$  except with  $\hat{f}[c] = e$ . This function can also be written:

$$[x \in \text{DOMAIN } f \mapsto \text{IF } x = c \text{ THEN } e \text{ ELSE } f[x]]$$

assuming that the symbol  $x$  does not occur in any of the expressions  $f$ ,  $c$ , and  $e$ . For example,  $[succ \text{ EXCEPT } ![42] = 86]$  is the function  $g$  that is the same as *succ* except that  $g[42]$  equals 86 instead of 43.

As in the EXCEPT construct for records, the expression  $e$  in

$$[f \text{ EXCEPT } ![c] = e]$$

can contain the symbol @, where it means  $f[c]$ . For example,

$$[succ \text{ EXCEPT } ![42] = 2 * @] = [succ \text{ EXCEPT } ![42] = 2 * succ[42]]$$

In general,

$$[f \text{ EXCEPT } ![c_1] = e_1, \dots, ![c_n] = e_n]$$

---

<sup>3</sup>The  $\in$  in  $[x \in S \mapsto e]$  is just part of the syntax; TLA<sup>+</sup> uses that particular symbol to help you remember what the construct means. Computer scientists write  $\lambda x : S. e$  to represent something similar to  $[x \in S \mapsto e]$ , except that their  $\lambda$  expressions aren't quite the same as the functions of ordinary mathematics that are used in TLA<sup>+</sup>.

is the function  $\hat{f}$  that is the same as  $f$  except with  $\hat{f}[c_i] = e_i$  for each  $i$ . More precisely, this expression equals

$$[\dots [f \text{ EXCEPT } ![c_1] = e_1] \text{ EXCEPT } ![c_2] = e_2] \dots \text{ EXCEPT } ![c_n] = e_n]$$

Functions correspond to the arrays of programming languages. The domain of a function corresponds to the index set of an array. The function  $[f \text{ EXCEPT } ![c] = e]$  corresponds to the array obtained from  $f$  by assigning  $e$  to  $f[c]$ . A function whose range is a set of functions corresponds to an array of arrays. TLA<sup>+</sup> defines  $[f \text{ EXCEPT } ![c][d] = e]$  to be the function corresponding to the array obtained by assigning  $e$  to  $f[c][d]$ . It can be written as

$$[f \text{ EXCEPT } ![c] = [@ \text{ EXCEPT } ![d] = e]]$$

The generalization to  $[f \text{ EXCEPT } ![c_1] \dots [c_n] = e]$  for any  $n$  should be obvious. Since a record is a function, this notation can be used for records as well. TLA<sup>+</sup> uniformly maintains the notation that  $\sigma.c$  is an abbreviation for  $\sigma["c"]$ . For example, this implies:

$$[f \text{ EXCEPT } ![c].d = e] = [f \text{ EXCEPT } ![c] = [@ \text{ EXCEPT } !.d = e]]$$

The TLA<sup>+</sup> definition of records as functions makes it possible to manipulate them in ways that have no counterparts in programming languages. For example, we can define an operator  $R$  such that  $R(r, s)$  is the record obtained from  $r$  by replacing the value of each field  $c$  that is also a field of the record  $s$  with  $s.c$ . In other words, for every field  $c$  of  $r$ , if  $c$  is a field of  $s$  then  $R(r, s).c = s.c$ ; otherwise  $R(r, s).c = r.c$ . The definition is:

$$R(r, s) \triangleq [c \in \text{DOMAIN } r \mapsto \text{IF } c \in \text{DOMAIN } s \text{ THEN } s[c] \text{ ELSE } r[c]]$$

So far, I have described only functions of a single argument. TLA<sup>+</sup> also allows functions of multiple arguments. Section 15.1.7 on page 285 describes the general versions of the TLA<sup>+</sup> function constructs for functions with any number of arguments. However, functions of a single argument are all you really need. You can always replace a function of multiple arguments with a function-valued function—for example, writing  $f[a][b]$  instead of  $f[a, b]$ .

### 5.3 A Linearizable Memory Specification

We specify a very simple memory system in which a processor  $p$  issues a memory request and then waits for a response before issuing the next request. In our specification, the request is executed by accessing (reading or modifying) a variable  $mem$ , which represents the current state of the memory. Because the memory can receive requests from other processors before responding to processor  $p$ , it matters when  $mem$  is accessed. We let the access of  $mem$  occur

any time between the request and the response. This specifies what is called a *linearizable* memory. Less restrictive, more practical memory specifications are described in Section 11.2.

In addition to *mem*, the specification has the internal variables *ctl* and *buf*, where *ctl*[*p*] describes the status of processor *p*'s request and *buf*[*p*] contains either the request or the response. Consider the request *req* that equals

$$[op \mapsto \text{"Wr"}, \text{adr} \mapsto a, \text{val} \mapsto v]$$

It is a request to write *v* to memory address *a*, and it generates the response *NoVal*. The processing of this request is represented by the following three steps:

$$\begin{array}{c} \left[ \begin{array}{l} \text{ctl}[p] = \text{"rdy"} \\ \text{buf}[p] = \dots \\ \text{mem}[a] = \dots \end{array} \right] \xrightarrow{\text{Req}(p)} \left[ \begin{array}{l} \text{ctl}[p] = \text{"busy"} \\ \text{buf}[p] = \text{req} \\ \text{mem}[a] = \dots \end{array} \right] \\ \xrightarrow{\text{Do}(p)} \left[ \begin{array}{l} \text{ctl}[p] = \text{"done"} \\ \text{buf}[p] = \text{NoVal} \\ \text{mem}[a] = v \end{array} \right] \xrightarrow{\text{Rsp}(p)} \left[ \begin{array}{l} \text{ctl}[p] = \text{"rdy"} \\ \text{buf}[p] = \text{NoVal} \\ \text{mem}[a] = v \end{array} \right] \end{array}$$

A *Req*(*p*) step represents the issuing of a request by processor *p*. It is enabled when *ctl*[*p*] = "rdy"; it sets *ctl*[*p*] to "busy" and sets *buf*[*p*] to the request. A *Do*(*p*) step represents the memory access; it is enabled when *ctl*[*p*] = "busy" and it sets *ctl*[*p*] to "done" and *buf*[*p*] to the response. A *Rsp*(*p*) step represents the memory's response to *p*; it is enabled when *ctl*[*p*] = "done" and it sets *ctl*[*p*] to "rdy".

Writing the specification is a straightforward exercise in representing these changes to the variables in TLA<sup>+</sup> notation. The internal specification, with *mem*, *ctl*, and *buf* visible (free variables), appears in module *InternalMemory* on the following two pages. The memory specification, which hides the three internal variables, is module *Memory* in Figure 5.4 on page 53.

## 5.4 Tuples as Functions

Before writing our caching memory specification, let's take a closer look at tuples. Recall that  $\langle a, b, c \rangle$  is the 3-tuple with components *a*, *b*, and *c*. In TLA<sup>+</sup>, this 3-tuple is actually the function with domain  $\{1, 2, 3\}$  that maps 1 to *a*, 2 to *b*, and 3 to *c*. Thus,  $\langle a, b, c \rangle[2]$  equals *b*.

TLA<sup>+</sup> provides the Cartesian product operator  $\times$  of ordinary mathematics, where  $A \times B \times C$  is the set of all 3-tuples  $\langle a, b, c \rangle$  such that *a* ∈ *A*, *b* ∈ *B*, and *c* ∈ *C*. Note that  $A \times B \times C$  is different from  $A \times (B \times C)$ , which is the set of pairs  $\langle a, p \rangle$  with *a* in *A* and *p* in the set of pairs  $B \times C$ .

The *Sequences* module defines finite sequences to be tuples. Hence, a sequence of length *n* is a function with domain  $1 \dots n$ . In fact, *s* is a sequence iff

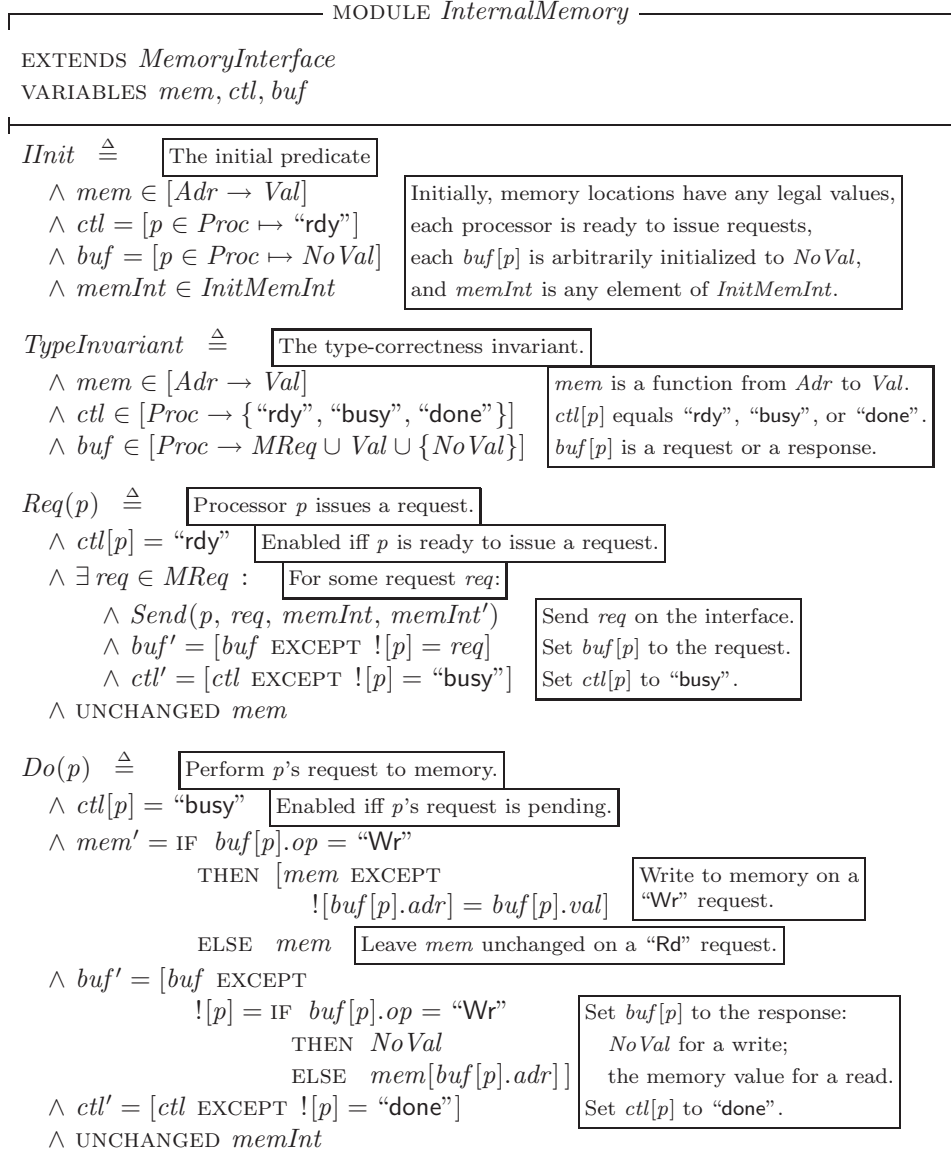


Figure 5.2: The internal memory specification (beginning).

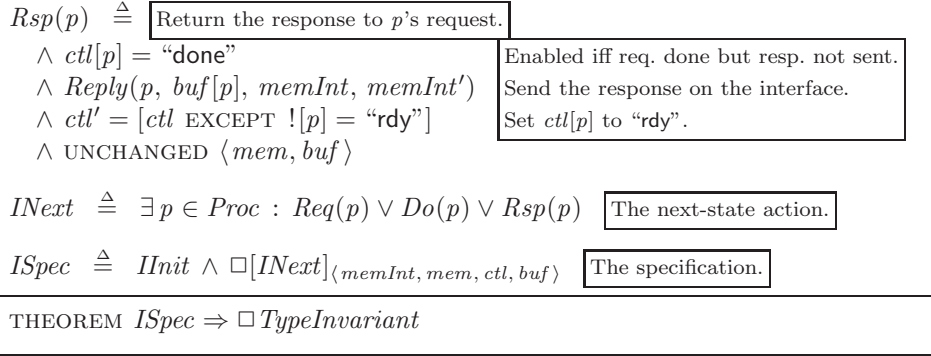


Figure 5.3: The internal memory specification (end).

it equals  $[i \in 1 \dots Len(s) \mapsto s[i]]$ . Below are a few operator definitions from the *Sequences* module. (The meanings of the operators are described in Section 4.1.)

$$\begin{aligned}
Head(s) &\triangleq s[1] \\
Tail(s) &\triangleq [i \in 1 \dots (Len(s) - 1) \mapsto s[i + 1]] \\
s \circ t &\triangleq [i \in 1 \dots (Len(s) + Len(t)) \mapsto \\
&\quad \text{IF } i \leq Len(s) \text{ THEN } s[i] \text{ ELSE } t[i - Len(s)]]
\end{aligned}$$

## 5.5 Recursive Function Definitions

We need one more tool to write the caching memory specification: recursive function definitions. Recursively defined functions are familiar to programmers. The classic example is the factorial function, which I'll call *fact*. It's usually defined by writing:

$$fact[n] = \text{IF } n = 0 \text{ THEN } 1 \text{ ELSE } n * fact[n - 1]$$

for all  $n \in Nat$ . The  $TLA^+$  notation for writing functions suggests trying to define *fact* by

$$fact \triangleq [n \in Nat \mapsto \text{IF } n = 0 \text{ THEN } 1 \text{ ELSE } n * fact[n - 1]]$$

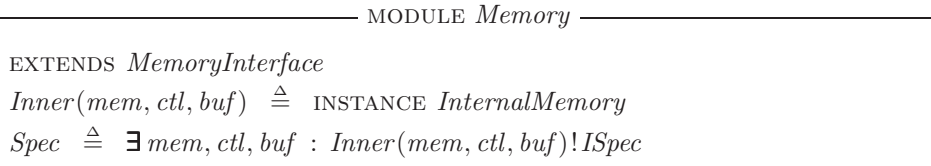


Figure 5.4: The memory specification.

This definition is illegal because the occurrence of *fact* to the right of the  $\triangleq$  is undefined—*fact* is defined only after its definition.

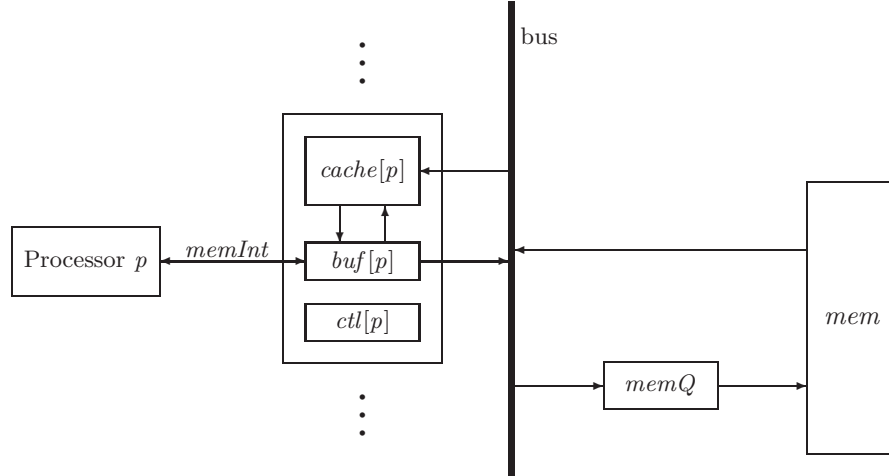
TLA<sup>+</sup> does allow the apparent circularity of recursive function definitions. We can define the factorial function *fact* by:

$$fact[n \in Nat] \triangleq \text{IF } n = 0 \text{ THEN } 1 \text{ ELSE } n * fact[n - 1]$$

In general, a definition of the form  $f[x \in S] \triangleq e$  can be used to define recursively a function *f* with domain *S*. Section 6.3 explains exactly what such a definition means. For now, we will just write recursive definitions without worrying about their meaning.

## 5.6 A Write-Through Cache

We now specify a simple write-through cache that implements the memory specification. The system is described by the following picture:



Each processor *p* communicates with a local controller, which maintains three state components: *buf[p]*, *ctl[p]*, and *cache[p]*. The value of *cache[p]* represents the processor's cache; *buf[p]* and *ctl[p]* play the same role as in the internal memory specification (module *InternalMemory*). (However, as we will see below, *ctl[p]* can assume an additional value “waiting”.) These local controllers communicate with the main memory *mem*,<sup>4</sup> and with one another, over a bus. Requests from the processors to the main memory are in the queue *memQ* of maximum length *QLen*.

<sup>4</sup>Although the write-through cache implements the linearizable memory, its main memory does not directly implement the variable *mem* of the specification in module *InternalMemory*.

A write request by processor  $p$  is performed by the action  $DoWr(p)$ . This is a write-through cache, meaning that every write request updates main memory. So, the  $DoWr(p)$  action writes the value into  $cache[p]$  and adds the write request to the tail of  $memQ$ . It also updates  $cache[q]$  for any other processor  $q$  that has a copy of the address in its cache. When the request reaches the head of  $memQ$ , the action  $MemQWr$  stores the value in  $mem$ .

A read request by processor  $p$  is performed by the action  $DoRd(p)$ , which obtains the value from the cache. If the value is not in the cache, the action  $RdMiss(p)$  adds the request to the tail of  $memQ$  and sets  $ctl[p]$  to “waiting”. When the enqueued request reaches the head of  $memQ$ , the action  $MemQRd$  reads the value and puts it in  $cache[p]$ , enabling the  $DoRd(p)$  action.

We might expect the  $MemQRd$  action to read the value from  $mem$ . However, this could cause an error if there is a write to that address enqueued in  $memQ$  behind the read request. In that case, reading the value from memory could lead to two processors having different values for the address in their caches: the one that issued the read request, and the one that issued the write request that followed the read in  $memQ$ . So, the  $MemQRd$  action must read the value from the last write to that address in  $memQ$ , if there is such a write; otherwise, it reads the value from  $mem$ .

Eviction of an address from processor  $p$ ’s cache is represented by a separate  $Evict(p)$  action. Since all cached values have been written to memory, eviction does nothing but remove the address from the cache. There is no reason to evict an address until the space is needed, so in an implementation, this action would be executed only when a request for an uncached address is received from  $p$  and  $p$ ’s cache is full. But that’s a performance optimization; it doesn’t affect the correctness of the algorithm, so it doesn’t appear in the specification. We allow a cached address to be evicted from  $p$ ’s cache at any time—except if the address was just put there by a  $MemQRd$  action for the current request. This is the case when  $ctl[p]$  equals “waiting” and  $buf[p].adr$  equals the cached address.

The actions  $Req(p)$  and  $Rsp(p)$ , which represent processor  $p$  issuing a request and the memory issuing a reply to  $p$ , are the same as the corresponding actions of the memory specification, except that they also leave the new variables  $cache$  and  $memQ$  unchanged.

To specify all these actions, we must decide how the processor caches and the queue of requests to memory are represented by the variables  $memQ$  and  $cache$ . We let  $memQ$  be a sequence of pairs of the form  $\langle p, req \rangle$ , where  $req$  is a request and  $p$  is the processor that issued it. For any memory address  $a$ , we let  $cache[p][a]$  be the value in  $p$ ’s cache for address  $a$  (the “copy” of  $a$  in  $p$ ’s cache). If  $p$ ’s cache does not have a copy of  $a$ , we let  $cache[p][a]$  equal  $NoVal$ .

The specification appears in module *WriteThroughCache* on pages 56–58. I’ll now go through this specification, explaining some of the finer points and some notation that we haven’t encountered before.

|                                                                                                      |                                                                                                             |
|------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------|
| MODULE <i>WriteThroughCache</i>                                                                      |                                                                                                             |
| EXTENDS <i>Naturals</i> , <i>Sequences</i> , <i>MemoryInterface</i>                                  |                                                                                                             |
| VARIABLES <i>mem</i> , <i>ctl</i> , <i>buf</i> , <i>cache</i> , <i>memQ</i>                          |                                                                                                             |
| CONSTANT <i>QLen</i>                                                                                 |                                                                                                             |
| ASSUME $(QLen \in Nat) \wedge (QLen > 0)$                                                            |                                                                                                             |
| $M \triangleq$ INSTANCE <i>InternalMemory</i>                                                        |                                                                                                             |
| <i>Init</i> $\triangleq$                                                                             | The initial predicate                                                                                       |
| $\wedge M!Init$                                                                                      | <i>mem</i> , <i>buf</i> , and <i>ctl</i> are initialized as in the internal memory spec.                    |
| $\wedge cache =$                                                                                     | All caches are initially empty ( $cache[p][a] = NoVal$ for all $p, a$ ).                                    |
| $[p \in Proc \mapsto [a \in Adr \mapsto NoVal]]$                                                     |                                                                                                             |
| $\wedge memQ = \langle \rangle$                                                                      | The queue <i>memQ</i> is initially empty.                                                                   |
| <i>TypeInvariant</i> $\triangleq$                                                                    | The type invariant.                                                                                         |
| $\wedge mem \in [Adr \rightarrow Val]$                                                               |                                                                                                             |
| $\wedge ctl \in [Proc \rightarrow \{\text{"rdy"}, \text{"busy"}, \text{"waiting"}, \text{"done"}\}]$ |                                                                                                             |
| $\wedge buf \in [Proc \rightarrow MReq \cup Val \cup \{NoVal\}]$                                     |                                                                                                             |
| $\wedge cache \in [Proc \rightarrow [Adr \rightarrow Val \cup \{NoVal\}]]$                           |                                                                                                             |
| $\wedge memQ \in Seq(Proc \times MReq)$                                                              | <i>memQ</i> is a sequence of $\langle proc., request \rangle$ pairs.                                        |
| <i>Coherence</i> $\triangleq$                                                                        | Asserts that if two processors' caches both have copies of an address, then those copies have equal values. |
| $\forall p, q \in Proc, a \in Adr :$                                                                 |                                                                                                             |
| $(NoVal \notin \{cache[p][a], cache[q][a]\}) \Rightarrow (cache[p][a] = cache[q][a])$                |                                                                                                             |
| <i>Req</i> ( <i>p</i> ) $\triangleq$                                                                 | Processor <i>p</i> issues a request.                                                                        |
| $M!Req(p) \wedge UNCHANGED \langle cache, memQ \rangle$                                              |                                                                                                             |
| <i>Rsp</i> ( <i>p</i> ) $\triangleq$                                                                 | The system issues a response to processor <i>p</i> .                                                        |
| $M!Rsp(p) \wedge UNCHANGED \langle cache, memQ \rangle$                                              |                                                                                                             |
| <i>RdMiss</i> ( <i>p</i> ) $\triangleq$                                                              | Enqueue a request to write value from memory to <i>p</i> 's cache.                                          |
| $\wedge (ctl[p] = \text{"busy"}) \wedge (buf[p].op = \text{"Rd"})$                                   | Enabled on a read request when the address is not in <i>p</i> 's cache and <i>memQ</i> is not full.         |
| $\wedge cache[p][buf[p].adr] = NoVal$                                                                |                                                                                                             |
| $\wedge Len(memQ) < QLen$                                                                            | Append $\langle p, request \rangle$ to <i>memQ</i> .<br>Set <i>ctl</i> [ <i>p</i> ] to "waiting".           |
| $\wedge memQ' = Append(memQ, \langle p, buf[p] \rangle)$                                             |                                                                                                             |
| $\wedge ctl' = [ctl \text{ EXCEPT } ![p] = \text{"waiting"}]$                                        |                                                                                                             |
| $\wedge UNCHANGED \langle memInt, mem, buf, cache \rangle$                                           |                                                                                                             |

Figure 5.5: The write-through cache specification (beginning).

$DoRd(p) \triangleq$  Perform a read by  $p$  of a value in its cache.  
 $\wedge ctl[p] \in \{\text{"busy"}, \text{"waiting"}\}$   
 $\wedge buf[p].op = \text{"Rd"}$   
 $\wedge cache[p][buf[p].adr] \neq NoVal$   
 $\wedge buf' = [buf \text{ EXCEPT } ![p] = cache[p][buf[p].adr]]$   
 $\wedge ctl' = [ctl \text{ EXCEPT } ![p] = \text{"done"}]$   
 $\wedge \text{UNCHANGED } \langle memInt, mem, cache, memQ \rangle$ 

 Enabled if a read request is pending and address is in cache.  
 Get result from cache.  
 Set  $ctl[p]$  to "done".

$DoWr(p) \triangleq$  Write to  $p$ 's cache, update other caches, and enqueue memory update.  
 $\text{LET } r \triangleq buf[p]$  Processor  $p$ 's request.  
 $\text{IN } \wedge (ctl[p] = \text{"busy"}) \wedge (r.op = \text{"Wr"})$  
 Enabled if write request pending and  $memQ$  is not full.
   
 $\wedge Len(memQ) < QLen$   
 $\wedge cache' =$  Update  $p$ 's cache and any other cache that has a copy.  
 $\quad [q \in Proc \mapsto \text{IF } (p = q) \vee (cache[q][r.adr] \neq NoVal)$   
 $\quad \quad \text{THEN } [cache[q] \text{ EXCEPT } ![r.adr] = r.val]$   
 $\quad \quad \text{ELSE } cache[q]]$   
 $\wedge memQ' = Append(memQ, \langle p, r \rangle)$  
 Enqueue write at tail of  $memQ$ .  
 Generate response.  
 Set  $ctl$  to indicate request is done.
   
 $\wedge buf' = [buf \text{ EXCEPT } ![p] = NoVal]$   
 $\wedge ctl' = [ctl \text{ EXCEPT } ![p] = \text{"done"}]$   
 $\wedge \text{UNCHANGED } \langle memInt, mem \rangle$

$vmem \triangleq$  The value  $mem$  will have after all the writes in  $memQ$  are performed.  
 $\text{LET } f[i \in 0 \dots Len(memQ)] \triangleq$  The value  $mem$  will have after the first  $i$  writes in  $memQ$  are performed.  
 $\quad \text{IF } i = 0 \text{ THEN } mem$   
 $\quad \quad \text{ELSE IF } memQ[i][2].op = \text{"Rd"}$   
 $\quad \quad \quad \text{THEN } f[i - 1]$   
 $\quad \quad \quad \text{ELSE } [f[i - 1] \text{ EXCEPT } ![memQ[i][2].adr] =$   
 $\quad \quad \quad \quad memQ[i][2].val]$   
 $\text{IN } f[Len(memQ)]$

$MemQWr \triangleq$  Perform write at head of  $memQ$  to memory.  
 $\text{LET } r \triangleq Head(memQ)[2]$  The request at the head of  $memQ$ .  
 $\text{IN } \wedge (memQ \neq \langle \rangle) \wedge (r.op = \text{"Wr"})$  
 Enabled if  $Head(memQ)$  a write.  
 Perform the write to memory.
   
 $\wedge mem' =$   
 $\quad [mem \text{ EXCEPT } ![r.adr] = r.val]$   
 $\wedge memQ' = Tail(memQ)$  
 Remove the write from  $memQ$ .
   
 $\wedge \text{UNCHANGED } \langle memInt, cache, buf, ctl, cache \rangle$

Figure 5.6: The write-through cache specification (middle).

|                                                                               |                                                             |
|-------------------------------------------------------------------------------|-------------------------------------------------------------|
| $MemQRd \triangleq$                                                           | Perform an enqueued read to memory.                         |
| LET $p \triangleq Head(memQ)[1]$                                              | The requesting processor.                                   |
| $r \triangleq Head(memQ)[2]$                                                  | The request at the head of $memQ$ .                         |
| IN $\wedge (memQ \neq \langle \rangle) \wedge (r.op = \text{"Rd"})$           | Enabled if $Head(memQ)$ is a read.                          |
| $\wedge memQ' = Tail(memQ)$                                                   | Remove the head of $memQ$ .                                 |
| $\wedge cache' =$                                                             | Put value from memory or $memQ$ in $p$ 's cache.            |
| $[cache \text{ EXCEPT } ![p][r.adr] = vmem[r.adr]]$                           |                                                             |
| $\wedge \text{UNCHANGED } \langle memInt, mem, buf, ctl \rangle$              |                                                             |
| $Evict(p, a) \triangleq$                                                      | Remove address $a$ from $p$ 's cache.                       |
| $\wedge (ctl[p] = \text{"waiting"}) \Rightarrow (buf[p].adr \neq a)$          | Can't evict $a$ if it was just read into cache from memory. |
| $\wedge cache' = [cache \text{ EXCEPT } ![p][a] = NoVal]$                     |                                                             |
| $\wedge \text{UNCHANGED } \langle memInt, mem, buf, ctl, memQ \rangle$        |                                                             |
| $Next \triangleq$                                                             |                                                             |
| $\vee \exists p \in Proc : \vee Req(p) \vee Rsp(p)$                           |                                                             |
| $\vee RdMiss(p) \vee DoRd(p) \vee DoWr(p)$                                    |                                                             |
| $\vee \exists a \in Adr : Evict(p, a)$                                        |                                                             |
| $\vee MemQWr \vee MemQRd$                                                     |                                                             |
| $Spec \triangleq$                                                             |                                                             |
| $Init \wedge \Box[Next]_{\langle memInt, mem, buf, ctl, cache, memQ \rangle}$ |                                                             |
| THEOREM $Spec \Rightarrow \Box(TypeInvariant \wedge Coherence)$               |                                                             |
| $LM \triangleq$                                                               |                                                             |
| INSTANCE $Memory$                                                             | The memory spec. with internal variables hidden.            |
| THEOREM $Spec \Rightarrow LM!Spec$                                            | Formula $Spec$ implements the memory spec.                  |

Figure 5.7: The write-through cache specification (end).

The EXTENDS, declaration statements, and ASSUME are familiar. We can re-use some of the definitions from the *InternalMemory* module, so an INSTANCE statement instantiates a copy of that module. (The parameters of module *InternalMemory* are instantiated by the parameters of the same name in module *WriteThroughCache*.)

The initial predicate *Init* contains the conjunct  $M!Init$ , which asserts that *mem*, *ctl*, and *buf* have the same initial values as in the internal memory specification. The write-through cache allows  $ctl[p]$  to have the value “waiting” that it didn’t in the internal memory specification, so we can’t re-use the internal memory’s type invariant  $M!TypeInvariant$ . Formula *TypeInvariant* therefore explicitly describes the types of *mem*, *ctl*, and *buf*. The type of *memQ* is the set of sequences of  $\langle processor, request \rangle$  pairs.

The module next defines the predicate *Coherence*, which asserts the basic cache coherence property of the write-through cache: for any processors  $p$  and

$q$  and any address  $a$ , if  $p$  and  $q$  both have copies of address  $a$  in their caches, then those copies are equal. Note the trick of writing  $x \notin \{y, z\}$  instead of the equivalent but longer formula  $(x \neq y) \wedge (x \neq z)$ .

The actions  $Req(p)$  and  $Rsp(p)$ , which represent a processor sending a request and receiving a reply, are essentially the same as the corresponding actions in module *InternalMemory*. The only difference is that they must specify that the variables *cache* and *memQ*, not present in module *InternalMemory*, are left unchanged.

In the definition of *RdMiss*, the expression  $Append(memQ, \langle p, buf[p] \rangle)$  is the sequence obtained by appending the element  $\langle p, buf[p] \rangle$  to the end of *memQ*.

The *DoRd*( $p$ ) action represents the performing of the read from  $p$ 's cache. If  $ctl[p] = \text{"busy"}$ , then the address was originally in the cache. If  $ctl[p] = \text{"waiting"}$ , then the address was just read into the cache from memory.

The *DoWr*( $p$ ) action writes the value to  $p$ 's cache and updates the value in any other caches that have copies. It also enqueues a write request in *memQ*. In an implementation, the request is put on the bus, which transmits it to the other caches and to the *memQ* queue. In our high-level view of the system, we represent all this as a single step.

The definition of *DoWr* introduces the  $TLA^+$  LET/IN construct. The LET clause consists of a sequence of definitions, whose scope extends until the end of the IN clause. In the definition of *DoWr*, the LET clause defines  $r$  to equal  $buf[p]$  within the IN clause. Observe that the definition of  $r$  contains the parameter  $p$  of the definition of *DoWr*. Hence, we could not move the definition of  $r$  outside the definition of *DoWr*.

A definition in a LET is just like an ordinary definition in a module; in particular, it can have parameters. These local definitions can be used to shorten an expression by replacing common subexpressions with an operator. In the definition of *DoWr*, I replaced five instances of  $buf[p]$  by the single symbol  $r$ . This was a silly thing to do, because it makes almost no difference in the length of the definition and it requires the reader to remember the definition of the new symbol  $r$ . But using a LET to eliminate common subexpressions can often greatly shorten and simplify an expression.

A LET can also be used to make an expression easier to read, even if the operators it defines appear only once in the IN expression. We write a specification with a sequence of definitions, instead of just defining a single monolithic formula, because a formula is easier to understand when presented in smaller chunks. The LET construct allows the process of splitting a formula into smaller parts to be done hierarchically. A LET can appear as a subexpression of an IN expression. Nested LETs are common in large, complicated specifications.

Next comes the definition of the state function *vmem*, which is used in defining action *MemQRd* below. It equals the value that the main memory *mem* will have after all the write operations currently in *memQ* have been performed. Recall that the value read by *MemQRd* must be the most recent one written

to that address—a value that may still be in *memQ*. That value is the one in *vmem*. The function *vmem* is defined in terms of the recursively defined function *f*, where *f*[*i*] is the value *mem* will have after the first *i* operations in *memQ* have been performed. Note that *memQ*[*i*][2] is the second component (the request) of *memQ*[*i*], the *i*th element in the sequence *memQ*.

The next two actions, *MemQWr* and *MemQRd*, represent the processing of the request at the head of the *memQ* queue—*MemQWr* for a write request, and *MemQRd* for a read request. These actions also use a LET to make local definitions. Here, the definitions of *p* and *r* could be moved before the definition of *MemQWr*. In fact, we could save space by replacing the two local definitions of *r* with one global (within the module) definition. However, making the definition of *r* global in this way would be somewhat distracting, since *r* is used only in the definitions of *MemQWr* and *MemQRd*. It might be better instead to combine these two actions into one. Whether you put a definition into a LET or make it more global should depend on what makes the specification easier to read. Writing specifications is a craft whose mastery requires talent and hard work.

The *Evict*(*p*, *a*) action represents the operation of removing address *a* from processor *p*'s cache. As explained above, we allow an address to be evicted at any time—unless the address was just written to satisfy a pending read request, which is the case iff *ctl*[*p*] = “waiting” and *buf*[*p*].*adr* = *a*. Note the use of the “double subscript” in the EXCEPT expression of the action's second conjunct. This conjunct “assigns *NoVal* to *cache*[*p*][*a*]”. If address *a* is not in *p*'s cache, then *cache*[*p*][*a*] already equals *NoVal* and an *Evict*(*p*, *a*) step is a stuttering step.

The definitions of the next-state action *Next* and of the complete specification *Spec* are straightforward. The module closes with two theorems that are discussed below.

## 5.7 Invariance

Module *WriteThroughCache* contains the theorem

$$\text{THEOREM } \textit{Spec} \Rightarrow \Box(\textit{TypeInvariant} \wedge \textit{Coherence})$$

which asserts that *TypeInvariant*  $\wedge$  *Coherence* is an invariant of *Spec*. A state predicate *P*  $\wedge$  *Q* is always true iff both *P* and *Q* are always true, so  $\Box(P \wedge Q)$  is equivalent to  $\Box P \wedge \Box Q$ . This implies that the theorem above is equivalent to the two theorems:

$$\text{THEOREM } \textit{Spec} \Rightarrow \Box \textit{TypeInvariant}$$

$$\text{THEOREM } \textit{Spec} \Rightarrow \Box \textit{Coherence}$$

The first theorem is the usual type-invariance assertion. The second, which asserts that *Coherence* is an invariant of *Spec*, expresses an important property of the algorithm.

Although *TypeInvariant* and *Coherence* are both invariants of the temporal formula *Spec*, they differ in a fundamental way. If  $s$  is any state satisfying *TypeInvariant*, then any state  $t$  such that  $s \rightarrow t$  is a *Next* step also satisfies *TypeInvariant*. This property is expressed by:

THEOREM  $TypeInvariant \wedge Next \Rightarrow TypeInvariant'$

(Recall that *TypeInvariant'* is the formula obtained by priming all the variables in formula *TypeInvariant*.) In general, when  $P \wedge N \Rightarrow P'$  holds, we say that predicate  $P$  is an invariant of action  $N$ . Predicate *TypeInvariant* is an invariant of *Spec* because it is an invariant of *Next* and it is implied by the initial predicate *Init*.

Predicate *Coherence* is not an invariant of the next-state action *Next*. For example, suppose  $s$  is a state in which

- $cache[p1][a] = 1$
- $cache[q][b] = NoVal$ , for all  $\langle q, b \rangle$  different from  $\langle p1, a \rangle$
- $mem[a] = 2$
- $memQ$  contains the single element  $\langle p2, [op \mapsto \text{“Rd”}, adr \mapsto a] \rangle$

for two different processors  $p1$  and  $p2$  and some address  $a$ . Such a state  $s$  (an assignment of values to variables) exists, assuming that there are at least two processors and at least one address. Then *Coherence* is true in state  $s$ . Let  $t$  be the state obtained from  $s$  by taking a *MemQRd* step. In state  $t$ , we have  $cache[p2][a] = 2$  and  $cache[p1][a] = 1$ , so *Coherence* is false. Hence *Coherence* is not an invariant of the next-state action.

*Coherence* is an invariant of formula *Spec* because states like  $s$  cannot occur in a behavior satisfying *Spec*. Proving its invariance is not so easy. We must find a predicate *Inv* that is an invariant of *Next* such that *Inv* implies *Coherence* and is implied by the initial predicate *Init*.

Important properties of a specification can often be expressed as invariants. Proving that a state predicate  $P$  is an invariant of a specification means proving a formula of the form

$$Init \wedge \Box[Next]_v \Rightarrow \Box P$$

This is done by finding an appropriate state predicate *Inv* and proving

$$Init \Rightarrow Inv, \quad Inv \wedge [Next]_v \Rightarrow Inv', \quad Inv \Rightarrow P$$

Since our subject is specification, not proof, I won't discuss how to find *Inv*.

An invariant of a specification  $S$  that is also an invariant of its next-state action is sometimes called an *inductive invariant* of  $S$ .

## 5.8 Proving Implementation

Module *WriteThroughCache* ends with the theorem

$$\text{THEOREM } Spec \Rightarrow LM!Spec$$

where  $LM!Spec$  is formula  $Spec$  of module *Memory*. By definition of this formula (page 53), we can restate the theorem as

$$\text{THEOREM } Spec \Rightarrow \exists mem, ctl, buf : LM!Inner(mem, ctl, buf)!ISpec$$

where  $LM!Inner(mem, ctl, buf)!ISpec$  is formula  $ISpec$  of the *InternalMemory* module. The rules of logic tell us that to prove such a theorem, we must find “witnesses” for the quantified variables  $mem$ ,  $ctl$ , and  $buf$ . These witness are state functions (ordinary expressions with no primes), which I’ll call  $omem$ ,  $octl$ , and  $obuf$ , that satisfy:

$$(5.1) \quad Spec \Rightarrow LM!Inner(omem, octl, obuf)!ISpec$$

The tuple  $\langle omem, octl, obuf \rangle$  of witness functions is called a *refinement mapping*, and we describe (5.1) as the assertion that  $Spec$  implements formula  $ISpec$  (of module *InternalMemory*) under this refinement mapping. Intuitively, this means  $Spec$  implies that the value of the tuple of state functions  $\langle memInt, omem, octl, obuf \rangle$  changes the way  $ISpec$  asserts that the tuple of variables  $\langle memInt, mem, ctl, buf \rangle$  should change.

I will now briefly describe how we prove (5.1); for details, see [2]. Let me first introduce a bit of non-TLA<sup>+</sup> notation. For any formula  $F$  of module *InternalMemory*, let  $\overline{F}$  equal  $LM!Inner(omem, octl, obuf)!F$ , which is formula  $F$  with  $omem$ ,  $octl$ , and  $obuf$  substituted for  $mem$ ,  $ctl$ , and  $buf$ . In particular,  $\overline{mem}$ ,  $\overline{ctl}$ , and  $\overline{buf}$  equal  $omem$ ,  $octl$ , and  $obuf$ , respectively.

Replacing  $Spec$  and  $ISpec$  by their definitions transforms (5.1) to

$$\begin{aligned} & Init \wedge \Box[Next]_{\langle memInt, mem, buf, ctl, cache, memQ \rangle} \\ & \Rightarrow \overline{Init} \wedge \Box[\overline{Next}]_{\langle memInt, \overline{mem}, \overline{ctl}, \overline{buf} \rangle} \end{aligned}$$

This is proved by finding an invariant  $Inv$  of  $Spec$  such that

$$\begin{aligned} & \wedge Init \Rightarrow \overline{Init} \\ & \wedge Inv \wedge Next \Rightarrow \vee \overline{Next} \\ & \quad \vee \text{UNCHANGED } \langle memInt, \overline{mem}, \overline{ctl}, \overline{buf} \rangle \end{aligned}$$

The second conjunct is called *step simulation*. It asserts that a  $Next$  step starting in a state satisfying the invariant  $Inv$  is either an  $\overline{Next}$  step—a step that changes the 4-tuple  $\langle memInt, omem, octl, obuf \rangle$  the way an  $\overline{Next}$  step changes  $\langle memInt, mem, ctl, buf \rangle$ —or else it leaves that 4-tuple unchanged.

The mathematics of an implementation proof is simple, so the proof is straightforward—in theory. For specifications of real systems, such proofs can

be quite difficult. Going from the theory to practice requires turning the mathematics of proofs into an engineering discipline—a subject that deserves a book to itself. However, when writing specifications, it helps to understand refinement mappings and step simulation.

We now return to the question posed in Section 3.2: what is the relation between the specifications of the asynchronous interface in modules *AsynchInterface* and *Channel*? Recall that module *AsynchInterface* describes the interface in terms of the three variables *val*, *rdy*, and *ack*, while module *Channel* describes it with a single variable *chan* whose value is a record with *val*, *rdy*, and *ack* components. In what sense are those two specifications of the interface equivalent?

One answer that now suggests itself is that each of the specifications should implement the other under a refinement mapping. We expect formula *Spec* of module *Channel* to imply the formula obtained from *Spec* of module *AsynchInterface* by substituting for its variables *val*, *rdy*, and *ack* the *val*, *rdy*, and *ack* components of the variable *chan* of module *Channel*. This assertion is expressed precisely by the theorem in the following module.

|                                                                                                                                                                  |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| MODULE <i>ChannelImplAsynch</i>                                                                                                                                  |
| EXTENDS <i>Channel</i><br>$AInt(val, rdy, ack) \triangleq$ INSTANCE <i>AsynchInterface</i><br>THEOREM $Spec \Rightarrow AInt(chan.val, chan.rdy, chan.ack)!Spec$ |

Here, the refinement mapping substitutes  $\langle chan.val, chan.rdy, chan.ack \rangle$  for the tuple  $\langle val, rdy, ack \rangle$  of variables in the formula *Spec* of module *AsynchInterface*.

Similarly, formula *Spec* of module *AsynchInterface* implies formula *Spec* of module *Channel* with *chan* replaced by the record-valued expression:

$$[val \mapsto val, rdy \mapsto rdy, ack \mapsto ack]$$

(The first *val* in  $val \mapsto val$  is the field name in the record constructor, while the second *val* is the variable of module *AsynchInterface*.)



## Chapter 6

# Some More Math

Our mathematics is built on a small, simple collection of concepts. You've already seen most of what's needed to describe almost any kind of mathematics. All you lack are a handful of operators on sets that are described below in Section 6.1. After learning about them, you will be able to define all the data structures and operations that occur in specifications.

While our mathematics is simple, its foundations are nonobvious—for example, the meanings of recursive function definitions and the CHOOSE operator are subtle. This section discusses some of those foundations. Understanding them will help you use mathematics more effectively.

### 6.1 Sets

The simple operations on sets described in Section 1.2 are all you'll need for writing most system specifications. However, you may occasionally have to use more sophisticated operators—especially if you need to define data structures beyond tuples, records, and simple functions.

Two powerful operators of set theory are the unary operators UNION and SUBSET, defined as follows.

UNION  $S$  The union of the elements of  $S$ . In other words, a value  $e$  is an element of UNION  $S$  iff it is an element of an element of  $S$ . For example:

$$\text{UNION } \{\{1, 2\}, \{2, 3\}, \{3, 4\}\} = \{1, 2, 3, 4\}$$

SUBSET  $S$  The set of all subsets of  $S$ . In other words,  $T \in \text{SUBSET } S$  iff  $T \subseteq S$ . For example:

$$\text{SUBSET } \{1, 2\} = \{\{\}, \{1\}, \{2\}, \{1, 2\}\}$$

|                                                 |
|-------------------------------------------------|
| Mathematicians write UNION $S$ as $\bigcup S$ . |
|-------------------------------------------------|

|                                                                                                    |
|----------------------------------------------------------------------------------------------------|
| Mathematicians call UNION $S$ the <i>power set</i> of $S$ and write it $\mathcal{P}(S)$ or $2^S$ . |
|----------------------------------------------------------------------------------------------------|

Mathematicians often describe a set as “the set of all ... such that ...”.  $\text{TLA}^+$  has two constructs that formalize such a description:

- $\{x \in S : P\}$  The subset of  $S$  consisting of all elements  $x$  satisfying property  $P$ . For example, the set of odd natural numbers can be written  $\{n \in \text{Nat} : n \% 2 = 1\}$ . The identifier  $x$  is bound in  $P$ ; it may not occur in  $S$ .
- $\{e : x \in S\}$  The set of elements of the form  $e$ , for all  $x$  in the set  $S$ . For example,  $\{2 * n + 1 : n \in \text{Nat}\}$  is the set of all odd natural numbers. The identifier  $x$  is bound in  $e$ ; it may not occur in  $S$ .

The modulus operator  $\%$  is described in Section 2.5.

The construct  $\{e : x \in S\}$  has the same generalizations as  $\exists x \in S : F$ . For example,  $\{e : x \in S, y \in T\}$  is the set of all elements of the form  $e$ , for  $x$  in  $S$  and  $y$  in  $T$ . In the construct  $\{x \in S : P\}$ , we can let  $x$  be a tuple. For example,  $\{\langle y, z \rangle \in S : P\}$  is the set of all pairs  $\langle y, z \rangle$  in the set  $S$  that satisfy  $P$ . The BNF grammar of  $\text{TLA}^+$  in Section 14.1.2 specifies precisely what set expressions you can write.

All the set operators we’ve seen so far are built-in operators of  $\text{TLA}^+$ . There is also a standard module *FiniteSets* that defines two operators:

- Cardinality*( $S$ ) The number of elements in set  $S$ , if  $S$  is a finite set.
- IsFiniteSet*( $S$ ) True iff  $S$  is a finite set.

Careless reasoning about sets can lead to problems. The classic example of this is Russell’s paradox:

Let  $\mathcal{R}$  be the set of all sets  $S$  such that  $S \notin S$ . The definition of  $\mathcal{R}$  implies that  $\mathcal{R} \in \mathcal{R}$  is true iff  $\mathcal{R} \notin \mathcal{R}$  is true.

The formula  $\mathcal{R} \notin \mathcal{R}$  is simply the negation of  $\mathcal{R} \in \mathcal{R}$ , and a formula and its negation can neither both be true nor both be false. The source of the paradox is that  $\mathcal{R}$  isn’t a set. There’s no way to write it in  $\text{TLA}^+$ . Intuitively,  $\mathcal{R}$  is too big to be a set. A collection  $\mathcal{C}$  is too big to be a set if it is as big as the collection of all sets—meaning that we can assign to every set a different element of  $\mathcal{C}$ . That is,  $\mathcal{C}$  is too big to be a set if we can define an operator *SMap* such that:

- *SMap*( $S$ ) is in  $\mathcal{C}$ , for any set  $S$ .
- If  $S$  and  $T$  are two different sets, then *SMap*( $S$ )  $\neq$  *SMap*( $T$ ).

For example, the collection of all sequences of length 2 is too big to be a set; we can define the operator *SMap* by

$$\text{SMap}(S) \triangleq \langle 1, S \rangle$$

## 6.2 Silly Expressions

Most modern programming languages introduce some form of type checking to prevent you from writing silly expressions like  $3/\text{“abc”}$ .  $\text{TLA}^+$  is based on the usual formalization of mathematics, which doesn’t have types. In an untyped formalism, every syntactically well-formed expression has a meaning—even a silly expression like  $3/\text{“abc”}$ . Mathematically, the expression  $3/\text{“abc”}$  is no sillier than the expression  $3/0$ , and mathematicians implicitly write that silly expression all the time. For example, consider the valid formula

$$\forall x \in \text{Real} : (x \neq 0) \Rightarrow (x * (3/x) = 3)$$

where *Real* is the set of all real numbers. This asserts that  $(x \neq 0) \Rightarrow (x * (3/x) = 3)$  is true for all real numbers  $x$ . Substituting 0 for  $x$  yields the valid formula  $(0 \neq 0) \Rightarrow (0 * (3/0) = 3)$  that contains the silly expression  $3/0$ . It’s valid because  $0 \neq 0$  equals FALSE, and  $\text{FALSE} \Rightarrow P$  is true for any formula  $P$ .

A correct formula can contain silly expressions. For example,  $3/0 = 3/0$  is a correct formula because any value equals itself. However, the validity of a correct formula cannot depend on the meaning of a silly expression. If an expression is silly, then its meaning is probably unspecified. The definitions of 0, 3, /, and \* (which are in the standard module *Reals*) don’t specify the value of  $0 * (3/0)$ , so there’s no way of knowing whether that value equals 3.

No sensible syntactic rules can prevent you from writing  $3/0$  without also preventing you from writing perfectly reasonable expressions. The typing rules of programming languages introduce complexity and limitations on what you can write that don’t exist in ordinary mathematics. In a well-designed programming language, the costs of types are balanced by benefits: types allow a compiler to produce more efficient code, and type checking catches errors. For programming languages, the benefits seem to outweigh the costs. For writing specifications, I have found that the costs outweigh the benefits.

If you’re used to the constraints of programming languages, it may be a while before you start taking advantage of the freedom afforded by mathematics. At first, you won’t think of defining anything like the operator  $R$  defined on page 50 of Section 5.2, which couldn’t be written in a typed programming language.

## 6.3 Recursive Function Definitions Revisited

Section 5.5 introduced recursive function definitions. Let’s now examine what such definitions mean mathematically. Mathematicians usually define the factorial function *fact* by writing:

$$\text{fact}[n] = \text{IF } n = 0 \text{ THEN } 1 \text{ ELSE } n * \text{fact}[n - 1], \text{ for all } n \in \text{Nat}$$

This definition can be justified by proving that it defines a unique function  $fact$  with domain  $Nat$ . In other words,  $fact$  is the unique value satisfying:

$$(6.1) \quad fact = [n \in Nat \mapsto \text{IF } n = 0 \text{ THEN } 1 \text{ ELSE } n * fact[n - 1]]$$

The CHOOSE operator, introduced on pages 47–48 of Section 5.1, allows us to express “the value satisfying property  $P$ ” as  $\text{CHOOSE } x : P$ . We can therefore define  $fact$  as follows to be the value satisfying (6.1):

$$(6.2) \quad fact \triangleq \text{CHOOSE } fact : \\ fact = [n \in Nat \mapsto \text{IF } n = 0 \text{ THEN } 1 \\ \text{ELSE } n * fact[n - 1]]$$

(Since the symbol  $fact$  is not yet defined in the expression to the right of the  $\triangleq$ , we can use it as the bound identifier in the CHOOSE expression.) The TLA<sup>+</sup> definition

$$fact[n \in Nat] \triangleq \text{IF } n = 0 \text{ THEN } 1 \text{ ELSE } n * fact[n - 1]$$

is simply an abbreviation for (6.2). In general,  $f[x \in S] \triangleq e$  is an abbreviation for:

$$(6.3) \quad f \triangleq \text{CHOOSE } f : f = [x \in S \mapsto e]$$

TLA<sup>+</sup> allows you to write silly definitions. For example, you can write

$$(6.4) \quad circ[n \in Nat] \triangleq \text{CHOOSE } y : y \neq circ[n]$$

This appears to define  $circ$  to be a function such that  $circ[n] \neq circ[n]$  for any natural number  $n$ . There obviously is no such function, so  $circ$  can’t be defined to equal it. A recursive function definition doesn’t necessarily define a function. If there is no  $f$  that equals  $[x \in S \mapsto e]$ , then (6.3) defines  $f$  to be some unspecified value. Thus, the nonsensical definition (6.4) defines  $circ$  to be some unknown value.

If we want to reason about a function  $f$  defined by  $f[x \in S] \triangleq e$ , we need to prove that there exists an  $f$  that equals  $[x \in S \mapsto e]$ . The existence of  $f$  is obvious if  $f$  does not occur in  $e$ . If it does, so this is a recursive definition, then there is something to prove. Since I’m not discussing proofs, I won’t describe how to prove it. Intuitively, you have to check that, as in the case of the factorial function, the definition uniquely determines the value of  $f[x]$  for every  $x$  in  $S$ .

Recursion is a common programming technique because programs must compute values using a small repertoire of simple elementary operations. It’s not used as often in mathematical definitions, where we needn’t worry about how to compute the value and can use the powerful operators of logic and set theory. For example, the operators *Head*, *Tail*, and  $\circ$  are defined in Section 5.4 without recursion, even though computer scientists usually define them recursively. Still, there are some things that are best defined inductively, using a recursive function definition.

## 6.4 Functions versus Operators

Consider these definitions, which we've seen before

$$\begin{aligned} \text{Tail}(s) &\triangleq [i \in 1 \dots (\text{Len}(s) - 1) \mapsto s[i + 1]] \\ \text{fact}[n \in \text{Nat}] &\triangleq \text{IF } n = 0 \text{ THEN } 1 \text{ ELSE } n * \text{fact}[n - 1] \end{aligned}$$

They define two very different kinds of objects: *fact* is a function, and *Tail* is an operator. Functions and operators differ in a few basic ways.

Their most obvious difference is that a function like *fact* by itself is a complete expression that denotes a value, but an operator like *Tail* is not. Both  $\text{fact}[n] \in S$  and  $\text{fact} \in S$  are syntactically correct expressions. But, while  $\text{Tail}(n) \in S$  is syntactically correct,  $\text{Tail} \in S$  is not. It is gibberish—a meaningless string of symbols, like  $x + = 0$ .

Their second difference is more profound. The definition of *Tail* defines  $\text{Tail}(s)$  for all values of  $s$ . For example, it defines  $\text{Tail}(1/2)$  to equal

$$(6.5) \quad [i \in 1 \dots (\text{Len}(1/2) - 1) \mapsto (1/2)[i + 1]]$$

We have no idea what this expression means, because we don't know what  $\text{Len}(1/2)$  or  $(1/2)[i + 1]$  mean. But, whatever (6.5) means, it equals  $\text{Tail}(1/2)$ .

The definition of *fact* defines  $\text{fact}[n]$  only for  $n \in \text{Nat}$ . It tells us nothing about the value of  $\text{fact}[1/2]$ . The expression  $\text{fact}[1/2]$  is syntactically well-formed, so it denotes some value. But the definition of *fact* tells us nothing about what that value is.

Unlike an operator, a function must have a domain, which is a set. We cannot define a function *Tail* so that  $\text{Tail}[s]$  is the tail of any nonempty sequence  $s$ ; the domain of such a function would have to include all nonempty sequences, and the collection of all such sequences is too big to be a set. (The operator *SMap* defined by  $\text{SMap}(S) \triangleq \langle S \rangle$  maps every set to a different nonempty sequence.) Hence, we can't define *Tail* to be a function.

Unlike a function, an operator cannot be defined recursively. However, we can usually transform an illegal recursive operator definition into a nonrecursive one using a recursive function definition. For example, let's try to define the *Cardinality* operator on the set of finite sets. (Recall that the cardinality of a finite set  $S$  is the number of elements in  $S$ .) The collection of all finite sets is too big to be a set. (The operator  $\text{SMap}(S) \triangleq \{S\}$  maps every set  $S$  to a different set  $\{S\}$  of cardinality 1.) The *Cardinality* operator has a simple intuitive definition:

- $\text{Cardinality}(\{\}) = 0$ .
- If  $S$  is a nonempty finite set, then

$$\text{Cardinality}(S) = 1 + \text{Cardinality}(S \setminus \{x\})$$

where  $x$  is an arbitrary element of  $S$ . (The set  $S \setminus \{x\}$  contains all the elements of  $S$  except  $x$ .)

Using the CHOOSE operator to describe an arbitrary element of  $S$ , we can write this as the more formal-looking, but still illegal, definition:

$$\begin{aligned} \text{Cardinality}(S) &\triangleq \\ &\text{IF } S = \{\} \text{ THEN } 0 \\ &\quad \text{ELSE } 1 + \text{Cardinality}(S \setminus \{\text{CHOOSE } x : x \in S\}) \end{aligned}$$

This definition is illegal because it's circular—only in a recursive function definition can the symbol being defined appear to the right of the  $\triangleq$ .

To turn this into a legal definition, observe that, for a given set  $S$ , we can define a function  $CS$  such that  $CS[T]$  equals the cardinality of  $T$  for every subset  $T$  of  $S$ . The definition is

$$\begin{aligned} CS[T \in \text{SUBSET } S] &\triangleq \\ &\text{IF } T = \{\} \text{ THEN } 0 \\ &\quad \text{ELSE } 1 + CS[T \setminus \{\text{CHOOSE } x : x \in T\}] \end{aligned}$$

Since  $S$  is a subset of itself, this defines  $CS[S]$  to equal  $\text{Cardinality}(S)$ , if  $S$  is a finite set. (We don't know or care what  $CS[S]$  equals if  $S$  is not finite.) So, we can define the *Cardinality* operator by:

$$\begin{aligned} \text{Cardinality}(S) &\triangleq \\ &\text{LET } CS[T \in \text{SUBSET } S] \triangleq \\ &\quad \text{IF } T = \{\} \text{ THEN } 0 \\ &\quad \quad \text{ELSE } 1 + CS[T \setminus \{\text{CHOOSE } x : x \in T\}] \\ &\text{IN } CS[S] \end{aligned}$$

Operators also differ from functions in that an operator can take an operator as an argument. For example, we can define an operator *IsPartialOrder* so that  $\text{IsPartialOrder}(R, S)$  equals true iff the operator  $R$  defines an irreflexive partial order on  $S$ . The definition is

$$\begin{aligned} \text{IsPartialOrder}(R(\_ , \_), S) &\triangleq \\ &\wedge \forall x, y, z \in S : R(x, y) \wedge R(y, z) \Rightarrow R(x, z) \\ &\wedge \forall x \in S : \neg R(x, x) \end{aligned}$$

We could also use an infix-operator symbol like  $\prec$  instead of  $R$  as the parameter of the definition, writing:

$$\begin{aligned} \text{IsPartialOrder}(\_ \prec \_, S) &\triangleq \\ &\wedge \forall x, y, z \in S : (x \prec y) \wedge (y \prec z) \Rightarrow (x \prec z) \\ &\wedge \forall x \in S : \neg(x \prec x) \end{aligned}$$

The first argument of *IsPartialOrder* is an operator that takes two arguments; its second argument is an expression. Since  $>$  is an operator that takes two arguments, the expression  $\text{IsPartialOrder}(>, \text{Nat})$  is syntactically correct. In fact, it's valid—if  $>$  is defined to be the usual operator on numbers. The expression

If you don't know what an irreflexive partial order is, read this definition of *IsPartialOrder* to find out.

$IsPartialOrder(+, 3)$  is also syntactically correct, but it's silly and we have no idea whether or not it's valid.

The last difference between operators and functions has nothing to do with mathematics and is an idiosyncrasy of  $TLA^+$ : the language doesn't permit you to define infix functions. So, if we want to define  $/$ , we have no choice but to make it an operator.

One can write equally nonsensical things using functions or operators. However, whether you use functions or operators may determine whether the nonsense you write is nonsyntactic gibberish or syntactically correct but semantically silly. The string of symbols  $2("a")$  is not a syntactically correct formula because  $2$  is not an operator. However,  $2["a"]$ , which can also be written  $2.a$ , is a syntactically correct expression. It's nonsensical because  $2$  isn't a function, so we don't know what  $2["a"]$  means. Similarly,  $Tail(s, t)$  is syntactically incorrect because  $Tail$  is an operator that takes a single argument. However, as explained in Section 15.1.7 (page 285),  $fact[m, n]$  is syntactic sugar for  $fact[\langle m, n \rangle]$ , so it is a syntactically correct, semantically silly formula. Whether an error is syntactic or semantic determines what kind of tool can catch it. In particular, the parser described in Chapter 12 catches syntactic errors, but not semantic silliness.

The distinction between functions and operators seems to confuse some people. One reason is that, although this distinction exists in ordinary math, it usually goes unnoticed by mathematicians. If you point out to them that  $SUBSET$  can't be a function because its domain couldn't be a set, mathematicians will realize that they use operators like  $SUBSET$  and  $\in$  all the time. But they never think of them as forming a distinct kind of entity. Logicians will observe that the distinction between operators and values, including functions, arises because  $TLA^+$  is a first-order logic rather than a higher-order logic.

When defining an object  $V$ , you may have to decide whether to make  $V$  an operator or a function. The differences between operators and functions will often determine the decision. For example, if a variable may have  $V$  as its value, then  $V$  must be a function. Thus, in the memory specification of Section 5.3, we had to represent the state of the memory by a function rather than an operator, since the variable  $mem$  couldn't equal an operator. If these differences don't determine whether to use an operator or a function, then it's a matter of taste. I usually prefer operators.

## 6.5 Using Functions

Consider the following two formulas:

$$(6.6) \quad f' = [i \in Nat \mapsto i + 1]$$

$$(6.7) \quad \forall i \in Nat : f'[i] = i + 1$$

These formulas both imply that  $f'[i] = i + 1$  for every natural number  $i$ , but they are not equivalent. Formula (6.6) uniquely determines  $f'$ , asserting that it's a function with domain  $Nat$ . Formula (6.7) is satisfied by lots of different values of  $f'$ —for example, by the function

$$[i \in Real \mapsto \text{IF } i \in Nat \text{ THEN } i + 1 \text{ ELSE } \sqrt{i}]$$

In fact, from (6.7), we can't even deduce that  $f'$  is a function. Formula (6.6) implies formula (6.7), but not vice-versa.

When writing specifications, we almost always want to specify the new value of a variable  $f$  rather than the new values of  $f[i]$  for all  $i$  in some set. We therefore usually write (6.6) rather than (6.7).

## 6.6 Choose

The CHOOSE operator was introduced in the memory interface of Section 5.1 in the simple idiom  $\text{CHOOSE } v : v \notin S$ , which is an expression whose value is not an element of  $S$ . In Section 6.3 above, we saw that it is a powerful tool that can be used in rather subtle ways.

The CHOOSE operator is known to logicians as Hilbert's  $\varepsilon$  [3].

The most common use for the CHOOSE operator is to “name” a uniquely specified value. For example,  $a/b$  is the unique real number that satisfies the formula  $a = b * (a/b)$ , if  $a$  and  $b$  are real numbers and  $b \neq 0$ . So, the standard module *Reals* defines division on the set *Real* of real numbers by

$$a/b \triangleq \text{CHOOSE } c \in Real : a = b * c$$

(The expression  $\text{CHOOSE } x \in S : P$  means  $\text{CHOOSE } x : (x \in S) \wedge P$ .) If  $a$  is a nonzero real number, then there is no real number  $c$  such that  $a = 0 * c$ . Therefore,  $a/0$  has an unspecified value. We don't know what a real number times a string equals, so we cannot say whether or not there is a real number  $c$  such that  $a$  equals “xyz”. Hence, we don't know what the value of  $a/\text{“xyz”}$  is.

People who do a lot of programming and not much mathematics often think that CHOOSE must be a nondeterministic operator. In mathematics, there is no such thing as a nondeterministic operator or a nondeterministic function. If some expression equals 42 today, then it will equal 42 tomorrow, and it will still equal 42 a million years from tomorrow. The specification

$$(x = \text{CHOOSE } n : n \in Nat) \wedge \Box[x' = \text{CHOOSE } n : n \in Nat]_x$$

allows only a single behavior—one in which  $x$  always equals  $\text{CHOOSE } n : n \in Nat$ , which is some particular, unspecified natural number. It is very different from the specification

$$(x \in Nat) \wedge \Box[x' \in Nat]_x$$

---

that allows all behaviors in which  $x$  is always a natural number—possibly a different number in each state. This specification is highly nondeterministic, allowing lots of different behaviors.



## Chapter 7

# Writing a Specification—Some Advice

You have now learned all you need to know about TLA<sup>+</sup> to write your own specifications. Here are a few additional hints to help you get started.

### 7.1 Why to Specify

Specifications are written to help eliminate errors. Writing a specification requires effort; the benefits it provides must be worth that effort. There are several benefits:

- Writing a TLA<sup>+</sup> specification can help the design process. Having to describe a design precisely often reveals problems—subtle interactions and “corner cases” that are easily overlooked. These problems are easier to correct when discovered in the design phase rather than after implementation has begun.
- A TLA<sup>+</sup> specification can provide a clear, concise way of communicating a design. It helps ensure that the designers agree on what they have designed, and it provides a valuable guide to the engineers who implement and test the system. It may also help users understand the system.
- A TLA<sup>+</sup> specification is a formal description to which tools can be applied to help find errors in the design and to help in testing the system. Some tools for TLA<sup>+</sup> specifications are being built.

Whether these benefits justify the effort of writing the specification depends on the nature of the project. Specification is not an end in itself; it is just one of many tools that an engineer should be able to use when appropriate.

## 7.2 What to Specify

Although we talk about specifying a system, that's not what we do. A specification is a mathematical model of a particular view of some part of a system. When writing a specification, the first thing you must choose is exactly what part of the system you want to describe. Sometimes the choice is obvious; often it isn't. The cache-coherence protocol of a real multiprocessor computer may be intimately connected with how the processors execute instructions. Finding an abstraction that describes the coherence protocol while suppressing the details of instruction execution may be difficult. It may require defining an interface between the processor and the memory that doesn't exist in the actual system design.

Remember that the purpose of a specification is to help avoid errors. You should specify those parts of the system for which a specification is most likely to reveal errors.  $\text{TLA}^+$  is particularly effective at revealing concurrency errors—ones that arise through the interaction of asynchronous components. So, you should concentrate your efforts on the parts of the system that are most likely to have such errors.

## 7.3 The Grain of Atomicity

After choosing what part of the system to specify, you must choose the specification's level of abstraction. The most important aspect of the level of abstraction is the grain of atomicity, the choice of what system changes are represented as a single step of a behavior. Sending a message in an actual system involves multiple suboperations, but we usually represent it as a single step. On the other hand, the sending of a message and its receipt are usually represented as separate steps when specifying a distributed system.

The same sequence of system operations is represented by a shorter sequence of steps in a coarser-grained representation than in a finer-grained one. This almost always makes the coarser-grained specification simpler than the finer-grained one. However, the finer-grained specification more accurately describes the behavior of the actual system. A coarser-grained specification may fail to reveal important details of the system.

There is no simple rule for deciding on the grain of atomicity. However, there is one way of thinking about the granularity that can help. To describe it, we need the  $\text{TLA}^+$  action-composition operator “ $\cdot$ ”. If  $A$  and  $B$  are actions, then the action  $A \cdot B$  is executed by first executing  $A$  then  $B$  in a single step. More precisely,  $A \cdot B$  is the action defined by letting  $s \rightarrow t$  be an  $A \cdot B$  step iff there exists a state  $u$  such that  $s \rightarrow u$  is an  $A$  step and  $u \rightarrow t$  is a  $B$  step.

When determining the grain of atomicity, we must decide whether to represent the execution of an operation as a single step or as a sequence of steps, each

corresponding to the execution of a suboperation. Let's consider the simple case of an operation consisting of two suboperations that are executed sequentially, where those suboperations are described by the two actions  $R$  and  $L$ . (Executing  $R$  enables  $L$  and disables  $R$ .) When the operation's execution is represented by two steps, each of those steps is an  $R$  step or an  $L$  step. The operation is then described with the action  $R \vee L$ . When its execution is represented by a single step, the operation is described with the action  $R \cdot L$ .<sup>1</sup> Let  $S2$  be the finer-grained specification in which the operation is executed in two steps, and let  $S1$  be the coarser-grained specification in which it is executed as a single  $R \cdot L$  step. To choose the grain of atomicity, we must choose whether to take  $S1$  or  $S2$  as the specification. Let's examine the relation between the two specifications.

We can transform any behavior  $\sigma$  satisfying  $S1$  into a behavior  $\hat{\sigma}$  satisfying  $S2$  by replacing each step  $s \xrightarrow{R \cdot L} t$  with the pair of steps  $s \xrightarrow{R} u \xrightarrow{L} t$ , for some state  $u$ . If we regard  $\sigma$  as being equivalent to  $\hat{\sigma}$ , then we can regard  $S1$  as being a strengthened version of  $S2$ —one that allows fewer behaviors. Specification  $S1$  requires that each  $R$  step be followed immediately by an  $L$  step, while  $S2$  allows behaviors in which other steps come between the  $R$  and  $L$  steps. To choose the appropriate grain of atomicity, we must decide whether those additional behaviors allowed by  $S2$  are important.

The additional behaviors allowed by  $S2$  are not important if the actual system executions they describe are also described by behaviors allowed by  $S1$ . So, we can ask whether each behavior  $\tau$  satisfying  $S2$  has a corresponding behavior  $\tilde{\tau}$  satisfying  $S1$  that is, in some sense, equivalent to  $\tau$ . One way to construct  $\tilde{\tau}$  from  $\tau$  is to transform a sequence of steps

$$(7.1) \quad s \xrightarrow{R} u_1 \xrightarrow{A_1} u_2 \xrightarrow{A_2} u_3 \dots u_n \xrightarrow{A_n} u_{n+1} \xrightarrow{L} t$$

into the sequence

$$(7.2) \quad s \xrightarrow{A_1} v_1 \dots v_{k-2} \xrightarrow{A_k} v_{k-1} \xrightarrow{R} v_k \xrightarrow{L} v_{k+1} \xrightarrow{A_{k+1}} v_{k+2} \dots v_{n+1} \xrightarrow{A_n} t$$

where the  $A_i$  are other system actions that can be executed between the  $R$  and  $L$  steps. Both sequences start in state  $s$  and end in state  $t$ , but the intermediate states may be different.

When is such a transformation possible? An answer can be given in terms of commutativity relations. We say that actions  $A$  and  $B$  commute if performing them in either order produces the same result. Formally,  $A$  and  $B$  commute iff  $A \cdot B$  is equivalent to  $B \cdot A$ . A simple sufficient condition for commutativity is that two actions commute if (i) each one leaves unchanged any variable whose value may be changed by the other, and (ii) neither enables or disables the other.

<sup>1</sup>We actually describe the operation with an ordinary action, like the ones we've been writing, that is equivalent to  $R \cdot L$ . The operator “ $\cdot$ ” rarely appears in an actual specification. If you're ever tempted to use it, look for a better way to write the specification; you can probably find one.

It's not hard to see that we can transform (7.1) to (7.2) in the following two cases:

- $R$  commutes with each  $A_i$ . (In this case,  $k = n$ .)
- $L$  commutes with each  $A_i$ . (In this case,  $k = 0$ .)

In general, if an operation consists of a sequence of  $m$  subactions, we must decide whether to choose the finer-grained representation  $O_1 \vee O_2 \vee \dots \vee O_m$  or the coarser-grained one  $O_1 \cdot O_2 \cdots O_m$ . The generalization of the transformation from (7.1) to (7.2) is one that transforms an arbitrary behavior satisfying the finer-grained specification into one in which the sequence of  $O_1, O_2, \dots, O_m$  steps come one right after the other. Such a transformation is possible if all but one of the actions  $O_i$  commute with every other system action. Commutativity can be replaced by weaker conditions, but it is the most common case.

By commuting actions and replacing a sequence  $s \xrightarrow{O_1} \dots \xrightarrow{O_m} t$  of steps by a single  $O_1 \cdots O_m$  step, you may be able to transform any behavior of a finer-grained specification into a corresponding behavior of a coarser-grained one. But that doesn't mean that the coarser-grained specification is just as good as the finer-grained one. The sequences (7.1) and (7.2) are not the same, and a sequence of  $O_i$  steps is not the same as a single  $O_1 \cdots O_m$  step. Whether you can consider the transformed behavior to be equivalent to the original one, and use the coarser-grained specification, depends on the particular system you are specifying and on the purpose of the specification. Understanding the relation between finer- and coarser-grained specifications can help you choose between them; it won't make the choice for you.

## 7.4 The Data Structures

Another aspect of a specification's level of abstraction is the accuracy with which it describes the system's data structures. For example, should the specification of a programming interface describe the actual layout of a procedure's arguments in memory, or should the arguments be represented more abstractly?

To answer such a question, you must remember that the purpose of the specification is to help catch errors. A precise description of the layout of procedure arguments will help prevent errors caused by misunderstandings about that layout, but at the cost of complicating the programming interface's specification. The cost is justified only if such errors are likely to be a real problem and the TLA<sup>+</sup> specification provides the best way to avoid them.

If the purpose of the specification is to catch errors caused by the asynchronous interaction of concurrently executing components, then detailed descriptions of data structures will be a needless complication. So, you will probably want to use high-level, abstract descriptions of the system's data structures

in the specification. For example, to specify a program interface, you might introduce constant parameters to represent the actions of calling and returning from a procedure—parameters analogous to *Send* and *Reply* of the memory interface described in Section 5.1 (page 45).

## 7.5 Writing the Specification

Once you've chosen the part of the system to specify and the level of abstraction, you're ready to start writing the  $\text{TLA}^+$  specification. We've already seen how this is done; let's review the steps.

First, pick the variables and define the type invariant and initial predicate. In the course of doing this, you will determine the constant parameters and assumptions about them that you need. You may also have to define some additional constants.

Next, write the next-state action, which forms the bulk of the specification. Sketching a few sample behaviors may help you get started. You must first decide how to decompose the next-state action as the disjunction of actions describing the different kinds of system operations. You then define those actions. The goal is to make the action definitions as compact and easy to read as possible. This requires carefully structuring them. One way to reduce the size of a specification is to define state predicates and state functions that are used in several different action definitions. When writing the action definitions, you will determine which of the standard modules you will need to use and add the appropriate `EXTENDS` statement. You may also have to define some constant operators for the data structures that you are using.

You must now write the temporal part of the specification. (If you want to specify liveness properties, you have to choose the fairness conditions, as described below in Chapter 8. You then combine the initial predicate, next-state action, and any fairness conditions you've chosen into the definition of a single temporal formula that is the specification.

Finally, you can assert theorems about the specification. If nothing else, you may want to add a type-correctness theorem.

## 7.6 Some Further Hints

Here are a few miscellaneous suggestions that may help you write better specifications.

**Don't be too clever.**

Cleverness can make a specification hard to read—and even wrong. The formula  $q = \langle h' \rangle \circ q'$  may look like a nice, short way of writing:

$$(7.3) \quad (h' = \text{Head}(q)) \wedge (q' = \text{Tail}(q))$$

But not only is  $q = \langle h' \rangle \circ q'$  harder to understand than (7.3), it's also wrong. We don't know what  $a \circ b$  equals if  $a$  and  $b$  are not both sequences, so we don't know whether  $h' = \text{Head}(q)$  and  $q' = \text{Tail}(q)$  are the only values of  $h'$  and  $q'$  that satisfy  $q = \langle h' \rangle \circ q'$ . There could be other values of  $h'$  and  $q'$ , which are not sequences, that satisfy the formula.

**A type invariant is not an assumption.**

Type invariance is a property of a specification, not an assumption. When writing a specification, we usually define a type invariant. But that's just a definition; a definition is not an assumption. Suppose you define a type invariant that asserts that a variable  $n$  is of type *Nat*. You may be tempted to then think that a conjunct  $n' > 7$  in an action asserts that  $n'$  is a natural number greater than 7. It doesn't. The formula  $n' > 7$  asserts only that  $n' > 7$ . It is satisfied if  $n' = \sqrt{96}$  as well as if  $n' = 8$ . Since we don't know whether or not “abc”  $> 7$  is true, it might be satisfied if  $n' = \text{“abc”}$ . The meaning of the formula is not changed just because you've defined a type invariant that asserts  $n \in \text{Nat}$ .

In general, you may want to describe the new value of a variable  $x$  by asserting some property of  $x'$ . However, the next-state relation should imply that  $x'$  is an element of some suitable set. For example, a specification might define:<sup>2</sup>

$$\begin{aligned} \text{Action1} &\triangleq (n' > 7) \wedge \dots \\ \text{Action2} &\triangleq (n' \leq 6) \wedge \dots \\ \text{Next} &\triangleq (n' \in \text{Nat}) \wedge (\text{Action1} \vee \text{Action2}) \end{aligned}$$

**Don't be too abstract**

Suppose a user interacts with the system by typing on a keyboard. We could describe the interaction abstractly with a variable *typ* and an operator parameter *KeyStroke*, where the action *KeyStroke*(“a”, *typ*, *typ'*) represents the user typing an “a”. This is the approach we took in describing the communication between the processors and the memory in the *MemoryInterface* module (page 48).

A more concrete description would be to let *kbd* represent the state of the keyboard, perhaps letting  $kbd = \{\}$  mean that no key is depressed, and  $kbd = \{\text{“a”}\}$  mean that the *a* key is depressed. The typing of an *a* is represented by

<sup>2</sup>An alternative approach is to define *Next* to equal  $\text{Action1} \vee \text{Action2}$  and to let the specification be  $\text{Init} \wedge \Box[\text{Next}] \dots \wedge \Box(n \in \text{Nat})$ . But it's usually better to stick to the simple form  $\text{Init} \wedge \Box[\text{Next}] \dots$  for specifications.

two steps, a  $[kbd = \{\}] \rightarrow [kbd = \{\text{"a"}\}]$  step represents the pressing of the  $a$  key, and a  $[kbd = \{\text{"a"}\}] \rightarrow [kbd = \{\}]$  step represents its release. This is the approach we took in the asynchronous interface specifications of Chapter 3.

The abstract interface is simpler; typing an  $a$  is represented by a single  $KeyStroke(\text{"a"}, typ, typ')$  step instead of a pair of steps. However, using the concrete representation leads us naturally to ask: what if the user presses the  $a$  key and, before releasing it, presses the  $b$  key? That's easy to describe with the concrete representation. The state with both keys depressed is  $kbd = \{a, b\}$ . Pressing and releasing a key are represented simply by the two actions

$$Press(k) \triangleq kbd' = kbd \cup \{\text{"k"}\} \quad Release(k) \triangleq kbd' = kbd \setminus \{\text{"k"}\}$$

This possibility cannot be expressed with the simple abstract interface. To express it abstractly, we would have to replace the parameter  $KeyStroke$  with two parameters  $Press$  and  $Release$ , and we would have to express explicitly the property that a key can't be released until it has been depressed, and vice-versa. The more concrete representation is simpler.

We might decide that we don't want to consider the possibility of the user pressing two keys at once, and that we prefer the abstract representation. But that should be a conscious decision. Our abstraction should not blind us to what can happen in the actual system. When in doubt, it's safer to use a concrete representation that more accurately describes the real system. That way, you are less likely to overlook problems that can arise in the actual system.

### Don't assume values that look different are unequal.

The rules of  $TLA^+$  do not imply that  $1 \neq \text{"a"}$ . If the system can send a message that is either a string or a number, represent the message as a record with a *type* and *value* field—for example,

$$[type \mapsto \text{"String"}, value \mapsto \text{"a"}] \quad \text{or} \quad [type \mapsto \text{"Nat"}, value \mapsto 1]$$

We know that these two values are different because they have different *type* fields.

### Move quantification to the outside.

Specifications are usually easier to read if  $\exists$  is moved outside disjunctions and  $\forall$  is moved outside conjunctions. For example, instead of:

$$\begin{aligned} Up &\triangleq \exists e \in Elevator : \dots \\ Down &\triangleq \exists e \in Elevator : \dots \\ Move &\triangleq Up \vee Down \end{aligned}$$

it's usually better to write:

$$\begin{aligned} Up(e) &\triangleq \dots \\ Down(e) &\triangleq \dots \\ Move &\triangleq \exists e \in Elevator : Up(e) \vee Down(e) \end{aligned}$$

### Prime only what you mean to prime.

When writing an action, be careful where you put your primes. The expression  $f[e]'$  equals  $f'[e']$ ; it equals  $f'[e]$  only if  $e' = e$ , which need not be true if the expression  $e$  contains variables. Be especially careful when priming an operator whose definition contains a variable. For example, suppose  $x$  is a variable and  $op$  is defined by

$$op(a) \triangleq x + a$$

Then  $op(y)'$  equals  $(x+y)'$ , which equals  $x'+y'$ , while  $op(y')$  equals  $x+y'$ . There is no way to use  $op$  and  $'$  to write the formula  $x' + y$ . (Writing  $op'(y)$  doesn't work because it's illegal—you can only prime an expression, not an operator.)

### Write comments as comments.

Don't put comments into the specification itself. I have seen people write things like the following action definition:

$$\begin{aligned} A &\triangleq \vee \wedge x \geq 0 \\ &\quad \wedge \dots \\ &\quad \vee \wedge x < 0 \\ &\quad \wedge \text{FALSE} \end{aligned}$$

The second disjunct is meant to indicate that the writer intended  $A$  not to be enabled when  $x < 0$ . But that disjunct is completely redundant, since  $F \wedge \text{FALSE}$  equals  $\text{FALSE}$ , and  $F \vee \text{FALSE}$  equals  $F$ , for any formula  $F$ . So the second disjunct of the definition serves only as a form of comment. It's better to write:

$$A \triangleq \wedge x \geq 0 \quad \boxed{A \text{ is not enabled if } x < 0} \\ \wedge \dots$$

## 7.7 When and How to Specify

Specifications are often written later than they should be. Engineers are usually under severe time constraints, and they may feel that writing a specification will slow them down. Only after a design has become so complex that they need help understanding it do engineers think about writing a precise specification.

Writing a specification helps you think clearly. Thinking clearly is hard; we can use all the help we can get. Making specification part of the design process can improve the design.

I have described how to write a specification assuming that the system design already exists. But it's better to write the specification as the system is being designed. The specification will start out being incomplete and probably incorrect. For example, an initial specification of the write-through cache of Section 5.6 (page 54) might include the definition:

$$\begin{array}{l}
 RdMiss(p) \triangleq \boxed{\text{Enqueue a request to write value from memory to } p\text{'s cache.}} \\
 \boxed{\text{Some enabling condition must be conjoined here.}} \\
 \wedge memQ' = Append(memQ, buf[p]) \quad \boxed{\text{Append request to } memQ.} \\
 \wedge ctl' = [ctl \text{ EXCEPT } ![p] = "?"] \quad \boxed{\text{Set } ctl[p] \text{ to value to be determined later.}} \\
 \wedge \text{UNCHANGED } \langle memInt, mem, buf, cache \rangle
 \end{array}$$

Some system functionality will at first be omitted; it can be included later by adding new disjuncts to the next-state action. Tools can be applied to these preliminary specifications to help find design errors.



# Part II

## More Advanced Topics



## Chapter 8

# Liveness and Fairness

The specifications we have written so far say what a system must *not* do. The clock must not advance from 11 to 9; the receiver must not receive a message if the FIFO is empty. They don't require that the system ever actually do anything. The clock need never tick; the sender need never send any messages. Our specifications have described what are called *safety properties*. If a safety property is violated, it is violated at some particular point in the behavior—by a step that advances the clock from 11 to 9, or that reads the wrong value from memory. Therefore, we can talk about a safety property being satisfied by a finite behavior, which means that it has not been violated by any step so far.

We now learn how to specify that something *does* happen—that the clock keeps ticking, or that a value is eventually read from memory. We specify *liveness properties*—ones that cannot be violated at any particular instant. Only by examining an entire infinite behavior can we tell that the clock has stopped ticking, or that a message is never sent.

We express liveness properties as temporal formulas. This means that, to add liveness conditions to your specifications, you have to understand temporal logic—the logic of temporal formulas. The chapter begins, in Section 8.1, with a more rigorous look at what a temporal formula means. To understand a logic, you have to understand what its true formulas are. Section 8.2 is about temporal tautologies, the true formulas of temporal logic. Sections 8.3–8.6 describe how to use temporal formulas to specify liveness properties. Section 8.7 completes our study of temporal logic by examining the temporal quantifier  $\exists$ . Finally, Section 8.8 reviews what we've done and explains why the undisciplined use of temporal logic is dangerous.

This chapter is the only one that contains proofs. It would be nice if you learned to write similar proofs yourself, but it doesn't matter if you don't. The proofs are here because studying them can help you developing the intuitive understanding of temporal formulas that you need to write specifications—

an understanding that makes the truth of a simple temporal tautology like  $\Box\Box F \equiv \Box F$  as obvious as the truth of a simple theorem about numbers like  $\forall n \in \text{Nat} : 2 * n \geq n$ .

## 8.1 Temporal Formulas

Recall that a state assigns a value to every variable, and a behavior is an infinite sequence of states. A temporal formula is true or false of a behavior. Formally, a temporal formula  $F$  assigns a Boolean value, which we write  $\sigma \models F$ , to a behavior  $\sigma$ . We say that  $F$  is true of  $\sigma$ , or that  $\sigma$  satisfies  $F$ , iff  $\sigma \models F$  equals TRUE. To define the meaning of a temporal formula  $F$ , we have to explain how to determine the value of  $\sigma \models F$  for any behavior  $\sigma$ . For now, we consider only temporal formulas that don't contain the temporal existential quantifier  $\exists$ .

It's easy to define the meaning of a Boolean combination of temporal formulas in terms of the meanings of those formulas. The formula  $F \wedge G$  is true of a behavior  $\sigma$  iff both  $F$  and  $G$  are true of  $\sigma$ , and  $\neg F$  is true of  $\sigma$  iff  $F$  is not true of  $\sigma$ . These definitions are written more formally as:

$$\sigma \models (F \wedge G) \triangleq (\sigma \models F) \wedge (\sigma \models G) \quad \sigma \models \neg F \triangleq \neg(\sigma \models F)$$

These are the definitions of the meaning of  $\wedge$  and of  $\neg$  as operators on temporal formulas. The meanings of the other Boolean operators are similarly defined. We can also define in this way the ordinary predicate-logic quantifiers  $\forall$  and  $\exists$  as operators on temporal formulas—for example:

$$\sigma \models (\exists r : F) \triangleq \exists r : (\sigma \models F)$$

Ordinary quantification over constant sets is defined the same way. For example, if  $S$  is an ordinary constant expression—that is, one containing no variables—then:

$$\sigma \models (\forall r \in S : F) \triangleq \forall r \in S : (\sigma \models F)$$

Quantifiers are discussed further in Section 8.7.

All the unquantified temporal formulas that we've seen have been Boolean combinations of three simple kinds of formulas, which have the following meanings:

- A state predicate, viewed as a temporal formula, is true of a behavior iff it is true in the first state of the behavior.
- A formula  $\Box P$ , where  $P$  is a state predicate, is true of a behavior iff  $P$  is true in every state of the behavior.
- A formula  $\Box[N]_v$ , where  $N$  is an action and  $v$  is a state function, is true of a behavior iff every successive pair of steps in the behavior is a  $[N]_v$  step.

*State function and state predicate are defined on page 25.*

Since a state predicate is an action that contains no primed variables, we can both combine and generalize these three kinds of temporal formulas into the two kinds of formulas  $A$  and  $\Box A$ , where  $A$  is an action. I'll first explain the meanings of these two kinds of formulas, and then define the operator  $\Box$  in general. To do this, I will use the notation that  $\sigma_i$  is the  $(i + 1)^{\text{st}}$  state of the behavior  $\sigma$ , for any natural number  $i$ , so  $\sigma$  is the behavior  $\sigma_0 \rightarrow \sigma_1 \rightarrow \sigma_2 \rightarrow \dots$ .

We interpret an arbitrary action  $A$  as a temporal formula by defining  $\sigma \models A$  to be true iff the first two states of  $\sigma$  are an  $A$  step. That is, we define  $\sigma \models A$  to be true iff  $\sigma_0 \rightarrow \sigma_1$  is an  $A$  step. In the special case when  $A$  is a state predicate,  $\sigma_0 \rightarrow \sigma_1$  is an  $A$  step iff  $A$  is true in state  $\sigma_0$ , so this definition of  $\sigma \models A$  generalizes our interpretation of a state predicate as a temporal formula.

We have already seen that  $\Box[N]_v$  is true of a behavior iff each step is a  $[N]_v$  step. This leads us to define  $\sigma \models \Box A$  to be true iff  $\sigma_n \rightarrow \sigma_{n+1}$  is an  $A$  step, for all natural numbers  $n$ .

We now generalize from the definition of  $\sigma \models \Box A$  for an action  $A$  to the definition of  $\sigma \models \Box F$  for an arbitrary temporal formula  $F$ . We defined  $\sigma \models \Box A$  to be true iff  $\sigma_n \rightarrow \sigma_{n+1}$  is an  $A$  step for all  $n$ . This is true iff  $A$ , interpreted as a temporal formula, is true of a behavior whose first step is  $\sigma_n \rightarrow \sigma_{n+1}$ , for all  $n$ . Let's define  $\sigma^{+n}$  to be the suffix of  $\sigma$  obtained by deleting its first  $n$  states:

$$\sigma^{+n} \triangleq \sigma_n \rightarrow \sigma_{n+1} \rightarrow \sigma_{n+2} \rightarrow \dots$$

Then  $\sigma_n \rightarrow \sigma_{n+1}$  is the first step of  $\sigma^{+n}$ , so  $\sigma \models \Box A$  is true iff  $\sigma^{+n} \models A$  is true for all  $n$ . In other words:

$$\sigma \models \Box A \equiv \forall n \in \text{Nat} : \sigma^{+n} \models A$$

The obvious generalization is:

$$\sigma \models \Box F \triangleq \forall n \in \text{Nat} : \sigma^{+n} \models F$$

for any temporal formula  $F$ . In other words,  $\sigma$  satisfies  $\Box F$  iff every suffix  $\sigma^{+n}$  of  $\sigma$  satisfies  $F$ . This defines the meaning of the temporal operator  $\Box$ .

We have now defined the meaning of any temporal formula built from actions (including state predicates), Boolean operators, and the  $\Box$  operator. For example:

$$\begin{aligned} \sigma \models \Box((x = 1) \Rightarrow \Box(y > 0)) \\ &\equiv \forall n \in \text{Nat} : \sigma^{+n} \models ((x = 1) \Rightarrow \Box(y > 0)) && \boxed{\text{By the meaning of } \Box.} \\ &\equiv \forall n \in \text{Nat} : (\sigma^{+n} \models (x = 1)) \Rightarrow (\sigma^{+n} \models \Box(y > 0)) && \boxed{\text{By the meaning of } \Rightarrow.} \\ &\equiv \forall n \in \text{Nat} : (\sigma^{+n} \models (x = 1)) \Rightarrow && \boxed{\text{By the meaning of } \Box.} \\ &\quad (\forall m \in \text{Nat} : (\sigma^{+n})^{+m} \models (y > 0)) \end{aligned}$$

Thus,  $\sigma \models \Box((x = 1) \Rightarrow \Box(y > 0))$  is true iff, for all  $n \in \text{Nat}$ , if  $x = 1$  is true in state  $\sigma_n$ , then  $y > 0$  is true in all states  $\sigma_{n+m}$  with  $m \geq 0$ .

To understand temporal formulas intuitively, think of  $\sigma_n$  as the state of the universe at time instant  $n$  during the behavior  $\sigma$ .<sup>1</sup> For any state predicate  $P$ , the expression  $\sigma^{+n} \models P$  asserts that  $P$  is true at time  $n$ . Thus,  $\Box((x = 1) \Rightarrow \Box(y > 0))$  asserts that, any time  $x = 1$  is true,  $y > 0$  is true from then on. For an arbitrary temporal formula  $F$ , we also interpret  $\sigma^{+n} \models F$  as the assertion that  $F$  is true at time instant  $n$ . The formula  $\Box F$  then asserts that  $F$  is true at all times. We can therefore read  $\Box$  as *always* or *henceforth* or *from then on*.

We saw in Section 2.2 that a specification should allow stuttering steps—ones that leave unchanged all the variables appearing in the formula. A stuttering step represents a change only to some part of the system not described by the formula; adding it to the behavior should not affect the truth of the formula. We say that a formula  $F$  is *invariant under stuttering*<sup>2</sup> iff adding or deleting a stuttering step to a behavior  $\sigma$  does not affect whether  $\sigma$  satisfies  $F$ . A sensible formula should be invariant under stuttering. There's no point writing formulas that aren't sensible, so TLA allows you to write only temporal formulas that are invariant under stuttering.

A state predicate (viewed as a temporal formula) is invariant under stuttering, since its truth depends only on the first state of a behavior, and adding a stuttering step doesn't change the first state. An arbitrary action is not invariant under stuttering. For example, the action  $[x' = x + 1]_x$  is satisfied by a behavior  $\sigma$  in which  $x$  is left unchanged in the first step and incremented by 2 in the second step; it isn't satisfied by the behavior obtained by removing the initial stuttering step from  $\sigma$ . However, the formula  $\Box[x' = x + 1]_x$  is invariant under stuttering, since it is satisfied by a behavior iff every step that changes  $x$  is an  $x' = x + 1$  step—a condition not affected by adding or deleting stuttering steps.

In general, the formula  $\Box[A]_v$  is invariant under stuttering, for any action  $A$  and state function  $v$ . However,  $\Box A$  is not invariant under stuttering for an arbitrary action  $A$ . For example,  $\Box(x' = x + 1)$  can be made false by adding a step that does not change  $x$ . So, even though we have assigned a meaning to  $\Box(x' = x + 1)$ , it isn't a legal TLA formula.

Invariance under stuttering is preserved by  $\Box$  and by the Boolean operators—that is, if  $F$  and  $G$  are invariant under stuttering, then so are  $\Box F$ ,  $\neg F$ ,  $F \wedge G$ ,  $\forall x \in S : F$ , etc. So, state predicates, formulas of the form  $\Box[N]_v$ , and all formulas obtainable from them by applying  $\Box$  and Boolean operators are invariant under stuttering.

We now examine five especially important classes of formulas that are constructed from arbitrary temporal formulas  $F$  and  $G$ . We introduce new opera-

<sup>1</sup>It is because we think of  $\sigma_n$  as the state at time  $n$ , and because we usually measure time starting from 0, that I number the states of a behavior starting with 0 rather than 1.

<sup>2</sup>This is a completely new sense of the word *invariant*; it has nothing to do with the concept of invariance discussed already.

tors for expressing the first three.

$\Diamond F$  is defined to equal  $\neg\Box\neg F$ . It asserts that  $F$  is not always false, which means that  $F$  is true at some time:

$$\begin{aligned}
 \sigma \models \Diamond F & \\
 \equiv \sigma \models \neg\Box\neg F & \quad \text{By definition of } \Diamond. \\
 \equiv \neg(\sigma \models \Box\neg F) & \quad \text{By the meaning of } \neg. \\
 \equiv \neg(\forall n \in \text{Nat} : \sigma^{+n} \models \neg F) & \quad \text{By the meaning of } \Box. \\
 \equiv \neg(\forall n \in \text{Nat} : \neg(\sigma^{+n} \models F)) & \quad \text{By the meaning of } \neg. \\
 \equiv \exists n \in \text{Nat} : \sigma^{+n} \models F & \quad \text{Because } \neg\forall\neg \text{ is equivalent to } \exists.
 \end{aligned}$$

We usually read  $\Diamond$  as *eventually*, taking eventually to include now.

$F \rightsquigarrow G$  is defined to equal  $\Box(F \Rightarrow \Diamond G)$ . The same kind of calculation we just did for  $\sigma \models \Diamond F$  shows:

$$\begin{aligned}
 \sigma \models (F \rightsquigarrow G) & \equiv \\
 \forall n \in \text{Nat} : (\sigma^{+n} \models F) \Rightarrow (\exists m \in \text{Nat} : (\sigma^{+(n+m)} \models G)) &
 \end{aligned}$$

The formula  $F \rightsquigarrow G$  asserts that whenever  $F$  is true,  $G$  is eventually true—that is,  $G$  is true then or at some later time. We read  $\rightsquigarrow$  as *leads to*.

$\Diamond\langle A \rangle_v$  is defined to equal  $\neg\Box[\neg A]_v$ , where  $A$  is an action and  $v$  a state function. It asserts that not every step is a  $(\neg A) \vee (v' = v)$  step, so some step is a  $\neg((\neg A) \vee (v' = v))$  step. Since  $\neg(P \vee Q)$  is equivalent to  $(\neg P) \wedge (\neg Q)$ , for any  $P$  and  $Q$ , action  $\neg((\neg A) \vee (v' = v))$  is equivalent to  $A \wedge (v' \neq v)$ . Hence,  $\Diamond\langle A \rangle_v$  asserts that some step is an  $A \wedge (v' \neq v)$  step—that is, an  $A$  step that changes  $v$ . We define the action  $\langle A \rangle_v$  by

$$\langle A \rangle_v \triangleq A \wedge (v' \neq v)$$

I pronounce  $\langle A \rangle_v$   
as *angle A sub v*.

so  $\Diamond\langle A \rangle_v$  asserts that eventually an  $\langle A \rangle_v$  step occurs. We think of  $\Diamond\langle A \rangle_v$  as the formula obtained by applying the operator  $\Diamond$  to  $\langle A \rangle_v$ , although technically it's not because  $\langle A \rangle_v$  isn't a temporal formula.

$\Box\Diamond F$  asserts that at all times,  $F$  is true then or at some later time. For time 0, this implies that  $F$  is true at some time  $n_0 \geq 0$ . For time  $n_0 + 1$ , it implies that  $F$  is true at some time  $n_1 \geq n_0 + 1$ . For time  $n_1 + 1$ , it implies that  $F$  is true at some time  $n_2 \geq n_1 + 1$ . Continuing the process, we see that  $F$  is true at an infinite sequence of time instants  $n_0, n_1, n_2, \dots$ . So,  $\Box\Diamond F$  implies that  $F$  is true at infinitely many instants. Conversely, if  $F$  is true at infinitely many instants, then, at every instant,  $F$  must be true at some later instant, so  $\Box\Diamond F$  is true. Therefore,  $\Box\Diamond F$  asserts that  $F$  is *infinitely often* true. In particular,  $\Box\Diamond\langle A \rangle_v$  asserts that infinitely many  $\langle A \rangle_v$  steps occur.

$\Diamond\Box F$  asserts that eventually (at some time),  $F$  becomes true and remains true thereafter. In other words,  $\Diamond\Box F$  asserts that  $F$  is *eventually always* true. In particular,  $\Diamond\Box[N]_v$  asserts that eventually, every step is a  $[N]_v$  step.

The operators  $\Box$  and  $\Diamond$  have higher precedence (bind more tightly) than the Boolean operators, so  $\Diamond F \vee \Box G$  means  $(\Diamond F) \vee (\Box G)$ . The operator  $\leadsto$  has the same precedence as  $\Rightarrow$ .

## 8.2 Temporal Tautologies

A temporal theorem is a temporal formula that is satisfied by all behaviors. In other words,  $F$  is a theorem iff  $\sigma \models F$  equals TRUE for all behaviors  $\sigma$ . For example, the *HourClock* module asserts that  $HC \Rightarrow \Box HCini$  is a theorem, where  $HC$  and  $HCini$  are the formulas defined in the module. This theorem expresses a property of the hour clock.

The formula  $\Box HCini \Rightarrow HCini$  is also a theorem. However, it tells us nothing about the hour clock because it's true regardless of how  $HCini$  is defined. For example, substituting  $x > 7$  for  $HCini$  yields the theorem  $\Box(x > 7) \Rightarrow (x > 7)$ . A formula like  $\Box HCini \Rightarrow HCini$  that is true when any formulas are substituted for its identifiers is called a *tautology*. To distinguish them from the tautologies of ordinary logic, tautologies containing temporal operators are sometimes called a *temporal* tautologies.

Let's prove that  $\Box HCini \Rightarrow HCini$  is a temporal tautology. To avoid confusing the arbitrary identifier  $HCini$  in this tautology with the formula  $HCini$  defined in the *HourClock* module, let's replace it by  $F$ , so the tautology becomes  $\Box F \Rightarrow F$ . There are axioms and inference rules for temporal logic from which we can prove any temporal tautology that, like  $\Box F \Rightarrow F$ , contains no quantifiers. However, it's often easier and more instructive to prove them directly from the meanings of the operators. We prove that  $\Box F \Rightarrow F$  is a tautology by proving that  $\sigma \models (\Box F \Rightarrow F)$  equals TRUE, for any behavior  $\sigma$  and any formula  $F$ . The proof is simple:

$$\begin{aligned}
 \sigma \models (\Box F \Rightarrow F) &\equiv (\sigma \models \Box F) \Rightarrow (\sigma \models F) && \boxed{\text{By the meaning of } \Rightarrow.} \\
 &\equiv (\forall n \in Nat : \sigma^{+n} \models F) \Rightarrow (\sigma \models F) && \boxed{\text{By definition of } \Box.} \\
 &\equiv (\forall n \in Nat : \sigma^{+n} \models F) \Rightarrow (\sigma^{+0} \models F) && \boxed{\text{By definition of } \sigma^{+0}.} \\
 &\equiv \text{TRUE} && \boxed{\text{By predicate logic.}}
 \end{aligned}$$

The temporal tautology  $\Box F \Rightarrow F$  asserts the obvious fact that, if  $F$  is true at all times, then it's true at time 0. Such a simple tautology should be obvious once you become accustomed to thinking in terms of temporal formulas. Here are three more simple tautologies, along with their English translations.

$$\begin{aligned}
 \neg\Box F &\equiv \Diamond\neg F \\
 F &\text{ is not always true iff it is eventually false.}
 \end{aligned}$$

$$\Box(F \wedge G) \equiv (\Box F) \wedge (\Box G)$$

$F$  and  $G$  are both always true iff  $F$  is always true and  $G$  is always true.  
Another way of saying this is that  $\Box$  distributes over  $\wedge$ .

$$\Diamond(F \vee G) \equiv (\Diamond F) \vee (\Diamond G)$$

$F$  or  $G$  is eventually true iff  $F$  is eventually true or  $G$  is eventually true.  
Another way of saying this is that  $\Diamond$  distributes over  $\vee$ .

At the heart of the proof of each of these tautologies is a tautology of predicate logic. For example, the proof that  $\Box$  distributes over  $\wedge$  relies on the fact that  $\forall$  distributes over  $\wedge$ :

$$\begin{aligned} \sigma \models (\Box(F \wedge G) \equiv (\Box F) \wedge (\Box G)) & \\ \equiv (\sigma \models \Box(F \wedge G)) \equiv (\sigma \models (\Box F) \wedge (\Box G)) & \quad \text{By the meaning of } \equiv. \\ \equiv (\sigma \models \Box(F \wedge G)) \equiv (\sigma \models \Box F) \wedge (\sigma \models \Box G) & \quad \text{By the meaning of } \wedge. \\ \equiv (\forall n \in \text{Nat} : \sigma^{+n} \models \Box(F \wedge G)) \equiv & \quad \text{By definition of } \Box. \\ (\forall n \in \text{Nat} : \sigma^{+n} \models (\Box F)) \wedge (\forall n \in \text{Nat} : \sigma^{+n} \models \Box G) & \\ \equiv \text{TRUE} & \quad \text{By the predicate-logic tautology } (\forall x \in S : P \wedge Q) \equiv (\forall x \in S : P) \wedge (\forall x \in S : Q). \end{aligned}$$

The operator  $\Box$  doesn't distribute over  $\vee$ , nor does  $\Diamond$  distribute over  $\wedge$ . For example,  $\Box((n \geq 0) \vee (n < 0))$  is not equivalent to  $(\Box(n \geq 0) \vee \Box(n < 0))$ ; the first formula is true for any behavior in which  $n$  is always a number, but the second is false for a behavior in which  $n$  assumes both positive and negative values. However, the following two formulas are tautologies:

$$(\Box F) \vee (\Box G) \Rightarrow \Box(F \vee G) \quad \Diamond(F \wedge G) \Rightarrow (\Diamond F) \wedge (\Diamond G)$$

Either of these tautologies can be derived from the other by substituting  $\neg F$  for  $F$  and  $\neg G$  for  $G$ . Making this substitution in the second tautology yields:

$$\begin{aligned} \text{TRUE} & \equiv \Diamond((\neg F) \wedge (\neg G)) \Rightarrow (\Diamond \neg F) \wedge (\Diamond \neg G) & \text{By substitution in the second tautology.} \\ & \equiv \Diamond \neg(F \vee G) \Rightarrow (\Diamond \neg F) \wedge (\Diamond \neg G) & \text{Because } (\neg P \wedge \neg Q) \equiv (\neg(P \vee Q)). \\ & \equiv \neg \Box(F \vee G) \Rightarrow (\neg \Box F) \wedge (\neg \Box G) & \text{Because } \Diamond \neg H \equiv \neg \Box H. \\ & \equiv \neg \Box(F \vee G) \Rightarrow \neg(\Box F \wedge \Box G) & \text{Because } (\neg P \wedge \neg Q) \equiv (\neg(P \vee Q)). \\ & \equiv (\Box F) \wedge (\Box G) \Rightarrow \Box(F \vee G) & \text{Because } (\neg P \Rightarrow \neg Q) \equiv (P \Rightarrow Q). \end{aligned}$$

This pair of tautologies illustrates a general law: from any temporal tautology, one can obtain a *dual* tautology by making the replacements

$$\Box \leftarrow \Diamond \quad \Diamond \leftarrow \Box \quad \wedge \leftarrow \vee \quad \vee \leftarrow \wedge$$

and reversing the direction of all implications. (Any  $\equiv$  or  $\neg$  is left unchanged.) As in the example above, the dual tautology can be proved from the original by replacing each identifier with its negation and applying the (dual) tautologies  $\Diamond \neg F \equiv \neg \Box F$  and  $\neg \Diamond F \equiv \Box \neg F$  along with propositional-logic reasoning.

Another important pair of dual tautologies assert that  $\Box\Diamond$  distributes over  $\vee$  and  $\Diamond\Box$  distributes over  $\wedge$ :

$$(8.1) \quad \Box\Diamond(F \vee G) \equiv (\Box\Diamond F) \vee (\Box\Diamond G) \quad \Diamond\Box(F \wedge G) \equiv (\Diamond\Box F) \wedge (\Diamond\Box G)$$

The first asserts that  $F$  or  $G$  is true infinitely often iff  $F$  is true infinitely often or  $G$  is true infinitely often. Its truth should be fairly obvious, but let's prove it. To reason about  $\Box\Diamond$ , it helps to introduce the symbol  $\exists_\infty$ , which means *there exist infinitely many*. In particular,  $\exists_\infty i \in \text{Nat} : P(i)$  means that  $P(i)$  is true for infinitely many natural numbers  $i$ . On page 91, we showed that  $\Box\Diamond F$  asserts that  $F$  is true infinitely often. Using  $\exists_\infty$ , we can express this as:

$$(8.2) \quad (\sigma \models \Box\Diamond F) \equiv (\exists_\infty i \in \text{Nat} : \sigma^{+i} \models F)$$

The same reasoning proves the following more general result, where  $P$  is any operator.

$$(8.3) \quad (\forall n \in \text{Nat} : \exists m \in \text{Nat} : P(n + m)) \equiv \exists_\infty i \in \text{Nat} : P(i)$$

Here is another useful tautology involving  $\exists_\infty$ , where  $P$  and  $Q$  are arbitrary operators and  $S$  is an arbitrary set.

$$(8.4) \quad (\exists_\infty i \in S : P(i) \vee Q(i)) \equiv (\exists_\infty i \in S : P(i)) \vee (\exists_\infty i \in S : Q(i))$$

Using these results, it's now easy to prove that  $\Box\Diamond$  distributes over  $\vee$ :

$$\begin{aligned} \sigma \models \Box\Diamond(F \vee G) & \\ \equiv \exists_\infty i \in \text{Nat} : \sigma^i \models (F \vee G) & \quad \boxed{\text{By (8.2).}} \\ \equiv (\exists_\infty i \in \text{Nat} : \sigma^i \models F) \vee (\exists_\infty i \in \text{Nat} : \sigma^i \models G) & \quad \boxed{\text{By (8.4).}} \\ \equiv (\sigma \models \Box\Diamond F) \wedge (\sigma \models \Box\Diamond G) & \quad \boxed{\text{By (8.2).}} \end{aligned}$$

From this, we deduce the dual tautology, that  $\Diamond\Box$  distributes over  $\wedge$ .

In any TLA tautology, replacing a temporal formula by an action yields a tautology—a formula that is true for all behaviors—even if that formula isn't a legal TLA formula. (Remember that we have defined the meaning of nonTLA formulas like  $\Box(x' = x + 1)$ .) We can turn apply the rules of logic to transform those nonTLA tautologies into TLA tautologies. Among these rules are the following dual equivalences, which are easy to check.

$$[A \wedge B]_v \equiv [A]_v \wedge [B]_v \quad \langle A \vee B \rangle_v \equiv \langle A \rangle_v \vee \langle B \rangle_v$$

(The second asserts that an  $A \vee B$  step that leaves  $v$  unchanged is either an  $A$  step that leaves  $v$  unchanged or a  $B$  step that leaves  $v$  unchanged.)

As an example of substituting actions for temporal formulas in TLA tautologies, let's substitute  $\langle A \rangle_v$  and  $\langle B \rangle_v$  for  $F$  and  $G$  in the first tautology of (8.1) to get

$$(8.5) \quad \Box\Diamond(\langle A \rangle_v \vee \langle B \rangle_v) \equiv (\Box\Diamond\langle A \rangle_v) \vee (\Box\Diamond\langle B \rangle_v)$$

This isn't a TLA tautology, because  $\Box\Diamond(\langle A \rangle_v \vee \langle B \rangle_v)$  isn't a TLA formula. However, a general rule of logic tells us that replacing a subformula by an equivalent one yields an equivalent formula. Substituting  $\langle A \vee B \rangle_v$  for  $\langle A \rangle_v \vee \langle B \rangle_v$  in (8.5) gives us the following TLA tautology:

$$\Box\Diamond(\langle A \vee B \rangle_v) \equiv (\Box\Diamond\langle A \rangle_v) \vee (\Box\Diamond\langle B \rangle_v)$$

### 8.3 Weak Fairness

Using the temporal operators  $\Box$  and  $\Diamond$ , it's easy to specify liveness properties. For example, consider the hour-clock specification of module *HourClock* in Figure 2.1 on page 20. We can require that the clock never stops by asserting that there must be infinitely many *HCnext* steps. The obvious way to write this assertion is  $\Box\Diamond HCnext$ , but that's not a legal TLA formula because *HCnext* is an action, not a temporal formula. However, an *HCnext* step advances the value *hr* of the clock, so it changes *hr*. Therefore, an *HCnext* step is also an  $\langle HCnext \rangle_{hr}$  step. We can thus write the liveness property that the clock never stops as  $\Box\Diamond\langle HCnext \rangle_{hr}$ . So, we can take  $HC \wedge \Box\Diamond\langle HCnext \rangle_{hr}$  to be the specification of a clock that never stops.

Before continuing, I must make a confession and then lead you on a brief digression about subscripts. Let me first confess that the argument I just gave, that we can write  $\Box\Diamond\langle HCnext \rangle_{hr}$  in place of  $\Box\Diamond HCnext$ , was sloppy (a polite term for *wrong*). Not every *HCnext* step changes *hr*. Consider a state in which *hr* has some value that is not a number—perhaps a value  $\infty$ . An *HCnext* step that starts in such a state sets the new value of *hr* to  $\infty + 1$ . We don't know what  $\infty + 1$  equals; it might or might not equal  $\infty$ . If it does, then the *HCnext* step leaves *hr* unchanged, so it is not an  $\langle HCnext \rangle_{hr}$  step. Fortunately, states in which the value of *hr* is not a number are irrelevant. Because we are conjoining the liveness condition to the safety specification *HC*, we care only about behaviors that satisfy *HC*. In all such behaviors, *hr* is always a number, and every *HCnext* step is an  $\langle HCnext \rangle_{hr}$  step. Therefore,  $HC \wedge \Box\Diamond\langle HCnext \rangle_{hr}$  is equivalent to the nonTLA formula  $HC \wedge \Box\Diamond HCnext$ .<sup>3</sup>

When writing liveness properties, the syntax of TLA often forces us to write  $\langle A \rangle_v$  instead of *A*, for some action *A*. As in the case of *HCnext*, the safety specification usually implies that any *A* step changes some variable. To avoid having to think about which variables *A* actually changes, we generally take the subscript *v* to be the tuple of all variables, which is changed iff any variable changes. But what if *A* does allow stuttering steps? It's silly to assert that a stuttering step eventually occurs, since such an assertion is not invariant under stuttering. So, if *A* does allow stuttering steps, we want to require not that an

<sup>3</sup>Even though  $HC \wedge \Box\Diamond HCnext$  is not a TLA formula, its meaning has been defined, so we can determine whether it is equivalent to a TLA formula.

A step eventually occurs, but that a nonstuttering  $A$  step occurs—that is, an  $\langle A \rangle_v$  step, where  $v$  is the tuple of all the specification’s variables. The syntax of TLA forces us to say what we should mean.

I usually ignore the angle brackets and subscripts in informal discussions—for example, describing  $\Box \Diamond \langle HCnext \rangle_{hr}$  as the assertion that there are infinitely many  $HCnext$  steps. This finishes the digression; we now return to specifying liveness conditions.

Let’s modify specification *Spec* of module *Channel* (Figure 3.2 on page 30) to require that every value sent is eventually received. We do this by conjoining a liveness condition to *Spec*. The analog of the liveness condition for the clock is  $\Box \Diamond \langle Rcv \rangle_{chan}$ , which asserts that there are infinitely many  $Rcv$  steps. However, only a value that has been sent can be received, so this condition would also require that infinitely many values be sent—a requirement we don’t want to make. We want to permit behaviors in which no value is ever sent, so no value is ever received. We require only that, if a value is ever sent, then it is eventually received.

To assure that every value sent eventually is received, it suffices to require only that the next value to be received eventually is received. (When the next value has been received, the one after it becomes the next value to be received, which by the requirement must eventually be received, and so on.) More precisely, we need only require that it’s always the case that, if there is a value to be received, then the next value to be received eventually is received. The next value is received by a  $Rcv$  step, so the requirement is:<sup>4</sup>

$$\Box(\text{There is an unreceived value} \Rightarrow \Diamond \langle Rcv \rangle_{chan})$$

There is an unreceived value iff action  $Rcv$  is enabled, meaning that it is possible to take a  $Rcv$  step.  $TLA^+$  defines  $ENABLED A$  to be the predicate that is true iff action  $A$  is enabled. The liveness condition can then be written:

$$(8.6) \quad \Box(ENABLED \langle Rcv \rangle_{chan} \Rightarrow \Diamond \langle Rcv \rangle_{chan})$$

In the  $ENABLED$  formula, it doesn’t matter if we write  $Rcv$  or  $\langle Rcv \rangle_{chan}$ . We add the angle brackets so the two actions appearing in the formula are the same.

In any behavior satisfying the safety specification  $HC$ , it’s always possible to take an  $HCnext$  step that changes  $hr$ . Action  $\langle HCnext \rangle_{hr}$  is therefore always enabled, so  $ENABLED \langle HCnext \rangle_{hr}$  is true throughout such a behavior. Since  $TRUE \Rightarrow \Diamond \langle HCnext \rangle_{hr}$  is equivalent to  $\Diamond \langle HCnext \rangle_{hr}$ , we can replace the liveness condition  $\Box \Diamond \langle HCnext \rangle_{hr}$  for the hour clock with:

$$\Box(ENABLED \langle HCnext \rangle_{hr} \Rightarrow \Diamond \langle HCnext \rangle_{hr})$$

This suggests the following general liveness condition for an action  $A$ :

$$\Box(ENABLED \langle A \rangle_v \Rightarrow \Diamond \langle A \rangle_v)$$

---

<sup>4</sup> $\Box(F \Rightarrow \Diamond G)$  equals  $F \rightsquigarrow G$ , so we could write this formula more compactly with  $\rightsquigarrow$ . However, it’s more convenient to keep it in the form  $\Box(F \Rightarrow \Diamond G)$

This condition asserts that, if  $A$  ever becomes enabled, then an  $A$  step will eventually occur—even if  $A$  remains enabled for only a fraction of a nanosecond and is never again enabled. The obvious practical difficulty of implementing such a condition suggests that it's too strong. So, we replace it with the weaker formula  $\text{WF}_v(A)$ , defined to equal:

$$(8.7) \quad \Box(\Box\text{ENABLED } \langle A \rangle_v \Rightarrow \Diamond \langle A \rangle_v)$$

This formula asserts that if  $A$  ever becomes forever enabled, then an  $A$  step must eventually occur. WF stands for *Weak Fairness*, and the condition  $\text{WF}_v(A)$  is called *weak fairness on A*. We'll soon see that our liveness conditions for the clock and the channel can be written as WF formulas. But first, let's examine (8.7) and the following two formulas, which turn out to be equivalent to it:

$$(8.8) \quad \Box\Diamond(\neg\text{ENABLED } \langle A \rangle_v) \vee \Box\Diamond \langle A \rangle_v$$

$$(8.9) \quad \Diamond\Box(\text{ENABLED } \langle A \rangle_v) \Rightarrow \Box\Diamond \langle A \rangle_v$$

These three formulas can be expressed in English as:

(8.7) It's always the case that, if  $A$  is enabled forever, then an  $A$  step eventually occurs.

(8.8)  $A$  is infinitely often disabled, or infinitely many  $A$  steps occur.

(8.9) If  $A$  is eventually enabled forever, then infinitely many  $A$  steps occur.

The equivalence of these three formulas isn't obvious. Trying to deduce their equivalence from the English expressions often leads to confusion. The best way to avoid confusion is to use mathematics. We show that the three formulas are equivalent by proving that (8.7) is equivalent to (8.8) and that (8.8) is equivalent to (8.9). Instead of proving that they are equivalent for an individual behavior, we can use tautologies that we've already seen to prove their equivalence directly. Here's a proof that (8.7) is equivalent to (8.8). Studying it will help you learn to write liveness conditions.

$$\begin{aligned}
& \Box(\Box\text{ENABLED } \langle A \rangle_v \Rightarrow \Diamond \langle A \rangle_v) \\
& \equiv \Box(\neg\Box\text{ENABLED } \langle A \rangle_v \vee \Diamond \langle A \rangle_v) && \text{Because } (F \Rightarrow G) \equiv (\neg F \vee G). \\
& \equiv \Box(\Diamond\neg\text{ENABLED } \langle A \rangle_v \vee \Diamond \langle A \rangle_v) && \text{Because } \neg\Box F \equiv \Diamond\neg F. \\
& \equiv \Box\Diamond(\neg\text{ENABLED } \langle A \rangle_v \vee \langle A \rangle_v) && \text{Because } \Diamond(F \vee G) \equiv \Diamond F \vee \Diamond G. \\
& \equiv \Box\Diamond(\neg\text{ENABLED } \langle A \rangle_v) \vee \Box\Diamond \langle A \rangle_v && \text{Because } \Box\Diamond(F \vee G) \equiv \Box\Diamond F \vee \Box\Diamond G.
\end{aligned}$$

The equivalence of (8.8) and (8.9) is proved as follows.

$$\begin{aligned}
& \Box\Diamond(\neg\text{ENABLED } \langle A \rangle_v) \vee \Box\Diamond \langle A \rangle_v \\
& \equiv \neg\Diamond\Box\text{ENABLED } \langle A \rangle_v \vee \Box\Diamond \langle A \rangle_v && \text{Because } \Box\Diamond\neg F \equiv \Box\neg\Box F \equiv \neg\Diamond\Box F. \\
& \equiv \Diamond\Box\text{ENABLED } \langle A \rangle_v \Rightarrow \Box\Diamond \langle A \rangle_v && \text{Because } (F \Rightarrow G) \equiv (\neg F \vee G).
\end{aligned}$$

We now show that the liveness conditions for the hour clock and the channel can be written as weak fairness conditions.

First, consider the hour clock. In any behavior satisfying  $HC$ , an  $\langle HCnext \rangle_{hr}$  step is always enabled, so  $\Diamond \Box (\text{ENABLED } \langle HCnext \rangle_{hr})$  equals TRUE. Therefore,  $HC$  implies that  $WF_{hr}(HCnext)$ , which equals (8.9), is equivalent to  $\Box \Diamond \langle HCnext \rangle_{hr}$ , our liveness condition for the hour clock.

Now, consider the channel. I claim that the liveness condition (8.6) can be replaced by  $WF_{chan}(Rcv)$ . More precisely,  $Spec$  implies that these two formulas are equivalent, so conjoining either of them to  $Spec$  yields equivalent specifications. The proof rests on the observation that, in any behavior satisfying  $Spec$ , once  $Rcv$  becomes enabled (because a value has been sent), it can be disabled only by a  $Rcv$  step (which receives the value). In other words, it's always the case that if  $Rcv$  is enabled, then it is enabled forever or a  $Rcv$  step eventually occurs. Stated formally, this observation asserts that  $Spec$  implies

$$(8.10) \quad \Box (\text{ENABLED } \langle Rcv \rangle_{chan} \Rightarrow \Box (\text{ENABLED } \langle Rcv \rangle_{chan}) \vee \Diamond \langle Rcv \rangle_{chan})$$

We show that we can take  $WF_{chan}(Rcv)$  as our liveness condition by showing that (8.10) implies the equivalence of (8.6) and  $WF_{chan}(Rcv)$ .

The proof is by purely temporal reasoning; we need no other facts about the channel specification. Both for compactness and to emphasize the generality of our reasoning, let's replace  $\text{ENABLED } \langle Rcv \rangle_{chan}$  by  $E$  and  $\langle Rcv \rangle_{chan}$  by  $A$ . Using version (8.7) of the definition of  $WF$ , we must prove:

$$(8.11) \quad \Box (E \Rightarrow \Box E \vee \Diamond A) \Rightarrow (\Box (E \Rightarrow \Diamond A) \equiv \Box (\Box E \Rightarrow \Diamond A))$$

So far, all our proofs have been by calculation. That is, we have proved that two formulas are equivalent, or that a formula is equivalent to TRUE, by proving a chain of equivalences. That's a good way to prove simple things, but it's usually better to tackle a complicated formula like (8.11) by splitting its proof into pieces. We have to prove that one formula implies the equivalence of two others. The equivalence of two formulas can be proved by showing that each implies the other. More generally, to prove that  $P$  implies  $Q \wedge R$ , we prove that  $P \wedge Q$  implies  $R$  and that  $P \wedge R$  implies  $Q$ . So, we prove (8.11) by proving the two formulas:

$$(8.12) \quad \Box (E \Rightarrow \Box E \vee \Diamond A) \wedge \Box (E \Rightarrow \Diamond A) \Rightarrow \Box (\Box E \Rightarrow \Diamond A)$$

$$(8.13) \quad \Box (E \Rightarrow \Box E \vee \Diamond A) \wedge \Box (\Box E \Rightarrow \Diamond A) \Rightarrow \Box (E \Rightarrow \Diamond A)$$

Let's start with (8.13), which I find easier. In fact, we don't even need the first conjunct; we can prove that  $\Box (\Box E \Rightarrow \Diamond A)$  implies  $\Box (E \Rightarrow \Diamond A)$ . We know that  $\Box E$  implies  $E$ . Strengthening the hypothesis weakens an implication, so  $\Box E \Rightarrow \Diamond A$  implies  $E \Rightarrow \Diamond A$ , for any behavior. If  $F$  implies  $G$  for any behavior, then  $\Box F$  implies  $\Box G$ , so we can conclude that  $\Box (\Box E \Rightarrow \Diamond A)$  implies  $\Box (E \Rightarrow \Diamond A)$ .

Let's write a more formal proof of (8.13). To do this, we need the following temporal proof rule.

**Implies Generalization Rule** From  $F \Rightarrow G$  we can infer  $\Box F \Rightarrow \Box G$ , for any temporal formulas  $F$  and  $G$ .

Using it, we prove (8.13) with the following sequence of steps, where each step is a theorem that we prove en route to proving (8.13).

1.  $(P \Rightarrow Q) \wedge (Q \Rightarrow R) \Rightarrow (P \Rightarrow R)$ , for any formulas  $P$  and  $Q$ .  
PROOF: This is a tautology of propositional logic.
2.  $\Box((P \Rightarrow Q) \wedge (Q \Rightarrow R)) \Rightarrow \Box(P \Rightarrow R)$ , for any formulas  $P$  and  $Q$ .  
PROOF: By 1 and the Implies Generalization Rule, substituting  $(P \Rightarrow Q) \wedge (Q \Rightarrow R)$  for  $F$  and  $P \Rightarrow R$  for  $G$ .
3.  $\Box((\Box E \Rightarrow E) \wedge \Box(E \Rightarrow \Diamond A) \Rightarrow \Box(\Box E \Rightarrow \Diamond A))$   
PROOF: By 2, substituting  $\Box E$  for  $P$ ,  $E$  for  $Q$  and  $\Diamond A$  for  $R$ .
4.  $\Box(\Box E \Rightarrow E) \wedge \Box\Box(E \Rightarrow \Diamond A) \Rightarrow \Box(\Box E \Rightarrow \Diamond A)$   
PROOF: By applying the tautology  $\Box(F \wedge G) \equiv \Box F \wedge \Box G$  to the hypothesis of 3.
5.  $\text{TRUE} \wedge \Box(E \Rightarrow \Diamond A) \Rightarrow \Box(\Box E \Rightarrow \Diamond A)$   
PROOF: By 4, since  $\Box(\Box E \Rightarrow E)$  is a tautology (hence equivalent to TRUE), and  $\Box\Box(E \Rightarrow \Diamond A)$  is equivalent to  $\Box(E \Rightarrow \Diamond A)$  by the tautology  $\Box\Box F \equiv \Box F$ .
6. (8.13) holds  
PROOF: By propositional logic from 5.

We now prove (8.12). The basic idea is to prove the formula with the outer  $\Box$ s removed, and then apply the Implies Generalization Rule.

1.  $(E \Rightarrow \Box E \vee \Diamond A) \wedge (E \Rightarrow \Diamond A) \Rightarrow (\Box E \Rightarrow \Diamond A)$   
PROOF: This follows from the proposition-logic tautology  

$$(P \Rightarrow Q \vee R) \wedge (Q \Rightarrow R) \Rightarrow (P \Rightarrow R)$$
by substituting  $E$  for  $P$ ,  $\Box E$  for  $Q$  and  $\Diamond A$  for  $R$ .
2.  $\Box((E \Rightarrow \Box E \vee \Diamond A) \wedge (E \Rightarrow \Diamond A)) \Rightarrow \Box(\Box E \Rightarrow \Diamond A)$   
PROOF: From 1 and the Implies Generalization Rule.
3. (8.12) holds.  
PROOF: By 2, using the tautology  $\Box(F \wedge G) \equiv \Box F \wedge \Box G$  to distribute the first  $\Box$  across the  $\wedge$ .

This completes the proof that we can take  $\text{WF}_{\text{chan}}(\text{Rcv})$  instead of (8.7) as our liveness condition for the channel.

## 8.4 The Memory Specification

### 8.4.1 The Liveness Requirement

Let's now strengthen the memory specification with the liveness requirement that every request must receive a response. (We don't require that a request ever be issued.) The liveness requirement is conjoined to the internal memory specification, formula *ISpec* of module *InternalMemory* (Figures 5.2 and 5.3 on pages 52–53).

We want to express the liveness requirement in terms of weak fairness. To do this, we must understand when actions are enabled. The action  $Rsp(p)$  is enabled only if the action

$$(8.14) \text{ Reply}(p, buf[p], memInt, memInt')$$

is enabled. Recall that the operator *Reply* is a constant parameter, declared in the *MemoryInterface* module (Figure 5.1 on page 48). Without knowing more about this operator, we can't say when action (8.14) is enabled.

Let's assume that *Reply* actions are always enabled. That is, for any processor  $p$  and reply  $r$ , and any old value  $miOld$  of  $memInt$ , there is a new value  $miNew$  of  $memInt$  such that  $Repl(p, r, miOld, miNew)$  is true. For simplicity, we just assume that this is true for all  $p$  and  $r$ , and add the following assumption to the *MemoryInterface* module:

$$\text{ASSUME } \forall p, r, miOld : \exists miNew : \text{Reply}(p, r, miOld, miNew)$$

We should also make a similar assumption for *Send*, but we don't need it here.

We will subscript our weak-fairness formulas with the tuple of all variables, so let's give that tuple a name:

$$vars \triangleq \langle memInt, mem, ctl, buf \rangle$$

When processor  $p$  issues a request, it enables the  $Do(p)$  action, which remains enabled until a  $Do(p)$  step occurs. The weak-fairness condition  $WF_{vars}(Do(p))$  implies that this  $Do(p)$  step must eventually occur. A  $Do(p)$  step enables the  $Rsp(p)$  action, which remains enabled until a  $Rsp(p)$  step occurs. The weak-fairness condition  $WF_{vars}(Rsp(p))$  implies that this  $Rsp(p)$  step, which produces the desired response, must eventually occur. Hence, the requirement

$$(8.15) \text{ } WF_{vars}(Do(p)) \wedge WF_{vars}(Rsp(p))$$

assures that every request issued by processor  $p$  must eventually receive a reply. We want this condition to hold for every processor  $p$ , so we can take, as the liveness condition for the memory specification, the formula:

$$(8.16) \text{ Liveness} \triangleq \forall p \in Proc : WF_{vars}(Do(p)) \wedge WF_{vars}(Rsp(p))$$

The internal memory specification is then  $ISpec \wedge \text{Liveness}$ .

### 8.4.2 Another Way to Write It

I find a single fairness condition simpler than the conjunction of fairness conditions. Seeing the conjunction of the two weak fairness formulas in the definition of *Liveness* leads me to ask if it can be replaced by a single weak fairness condition on  $Do(p) \vee Rsp(p)$ . Such a replacement isn't always possible; in general, the formulas  $WF_v(A) \wedge WF_v(B)$  and  $WF_v(A \vee B)$  are not equivalent. However, in this case, we can replace the two fairness conditions with one. If we define

$$Liveness2 \triangleq \forall p \in Proc : WF_{vars}(Do(p) \vee Rsp(p))$$

then  $ISpec \wedge Liveness2$  is equivalent to  $ISpec \wedge Liveness$ . As we will see, this equivalence holds because any behavior satisfying *ISpec* satisfies the following two properties:

- DR1. Whenever  $Do(p)$  is enabled,  $Rsp(p)$  can never become enabled unless a  $Do(p)$  step eventually occurs.
- DR2. Whenever  $Rsp(p)$  is enabled,  $Do(p)$  can never become enabled unless a  $Rsp(p)$  step eventually occurs.

These properties are satisfied because a request to  $p$  is issued by a  $Req(p)$  step, executed by a  $Do(p)$  step, and responded to by a  $Rsp(p)$  step; and then, the next request to  $p$  can be issued by a  $Req(p)$  step. Each of these steps becomes possible (the action enabled) only after the preceding one occurs.

Let's now show that DR1 and DR2 imply that the conjunction of weak fairness of  $Do(p)$  and of  $Rsp(p)$  is equivalent to weak fairness of  $Do(p) \vee Rsp(p)$ . For compactness, and to emphasize the generality of what we're doing, let's replace  $Do(p)$ ,  $Rsp(p)$ , and  $vars$  by  $A$ ,  $B$ , and  $v$ , respectively.

First, we must restate DR1 and DR2 as temporal formulas. The basic form of DR1 and DR2 is:

Whenever  $F$  is true,  $G$  can never be true unless  $H$  is eventually true.

This is expressed in temporal logic as  $\Box(F \Rightarrow \Box \neg G \vee \Diamond H)$ . (The assertion " $P$  unless  $Q$ " just means  $P \vee Q$ .) Adding suitable subscripts, we can therefore write DR1 and DR2 in temporal logic as:

$$DR1 \triangleq \Box (\text{ENABLED } \langle A \rangle_v \Rightarrow \Box \neg \text{ENABLED } \langle B \rangle_v \vee \Diamond \langle A \rangle_v)$$

$$DR2 \triangleq \Box (\text{ENABLED } \langle B \rangle_v \Rightarrow \Box \neg \text{ENABLED } \langle A \rangle_v \vee \Diamond \langle B \rangle_v)$$

Our goal is to prove

$$(8.17) \quad DR1 \wedge DR2 \Rightarrow (WF_v(A) \wedge WF_v(B) \equiv WF_v(A \vee B))$$

This is complicated, so we split the proof into pieces. As in the proof of (8.11) in Section 8.3 above, we prove an equivalence by proving two implications. To prove (8.17), we prove the two theorems:

$$\begin{aligned} DR1 \wedge DR2 \wedge WF_v(A) \wedge WF_v(B) &\Rightarrow WF_v(A \vee B) \\ DR1 \wedge DR2 \wedge WF_v(A \vee B) &\Rightarrow WF_v(A) \wedge WF_v(B) \end{aligned}$$

We prove them by showing that they are true for an arbitrary behavior  $\sigma$ . In other words, we prove:

$$\begin{aligned} (8.18) \quad (\sigma \models DR1 \wedge DR2 \wedge WF_v(A) \wedge WF_v(B)) &\Rightarrow (\sigma \models WF_v(A \vee B)) \\ (8.19) \quad (\sigma \models DR1 \wedge DR2 \wedge WF_v(A \vee B)) &\Rightarrow (\sigma \models WF_v(A) \wedge WF_v(B)) \end{aligned}$$

These formulas seem daunting. Whenever you have trouble proving something, try a proof by contradiction; it gives you an extra hypothesis for free—namely, the negation of what you’re trying to prove. Proofs by contradiction are especially useful in temporal logic. To prove (8.18) and (8.19) by contradiction, we need to compute  $\neg(\sigma \models WF_v(C))$  for an action  $C$ . From the definition (8.7) of  $WF$ , we easily get

$$(8.20) \quad (\sigma \models WF_v(C)) \equiv \forall n \in Nat : (\sigma^{+n} \models \Box \text{ENABLED } \langle C \rangle_v) \Rightarrow (\sigma^{+n} \models \Diamond \langle C \rangle_v)$$

This and the tautology

$$\neg(\forall x \in S : P \Rightarrow Q) \equiv (\exists x \in S : P \wedge \neg Q)$$

of predicate logic yields:

$$(8.21) \quad \neg(\sigma \models WF_v(C)) \equiv \exists n \in Nat : (\sigma^{+n} \models \Box \text{ENABLED } \langle C \rangle_v) \wedge \neg(\sigma^{+n} \models \Diamond \langle C \rangle_v)$$

We also need two further results, both of which are derived from the tautology  $\langle A \vee B \rangle_v \equiv \langle A \rangle_v \vee \langle B \rangle_v$ . Combining this tautology with the temporal tautology  $\Diamond(F \vee G) \equiv \Diamond F \vee \Diamond G$  yields

$$(8.22) \quad \Diamond \langle A \vee B \rangle_v \equiv \Diamond \langle A \rangle_v \vee \Diamond \langle B \rangle_v$$

Combining it with the observation that an action  $C \vee D$  is enabled iff action  $C$  or action  $D$  is enabled yields

$$(8.23) \quad \text{ENABLED } \langle A \vee B \rangle_v \equiv \text{ENABLED } \langle A \rangle_v \vee \text{ENABLED } \langle B \rangle_v$$

We can now prove (8.18) and (8.19). To prove (8.18), we assume that  $\sigma$  satisfies  $DR1$ ,  $DR2$ ,  $WF_v(A)$ , and  $WF_v(B)$ , but it does not satisfy  $WF_v(A \vee B)$ , and we obtain a contradiction. By (8.21), the assumption that  $\sigma$  does not satisfy  $WF_v(A \vee B)$  means that there exists some number  $n$  such that:

$$(8.24) \quad \sigma^{+n} \models \Box \text{ENABLED } \langle A \vee B \rangle_v$$

$$(8.25) \quad \neg(\sigma^{+n} \models \Diamond \langle A \vee B \rangle_v)$$

We obtain a contradiction from (8.24) and (8.25) as follows.

1.  $\neg(\sigma^{+n} \models \Diamond \langle A \rangle_v) \wedge \neg(\sigma^{+n} \models \Diamond \langle B \rangle_v)$

PROOF: By (8.25) and (8.22), using the tautology  $\neg(P \vee Q) \equiv (\neg P \wedge \neg Q)$ .

2. (a)  $(\sigma^{+n} \models \text{ENABLED } \langle A \rangle_v) \Rightarrow (\sigma^{+n} \models \Box \neg \text{ENABLED } \langle B \rangle_v)$   
 (b)  $(\sigma^{+n} \models \text{ENABLED } \langle B \rangle_v) \Rightarrow (\sigma^{+n} \models \Box \neg \text{ENABLED } \langle A \rangle_v)$

PROOF: By definition of *DR1*, the assumption  $\sigma \models \text{DR1}$  implies

$$(\sigma^{+n} \models \text{ENABLED } \langle A \rangle_v) \Rightarrow (\sigma^{+n} \models \Box \neg \text{ENABLED } \langle B \rangle_v) \vee (\sigma^{+n} \models \Diamond \langle A \rangle_v)$$

and part (a) then follows from 1. The proof of (b) is similar.

3. (a)  $(\sigma^{+n} \models \text{ENABLED } \langle A \rangle_v) \Rightarrow (\sigma^{+n} \models \Box \text{ENABLED } \langle A \rangle_v)$   
 (b)  $(\sigma^{+n} \models \text{ENABLED } \langle B \rangle_v) \Rightarrow (\sigma^{+n} \models \Box \text{ENABLED } \langle B \rangle_v)$

PROOF: Part (a) follows from 2(a), (8.24), (8.23), and the temporal tautology

$$\Box(F \vee G) \wedge \Box \neg G \Rightarrow \Box F$$

The proof of part (b) is similar.

4. (a)  $(\sigma^{+n} \models \text{ENABLED } \langle A \rangle_v) \Rightarrow (\sigma^{+n} \models \Diamond \langle A \rangle_v)$   
 (b)  $(\sigma^{+n} \models \text{ENABLED } \langle B \rangle_v) \Rightarrow (\sigma^{+n} \models \Diamond \langle B \rangle_v)$

PROOF: The assumption  $\sigma \models \text{WF}_v(A)$  and (8.20) imply

$$(\sigma^{+n} \models \Box \text{ENABLED } \langle A \rangle_v) \Rightarrow (\sigma^{+n} \models \Diamond \langle A \rangle_v)$$

Part (a) follows from this and 3(a). The proof of part (b) is similar.

5.  $(\sigma^{+n} \models \Diamond \langle A \rangle_v) \vee (\sigma^{+n} \models \Diamond \langle B \rangle_v)$

PROOF: Since  $\Box F$  implies  $F$ , for any  $F$ , (8.24) and (8.23) imply

$$(\sigma^{+n} \models \text{ENABLED } \langle A \rangle_v) \vee (\sigma^{+n} \models \text{ENABLED } \langle B \rangle_v)$$

Step 5 then follows by propositional logic from step 4.

Steps 1 and 5 provide the required contradiction.

We can prove (8.19) by assuming that  $\sigma$  satisfies *DR1*, *DR2*, and  $\text{WF}_v(A \vee B)$ , and then proving that it satisfies  $\text{WF}_v(A)$  and  $\text{WF}_v(B)$ . We prove only that it satisfies  $\text{WF}_v(A)$ ; the proof for  $\text{WF}_v(B)$  is similar. The proof is by contradiction; we assume that  $\sigma$  does not satisfy  $\text{WF}_v(A)$  and obtain a contradiction. By (8.21), the assumption that  $\sigma$  does not satisfy  $\text{WF}_v(A)$  means that there exists some number  $n$  such that:

$$(8.26) \quad \sigma^{+n} \models \Box \text{ENABLED } \langle A \rangle_v$$

$$(8.27) \quad \neg(\sigma^{+n} \models \Diamond \langle A \rangle_v)$$

We obtain the contradiction as follows.

1.  $\sigma^{+n} \models \Diamond \langle A \vee B \rangle_v$

PROOF: From (8.26) and (8.23) we deduce  $\sigma^{+n} \models \Box \text{ENABLED } \langle A \vee B \rangle_v$ . By the assumption  $\sigma \models \text{WF}_v(A \vee B)$  and (8.20), this implies  $\sigma^{+n} \models \Diamond \langle A \vee B \rangle_v$ .

2.  $\sigma^{+n} \models \Box \neg \text{ENABLED } \langle B \rangle_v$

PROOF: From (8.26) we obtain  $\sigma^{+n} \models \text{ENABLED } \langle A \rangle_v$ , which by the assumption  $\sigma \models DR1$  and the definition of  $DR1$  implies

$$(\sigma^{+n} \models \Box \neg \text{ENABLED } \langle B \rangle_v) \vee (\sigma^{+n} \models \Diamond \langle A \rangle_v)$$

The assumption (8.27) then implies  $\sigma^{+n} \models \Box \neg \text{ENABLED } \langle B \rangle_v$ .

3.  $\neg(\sigma^{+n} \models \Diamond \langle B \rangle_v)$

PROOF: The definition of  $\text{ENABLED}$  implies  $\neg \text{ENABLED } \langle B \rangle_v \Rightarrow \neg \langle B \rangle_v$ . (A  $\langle B \rangle_v$  step can occur only when it is enabled.) From this, simple temporal reasoning implies

$$(\sigma^{+n} \models \Box \neg \text{ENABLED } \langle B \rangle_v) \Rightarrow \neg(\sigma^{+n} \models \Diamond \langle B \rangle_v)$$

(A formal proof uses the Implies Generalization Rule and the tautology  $\Box \neg F \equiv \neg \Diamond F$ .) We then deduce  $\neg(\sigma^{+n} \models \Diamond \langle B \rangle_v)$  from 2.

4.  $\neg \sigma^{+n} \models \Diamond \langle A \vee B \rangle_v$

PROOF: By (8.27), 3, and (8.22), using the tautology  $\neg P \wedge \neg Q \equiv \neg(P \vee Q)$ .

Steps 1 and 4 provide the necessary contradiction. This completes our proof of (8.19), which completes our proof of (8.17).

### 8.4.3 A Generalization

Formula (8.17) provides a rule for replacing the conjunction of weak fairness requirement on two actions with weak fairness of their disjunction. We now generalize it from two actions  $A$  and  $B$  to  $n$  actions  $A_1, \dots, A_n$ . The generalization of  $DR1$  and  $DR2$  is

$$DR(i, j) \triangleq \Box (\text{ENABLED } \langle A_i \rangle_v \Rightarrow \Box \neg \text{ENABLED } \langle A_j \rangle_v \vee \Diamond \langle A_i \rangle_v)$$

If we substitute  $A_1$  for  $A$  and  $A_2$  for  $B$ , then  $DR1$  becomes  $DR(1, 2)$  and  $DR2$  becomes  $DR(2, 1)$ . The generalization of (8.17) is:

$$(8.28) \quad (\forall i, j \in 1 \dots n : (i \neq j) \Rightarrow DR(i, j)) \Rightarrow$$

$$(\text{WF}_v(A_1) \wedge \dots \wedge \text{WF}_v(A_n) \equiv \text{WF}_v(A_1 \vee \dots \vee A_n))$$

To decide if you can replace the conjunction of weak fairness conditions by a single one in a specification, you will probably find it easier to use the following informal statement of (8.28).

**WF Conjunction Rule** If  $A_1, \dots, A_n$  are actions such that, for any distinct  $i$  and  $j$ , whenever  $\langle A_i \rangle_v$  is enabled,  $\langle A_j \rangle_v$  cannot become enabled unless an  $\langle A_i \rangle_v$  step occurs, then  $\text{WF}_v(A_1) \wedge \dots \wedge \text{WF}_v(A_n)$  is equivalent to  $\text{WF}_v(A_1 \vee \dots \vee A_n)$ .

Perhaps the best way to think of this rule is as an assertion about an arbitrary individual behavior  $\sigma$ . Its hypothesis is then that  $\sigma \models DR(i, j)$  holds for all distinct  $i$  and  $j$ ; its conclusion is:

$$\sigma \models (\text{WF}_v(A_1) \wedge \dots \wedge \text{WF}_v(A_n) \equiv \text{WF}_v(A_1 \vee \dots \vee A_n))$$

To replace  $\text{WF}_v(A_1) \wedge \dots \wedge \text{WF}_v(A_n)$  by  $\text{WF}_v(A_1 \vee \dots \vee A_n)$  in a specification, you have to check that any behavior satisfying the safety part of the specification also satisfies  $\text{DR}(i, j)$ , for all distinct  $i$  and  $j$ .

Conjunction and disjunction are special cases of quantification:

$$\begin{aligned} F_1 \vee \dots \vee F_n &\equiv \exists i \in 1 \dots n : F_i \\ F_1 \wedge \dots \wedge F_n &\equiv \forall i \in 1 \dots n : F_i \end{aligned}$$

We can therefore easily restate the WF Conjunction Rule as a condition on when  $\forall i \in S : \text{WF}_v(A_i)$  and  $\text{WF}_v(\exists i \in S : A_i)$  are equivalent, for a finite set  $S$ . The resulting rule is actually valid for any set  $S$ :

**WF Quantifier Rule** If, for all  $i \in S$ , the  $A_i$  are actions such that, for any distinct  $i$  and  $j$  in  $S$ , whenever  $\langle A_i \rangle_v$  is enabled,  $\langle A_j \rangle_v$  cannot become enabled unless an  $\langle A_i \rangle_v$  step occurs, then  $\forall i \in S : \text{WF}_v(A_i)$  is equivalent to  $\text{WF}_v(\exists i \in S : A_i)$ .

## 8.5 Strong Fairness

We define  $\text{SF}_v(A)$ , *strong fairness* of action  $A$ , to be either of the following two equivalent formulas.

$$(8.29) \quad \Diamond \Box (\neg \text{ENABLED } \langle A \rangle_v) \vee \Box \Diamond \langle A \rangle_v$$

$$(8.30) \quad \Box \Diamond \text{ENABLED } \langle A \rangle_v \Rightarrow \Box \Diamond \langle A \rangle_v$$

Intuitively, these two formulas assert:

(8.29)  $A$  is eventually disabled forever, or infinitely many  $A$  steps occur.

(8.30) If  $A$  is infinitely often enabled, then infinitely many  $A$  steps occur.

The proof that (8.29) and (8.30) are equivalent is similar to the proof on page 97 that the two formulations (8.8) and (8.9) of  $\text{WF}_v(A)$  are equivalent.

Definition (8.29) of  $\text{SF}_v(A)$  is obtained from definition (8.8) of  $\text{WF}_v(A)$  replacing  $\Box \Diamond (\neg \text{ENABLED } \langle A \rangle_v)$  with  $\Diamond \Box (\neg \text{ENABLED } \langle A \rangle_v)$ . Since  $\Diamond \Box F$  (eventually always  $F$ ) is stronger than (implies)  $\Box \Diamond F$  (infinitely always  $F$ ) for any formula  $F$ , strong fairness is stronger than weak fairness. We can express weak and strong fairness as follows.

- Weak fairness of  $A$  asserts that an  $A$  step must eventually occur if  $A$  is *continuously* enabled.
- Strong fairness of  $A$  asserts that an  $A$  step must eventually occur if  $A$  is *continually* enabled.

*Continuously* means without interruption. *Continually* means repeatedly, possibly with interruptions.

Strong fairness need not be strictly stronger than weak fairness. Weak and strong fairness of an action  $A$  are equivalent iff  $A$  infinitely often disabled implies that either  $A$  eventually becomes forever disabled, or infinitely many  $A$  steps occur. This is expressed formally by the tautology:

$$\begin{aligned} (\text{WF}_v(A) \equiv \text{SF}_v(A)) &\equiv \\ (\Box\Diamond(\neg \text{ENABLED } \langle A \rangle_v) \Rightarrow \Diamond\Box(\neg \text{ENABLED } \langle A \rangle_v) \vee \Box\Diamond\langle A \rangle_v) \end{aligned}$$

In the channel example, weak and strong fairness of  $Rcv$  are equivalent because  $Spec$  implies that, once enabled,  $Rcv$  can be disabled only by a  $Rcv$  step. Hence, if  $Rcv$  is disabled infinitely often, then it either eventually remains disabled forever, or else it is disabled infinitely often by  $Rcv$  steps.

The analogs of the WF Conjunction and WF Quantifier Rules (pages 104–105) hold for strong fairness—for example:

**SF Conjunction Rule** If  $A_1, \dots, A_n$  are actions such that, for any distinct  $i$  and  $j$ , whenever action  $A_i$  is enabled, action  $A_j$  cannot become enabled until an  $A_i$  step occurs, then  $\text{SF}_v(A_1) \wedge \dots \wedge \text{SF}_v(A_n)$  is equivalent to  $\text{SF}_v(A_1 \vee \dots \vee A_n)$ .

Strong fairness can be more difficult to implement than weak fairness, and it is a less common requirement. A strong fairness condition should be used in a specification only if it is needed. When strong and weak fairness are equivalent, the fairness property should be written as weak fairness.

Liveness properties can be subtle. Expressing them with *ad hoc* temporal formulas can lead to errors. We will specify liveness as the conjunction of weak and/or strong fairness properties whenever possible—and it almost always is possible. Having a uniform way of expressing liveness makes specifications easier to understand. Section 8.8.2 discusses an even more compelling reason for using fairness to specify liveness.

## 8.6 Liveness for the Write-Through Cache

Let's now add liveness to the write-through cache, specified in Figures 5.5–5.7 on pages 56–58. We want our specification to guarantee that every request eventually receives a response, without requiring that any requests are issued. This requires fairness on all the actions that make up the next-state action *Next* except for the following:

- A  $Req(p)$  action, which issues a request.
- An  $Evict(p, a)$  action, which evicts an address from the cache.

- A *MemQWr* action, if *memQ* contains only write requests and is not full (has fewer than *QLen* elements). Since a response to a write request can be issued before the value is written to memory, failing to execute a *MemQWr* action can prevent a response only if it prevents the dequeuing of a read operation in *memq* or the enqueueing of a read operation (because *memq* is full).

For simplicity, let's always require fairness for the *MemQWr* action too. We'll weaken this requirement later. Our liveness condition then has to assert fairness of the actions

$$MemQWr \quad MemQRd \quad Rsp(p) \quad RdMiss(p) \quad DoRd(p) \quad DoWr(p)$$

for all *p* in *Proc*. We now must decide whether to assert weak or strong fairness for these actions. Weak and strong fairness are equivalent for an action that, once enabled, remains enabled until it is executed. This is the case for all of these actions except *RdMiss(p)* and *DoWr(p)*. These two actions append a request to the *memQ* queue, and they are disabled if that queue is full. A *RdMiss(p)* or *DoWr(p)* could be enabled and then become disabled because a *RdMiss(q)* or *DoWr(q)*, for a different processor *q*, appends a request to *memQ*. We therefore need strong fairness for the *RdMiss(p)* and *DoWr(p)* actions. So, the fairness conditions we need are:

Weak Fairness    for *Rsp(p)*, *DoRd(p)*, *MemQWr*, and *MemQRd*

Strong Fairness    for *RdMiss(p)* and *DoWr(p)*.

As before, let's define *vars* to be the tuple of all variables.

$$vars \triangleq \langle memInt, mem, buf, ctl, cache, memQ \rangle$$

We could just write the liveness condition as

$$(8.31) \quad \wedge \forall p \in Proc : \wedge WF_{vars}(Rsp(p)) \wedge WF_{vars}(DoRd(p)) \\ \wedge SF_{vars}(RdMiss(p)) \wedge SF_{vars}(DoWr(p)) \\ \wedge WF_{vars}(MemQWr) \wedge WF_{vars}(MemQRd)$$

However, I prefer replacing the conjunction of fairness conditions by a single fairness condition on a disjunction, as we did in Section 8.4 for the memory specification. The WF and SF Conjunction Rules (page 104 and 106) imply that the liveness condition (8.31) can be rewritten as

$$(8.32) \quad \wedge \forall p \in Proc : \wedge WF_{vars}(Rsp(p) \vee DoRd(p)) \\ \wedge SF_{vars}(RdMiss(p) \vee DoWr(p)) \\ \wedge WF_{vars}(MemQWr \vee MemQRd)$$

We can now try to simplify (8.32) by applying the WF Quantifier Rule (page 105) to replace  $\wedge \forall p \in Proc : WF_{vars}(\dots)$  with  $WF_{vars}(\exists p \in Proc : \dots)$ . However, that

rule doesn't apply; it's possible for  $Rsp(p) \vee DoRd(p)$  to be enabled for two different processors  $p$  at the same time. In fact, the two liveness conditions are not equivalent. Weak fairness of  $\exists p \in Proc : Rsp(p) \vee DoRd(p)$  is satisfied by any behavior in which infinitely many  $Rsp(p)$  and  $DoRd(p)$  actions occur for some processor  $p$ . In such a behavior,  $Rsp(q)$  could be enabled for some other processor  $q$  without an  $Rsp(q)$  step ever occurring, making  $WF_{vars}(Rsp(p) \vee DoRd(p))$  false. Similarly, the analogous rule for strong fairness doesn't apply. Formula (8.32) is as simple as we can make it.

Let's return to the observation that we don't have to execute  $MemQWr$  if the  $memQ$  queue contains only write requests and is not full. Let's define  $QCond$  to be the assertion that  $memQ$  is full and contains a read request:

$$QCond \triangleq \wedge Len(memQ) = QLen \\ \wedge \exists i \in 1 \dots Len(memQ) : memQ[i][2].op = \text{"Rd"}$$

We have to eventually execute a  $MemQWr$  action only when it's enabled and  $QCond$  is true, which is the case iff the action  $QCond \wedge MemQWr$  is enabled. In this case, a  $MemQWr$  step is a  $QCond \wedge MemQWr$  step. Hence, it suffices to require weak fairness of the action  $QCond \wedge MemQWr$ . We can therefore replace the second conjunct of (8.32) with

$$WF_{vars}((QCond \wedge MemQWr) \vee MemQRd)$$

We would do this if we wanted the specification to describe the weakest liveness condition that implements the memory specification's liveness condition. However, if the specification were a description of an actual device, then that device would probably implement weak fairness on all  $MemQWr$  actions, so we would take (8.32) as the liveness condition.

## 8.7 Quantification

Section 8.1 describes the meaning of ordinary quantification of temporal formulas. For example, the meaning of the formula  $\forall r : F$ , for any temporal formula  $F$ , is defined by

$$\sigma \models (\forall r : F) \triangleq \forall r : (\sigma \models F)$$

where  $\sigma$  is any behavior.

The symbol  $r$  in  $\exists r : F$  is usually called a bound variable. But we've been using the term *variable* to mean something else—something that's declared by a VARIABLE statement in a module. The bound "variable"  $r$  is actually a constant in these formulas—a value that is the same in every state of the behavior.<sup>5</sup> For

<sup>5</sup>Logicians use the term *flexible variable* for a TLA variable, and the term *rigid variable* for a symbol like  $r$  that represents a constant.

example, the formula  $\exists r : \Box(x = r)$  asserts that  $x$  has the same value in every state of a behavior.

Bounded quantification over a constant set  $S$  is defined by:

$$\begin{aligned}\sigma \models (\forall r \in S : F) &\triangleq (\forall r \in S : \sigma \models F) \\ \sigma \models (\exists r \in S : F) &\triangleq (\exists r \in S : \sigma \models F)\end{aligned}$$

The symbol  $r$  is declared to be a constant in formula  $F$ . The expression  $S$  lies outside the scope of the declaration of  $r$ , so the symbol  $r$  cannot occur in  $S$ . It's easy to define the meanings of these formulas even if  $S$  is not a constant—for example, by letting  $\exists r \in S : F$  equal  $\exists r : (r \in S) \wedge F$ . However, for nonconstant  $S$ , it's better to write  $\exists r : (r \in S) \wedge F$  explicitly.

It's also easy to define the meaning of CHOOSE as a temporal operator. We can just let  $\sigma \models (\text{CHOOSE } r : F)$  be an arbitrary constant value  $r$  such that  $\sigma \models F$  equals TRUE, if such an  $r$  exists. However, a temporal CHOOSE operator is not needed for writing specifications, so  $\text{CHOOSE } r : F$  is not a legal TLA formula if  $F$  is a temporal formula.

We now come to the temporal existential quantifier  $\exists$ . In the formula  $\exists x : F$ , the symbol  $x$  is declared to be a variable in  $F$ . Unlike  $\exists r : F$ , which asserts the existence of a single value  $r$ , the formula  $\exists x : F$  asserts the existence of a value for  $x$  in each state of a behavior. For example, if  $y$  is a variable, then the formula  $\exists x : \Box(x \in y)$  asserts that  $y$  always has some element  $x$ , so  $y$  is always a nonempty set. However, the element  $x$  could be different in different states, so the values of  $y$  in different states could be disjoint.

We have been using  $\exists$  as a hiding operator, thinking of  $\exists x : F$  as  $F$  with variable  $x$  hidden. The precise definition of  $\exists$  is a bit tricky because, as discussed in Section 8.1, the formula  $\exists x : F$  should be invariant under stuttering. Intuitively,  $\exists x : F$  is satisfied by a behavior  $\sigma$  iff  $F$  is satisfied by a behavior  $\tau$  that is obtained from  $\sigma$  by adding and/or deleting stuttering steps and changing the value of  $x$ . A precise definition appears in Section 15.2.4 (page 298). However, for writing specifications, you can simply think of  $\exists x : F$  as  $F$  with  $x$  hidden.

TLA also has a temporal universal quantifier  $\forall$ , defined by:

$$\forall x : F \triangleq \neg \exists x : \neg F$$

This operator is hardly ever used. TLA<sup>+</sup> does not allow bounded versions of the operators  $\exists$  and  $\forall$ .

## 8.8 Temporal Logic Examined

### 8.8.1 A Review

Let's look at the shapes of the specifications that we've written so far. We started with the simple form

$$(8.33) \text{ } Init \wedge \Box[Next]_{vars}$$

where *Init* is the initial predicate, *Next* the next-state action, and *vars* the tuple of all variables. This kind of specification is, in principle, quite straightforward. We then introduced hiding, using  $\exists$  to bind variables that should not appear in the specification. Those bound variables, also called *hidden* or *internal* variables, serve only to help describe how the values of the free variables (also called visible variables) change.

Hiding variables is easy enough, and it is mathematically elegant and philosophically satisfying. However, in practice, it doesn't make much difference to a specification. A comment can also tell a reader that a variable should be regarded as internal. Explicit hiding allows implementation to mean implication. A lower-level specification that describes an implementation can be expected to imply a specification only if the specification's internal variables, whose values don't really matter, are explicitly hidden. Otherwise, implementation means implementation under a refinement mapping. (See Section 5.8.) However, as explained in Section 10.8, implementation often involves a refinement of the visible variables as well.

To express liveness, the specification (8.33) is strengthened to the form

$$(8.34) \text{ } Init \wedge \Box[Next]_{vars} \wedge Liveness$$

where *Liveness* is the conjunction of formulas of the form  $WF_{vars}(A)$  and/or  $SF_{vars}(A)$ , for actions *A*. (I'm considering universal quantification to be a form of conjunction.)

### 8.8.2 Machine Closure

In the specifications of the form (8.34) we've written so far, the actions *A* whose fairness properties appear in formula *Liveness* have one thing in common: they are all *subactions* of the next-state action *Next*. An action *A* is a subaction of *Next* iff every *A* step is a *Next* step. Equivalently, *A* is a subaction of *Next* iff *A* implies *Next*.<sup>6</sup> In almost all specifications of the form (8.34), formula *Liveness* should be the conjunction of weak and/or strong fairness formulas for subactions of *Next*. I'll now explain why.

<sup>6</sup>A weaker notion of subaction is that *A* is a subaction of the formula (8.33) iff, in every state of every behavior satisfying (8.33), if *A* is enabled then *Next*  $\wedge$  *A* is enabled.

When we look at the specification (8.34), we expect *Init* to constrain the initial state, *Next* to constrain what steps may occur, and *Liveness* to describe only what must eventually happen. However, consider the following formula

$$(8.35) \quad (x = 0) \wedge \Box[x' = x + 1]_x \wedge \text{WF}_x((x > 99) \wedge (x' = x - 1))$$

The first two conjuncts of (8.35) assert that  $x$  is initially 0 and that any step either increments  $x$  by 1 or leaves it unchanged. Hence, they imply that if  $x$  ever exceeds 99, then it forever remains greater than 99. The weak fairness property asserts that, if this happens, then  $x$  must eventually be decremented by 1—contradicting the second conjunct. Hence, (8.35) implies that  $x$  can never exceed 99, so it is equivalent to

$$(x = 0) \wedge \Box[(x < 99) \wedge (x' = x + 1)]_x$$

Conjoining the weak fairness property to the first two conjuncts of (8.35) forbids an  $x' = x + 1$  step when  $x = 99$ .

A specification of the form (8.34) is called *machine closed* iff the conjunct *Liveness* constrains neither the initial state nor what steps may occur. A more general way to express this is as follows. Let a finite behavior be a finite sequence of states.<sup>7</sup> We say that a finite behavior  $\sigma$  satisfies a safety property  $S$  iff the behavior obtained by adding infinitely many stuttering steps to the end of  $\sigma$  satisfies  $S$ . If  $S$  is a safety property, then we define the pair of formulas  $S, L$  to be machine closed iff every finite behavior that satisfies  $S$  can be extended to an infinite behavior that satisfies  $S \wedge L$ . We call (8.34) machine closed if the pair of formulas  $\text{Init} \wedge \Box[\text{Next}]_{\text{vars}}, \text{Liveness}$  is machine closed.

We seldom want to write a specification that isn't machine closed. If we do write one, it's usually by mistake. Specification (8.34) is guaranteed to be machine closed if *Liveness* is the conjunction of weak and/or strong fairness properties for subactions of *Next*.<sup>8</sup> This condition doesn't hold for specification (8.35), which is not machine closed, because  $(x > 99) \wedge (x' = x - 1)$  is not a subaction of  $x' = x + 1$ .

Liveness requirements are philosophically satisfying. A specification of the form (8.33), which specifies only a safety property, allows behaviors in which the system does nothing. Therefore, the specification is satisfied by a system that does nothing. Expressing liveness requirements with fairness properties is less satisfying. These properties are subtle and it's easy to get them wrong. It requires some thought to determine that the liveness condition for the write-through cache, formula (8.32) on page 107, does imply that every request receives a reply.

<sup>7</sup>A finite behavior therefore isn't a behavior, which is an infinite sequence of states. Mathematicians often abuse language in this way.

<sup>8</sup>More precisely, this is the case for a finite or countably infinite conjunction of properties of the form  $\text{WF}_v(A)$  and/or  $\text{SF}_v(A)$ , where each  $\langle A \rangle_v$  is a subaction of *Next*. This result also holds for the weaker definition of subaction in the footnote on the adjacent page.

It's tempting to express liveness properties more directly, without using fairness properties. For example, it's easy to write a temporal formula asserting for the write-through cache that every request receives a response. When processor  $p$  issues a request, it sets  $ctl[p]$  to “rdy”. We just have to assert that, for every processor  $p$ , whenever a state in which  $ctl[p] = \text{“rdy”}$  is true occurs, there will eventually be a  $Rsp(p)$  step:

$$(8.36) \quad \forall p \in Proc : \Box((ctl[p] = \text{“rdy”}) \Rightarrow \Diamond \langle Rsp(p) \rangle_{vars})$$

While such formulas are appealing, they are dangerous. It's very easy to make a mistake and write a specification that isn't machine closed.

Except in unusual circumstances, you should express liveness with fairness properties for subactions of the next-state action. These are the most straightforward specifications, and hence the easiest to write and to understand. Most system specifications, even if very detailed and complicated, can be written in this straightforward manner. The exceptions are usually in the realm of subtle, high-level specifications that attempt to be very general. An example of such a specification appears in Section 11.2.

### 8.8.3 Machine Closure and Possibility

Machine closure can be thought of as a possibility condition. For example, machine closure of the pair  $S, \Box \Diamond \langle A \rangle_v$  asserts that for every finite behavior  $\sigma$  satisfying  $S$ , it is possible to extend  $\sigma$  to an infinite behavior satisfying  $S$  in which infinitely many  $\langle A \rangle_v$  actions occur. If we regard  $S$  as a system specification, so a behavior that satisfies  $S$  represents a possible execution of the system, then we can restate machine closure of  $S, \Box \Diamond \langle A \rangle_v$  as follows: in any system execution, it is always possible for infinitely many  $\langle A \rangle_v$  actions to occur.

TLA specifications express safety and liveness properties, not possibility properties. A safety property asserts that something is impossible—for example, the system cannot take a step that doesn't satisfy the next-state action. A liveness property asserts that something must eventually happen. System requirements are sometimes stated informally in terms of what is possible. Most of the time, when examined rigorously, these requirements can be expressed with liveness and/or safety properties. (The most notable exceptions are statistical properties, such as assertions about the probability that something happens.) We are never interested in specifying that something *might* happen. It's never useful to know that the system *might* produce the right answer. We never have to specify that the user *might* type an “a”; we must specify what happens if he does.

Machine closure is a property of a pair of formulas, not of a system. Although a possibility property is never a useful assertion about a system, it can be a useful assertion about a specification. A specification  $S$  of a system with keyboard input should always allow the user to type an “a”. So, we want every finite

behavior satisfying  $S$  to be extendable to an infinite behavior satisfying  $S$  in which infinitely many “a”s are typed. If the action  $\langle A \rangle_v$  represents the typing of an “a”, then saying that the user should always be able to type infinitely many “a”s is equivalent to saying that the pair  $S, \Box \Diamond \langle A \rangle_v$  should be machine closed. If  $S, \Box \Diamond \langle A \rangle_v$  isn’t machine closed, then it could become impossible for the user ever to type an “a”. Unless the system is allowed to lock the keyboard, this means that there is something wrong with our specification.

This kind of possibility property can be proved. For example, to prove that it’s always possible for the user to type infinitely many “a”s, we show that conjoining suitable fairness conditions on the input actions implies that the user *must* type infinitely many “a”s. However, proofs of this kind of simple property don’t seem to be worth the effort. When writing a specification, you should make sure that possibilities allowed by the real system are allowed by the specification. Once you are aware of what should be possible, you will usually have little trouble ensuring that the specification makes it possible. You should also make sure that what the system *must* do is implied by the specification’s fairness conditions. This can be more difficult.

### 8.8.4 The Unimportance of Liveness

While philosophically important, in practice the liveness property of (8.34) is not as important as the safety part,  $Init \wedge \Box[Next]_{vars}$ . The ultimate purpose of writing a specification is to avoid errors. Experience shows that most of the benefit from writing and using a specification comes from the safety part. On the other hand, the liveness property is usually easy enough to write. It typically constitutes less than five percent of a specification. So, you might as well write the liveness part. However, when looking for errors, most of your effort should be devoted to examining the safety part.

### 8.8.5 Temporal Logic Considered Confusing

The most general type of specification I’ve discussed so far has the form

$$(8.37) \quad \exists v_1, \dots, v_n : Init \wedge \Box[Next]_{vars} \wedge Liveness$$

where *Liveness* is the conjunction of fairness properties of subactions of *Next*. This is a very restricted class of temporal-logic formulas. Temporal logic is quite expressive, and one can combine its operators in all sorts of ways to express a wide variety of properties. This suggests the following approach to writing a specification: express each property that the system must satisfy with a temporal formula, and then conjoin all these formulas. For example, formula (8.36) above expresses the property of the write-through cache that every request eventually receives a response.

This approach is philosophically appealing. It has just one problem: it's practical for only the very simplest of specifications—and even for them, it seldom works well. The unbridled use of temporal logic produces formulas that are hard to understand. Conjoining several of these formulas produces a specification that is impossible to understand.

The basic form of a TLA specification is (8.37). Most specifications should have this form. We can also use this kind of specification as a building block. Chapters 9 and 10 describe situations in which we write a specification as a conjunction of such formulas. Section 10.7 introduces an additional temporal operator  $\overset{\pm}{\triangleright}$  and explains why we might want to write a specification  $F \overset{\pm}{\triangleright} G$ , where  $F$  and  $G$  have the form (8.34). But such specifications are of limited practical use. Most engineers need only know how to write specifications of the form (8.37). Indeed, they can get along quite well with specifications of the form (8.33) that express only safety properties and don't hide any variables.

## Chapter 9

# Real Time

### 9.1 The Hour Clock Revisited

Let's return to our specification of the simple hour clock in Chapter 2, which asserts that the variable *hr* cycles through the values 1 through 12. We now add the obvious requirement that the clock keep correct time. In science, one represents the real-time behavior of a system by introducing a variable, traditionally *t*, whose value is a real number that represents time. A state in which  $t = -17.51$  represents a state of the system at time  $-17.51$ , perhaps measured in seconds elapsed since 00:00 UT on 1 January 2000. In TLA<sup>+</sup> specifications, I prefer to use the variable *now* rather than *t*. For linguistic convenience, I will usually assume that the unit of time is the second, though we could just as well choose any other unit.

Remember that a state is an assignment of values to all relevant variables.

Unlike sciences such as physics and chemistry, computer science studies systems whose behavior can be described by a sequence of discrete states, rather than by states that vary continuously with time. We consider the hour clock's display to change directly from reading 12 to reading 1, and ignore the continuum of intermediate states that occur in the physical display. This means that we pretend that the change is instantaneous. So, a real-time specification of the clock might allow the step

$$\begin{bmatrix} hr & = & 12 \\ now & = & \sqrt{2.47} \end{bmatrix} \rightarrow \begin{bmatrix} hr & = & 1 \\ now & = & \sqrt{2.47} \end{bmatrix}$$

The value of *now* advances between changes to *hr*. If we wanted to specify how long it takes the display to change from 12 to 1, we would have to introduce an intermediate state that represents a changing display—perhaps by letting *hr* assume some intermediate value such as 12.5, or by adding a Boolean-valued variable *chg* whose value indicates whether the display is changing. We won't

do this, but will be content to specify an hour clock in which we consider the display to change instantaneously.

The value of *now* changes between changes to *hr*. Just as we represent a continuously varying clock display by a variable whose value changes in discrete steps, we let the value of *now* change in discrete steps. A behavior in which *now* increases in femtosecond increments would be an accurate enough description of continuously changing time for our specification of the hour clock. In fact, there's no need to specify any particular granularity of time; we can let *now* advance by arbitrary amounts between clock ticks. (Since the value of *hr* is unchanged by steps that change *now*, the requirement that the clock keep correct time will rule out behaviors in which *now* changes by too much in a single step.)

What real-time condition do we want to require for our hour clock? We might require that it always display the correct time. A more realistic requirement would be that it display the time correctly to within  $\rho$  seconds, for some real number  $\rho$ . However, this is not typical of the real-time requirements that arise in actual systems. Instead, we require that the clock tick approximately once per hour. More precisely, we require that the interval between ticks be one hour plus or minus  $\rho$  seconds, for some positive number  $\rho$ . Of course, this requirement allows the time displayed by the clock eventually to drift away from the actual time. But that's what most real clocks will do if not reset.

We could start our specification of the real-time clock from scratch. However, we still want the hour clock to satisfy the specification *HC* of module *HourClock* (Figure 2.1 on page 20). We just want to add an additional real-time requirement. So, we will write the specification as the conjunction of *HC* and a formula requiring that the clock tick every hour, plus or minus  $\rho$  seconds. This requirement is the conjunction of two separate requirements: that the clock tick at most once every  $3600 - \rho$  seconds, and at least once every  $3600 + \rho$  seconds.

To specify these requirements, we must introduce a variable that records how much time has elapsed since the last clock tick. Let's call it *t* for *timer*. The value of *t* is set to 0 by a step that represents a clock tick—namely, by an *HCnext* step. Any step that represents the passing of *s* seconds should advance *t* by *s*. A step represents the passing of time iff it changes *now*, and such step represents the passage of  $now' - now$  seconds. So, the change to the timer *t* is described by the action:

$$TNext \triangleq t' = \text{IF } HCnext \text{ THEN } 0 \text{ ELSE } t + (now' - now)$$

The specification of how *t* changes is a formula asserting that every step is a *TNext* step or else leaves *t* and *now* unchanged. Letting *t* initially equal 0 (so we start in a state in which the clock has just ticked), this formula is:

$$Timer \triangleq (t = 0) \wedge \Box [TNext]_{\langle t, hr \rangle}$$

The requirement that the clock tick at least once every  $3600 + \rho$  seconds means that it's always the case that at most  $3600 + \rho$  seconds have elapsed since the

last *HCnext* step. Since  $t$  always equals the elapsed time since the last *HCnext* step, this requirement is expressed by the formula:

$$MaxTime \triangleq \Box(t \leq 3600 + \rho)$$

(Since we can't measure time with perfect accuracy, it doesn't matter whether we use  $<$  or  $\leq$ . When we generalize from this example, it will be somewhat more convenient to use  $\leq$ .)

The requirement that the clock tick at most once every  $3600 - \rho$  seconds means that, whenever an *HCnext* step occurs, at least  $3600 - \rho$  seconds have elapsed since the previous *HCnext* step. This suggests the condition

$$(9.1) \quad \Box(HCnext \Rightarrow (t \geq 3600 - \rho))$$

However, (9.1) isn't a legal TLA formula because  $HCnext \Rightarrow \dots$  is an action (a formula containing primes), and a TLA formula asserting that an action is always true must have the form  $\Box[A]_v$ . We don't care about steps that leave  $hr$  unchanged, so we can replace (9.1) by the TLA formula:

$$MinTime \triangleq \Box[HCnext \Rightarrow (t \geq 3600 - \rho)]_{hr}$$

The desired real-time constraint on the clock is expressed by the conjunction of these three formulas:

$$HCTime \triangleq Timer \wedge MaxTime \wedge MinTime$$

However, formula *HCTime* contains the variable  $t$ , and the specification of the real-time clock should describe only the changes to  $hr$  (the clock display) and  $now$  (the time). So, we have to hide  $t$ . Hiding is expressed in TLA<sup>+</sup> by the temporal existential quantifier  $\exists$ , introduced in Section 4.3 (page 41). However, as explained in that section, we can't simply write  $\exists t : HCTime$ . We must define *HCTime* in a module that declares  $t$ , and then use a parametrized instantiation of that module. This is done in Figure 9.1 on page 119. Instead of defining *HCTime* in a completely separate module, I have defined it in a submodule named *Inner* of the module *RealTimeHourClock* containing the specification of the real-time hour clock. Note that all the symbols declared and defined in the main module up to that point can be used in the submodule. Submodule *Inner* is instantiated in the main module with the statement

$$I(t) \triangleq \text{INSTANCE } Inner$$

The  $t$  in *HCTime* can then be hidden by writing  $\exists t : I(t)!HCTime$ .

The formula  $HC \wedge (\exists t : I(t)!HCTime)$  describes the possible changes to the value of  $hr$ , and relates those changes to the value of  $now$ . But it says very little about how the value of  $now$  can change. For example, it allows the following behavior:

$$\begin{bmatrix} hr & = & 11 \\ now & = & 23.5 \end{bmatrix} \rightarrow \begin{bmatrix} hr & = & 11 \\ now & = & 23.4 \end{bmatrix} \rightarrow \begin{bmatrix} hr & = & 11 \\ now & = & 23.5 \end{bmatrix} \rightarrow \begin{bmatrix} hr & = & 11 \\ now & = & 23.4 \end{bmatrix} \rightarrow \dots$$

In the generalization,  $\geq$  will be more convenient than  $>$ .

Because time can't go backwards, such a behavior doesn't represent a physical possibility. Everyone knows that time only increases, so there's no need to forbid this behavior if the only purpose of our specification is to describe the hour clock. However, a specification should also allow us to reason about a system. If the clock ticks approximately once per hour, then it can't stop. However, as the behavior above shows, the formula  $HC \wedge (\exists t : I(t)!HCTime)$  by itself allows the clock to stop. To infer that it can't, we also need to state how *now* changes.

We define a formula  $RTnow$  that specifies the possible changes to *now*. This formula does not specify the granularity of the changes to *now*; it allows *now* to advance by a microsecond or by a century. However, we have decided that a step that changes *hr* should leave *now* unchanged, so a step that changes *now* should leave *hr* unchanged. Therefore, steps that change *now* are described by the action:

$$NowNext \triangleq \wedge now' \in \{r \in Real : r > now\} \\ \wedge UNCHANGED \ hr$$

Formula  $RTnow$  should also allow steps that leave *now* unchanged. The initial value of *now* is arbitrary (we can start the system at any time), so the safety part of  $RTnow$  is:

$$(now \in Real) \wedge \square[NowNext]_{now}$$

The liveness condition we want is that *now* should increase without bound. Simple weak fairness of the  $NowNext$  action isn't good enough, because it allows "Zeno" behaviors such as:

$$[now = .9] \rightarrow [now = .99] \rightarrow [now = .999] \rightarrow [now = .9999] \rightarrow \dots$$

in which the value of *now* remains bounded. Weak fairness of the action  $NowNext \wedge (now' > r)$  implies that eventually a  $NowNext$  step will occur in which the new value of *now* is greater than *r*. (This action is always enabled, so weak fairness implies that infinitely many such actions must occur.) Asserting this for all real numbers *r* implies that *now* grows without bound, so we take as the fairness condition:

$$\forall r \in Real : WF_{now}(NowNext \wedge (now' > r))$$

The complete specification of the real-time hour clock, including the definition of formula  $RTnow$ , appears in Figure 9.1 on the next page.

Fairness is discussed in Chapter 8.

## 9.2 Real-Time Specifications in General

In Section 8.3 (page 95), we saw that the appropriate generalization of the liveness requirement that the hour clock tick infinitely often is weak fairness of the

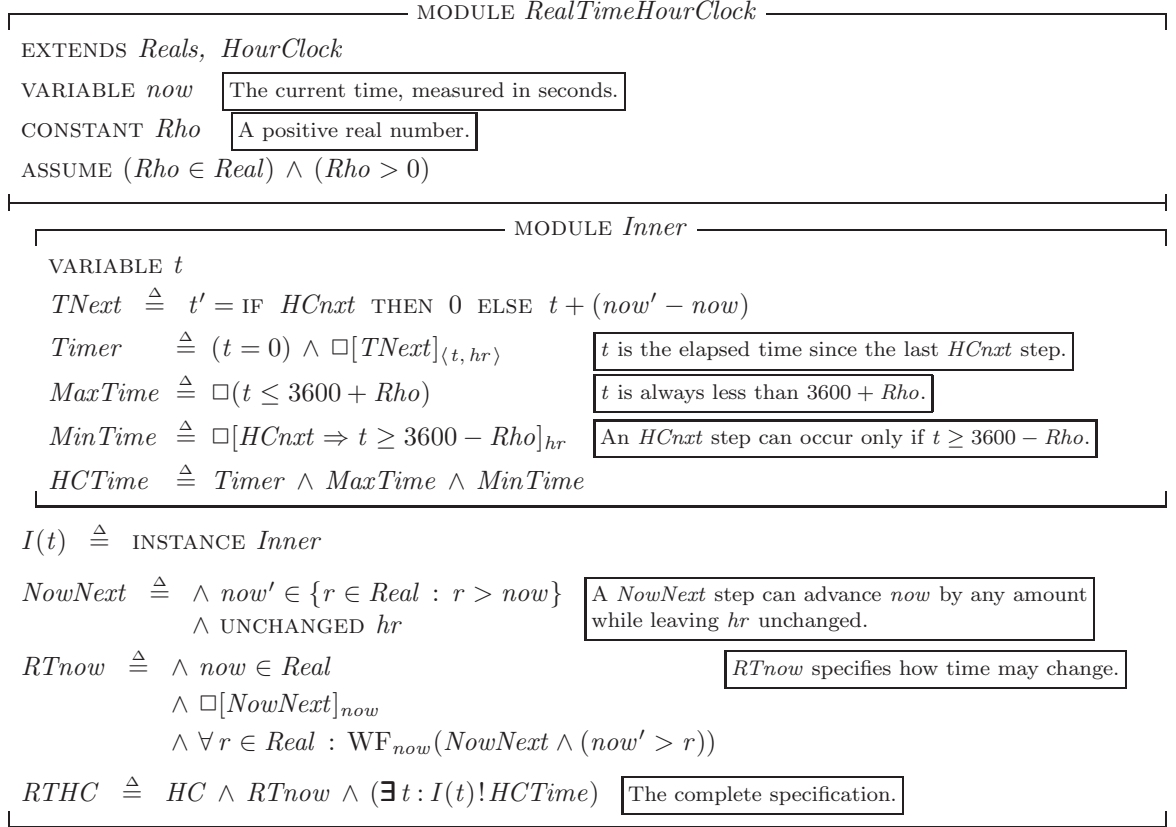


Figure 9.1: The real-time specification of an hour clock that ticks every hour, plus or minus *Rho* seconds.

clock-tick action. There is a similar generalization for real-time specifications. Weak fairness of an action *A* asserts that if *A* is continuously enabled, then an *A* step must eventually occur. The real-time analog is that if *A* is continuously enabled for  $\epsilon$  seconds, then an *A* step must occur. Since an *HCnext* action is always enabled, the requirement that the clock tick at least once every  $3600 + \rho$  seconds is can be expressed in this way by letting *A* be *HCnext* and  $\epsilon$  be  $3600 + \rho$ .

The requirement that an *HCnext* action occur at most once every  $3600 - \rho$  seconds can be similarly generalized to the condition that an action *A* must be continuously enabled for at least  $\delta$  seconds before an *A* step can occur.

The first condition, the upper bound  $\epsilon$  on how long *A* can be enabled without an *A* step occurring, is vacuously satisfied if  $\epsilon$  equals *Infinity*—a value defined in the *Reals* module to be greater than any real number. The second condition, the lower bound  $\delta$  on how long *A* must be enabled before an *A* step can occur, is

vacuously satisfied if  $\delta$  equals 0. So, nothing is lost by combining both of these conditions into a single formula containing  $\delta$  and  $\epsilon$  as parameters. I now define such a formula, which I call a *real-time bound condition*.

The weak-fairness formula  $WF_v(A)$  actually asserts weak fairness of the action  $\langle A \rangle_v$ , which equals  $A \wedge (v' \neq v)$ . The subscript  $v$  is needed to rule out stuttering steps. Since the truth of a meaningful formula can't depend on whether or not there are stuttering steps, it makes no sense to say that an  $A$  step did or did not occur if that step could be a stuttering step. For this reason, the corresponding real-time condition must also be a condition on an action  $\langle A \rangle_v$ , not on an arbitrary action  $A$ . In most cases of interest,  $v$  is the tuple of all variables that occur in  $A$ . I therefore define the real-time bound formula  $RTBound(A, v, \delta, \epsilon)$  to assert that:

- An  $\langle A \rangle_v$  step cannot occur until  $\langle A \rangle_v$  has been continuously enabled for at least  $\delta$  time units since the last  $\langle A \rangle_v$  step—or since the beginning of the behavior.
- $\langle A \rangle_v$  can be continuously enabled for at most  $\epsilon$  time units before an  $\langle A \rangle_v$  step occurs.

$RTBound(A, v, \delta, \epsilon)$  generalizes the formula  $\exists t : I(t)!HCTime$  of the real-time hour clock specification and it can be defined in the same way, using a submodule. However, the definition can be structured a little more compactly as:

$$RTBound(A, v, D, E) \triangleq \text{LET } Timer(t) \triangleq \dots \\ \dots \\ \text{IN } \exists t : Timer(t) \wedge \dots$$

For the TLA<sup>+</sup> specification, I have replaced  $\delta$  and  $\epsilon$  by  $D$  and  $E$ .

We first define  $Timer(t)$  to be a temporal formula asserting that  $t$  always equals the length of time that  $\langle A \rangle_v$  has been continuously enabled since the last  $\langle A \rangle_v$  step. The value of  $t$  should be set to 0 by an  $\langle A \rangle_v$  step or a step that disables  $\langle A \rangle_v$ . A step that advances *now* should increment  $t$  by  $now' - now$  iff  $\langle A \rangle_v$  is enabled. Changes to  $t$  are therefore described by the action:

$$TNext(t) \triangleq t' = \text{IF } \langle A \rangle_v \vee \neg(\text{ENABLED } \langle A \rangle_v)' \text{ THEN } 0 \\ \text{ELSE } t + (now' - now)$$

We are interested in the meaning of  $Timer(t)$  only when  $v$  is a tuple whose components include all the variables that may appear in  $A$ . In this case, a step that leaves  $v$  unchanged cannot enable or disable  $\langle A \rangle_v$ . So, the formula  $Timer(t)$  should allow steps that leave  $t$ ,  $v$ , and *now* unchanged. Letting the initial value of  $t$  be 0, we define:

$$Timer(t) \triangleq (t = 0) \wedge \Box [TNext(t)]_{\langle t, v, now \rangle}$$

Formulas *MaxTime* and *MinTime* of the real-time hour clock's specification have the obvious generalizations:

- $MaxTime(t)$  asserts that  $t$  is always less than or equal to  $E$ :

$$MaxTime(t) \triangleq \Box(t \leq E)$$

- $MinTime(t)$  asserts that an  $\langle A \rangle_v$  step can occur only if  $t \geq D$ :

$$MinTime(t) \triangleq \Box[A \Rightarrow (t \geq D)]_v$$

(A little propositional logic reasoning shows that  $[A \Rightarrow (t \geq D)]_v$  is equivalent to  $[\langle A \rangle_v \Rightarrow (t \geq D)]_v$ .)

We can then define  $RTBound(A, v, D, E)$  to equal

$$\exists t : Timer(t) \wedge MaxTime(t) \wedge MinTime(t)$$

We must also generalize formula  $RTnow$  of the real-time hour clock's specification. That formula describes how *now* changes, and it asserts that *hr* remains unchanged when *now* changes. The generalization is the formula  $RTnow(v)$ , which replaces *hr* with an arbitrary state function  $v$  that will usually be the tuple of all variables (other than *now*) appearing in the specification. Using these definitions, the specification  $RTHC$  of the real-time hour clock can be written:

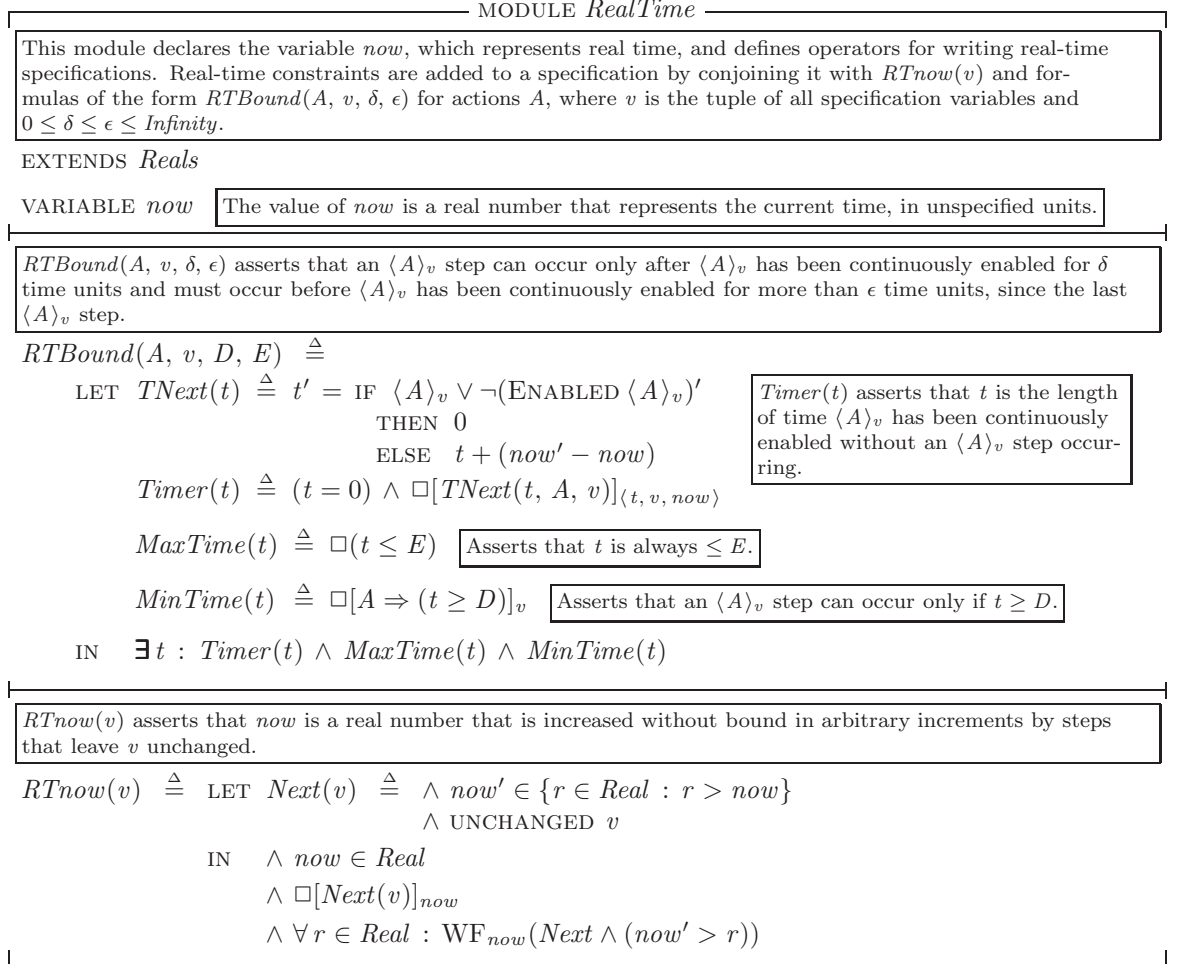
$$HC \wedge RTnow(hr) \wedge RTBound(HCnext, hr, 3600 - Rho, 3600 + Rho)$$

The *RealTime* module, with its definitions of  $RTBound$  and  $RTnow$ , appears in Figure 9.2 on the next page.

Strong fairness strengthens weak fairness by requiring an  $A$  step not only if action  $A$  is continuously enabled, but if it is repeatedly enabled. Being repeatedly enabled includes the possibility that it is also repeatedly disabled. We can similarly strengthen our real-time bound conditions by defining a stronger formula  $SRTBound(A, v, \delta, \epsilon)$  to assert that:

- An  $\langle A \rangle_v$  step cannot occur until  $\langle A \rangle_v$  has been enabled for a total of at least  $D$  time units since the last  $\langle A \rangle_v$  step—or since the beginning of the behavior.
- $\langle A \rangle_v$  can be enabled for a total of at most  $\epsilon$  time units before an  $\langle A \rangle_v$  step occurs.

If  $\epsilon < Infinity$ , then  $RTBound(A, v, \delta, \epsilon)$  and  $RTnow(v)$  imply  $WF_v(A)$ , since they imply that if  $\langle A \rangle_v$  is continuously enabled, then it must eventually be executed within  $\epsilon$  seconds. However,  $SRTBound(A, v, \delta, \epsilon)$  and  $RTnow(v)$  do not similarly imply  $SF_v(A)$ . They allow behaviors in which  $\langle A \rangle_v$  is enabled infinitely often but never executed—for example, it can be enabled for  $\epsilon/2$  seconds, then for  $\epsilon/4$  seconds, then for  $\epsilon/8$  seconds, and so on. For this reason,  $SRTBound$  does not seem to be of much practical use, so I won't bother defining it formally.

Figure 9.2: The *RealTime* module for writing real-time specifications.

### 9.3 The Real-Time Write-Through Cache

I now illustrate the use of the *RealTime* module for writing real-time specifications by writing real-time versions of the linearizable memory specification of Section 5.3 (page 50) and the write-through cache specification of Section 5.6 (page 54).

The real-time memory specification is obtained by strengthening the specification in Module *Memory* (Figure 5.4 on page 53) to require that the memory responds to a processor's requests within *Rho* seconds. The complete memory specification *Spec* of module *Memory* was obtained by hiding the variables *mem*,

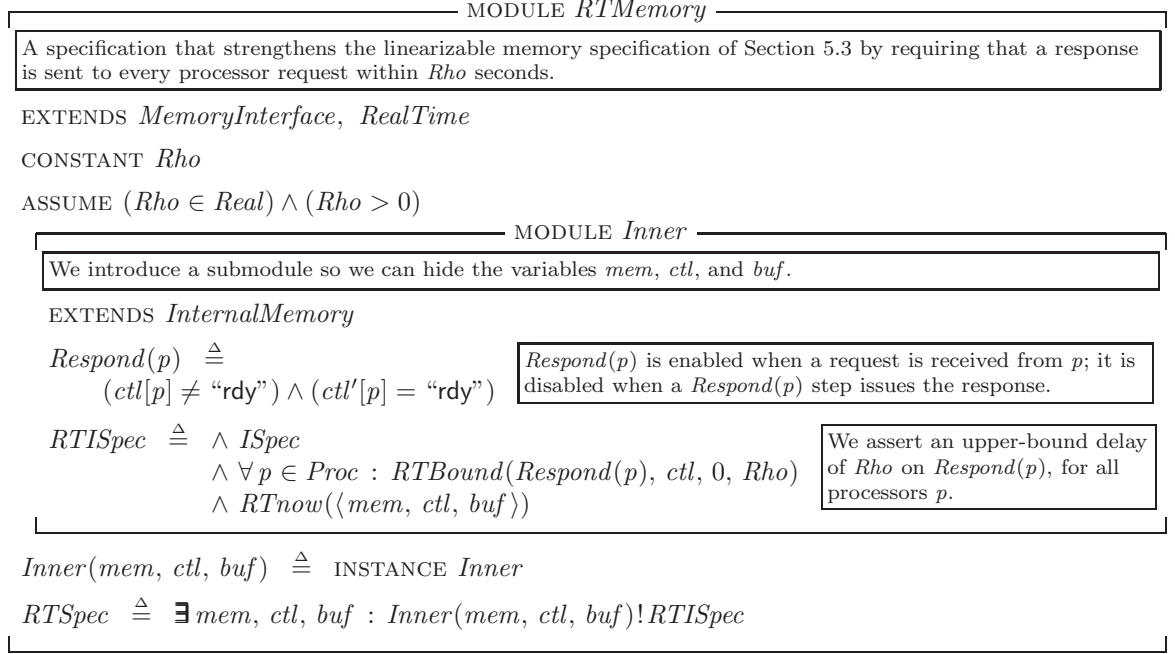


Figure 9.3: A real-time version of the linearizable memory specification.

*ctl*, and *buf* in the internal specification *ISpec* of module *InternalMemory*. It's usually easier to add a real-time constraint to an internal specification, where the constraints can mention the internal (hidden) variables. So, we first add the timing constraint to *ISpec* and then hide the internal variables.

To specify that the system must respond to a processor request within *Rho* seconds, we add an upper-bound timing constraint for an action that becomes enabled when a request is issued, and becomes disabled (possibly being executed) only when the processor responds to the request. In specification *ISpec*, responding to a request requires two actions—*Do*(*p*) to perform the operation internally, and *Rsp*(*p*) to issue the response. Neither of these actions is the one we want; we have to define a new action for the purpose. There is a pending request for processor *p* iff *ctl*[*p*] equals "rdy". So, we assert that the following action cannot be enabled for more than *Rho* seconds without being executed:

$$Respond(p) \triangleq (ctl[p] \neq \text{"rdy"}) \wedge (ctl'[p] = \text{"rdy"})$$

The complete specification is formula *RTSpec* of Module *RTMemory* in Figure 9.3 on this page. To permit the hiding of variables *mem*, *ctl*, and *buf*, Module *RTMemory* contains a submodule *Inner* that extends module *InternalMemory*.

Having added a real-time constraint to the specification of a linearizable memory, let's strengthen the specification of the write-through cache so it satisfies that constraint. The object is not just to add any real-time constraint that does the job—that's easy to do by using the same constraint that we added to the memory specification. We want to write a specification of a real-time algorithm—a specification that tells an implementor how to meet the real-time constraints. This is generally done by placing real-time bounds on the actual actions of the untimed specification, not by adding time bounds on a new action, as we did for the memory specification. An upper-bound constraint on the response time should be achieved by enforcing upper-bound constraints on the system's actions.

If we try to achieve a bound on response time by adding real-time bounds to the write-through cache specification's actions, we encounter the following problem. Operations by different processors “compete” with one another to enqueue operations on the finite queue *memQ*. For example, when servicing a write request for processor *p*, the system must execute a *DoWr(p)* action to enqueue the operation to the tail of *memQ*. That action is not enabled if *memQ* is full. The *DoWr(p)* action can be continually disabled by the system performing *DoWr* or *RdMiss* actions for other processors. That's why, to guarantee liveness—that each request eventually receives a response—in Section 8.6 (page 106) we had to assert strong fairness of *DoWr* and *RdMiss* actions. The only way to ensure that a *DoWr(p)* action is executed within some length of time is to use lower-bound constraints on the actions of other processors to ensure that they cannot perform *DoWr* or *RdMiss* actions too frequently. Although such a specification is possible, it is not the kind of approach anyone is likely to take in practice.

The usual method of enforcing real-time bounds on accesses to a shared resource is to schedule the use of the resource by different processors. So, let's modify the write-through cache to add a scheduling discipline to actions that enqueue operations on *memQ*. We use round robin scheduling, which is probably the easiest one to implement. Suppose processors are numbered from 0 through  $N - 1$ . Round-robin scheduling means that an operation for processor *p* is the next one to be enqueued after an operation for processor *q* iff there is not an operation for any of the processors  $(q + 1) \% N, (q + 2) \% N, \dots, (p - 1) \% N$  waiting to be put on *memQ*.

To express this formally, we first let the set *Proc* of processors equal the set  $0 \dots (N - 1)$  of integers. Normally, this is done by defining *Proc* to equal  $0 \dots (N - 1)$ . However, we want to reuse the parameters and definitions from the write-through cache specification. The easiest way to do this is by extending module *WriteThroughCache*. Since *Proc* is a parameter in that module, we can't define it. We therefore let *N* be a new constant parameter and let *Proc* =  $0 \dots (N - 1)$  be an assumption.<sup>1</sup>

<sup>1</sup>We could also instantiate module *WriteThroughCache* with  $0 \dots (N - 1)$  substituted for *Proc*; but that requires declaring the other parameters of *WriteThroughCache*, including the

To implement round-robin scheduling, we use a variable *lastP* whose value is the last processor whose operation is enqueued to *memQ*. We define the operator *position* so that *p* is the *position(p)*<sup>th</sup> processor after *lastP* in the round-robin order:

$$\text{position}(p) \triangleq \text{CHOOSE } i \in 1 \dots N : p = (\text{lastP} + i) \% N$$

An operation for processor *p* can be the next to access *memQ* iff there is no operation for a processor *q* with *position(q) < position(p)* ready to access it—that is, iff *canGoNext(p)* is true, where

$$\begin{aligned} \text{canGoNext}(p) &\triangleq \\ &\forall q \in \text{Proc} : (\text{position}(q) < \text{position}(p)) \Rightarrow \neg \text{ENABLED}(\text{RdMiss}(q) \vee \text{DoWr}(q)) \end{aligned}$$

We then define *RTRdMiss(p)* and *RTDoWr(p)* to be the same as *RdMiss(p)* and *DoWr(p)*, respectively, except that they have the additional enabling condition *canGoNext(p)*, and they set *lastP* to *p*. The other subactions of the next-state action are the same as before, except that they must also leave *lastP* unchanged.

For simplicity, we assume a single upper bound of *Epsilon* on the length of time any of the actions of processor *p* can remain enabled without being executed—except for the *Evict(p, a)* action, which we never require to happen. If two actions *A* and *B* are never simultaneously enabled, and one must be executed before the other can become enabled, then a single real-time constraint on *A*  $\vee$  *B* is equivalent to separate constraints on *A* and *B*. We can therefore place a single constraint on the disjunction of all the actions of processor *p*, except we can't use the same constraint for both *DoRd(p)* and *RTRdMiss(p)*, since an *Evict(p, a)* step could disable *DoRd(p)* and enable *RTRdMiss(p)*.

We assume an upper bound of *Delta* on the time *MemQWr* or *MemQRd* can be enabled without dequeuing an operation from *memQ*. The variable *memQ* represents a physical queue between the bus and the main memory, and *Delta* must be large enough so an operation inserted into an empty queue will reach the memory and be dequeued within *Delta* seconds.

We want the real-time write-through cache to implement the real-time memory specification. This requires an assumption relating *Delta*, *Epsilon*, and *Rho* to assure that the memory specification's timing constraint is satisfied—namely, that the delay between when the memory receives a request from processor *p* when it responds is at most *Rho*. Determining this assumption requires computing an upper bound on that delay. Finding the smallest upper bound is complicated, but it isn't too hard to show that

$$2 * (N + 1) * \text{Epsilon} + (N + \text{QLen}) * \text{Delta}$$

is an upper bound. So we assume that this value is less than or equal to *Rho*.

The complete specification appears in Figures 9.4 and 9.5 on the following two pages. The module also asserts as a theorem that the specification *RTSpec*

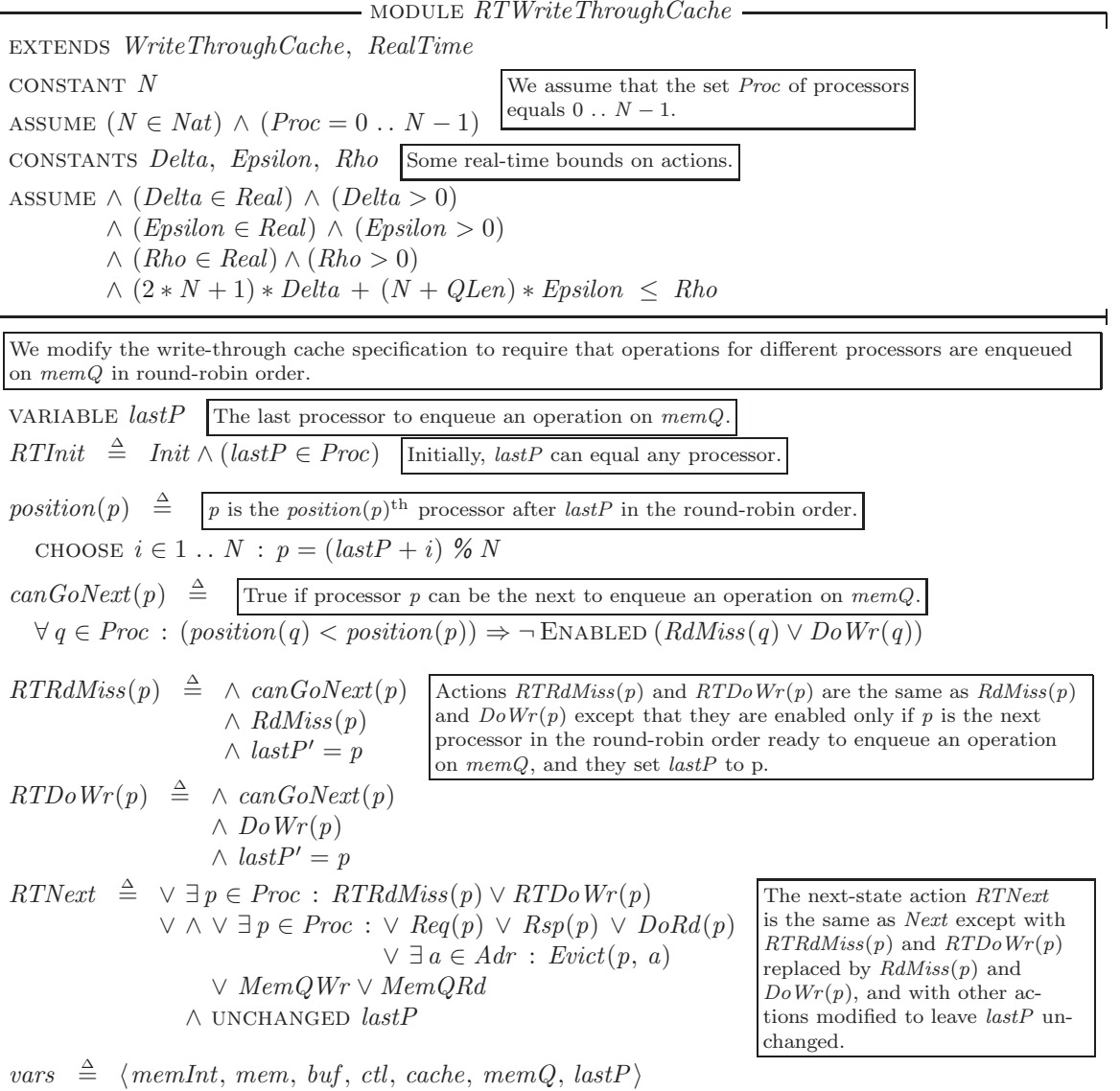


Figure 9.4: A real-time version of the write-through cache (beginning).

|                                                                                                                                                                                                                                                                                                                                                       |                                                                                                                                                                                                                         |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $ \begin{aligned} RTSpec &\triangleq \\ &\wedge RTInit \wedge \Box[RTNext]_{vars} \\ &\wedge RTBound(MemQWr \vee MemQRd, vars, 0, Delta) \\ &\wedge \forall p \in Proc : \wedge RTBound(RTDoWr(p) \vee DoRd(p) \\ &\quad \vee Rsp(p), vars, 0, Epsilon) \\ &\quad \wedge RTBound(RTRdMiss(p) vars, 0, Epsilon) \\ &\wedge RTnow(vars) \end{aligned} $ | <p>We put an upper-bound delay of <i>Delta</i> seconds on <i>MemQWr</i> and <i>MemQRd</i> actions (which dequeue operations from <i>memQ</i>), and an upper-bound delay of <i>Epsilon</i> seconds on other actions.</p> |
| $ \begin{aligned} RTM &\triangleq \text{INSTANCE } RTMemory \\ \text{THEOREM } RTSpec &\Rightarrow RTM!RTSpec \end{aligned} $                                                                                                                                                                                                                         |                                                                                                                                                                                                                         |

Figure 9.5: A real-time version of the write-through cache (end).

of the real-time write-through cache implements (implies) the real-time memory specification, formula  $RTSpec$  of module  $RTMemory$ .

## 9.4 Zeno Specifications

I have described the formula  $RTBound(HCnext, hr, \delta, \epsilon)$  as asserting that an  $HCnext$  step must occur within  $\epsilon$  seconds of the previous  $HCnext$  step. However, implicit in this description is a notion of causality that is not present in the formula. It would be just as accurate to describe the formula as asserting that *now* cannot advance by more than  $\epsilon$  seconds before the next  $HCnext$  step occurs. The formula doesn't tell us whether this condition is met by causing the clock to tick or by preventing time from advancing. Indeed, the formula is satisfied by a “Zeno” behavior<sup>2</sup>

$$\begin{bmatrix} hr &= 11 \\ now &= 0 \end{bmatrix} \rightarrow \begin{bmatrix} hr &= 11 \\ now &= \epsilon/2 \end{bmatrix} \rightarrow \begin{bmatrix} hr &= 11 \\ now &= 3\epsilon/4 \end{bmatrix} \rightarrow \begin{bmatrix} hr &= 11 \\ now &= 7\epsilon/8 \end{bmatrix} \rightarrow \dots$$

in which  $\epsilon$  seconds never pass. We rule out such Zeno behaviors by conjoining to our specification the formula  $RTnow(hr)$ —more precisely by conjoining its liveness conjunct

$$\forall r \in Real : \text{WF}_{now}(Next \wedge (now' > r))$$

which implies that time advances without bound. Let's call this formula  $NZ$  (for NonZeno).

---

ones from the *MemoryInterface* module.

<sup>2</sup>The Greek philosopher Zeno posed the paradox that an arrow first had to travel half the distance to its target, then the next quarter of the distance, then the next eighth, and so on; thus it should not be able to land within a finite length of time.

Zeno behaviors pose no problem; they are trivially forbidden by conjoining  $NZ$ . A problem does exist if a specification allows only Zeno behaviors. For example, suppose we conjoined to the untimed hour-clock's specification the condition  $RTBound(HCnext, hr, \delta, \epsilon)$  for some  $\delta$  and  $\epsilon$  with  $\delta > \epsilon$ . This would assert that the clock must wait at least  $\delta$  seconds before ticking, but must tick within a shorter length of time. In other words, the clock could never tick. Only a Zeno behavior, in which  $\epsilon$  seconds never elapsed, can satisfy this specification. Conjoining  $NZ$  to this specification yields a formula that allows no behaviors—that is, a formula equivalent to  $FALSE$ .

This example is an extreme case of what is called a *Zeno specification*. A Zeno specification is one for which there exists a finite behavior  $\sigma$  that satisfies the safety part but cannot be extended to an infinite behavior that satisfies both the safety part and  $NZ$ .<sup>3</sup> In other words, the only complete behaviors satisfying the safety part that extend  $\sigma$  are Zeno behaviors. Equivalently, a specification is nonZeno iff the pair of properties consisting of the safety part of the specification (the conjunction of the untimed specification, the real-time bound conditions, and the safety part of the  $RTnow$  formula) and  $NZ$  is machine closed.

A Zeno specification is one in which the requirement that time increases without bound rules out some finite behaviors that would otherwise be allowed. Such a specification is likely to be incorrect because the real-time bound conditions are probably constraining the system in unintended ways. In this respect, Zeno specifications are much like other non-machine closed specifications.

In Section 8.8.2 I told you that the conjunction of fairness conditions on subactions of the next-state relation produces a machine closed specification. There is an analogous result for real-time bound conditions and nonZeno specifications. A specification is nonZeno if it is the conjunction of a formula  $Init \wedge \Box[Next]_{vars}$ , the formula  $RTnow(vars)$ , and a finite number of formulas of the form  $RTBound(A_i, vars, \delta_i, \epsilon_i)$ , where for each  $i$ :

- $0 \leq \delta_i \leq \epsilon_i \leq Infinity$
- $A_i$  is a subaction of the next-state action  $Next$ .
- No step is both an  $A_i$  and an  $A_j$  step, for any  $A_j$  with  $j \neq i$ .

In particular, this implies that the specification  $RTSpec$  of the real-time write-through cache in module  $RTWriteThroughCache$  is nonZeno.

This result does not apply to the specification of the real-time memory in module  $RTMemory$  (Figure 9.3 on page 123) because the action  $Respond(p)$  is not a subaction of the next-state action of formula  $ISpec$ . ( $Respond(p)$  can't possibly be a subaction of the next-state action because it doesn't even mention any of the specification's variables except  $ctl$ .) The specification is nonetheless nonZeno, because any finite behavior  $\sigma$  that satisfies the specification can be

Machine closed specifications are defined in Section 8.8.2 on page 110.

<sup>3</sup>Recall that a finite behavior  $\sigma$  satisfies a safety property  $P$  iff adding infinitely many stuttering steps to the end of  $\sigma$  produces a behavior that satisfies  $P$ .

extended to one in which time advances without bound. One can first extend  $\sigma$  to respond to all pending requests immediately (in 0 time), and then extend it to a behavior performing nothing but clock ticks.

It's easy to construct an example in which conjoining an *RTBound* formula for an action that is not a subaction of the next-state action produces a Zeno specification. For example, consider the formula

$$(9.2) \quad HC \wedge RTBound(hr' = hr - 1, hr, 0, 3600) \wedge RTNow(hr)$$

where *HC* is the specification of the hour clock, whose next-state action *HCnext* asserts that *hr* is either incremented by 1 or changes from 12 to 1. The *RTBound* formula asserts that *now* cannot advance for 3600 or more seconds without an  $hr' = hr - 1$  step occurring. Since *HC* asserts that every step that changes *hr* is an *HCnext* step, the safety part of (9.2) is satisfied only by behaviors in which *now* increases by less than 3600 seconds. Since the complete specification (9.2) contains the conjunct *NZ*, which asserts that *now* increases without bound, it is equivalent to FALSE, and is thus a Zeno specification.

When a specification describes how a system is implemented, the real-time constraints are likely to be expressed as *RTBound* formulas for subactions of the next-state action. These are the kinds of formulas that correspond fairly directly to an implementation. More abstract, higher-level specifications—ones describing what a system is supposed to do rather than how to do it—are less likely to have real-time constraints expressed in this ways. Thus, module *RTWriteThroughCache* describes an algorithm for implementing a memory, and it has real-time bounds on subactions of the next-state action. On the other hand, the high-level specification of the real-time memory in module *RTMemory* contains an *RTBound* formula for an action that is not a subaction of the next-state action.

## 9.5 Hybrid System Specifications

A system described by a  $TLA^+$  specification is a physical entity. The specification's variables represent some part of the physical state—the display of a clock, or the distribution of charge in a piece of silicon that implements a memory cell. In a real-time specification, the variable *now* is different from the others because we are not abstracting away the continuous nature of time. The specification allows *now* to assume any of a continuum of values. The discrete states in a behavior mean that we are observing the state of the system, and hence the value of *now*, at a sequence of discrete instants.

There may be physical quantities other than time whose continuous nature we want to represent in a specification. For an air traffic control system, we might want to represent the positions and velocities of the aircraft. For a system controlling a nuclear reactor, we might want to represent the physical parameters

of the reactor itself. A specification that represents such continually varying quantities is called a *hybrid system specification*.

As an example, consider a system that, among other things, controls a switch that influences the one-dimensional motion of some object. Suppose the object's position  $p$  obeys one of the following laws, depending on whether the switch is off or on:

$$(9.3) \quad d^2p/dt^2 + c * dp/dt + f[t] = 0 \qquad d^2p/dt^2 + c * dp/dt + f[t] + k * p = 0$$

where  $c$  and  $k$  are constants,  $f$  is some function, and  $t$  represents time. At any instant, the future position of the object is determined by the object's current position and velocity. So, the state object is described by two variables—namely, its position  $p$  and its velocity  $w$ . These variables are related by  $w = dp/dt$ .

We describe this system with a TLA<sup>+</sup> specification in which the variables  $p$  and  $w$  are changed only by steps that change *now*—that is, steps representing the passage of time. We specify the changes to the discrete system state and any real-time constraints as before. However, we replace  $RTnow(v)$  with a formula having the following next-state action, where *Integrate* and  $D$  are explained below:

$$\begin{aligned} & \wedge now' \in \{r \in Real : r > now\} \\ & \wedge \langle p', w' \rangle = Integrate(D, now, now', \langle p, w \rangle) \\ & \wedge UNCHANGED \ v \quad \boxed{v \text{ is the tuple of all discrete variables, which change instantaneously.}} \end{aligned}$$

The second conjunct asserts that  $p'$  and  $w'$  equal the expressions obtained by solving the appropriate differential equation for the object's position and velocity at time  $now'$ , assuming that their values at time  $now$  are  $p$  and  $w$ . The differential equation is specified by  $D$ , while *Integrate* is a general operator for solving an arbitrary differential equation.

To specify the differential equation satisfied by the object, let's suppose that *switchOn* is a Boolean-valued state variable that describes the position of the switch. We can then rewrite the pair of equations (9.3) as

$$d^2p/dt^2 + c * dp/dt + f[t] + (\text{IF } switchOn \text{ THEN } -k * p \text{ ELSE } 0) = 0$$

We then define the function  $D$  so this equation can be written as

$$D[t, p, dp/dt, d^2p/dt^2] = 0$$

Using the notation explained in Section 15.1.7 on page 285, the definition is:

$$D[t, p0, p1, p2 \in Real] \triangleq p2 + c * p1 + f[t] + (\text{IF } switchOn \text{ THEN } -k * p \text{ ELSE } 0)$$

We obtain the desired specification if the operator *Integrate* is defined so that  $Integrate(D, t_0, t_1, \langle x_0, \dots, x_{n-1} \rangle)$  is the value at time  $t_1$  of the  $n$ -tuple

$$\langle x, dx/dt, \dots, d^{n-1}/dt^{n-1} \rangle$$

where  $x$  is a solution to the differential equation

$$D[t, x, dx/dt, \dots, d^n x/dt^n] = 0$$

whose 0<sup>th</sup> through  $(n-1)$ <sup>st</sup> derivatives at time  $t_0$  are  $x_0, \dots, x_{n-1}$ . The definition of *Integrate* appears in the *DifferentialEquations* module of Section 11.1.3.

In general, a hybrid-system specification is similar to a real-time specification, except that the formula  $RTNow(v)$  is replaced by one that describes the changes to all variables that represent continuously changing physical quantities. The *Integrate* operator will allow you to specify those changes for many hybrid systems. Some systems will require different operators. For example, describing the evolution of some physical quantities might require an operator for describing the solution to a partial differential equation. However, if you can describe the evolution mathematically, then it can be specified in  $TLA^+$ .

Hybrid system specifications still seem to be of only academic interest, so I won't say any more about them. If you do have occasion to write one, this brief discussion should indicate what you can do.

## 9.6 Remarks on Real Time

Real-time constraints are used most often to place an upper bound on how long it can take the system to do something. In this capacity, they can be considered a strong form of liveness, specifying not just that something must eventually happen, but when it must happen. In very simple specifications, such as the hour clock and the write-through cache, real-time constraints usually replace liveness conditions. More complicated specifications can assert both real-time constraints and liveness properties.

The real-time specifications I have seen have not required very complicated timing constraints. They have been specifications either of fairly simple algorithms in which timing constraints are crucial to correctness, or of more complicated systems in which real time appears only through the use of simple timeouts to ensure liveness. I suspect that people don't build systems with complicated real-time constraints because it's too hard to get them right.

I've described how to write a real-time specification by conjoining *RTnow* and *RTBound* formulas to an untimed specification. One can prove that all real-time specifications can be written in this form. In fact, it suffices to use *RTBound* formulas only for subactions of the next-state action. However, this result is of theoretical interest only because the resulting specification can be incredibly complicated. The operators *RTnow* and *RTBound* solve all the real-time specification problems that I have encountered; but I haven't encountered enough to say with confidence that they're all you will ever need. Still, I am quite confident that whatever real-time properties you have to specify, it will not be hard to express them in  $TLA^+$ .



## Chapter 10

# Composing Specifications

Systems are usually described in terms of their components. In the specifications we've written so far, the components have been represented as separate disjuncts of the next-state action. For example, the FIFO system pictured on page 35 is specified in module *InnerFIFO* on page 38 by representing the three components with the following disjuncts of the next-state action:

Sender:  $\exists msg \in Message : SSend(msg)$

Buffer:  $BufRcv \vee BufSend$

Receiver:  $RRcv$

In this chapter, we learn how to specify the components separately and compose their specifications to form a single system specification. Most of the time, there's no point to doing this. The two ways of writing the specification differ by only a few lines—a trivial difference in a specification of hundreds or thousands of lines. Still, you may encounter a situation in which it's better to specify a system as a composition.

First, we must understand what it means to compose specifications. We usually say that a TLA formula specifies the correct behavior of a system. However, as explained in Section 2.3 (page 18), a behavior actually represents a possible history of the entire universe, not just of the system. So, it would be more accurate to say that a TLA formula specifies a universe in which the system behaves correctly. Building a system that implements a specification  $F$  means constructing the universe so it satisfies  $F$ . (Fortunately, correctness of the system depends on the behavior of only a tiny part of the universe, so it's just that part that we must build.) Composing two systems whose specifications are  $F$  and  $G$  means making the universe satisfy both  $F$  and  $G$ , which is the same as making it satisfy  $F \wedge G$ . Thus, the specification of the composition of two systems is the conjunction of their specifications.

Writing a specification as the composition of its components therefore means writing the specification as a conjunction, each conjunct of which can be viewed as the specification of a component. While the basic idea is simple, the details are not always obvious. To simplify the exposition, I begin by considering only safety properties, ignoring liveness and largely ignoring hiding. Liveness and hiding are discussed in Section 10.6.

## 10.1 Composing Two Specifications

Let's return once again to the simple hour clock, with no liveness or real-time requirement. In Chapter 2, we specified such a clock whose display is represented by the variable  $hr$ . We can write that specification as

$$(hr \in 1 \dots 12) \wedge \Box[HCN(hr)]_{hr}$$

where  $HCN$  is defined by:

$$HCN(h) \triangleq h' = (h \% 12) + 1$$

Now let's write a specification  $TwoClocks$  of a system composed of two separate hour clocks, whose displays are represented by the variables  $x$  and  $y$ . (The two clocks are not synchronized and are completely independent of one another.) We can just define  $TwoClocks$  to be the conjunction of the two clock specifications:

$$\begin{aligned} TwoClocks &\triangleq \wedge (x \in 1 \dots 12) \wedge \Box[HCN(x)]_x \\ &\quad \wedge (y \in 1 \dots 12) \wedge \Box[HCN(y)]_y \end{aligned}$$

The following calculation shows how we can rewrite  $TwoClocks$  in the usual form as a “monolithic” specification with a single next-state action:<sup>1</sup>

$$\begin{aligned} TwoClocks &\equiv \wedge (x \in 1 \dots 12) \wedge (y \in 1 \dots 12) \\ &\quad \wedge \Box[HCN(x)]_x \wedge \Box[HCN(y)]_y \\ &\equiv \wedge (x \in 1 \dots 12) \wedge (y \in 1 \dots 12) && \boxed{\text{Because } \Box(F \wedge G) \equiv (\Box F) \wedge (\Box G).} \\ &\quad \wedge \Box([HCN(x)]_x \wedge [HCN(y)]_y) \\ &\equiv \wedge (x \in 1 \dots 12) \wedge (y \in 1 \dots 12) && \boxed{\text{By definition of } [\dots]_x \text{ and } [\dots]_y.} \\ &\quad \wedge \Box(\wedge HCN(x) \vee x' = x \\ &\quad \quad \wedge HCN(y) \vee y' = y) \end{aligned}$$

<sup>1</sup>This calculation is informal because it contains formulas that are not legal TLA—namely, ones of the form  $\Box A$  where  $A$  is an action that doesn't have the syntactic form  $[B]_v$ . However, it can be done rigorously.

$$\begin{aligned}
&\equiv \wedge (x \in 1 \dots 12) \wedge (y \in 1 \dots 12) \\
&\quad \wedge \square (\vee HCN(x) \wedge HCN(y) \\
&\quad \quad \vee HCN(x) \wedge (y' = y) \\
&\quad \quad \vee HCN(y) \wedge (x' = x) \\
&\quad \quad \vee (x' = x) \wedge (y' = y)) \\
&\equiv \wedge (x \in 1 \dots 12) \wedge (y \in 1 \dots 12) \\
&\quad \wedge \square [\vee HCN(x) \wedge HCN(y) \\
&\quad \quad \vee HCN(x) \wedge (y' = y) \\
&\quad \quad \vee HCN(y) \wedge (x' = x)]_{\langle x, y \rangle}
\end{aligned}$$

Because:

$$\begin{pmatrix} \wedge \vee A_1 \\ \vee A_2 \\ \wedge \vee B_1 \\ \vee B_2 \end{pmatrix} \equiv \begin{pmatrix} \vee A_1 \wedge B_1 \\ \vee A_1 \wedge B_2 \\ \vee A_2 \wedge B_1 \\ \vee A_2 \wedge B_2 \end{pmatrix}$$

By definition of  $[\dots]_{\langle x, y \rangle}$ .

Thus, *TwoClocks* is equivalent to  $Init \wedge \square[TCNxt]_{\langle x, y \rangle}$  where the next-state action *TCNxt* is:

$$\begin{aligned}
TCNxt &\triangleq \vee HCN(x) \wedge HCN(y) \\
&\quad \vee HCN(x) \wedge (y' = y) \\
&\quad \vee HCN(y) \wedge (x' = x)
\end{aligned}$$

This next-state action differs from the ones we are used to writing because of the disjunct  $HCN(x) \wedge HCN(y)$ , which represents the simultaneous advance of the two displays. In the specifications we have written so far, different components never act simultaneously.

Up until now, we have been writing what are called *interleaving* specifications. In an interleaving specification, each step represents an operation of only one component. For example, in our FIFO specification, a (nonstuttering) step represents an action of either the sender, the buffer, or the receiver. For want of a better term, we describe as *noninterleaving* a specification that, like *TwoClocks*, does permit simultaneous actions by two components.

Suppose we want to write an interleaving specification of the two-clock system as the conjunction of two component specifications. One way is to replace the next-state actions  $HCN(x)$  and  $HCN(y)$  of the two components by two actions  $HCNx$  and  $HCNy$  so that, when we perform the analogous calculation to the one above, we get

$$\begin{pmatrix} \wedge (x \in 1 \dots 12) \wedge \square[HCNx]_x \\ \wedge (y \in 1 \dots 12) \wedge \square[HCNy]_y \end{pmatrix} \equiv \begin{pmatrix} \wedge (x \in 1 \dots 12) \wedge (y \in 1 \dots 12) \\ \wedge \square [\vee HCNx \wedge (y' = y) \\ \vee HCNy \wedge (x' = x)]_{\langle x, y \rangle} \end{pmatrix}$$

From the calculation above, we see that this equivalence holds if the following three conditions are satisfied: (i)  $HCNx$  implies  $HCN(x)$ , (ii)  $HCNy$  implies  $HCN(y)$ , and (iii)  $HCNx \wedge HCNy$  implies  $x' = x$  or  $y' = y$ . (Condition (iii) implies that the disjunct  $HCNx \wedge HCNy$  of the next-state relation is subsumed by one of the disjuncts  $HCNx \wedge (y' = y)$  and  $HCNy \wedge (x' = x)$ .) The common

way of satisfying these conditions is to let the next-state relation of each clock assert that the other clock's display is unchanged. We do this by defining:

$$HCNx \triangleq HCN(x) \wedge (y' = y) \quad HCNy \triangleq HCN(y) \wedge (x' = x)$$

Another way to write an interleaving specification is simply to disallow simultaneous changes to both clock displays. We can do this by taking as our specification the formula:

$$TwoClocks \wedge \Box[(x' = x) \vee (y' = y)]_{\langle x, y \rangle}$$

The second conjunct asserts that any step must leave  $x$  or  $y$  (or both) unchanged.

Everything we have done for the two-clock system generalizes to any system comprising two components. The same calculation as above shows that if

$$(v_1' = v_1) \wedge (v_2' = v_2) \equiv (v' = v) \quad \boxed{\text{This asserts that } v \text{ is unchanged iff both } v_1 \text{ and } v_2 \text{ are.}}$$

then

$$(10.1) \quad \left( \begin{array}{l} \wedge I_1 \wedge \Box[N_1]_{v_1} \\ \wedge I_2 \wedge \Box[N_2]_{v_2} \end{array} \right) \equiv \left( \begin{array}{l} \wedge I_1 \wedge I_2 \\ \wedge \Box[\vee N_1 \wedge N_2 \\ \vee N_1 \wedge (v_2' = v_2) \\ \vee N_2 \wedge (v_1' = v_1)]_v \end{array} \right)$$

for any state predicates  $I_1$  and  $I_2$  and any actions  $N_1$  and  $N_2$ . The left-hand side of this equivalence represents the composition of two component specifications if  $v_k$  is a tuple containing the variables that describe the  $k^{\text{th}}$  component, for  $k = 1, 2$ , and  $v$  is the tuple of all the variables.

The equivalent formulas in (10.1) represent an interleaving specification if the first disjunct in the next-state action of the right-hand side is redundant, so it can be removed. This is the case if  $N_1 \wedge N_2$  implies that  $v_1$  or  $v_2$  is unchanged. The usual way to ensure that this condition is satisfied is by defining each  $N_k$  so it implies that the other component's tuple is left unchanged. Another way to obtain an interleaving specification by conjoining the formula  $\Box[(v_1' = v_1) \vee (v_2' = v_2)]_v$ .

## 10.2 Composing Many Specifications

We can generalize (10.1) to the composition of any set  $C$  of components. Because universal quantification generalizes conjunction, the following rule is a generalization of (10.1):

**Composition Rule** For any set  $C$ , if

$$(\forall k \in C : v_k' = v_k) \equiv (v' = v) \quad \boxed{\text{This asserts that } v \text{ is unchanged iff all the } v_k \text{ are.}}$$

then

$$\begin{aligned}
 (\forall k \in C : I_k \wedge \Box[N_k]_{v_k}) &\equiv \\
 \wedge \forall k \in C : I_k & \\
 \wedge \Box \left[ \begin{array}{l} \vee \exists k \in C : N_k \wedge (\forall i \in C \setminus \{k\} : v_i' = v_i) \\ \vee \exists i, j \in C : (i \neq j) \wedge N_i \wedge N_j \wedge F_{ij} \end{array} \right]_v &
 \end{aligned}$$

for some action  $F_{ij}$ .

The disjunct containing  $F_{ij}$  is redundant, and we have an interleaving specification, if  $N_i \wedge N_j$  implies that  $v_i$  or  $v_j$  is unchanged, for all  $i$  and  $j$  in  $C$  with  $i \neq j$ . Typically, this is made true by letting each  $N_k$  imply that  $v_j$  is unchanged for all  $j$  in  $C$  other than  $k$ . However, that means that  $N_k$  must mention  $v_j$  for components  $j$  other than  $k$ . You might object to this approach—either on philosophical grounds, because you feel that the specification of one component should not mention the state of another component, or because mentioning other component’s variables complicates the component’s specification. An alternative approach is simply to assert interleaving. You can do this by conjoining the following formula, which states that no step changes both  $v_i$  and  $v_j$ , for any  $i$  and  $j$  with  $i \neq j$ :

$$\Box[\exists k \in C : \forall i \in C \setminus \{k\} : v_i' = v_i]_v$$

This conjunct can be viewed as a global condition, not attached to any component’s specification.

For the left-hand side of the conclusion of the Composition Rule to represent the composition of separate components, the  $v_k$  need not be composed of separate variables. They could contain different “parts” of the same variable that describe different components. For example, our system might consist of a set *Clock* of separate, independent clocks, where clock  $k$ ’s display is described by the value of  $hr[k]$ . Then  $v_k$  would equal  $hr[k]$ . It’s easy to specify such an array of clocks as a composition. Using the definition of *HCN* on page 134 above, we can write the specification as:

$$(10.2) \text{ ClockArray} \triangleq \forall k \in \text{Clock} : (hr[k] \in 1 \dots 12) \wedge \Box[HCN(hr[k])]_{hr[k]}$$

This is a noninterleaving specification, since it allows simultaneous steps by different clocks.

Suppose we wanted to use the Composition Rule to express *ClockArray* as a monolithic specification. What would we substitute for  $v$ ? Our first thought is to substitute  $hr$  for  $v$ . However, the hypothesis of the rule requires that  $v$  must be left unchanged iff  $hr[k]$  is left unchanged, for all  $k \in \text{Clock}$ . However, as explained in Section 6.5 on page 71, specifying the values of  $hr[k]'$  for all

$k \in \text{Clock}$  does not specify the value of  $hr$ . It doesn't even imply that  $hr$  is a function. We must substitute for  $v$  the function  $hrfcn$  defined by

$$(10.3) \quad hrfcn \triangleq [k \in \text{Clock} \mapsto hr[k]]$$

The function  $hrfcn$  equals  $hr$  iff  $hr$  is a function with domain  $\text{Clock}$ . Formula  $\text{ClockArray}$  does not imply that  $hr$  is always a function. It specifies the possible values of  $hr[k]$ , for all  $k \in \text{Clock}$ , but it doesn't specify the value of  $hr$ . Even if we changed the initial condition to imply that  $hr$  is initially a function with domain  $\text{Clock}$ , formula  $\text{ClockArray}$  would not imply that  $hr$  is always a function. For example, it would still allow “stuttering” steps that leave each  $hr[k]$  unchanged, but can change  $hr$  in unknown ways.

We might prefer to write a specification in which  $hr$  is a function with domain  $\text{Clock}$ . (For example, some tool might require that the value of  $hr$  be completely specified.) One way of doing this is to conjoin to the specification the formula  $\Box \text{IsFcnOn}(hr, \text{Clock})$ , where  $\text{IsFcnOn}(hr, \text{Clock})$  asserts that  $hr$  is an arbitrary function with domain  $\text{Clock}$ . The operator  $\text{IsFcnOn}$  is defined by

$$\text{IsFcnOn}(f, S) \triangleq f = [x \in S \mapsto f[x]]$$

We can view the formula  $\Box \text{IsFcnOn}(hr, \text{Clock})$  as a global constraint on  $hr$ , while the value of  $hr[k]$  for each component  $k$  is described by that component's specification.

Now suppose we want to write an interleaving specification of the array of clocks, again as the composition of specifications of the individual clocks. In general, for the conjunction in the Composition Rule to be an interleaving specification,  $N_i \wedge N_j$  should imply that  $v_i$  or  $v_j$  is unchanged. We can do this by letting the next-state relation  $N_k$  of clock  $k$  imply that  $hr[i]$  is unchanged for every clock  $i$  other than  $k$ . The most obvious way to do this is to define  $N_k$  to equal:

$$\begin{aligned} \wedge \quad & hr'[k] = (hr[k] \% 12) + 1 \\ \wedge \quad & \forall i \in \text{Clock} \setminus \{k\} : hr'[i] = hr[i] \end{aligned}$$

We can express this formula more compactly using the EXCEPT construct. This construct applies only to functions, so we must choose whether or not to require  $hr$  to be a function. If  $hr$  is a function, then we can let  $N_k$  equal

$$(10.4) \quad hr' = [hr \text{ EXCEPT } ![k] = (hr[k] \% 12) + 1]$$

As noted above, we can ensure that  $hr$  is a function by conjoining the formula  $\Box \text{IsFcnOn}(hr, \text{Clock})$  to the specification. Another way is to define the state function  $hrfcn$  by (10.3) on this page and let  $N(k)$  equal

$$hrfcn' = [hrfcn \text{ EXCEPT } ![k] = (hr[k] \% 12) + 1]$$

A specification is just a mathematical formula; and, as we've seen before, there are often many equivalent ways of writing a formula. Which one you choose is usually a matter of taste.

The EXCEPT construct is explained in Section 5.2 on page 48.

## 10.3 The FIFO

Let's now specify the FIFO, described in Chapter 4, as the composition of its three components—the Sender, the Buffer, and the Receiver. We start with the internal specification, in which the variable  $q$  occurs—that is,  $q$  is not hidden. First, we decide what part of the state describes each component. The variables  $in$  and  $out$  are channels. Recall that the *Channel* module (page 30) specifies a channel  $chan$  to be a record with  $val$ ,  $rdy$ , and  $ack$  components. The *Send* action, which sends a value, modifies the  $val$  and  $rdy$  components; the *Rcv* action, which receives a value, modifies the  $ack$  component. So, the components' states are described by the following state functions:

Sender:  $\langle in.val, in.rdy \rangle$

Buffer:  $\langle in.ack, q, out.val, out.rdy \rangle$

Receiver:  $out.ack$

Unfortunately, we can't reuse the definitions from the *InnerFIFO* module on page 38 for the following reason. The variable  $q$ , which is hidden in the final specification, is part of the Buffer component's internal state. Therefore, it should not appear in the specifications of the Sender or Receiver component. The Sender and Receiver actions defined in the *InnerFIFO* module all mention  $q$ , so we can't use them. We therefore won't bother reusing that module. However, instead of starting completely from scratch, we can make use of the *Send* and *Rcv* actions from the *Channel* module on page 30 to describe the changes to  $in$  and  $out$ .

Let's write a noninterleaving specification. The next-state actions of the components are then the same as the corresponding disjuncts of the *Next* action in module *InnerFIFO*, except that they do not mention the parts of the states belonging to the other components. These contain *Send* and *Rcv* actions, instantiated from the *Channel* module, which use the EXCEPT construct. As noted above, we can apply EXCEPT only to functions—and to records, which are functions. We therefore add to our specification the conjunct

$$\Box(IsChannel(in) \wedge IsChannel(out))$$

where  $IsChannel(c)$  asserts that  $c$  is a channel—that is a record with  $val$ ,  $ack$ , and  $rdy$  fields. Since a record with  $val$ ,  $ack$ , and  $rdy$  fields is a function whose domain is  $\{\text{"val"}, \text{"ack"}, \text{"rdy"}\}$ , we can define  $IsChannel(c)$  to equal  $IsFcnOn(c, \{\text{"val"}, \text{"ack"}, \text{"rdy"}\})$ . However, it's just as easy to define  $IsChannel(c)$  directly by

$$IsChannel(c) \triangleq c = [ack \mapsto c.ack, val \mapsto c.val, rdy \mapsto c.rdy]$$

In writing this specification, we face the same problem as in our original FIFO specification of introducing the variable  $q$  and then hiding it. In Chapter 4, we

Section 5.2 on page 48 explains why records are functions.

solved this problem by introducing  $q$  in a separate *InnerFIFO* module, which is instantiated by the *FIFO* module that defines the final specification. We do essentially the same thing here, except that we introduce  $q$  in a submodule instead of in a completely separate module. All the symbols declared and defined at the point where the submodule appears can be used within it. The submodule itself can be instantiated in the containing module anywhere after it appears. (Submodules are used in the *RealTimeHourClock* and *RTMemory* specifications on pages 119 and 123 of Chapter 9.)

There is one small problem to be solved before we can write a composite specification of the FIFO—how to specify the initial predicates. It makes sense for the initial predicate of each component's specification to specify the initial values of its part of the state. However the initial condition includes the requirements  $in.ack = in.rdy$  and  $out.ack = out.rdy$ , each of which relates the initial states of two different components. (These requirements are stated in module *InnerFIFO* by the conjuncts *InChan!Init* and *OutChan!Init* of the initial predicate *Init*.) There are three ways of expressing a requirement that relates the initial states of multiple components:

- Assert it in the initial conditions of all the components. Although symmetric, this seems needlessly redundant.
- Arbitrarily assign the requirement to one of the components. This intuitively suggests that we are assigning to that component the responsibility of ensuring that the requirement is met.
- Assert the requirement as a conjunct separate from either of the component specifications. This intuitively suggests that it is an assumption about how the components are put together, rather than a requirement of either component.

When we write an open-system specification, as described in Section 10.7 below, the intuitive suggestions of the last two approaches can be turned into formal requirements. I've taken the last approach and added

$$(in.ack = in.rdy) \wedge (out.ack = out.rdy)$$

as a separate condition. The complete specification is in module *CompositeFIFO* of Figure 10.1 on the next page. Formula *Spec* of this module is a noninterleaving specification; for example, it allows a single step that is both an *InChan!Send* step (the sender sends a value) and an *OutChan!Rcv* step (the receiver acknowledges a value). Hence, it is not equivalent to the interleaving specification *Spec* of the *FIFO* module on page 41, which does not allow such a step.

|                                                                                                                                                                                                                    |                                                             |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------|
| MODULE <i>CompositeFIFO</i>                                                                                                                                                                                        |                                                             |
| EXTENDS <i>Naturals, Sequences</i>                                                                                                                                                                                 |                                                             |
| CONSTANT <i>Message</i>                                                                                                                                                                                            |                                                             |
| VARIABLES <i>in, out</i>                                                                                                                                                                                           |                                                             |
| $InChan \triangleq \text{INSTANCE } Channel \text{ WITH } Data \leftarrow Message, chan \leftarrow in$<br>$OutChan \triangleq \text{INSTANCE } Channel \text{ WITH } Data \leftarrow Message, chan \leftarrow out$ |                                                             |
| $SenderInit \triangleq (in.rdy \in Boolean) \wedge (in.val \in Message)$<br>$Sender \triangleq SenderInit \wedge \square [\exists msg \in Message : InChan!Send(msg)]_{\langle in.val, in.rdy \rangle}$            | The Sender's specification.                                 |
| MODULE <i>InnerBuf</i>                                                                                                                                                                                             |                                                             |
| VARIABLE <i>q</i>                                                                                                                                                                                                  |                                                             |
| $BufferInit \triangleq \wedge in.ack \in Boolean$<br>$\wedge q = \langle \rangle$<br>$\wedge (out.rdy \in Boolean) \wedge (out.val \in Message)$                                                                   | The Buffer's internal specification, with <i>q</i> visible. |
| $BufRcv \triangleq \wedge InChan!Rcv$<br>$\wedge q' = Append(q, in.val)$<br>$\wedge \text{UNCHANGED } \langle out.val, out.rdy \rangle$                                                                            |                                                             |
| $BufSend \triangleq \wedge q \neq \langle \rangle$<br>$\wedge OutChan!Send(Head(q))$<br>$\wedge q' = Tail(q)$<br>$\wedge \text{UNCHANGED } in.ack$                                                                 |                                                             |
| $InnerBuffer \triangleq BufferInit \wedge \square [BufRcv \vee BufSend]_{\langle in.ack, q, out.val, out.rdy \rangle}$                                                                                             |                                                             |
| $Buf(q) \triangleq \text{INSTANCE } InnerBuf$<br>$Buffer \triangleq \exists q : Buf(q)!InnerBuffer$                                                                                                                | The buffer's external specification with <i>q</i> hidden.   |
| $ReceiverInit \triangleq out.ack \in Boolean$<br>$Receiver \triangleq ReceiverInit \wedge \square [OutChan!Rcv]_{\langle in.val, in.rdy \rangle}$                                                                  | The Receiver's specification.                               |
| $IsChannel(c) \triangleq c = [ack \mapsto c.ack, val \mapsto c.val, rdy \mapsto c.rdy]$                                                                                                                            |                                                             |
| $Spec \triangleq \wedge \square (IsChannel(in) \wedge IsChannel(out))$<br>$\wedge (in.ack = in.rdy) \wedge (out.ack = out.rdy)$<br>$\wedge Sender \wedge Buffer \wedge Receiver$                                   | Asserts that <i>in</i> and <i>out</i> are always records.   |
|                                                                                                                                                                                                                    | Relates different components' initial states.               |
|                                                                                                                                                                                                                    | Conjoins the three specifications.                          |

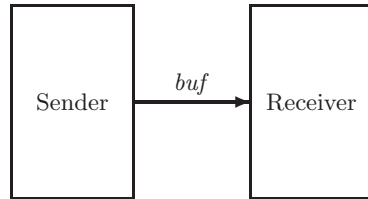
Figure 10.1: A noninterleaving composite specification of the FIFO.

## 10.4 Composition with Shared State

Thus far, we have been considering *disjoint-state compositions*—ones in which the components are represented by disjoint parts of the state, and a component’s next-state action describes changes only to its part of the state.<sup>2</sup> We now consider the case when this may not be possible.

### 10.4.1 Explicit State Changes

We first examine the situation in which some part of the state cannot be partitioned among the different components, but the state change that each component performs is completely described by the specification. As an example, let’s again consider a Sender and a Receiver that communicate with a FIFO buffer. In the system we studied in Chapter 4, sending or receiving a value required two steps. For example, the Sender executes a *Send* step to send a value, and it must then wait until the buffer executes a *Rcv* step before it can send another value. We simplify the system by replacing the Buffer component with a variable *buf* whose value is the sequence of values sent by the Sender but not yet received by the Receiver. This replaces the three-component system pictured on page 35 with this two-component one:



The Sender sends a value by appending it to the end of *buf*; the Receiver receives a value by removing it from the head of *buf*.

In general, the Sender performs some computation to produce the values that it sends, and the Receiver does some computation on the values that it receives. The system state consists of *buf* and two tuples *s* and *r* of variables that describe the Sender and Receiver states. In a monolithic specification, the system’s next-state action is a disjunction  $Sndr \vee Rcvr$ , where *Sndr* and *Rcvr* describe steps taken by the Sender and Receiver, respectively. These actions are

---

<sup>2</sup>In an interleaving composition, a component specification may assert that the state of other components is *not* changed.

defined by

$$\begin{array}{ll}
 Sndr \triangleq & \vee \wedge buf' = Append(buf, \dots) \\
 & \wedge SComm \\
 & \wedge UNCHANGED \ r \\
 & \vee \wedge SCompute \\
 & \wedge UNCHANGED \langle buf, r \rangle \\
 Rcvr \triangleq & \vee \wedge buf \neq \langle \rangle \\
 & \wedge buf' = Tail(buf) \\
 & \wedge RComm \\
 & \wedge UNCHANGED \ s \\
 & \vee \wedge RCompute \\
 & \wedge UNCHANGED \langle buf, s \rangle
 \end{array}$$

for some actions  $SComm$ ,  $SCompute$ ,  $RComm$ , and  $RCompute$ . For simplicity, we assume that neither  $Sndr$  nor  $Rcvr$  allows stuttering actions, so  $SCompute$  changes  $s$  and  $RCompute$  changes  $r$ . We now write the specification as the composition of separate specifications of the Sender and Receiver.

Splitting the initial predicate is straightforward. The initial conditions on  $s$  belong to the Sender's initial predicate; those on  $r$  belong to the Receiver's initial predicate; and the initial condition  $buf = \langle \rangle$  can be assigned arbitrarily to either of them.

Now let's consider the next-state actions  $NS$  and  $NR$  of the Sender and Receiver components. The trick is to define them by

$$NS \triangleq Sndr \vee (\sigma \wedge (s' = s)) \quad NR \triangleq Rcvr \vee (\rho \wedge (r' = r))$$

where  $\sigma$  and  $\rho$  are actions containing only the variable  $buf$ . Think of  $\sigma$  as describing possible changes to  $buf$  that are not caused by the Sender, and  $\rho$  as describing possible changes to  $buf$  that are not caused by the Receiver. Thus,  $NS$  permits any step that is either a  $Sndr$  step or one that leaves  $s$  unchanged and is a change to  $buf$  that can't be "blamed" on the Sender.

Suppose  $\sigma$  and  $\rho$  satisfy the following three conditions:

- $\forall d : (buf' = Append(buf, d)) \Rightarrow \rho$   
A step that appends a value to  $buf$  is not caused by the Receiver.
- $(buf \neq \langle \rangle) \wedge (buf' = Tail(buf)) \Rightarrow \sigma$   
A step that removes a value from the head of  $buf$  is not caused by the Sender.
- $(\sigma \wedge \rho) \Rightarrow (buf' = buf)$   
A step that is caused by neither the Sender nor the Receiver cannot change  $buf$ .

Using the relation<sup>3</sup>

$$(buf' = Append(buf, \dots)) \wedge (buf \neq \langle \rangle) \wedge (buf' = Tail(buf)) \equiv \text{FALSE}$$

<sup>3</sup>This relation is valid only if  $buf$  is a sequence. A rigorous calculation requires the use of an invariant to assert that  $buf$  is indeed a sequence.

a computation like the one by which we derived (10.1) shows

$$\Box[NS]_{\langle buf, s \rangle} \wedge \Box[NR]_{\langle buf, r \rangle} \equiv \Box[Sndr \vee Rcvr]_{\langle buf, s, r \rangle}$$

Thus,  $NS$  and  $NR$  are suitable next-state relations for the components, if we choose  $\sigma$  and  $\rho$  to satisfy the three conditions above. There is considerable freedom in that choice. The strongest possible choices of  $\sigma$  and  $\rho$  are ones that describe exactly the changes permitted by the other component:

$$\begin{aligned} \sigma &\triangleq (buf \neq \langle \rangle) \wedge (buf' = Tail(buf)) \\ \rho &\triangleq \exists d : buf' = Append(buf, d) \end{aligned}$$

We can weaken these definitions any way we want, so long as we maintain the condition that  $\sigma \wedge \rho$  implies that  $buf$  is unchanged. For example, we can define  $\sigma$  as above and let  $\rho$  equal  $\neg\sigma$ . The choice is a matter of taste.

I've been describing an interleaving specification of the Sender/Receiver system. Now let's consider a noninterleaving specification—one that allows steps in which both the Sender and the Receiver are computing. In other words, we want the specification to allow  $SCompute \wedge RCompute$  steps that leave  $buf$  unchanged. Let  $SSndr$  be the action that is the same as  $Sndr$  except it doesn't mention  $r$ , and let  $RRcvr$  be defined analogously. Then we have:

$$Sndr \equiv SSndr \wedge (r' = r) \quad Rcvr \equiv RRcvr \wedge (s' = s)$$

A monolithic noninterleaving specification has the next-state relation

$$Sndr \vee Rcvr \vee (SSndr \wedge SRcvr \wedge (buf' = buf))$$

It is the conjunction of component specifications having the next-state relations  $NS$  and  $NR$  defined by

$$NS \triangleq SSndr \vee (\sigma \wedge (s' = s)) \quad NR \triangleq RRcvr \vee (\rho \wedge (r' = r))$$

where  $\sigma$  and  $\rho$  are as above.

This two-process situation generalizes to the composition of any set  $C$  of components that share a variable or tuple of variables  $w$ . The interleaving case generalizes to the following rule, in which  $N_k$  is the next-state action of component  $k$ , the action  $\mu_k$  describes all changes to  $w$  that are attributed to some component other than  $k$ , the tuple  $v_k$  describes the private state of  $k$ , and  $v$  is the tuple formed by all the  $v_k$ .

**Shared-State Composition Rule** The four conditions

1.  $(\forall k \in C : v_k' = v_k) \equiv (v' = v)$

$v$  is unchanged iff the private state  $v_k$  of every component is unchanged.

2.  $\forall i, k \in C : N_k \wedge (i \neq k) \Rightarrow (v_i' = v_i)$

The next-state action of any component  $k$  leaves the private state  $v_i$  of all other components  $i$  unchanged.

$$3. \forall i, k \in C : N_k \wedge (w' \neq w) \wedge (i \neq k) \Rightarrow \mu_i$$

A step of any component  $k$  that changes  $w$  is a  $\mu_i$  step, for any other component  $i$ .

$$4. (\forall k \in C : \mu_k) \Rightarrow (w' = w)$$

A step that is caused by no component does not change  $w$ .

imply

$$\begin{aligned} & (\forall k \in C : I_k \wedge \Box[N_k \vee (\mu_k \wedge (v_k' = v_k))])_{\langle w, v_k \rangle} \\ & \equiv (\forall k \in C : I_k) \wedge \Box[\exists k \in C : N_k]_{\langle w, v \rangle} \end{aligned}$$

Assumption 2 asserts that we have an interleaving specification. If we drop that assumption, then the right-hand side of the conclusion may not be a sensible specification, since a disjunct  $N_k$  may allow steps in which a variable of some other component assumes arbitrary values. However, if each  $N_k$  correctly determines the new values of component  $k$ 's private state  $v_k$ , then the left-hand side will be a reasonable specification, though possibly a noninterleaving one (and not necessarily equivalent to the right-hand side).

## 10.4.2 Composition with Joint Actions

Consider the linearizable memory of Chapter 5. As shown in the picture on page 45, it is a system consisting of a collection of processors, a memory, and an interface represented by the variable *memInt*. We take it to be a two-component system, where the set of processors forms one component, called the *environment*, and the memory is the other component. Let's neglect hiding for now and consider only the internal specification, with all variables visible. We want to write the specification in the form

$$(10.5) \quad (IE \wedge \Box[NE]_{vE}) \wedge (IM \wedge \Box[NM]_{vM})$$

where  $E$  refers to the environment component (the processors) and  $M$  to the memory component. The tuple  $vE$  of variables includes *memInt* and the variables of the environment component; the tuple  $vM$  includes *memInt* and the variables of the memory component. We must choose the formulas  $IE$ ,  $NE$ , etc. so that (10.5), with internal variables hidden, is equivalent to the memory specification *Spec* of module *Memory* on page 53.

In the memory specification, communication between the environment and the memory is described by an action of the form

$$Send(p, d, memInt, memInt') \quad \text{or} \quad Reply(p, d, memInt, memInt')$$

where *Send* and *Reply* are unspecified operators declared in the *MemoryInterface* module (page 48). The specification says nothing about the actual value of

*memInt*. So, not only do we not know how to split *memInt* into two parts that are each changed by only one of the components, we don't even know exactly how *memInt* changes.

The trick to writing the specification as a composition is to put the *Send* and *Reply* actions in the next-state relations of both components. We represent the sending of a value over *memInt* as a *joint action* performed by both the memory and the environment. The next-state relations have the following form:

$$\begin{aligned} NM &\triangleq \exists p \in Proc : MReq(p) \vee MRsp(p) \vee MInternal(p) \\ NE &\triangleq \exists p \in Proc : EReq(p) \vee ERsp(p) \end{aligned}$$

where an  $MReq(p) \wedge EReq(p)$  step represents the sending of a request by processor  $p$  (part of the environment) to the memory, an  $MRsp(p) \wedge ERsp(p)$  step represents the sending of a reply by the memory to processor  $p$ , and an  $MInternal(p)$  step is an internal step of the memory component that performs the request. (There are no internal steps of the environment.)

The sending of a reply is controlled by the memory, which chooses what value is sent and when it is sent. The enabling condition and the value sent are therefore specified by the  $MRsp(p)$  action. Let's take the internal variables of the memory component to be the same variables *mem*, *ctl*, and *buf* as in the internal monolithic memory specification of module *InternalMemory* on pages 52 and 53. We can then let  $MRsp(p)$  be the same as the action  $Rsp(p)$  defined in that module. The  $ERsp(p)$  action should always be enabled, and it should allow any legal response to be sent. A legal response is an element of *Val* or the special value *NoVal*, so we can define  $ERsp(p)$  to equal:<sup>4</sup>

$$\begin{aligned} &\wedge \exists rsp \in Val \cup \{NoVal\} : Reply(p, rsp, memInt, memInt') \\ &\wedge \dots \end{aligned}$$

where the “...” describes the new values of the environment's internal variables.

The sending of a request is controlled by the environment, which chooses what value is sent and when it is sent. Hence, the enabling condition should be part of the  $EReq(p)$  action. In the monolithic specification of the *InternalMemory* module, that enabling condition was  $ctl[p] = \text{“rdy”}$ . However, if *ctl* is an internal variable of the memory, it can't also appear in the environment specification. We therefore have to add a new variable whose value indicates whether a processor is allowed to send a new request. Let's use a Boolean variable *rdy*, where  $rdy[p]$  is true iff processor  $p$  can send a request. The value of  $rdy[p]$  is set false when  $p$  sends a request and is set true again when the corresponding response to  $p$  is sent. We can therefore define  $EReq(p)$ , and

<sup>4</sup>The bound on the  $\exists$  isn't necessary. We can let the processor accept any value, not just a legal one, by taking  $\exists rsp : Reply(p, rsp, memInt, memInt')$  as the first conjunct. However, it's generally better to use bounded quantifiers when possible.

complete the definition of  $ERsp(p)$ , as follows:

$$\begin{aligned}
 EReq(p) &\triangleq \wedge rdy[p] \\
 &\quad \wedge \exists req \in MReq : Send(p, req, memInt, memInt') \\
 &\quad \wedge rdy' = [rdy \text{ EXCEPT } ![p] = \text{FALSE}] \\
 ERsp(p) &\triangleq \wedge \exists rsp \in Val \cup \{NoVal\} : Reply(p, rsp, memInt, memInt') \\
 &\quad \wedge rdy' = [rdy \text{ EXCEPT } ![p] = \text{TRUE}]
 \end{aligned}$$

The memory's  $MReq(p)$  action is the same as the  $Req(p)$  action of the *InternalMemory* module, except without the enabling condition  $ctl[p] = \text{"rdy"}$ .

Finally, the memory's internal action  $MInternal(p)$  is the same as the  $Do(p)$  action of the *InternalMemory* module.

The rest of the specification is easy. The tuples  $vE$  and  $vM$  are  $\langle memInt, rdy \rangle$  and  $\langle memInt, mem, ctl, buf \rangle$ , respectively. Defining the initial predicates  $IE$  and  $IM$  is straightforward, except for the decision of where to put the initial condition  $memInt \in InitMemInt$  for  $memInt$ . We can put it in either  $IE$  or  $IM$ , in both, or else in a separate conjunct that belongs to neither component's specification. Let's put it in  $IM$ , which then equals the initial predicate  $IInit$  from the *InternalMemory* module. The final environment specification is obtained by hiding  $rdy$  in its internal specification; the final memory component specification is obtained by hiding  $mem$ ,  $ctl$ , and  $buf$  in its internal specification. The complete specification appears in Figure 10.2 on the next page. I have not bothered to define  $IM$ ,  $MRsp(p)$ , or  $MInternal(p)$ , since they equal  $IInit$ ,  $Rsp(p)$ , and  $Do(p)$  from the *InternalMemory* module, respectively.

What we've just done for the environment-memory system generalizes naturally to joint-action specifications of any two-component system in which part of the state cannot be considered to belong to either component. It also generalizes to systems in which any number of components share some part of the state. For example, suppose we want to write a composite specification of the linearizable memory system in which each processor is a separate component. The specification of the memory component would be the same as before. The next-state relation of processor  $p$  now be

$$EReq(p) \vee ERsp(p) \vee OtherProc(p)$$

where  $EReq(p)$  and  $ERsp(p)$  are the same as above, and an  $OtherProc(p)$  step represents the sending of a request by, or a response to, some processor other than  $q$ . Action  $OtherProc(p)$  represents  $p$ 's participation in the joint action by which another processor  $q$  communicates with the memory component. It is defined to equal:

$$\begin{aligned}
 \exists q \in Proc \setminus \{p\} : &\vee \exists req \in MReq : Send(p, req, memInt, memInt') \\
 &\vee \exists rsp \in Val \cup \{NoVal\} : Reply(q, rsp, memInt, memInt')
 \end{aligned}$$

This example is rather silly because each processor must participate in communication actions that concern only other components. It would be better to

```

┌────────────────── MODULE JointActionMemory ───────────────────┐
EXTENDS MemoryInterface
┌────────────────── MODULE InnerEnvironmentComponent ───────────────────┐
  VARIABLE rdy
   $IE \triangleq rdy = [p \in Proc \mapsto \text{TRUE}]$ 
   $EReq(p) \triangleq \wedge rdy[p]$ 
   $\wedge \exists req \in MReq : Send(p, req, memInt, memInt')$ 
   $\wedge rdy' = [rdy \text{ EXCEPT } ![p] = \text{FALSE}]$ 
   $ERsp(p) \triangleq \wedge \exists rsp \in Val \cup \{NoVal\} : Reply(p, rsp, memInt, memInt')$ 
   $\wedge rdy' = [rdy \text{ EXCEPT } ![p] = \text{TRUE}]$ 
   $NE \triangleq \exists p \in Proc : EReq(p) \vee ERsp(p)$ 
   $IESpec \triangleq IE \wedge \Box[MN]_{\langle memInt, rdy \rangle}$ 
└──────────────────┐

┌────────────────── MODULE InnerMemoryComponent ───────────────────┐
EXTENDS InternalMemory
   $MReq(p) \triangleq \wedge \exists req \in MReq : \wedge Send(p, req, memInt, memInt')$ 
   $\wedge buf' = [buf \text{ EXCEPT } ![p] = req]$ 
   $\wedge ctl' = [ctl \text{ EXCEPT } ![p] = \text{"busy"}]$ 
   $\wedge \text{UNCHANGED } mem$ 
   $NM \triangleq \exists p \in Proc : MReq(p) \vee Do(p) \vee Rsp(p)$ 
   $IMSpec \triangleq IInit \wedge \Box[IMNext]_{\langle memInt, mem, ctl, buf \rangle}$ 
└──────────────────┐

 $IE(rdy) \triangleq \text{INSTANCE } InnerEnvironmentComponent$ 
 $IM(mem, ctl, buf) \triangleq \text{INSTANCE } InnerMemoryComponent$ 
 $Spec \triangleq \wedge \exists rdy : IE(rdy)!IMSpec$ 
 $\wedge \exists mem, ctl, buf : IM(mem, ctl, buf)!IESpec$ 
└──────────────────┐

```

Figure 10.2: A joint-action compositional specification of the linearizable memory.

change the interface to make *memInt* an array, with communication between processor *p* and the memory represented by a change to *memInt*[*p*]. A sensible example would require that a joint action represent a true interaction between all the components—for example, a barrier synchronization operation in which the components wait until they are all ready and then perform a synchronization step together.

## 10.5 A Brief Review

The basic idea of composing specifications is simple: a composite specification is the conjunction of formulas, each of which can be considered to be the specification of a separate component. I have presented several techniques for writing a specification as a composition. Before going further, let's put these techniques in perspective.

### 10.5.1 A Taxonomy of Composition

I have introduced three different ways of categorizing composite specifications:

**Interleaving versus noninterleaving.** An interleaving specification is one in which each (nonstuttering) step can be attributed to exactly one component. A noninterleaving specification allows steps that represent simultaneous operations of two or more different components.

**Disjoint-state versus shared-state.** A disjoint-state specification is one in which the state can be partitioned, with each part belonging to a separate component. A part of the state can be a variable  $v$ , or a “piece” of that variable such as  $v.c$  or  $v[c]$  for some fixed  $c$ . Any change to a component's part of the state is attributed to that component. In a shared-state specification, some part of the state can be changed by steps attributed to more than one component.

**Joint-action versus separate-action.** A joint-action specification is a noninterleaving one in which some step attributed to one component must occur simultaneously with a step attributed to another component. A separate-action specification is simply one that is not a joint-action specification.

These categories are orthogonal, except that a joint-action specification must be noninterleaving.

### 10.5.2 Interleaving Reconsidered

Should we write interleaving or noninterleaving specifications? We might try to answer this question by asking, can different components really take simultaneous steps? However, this question makes no sense. A step is a mathematical abstraction; real components perform operations that take a finite amount of time. Operations performed by two different components could overlap in time. We are free to represent this physical situation either with a single simultaneous step of the two components, or with two separate steps. In the latter case, the specification usually allows the two steps to occur in either order. (If the two

operations must occur simultaneously, then we have written a joint-action specification.) It's up to you whether to write an interleaving or a noninterleaving specification. You should choose whichever is more convenient.

The choice is not completely arbitrary if you want one specification to implement another. A noninterleaving specification will not, in general, implement an interleaving one because the noninterleaving specification will allow simultaneous actions that the interleaving specification prohibits. So, if you want to write a low-level specification that implements a high-level interleaving specification, then you'll have to use an interleaving specification. As we've seen, it's easy to turn a noninterleaving specification into an interleaving one by conjoining an interleaving assumption.

### 10.5.3 Joint Actions Reconsidered

The reason for writing a composite specification is to separate the specifications of the different components. The mixing of actions from different components in a joint-action specification destroys this separation. So, why should we write such a specification?

Joint-action specifications arise most often in highly abstract descriptions of inter-component communication. In writing a composite specification of the linearizable memory, we were led to use joint actions because of the abstract nature of the interface. In real systems, communication occurs when one component changes the state and another component later observes that change. The interface described by the *MemoryInterface* module abstracts away those two steps, replacing them with a single one that represents instantaneous communication—a fiction that does not exist in the real world. Since each component must remember that the communication has occurred, the single communication step has to change the private state of both components. That's why we couldn't use the approach of Section 10.4.1, which requires that any change to the shared interface change the nonshared state of just one component.

The abstract memory interface simplifies the specification, allowing communication to be represented as one step instead of two. But this simplification comes at the cost of blurring the distinction between the two components. If we blur this distinction, it may not make sense to write the specification as the conjunction of separate component specifications. As the memory system example illustrates, decomposing the system into separate components communicating with joint actions may require the introduction of extra variables. There may occasionally be a good reason for adding this kind of complexity to a specification, but it should not be done as a matter of course.

## 10.6 Liveness and Hiding

### 10.6.1 Liveness and Machine Closure

Thus far, the discussion of composition has neglected liveness. In composite specifications, it is usually easy to specifying liveness by placing fairness conditions on the actions of individual components. For example, to specify an array of clocks that all keep ticking forever, we would modify the specification *ClockArray* of (10.2) on page 137 to equal:

$$\forall k \in \text{Clock} : (hr[k] \in 1 \dots 12) \wedge \Box[HCN(hr[k])]_{hr[k]} \wedge WF_{hr[k]}(HCN(hr[k]))$$

When writing a weak or strong fairness formula for an action  $A$  of component  $c$ , there arises the question of what the subscript should be. The obvious choices are (i) the tuple  $v$  describing the entire specification state, and (ii) the tuple  $v_c$  describing that component's state.<sup>5</sup> The choice matters only if the safety part of the specification allows the system to reach some state in which an  $A$  step could leave  $v_c$  unchanged while changing  $v$ . In practice, this is seldom the case. If it is, we probably don't want the fairness condition to be satisfied by a step that leaves the component's state unchanged, so we would use the subscript  $v_c$ .

Fairness conditions for composite specifications do raise one important question: if each component specification is machine closed, is the composite specification necessarily machine closed? Suppose we write the specification as  $\forall k \in C : S_k \wedge L_k$ , where each pair  $S_k, L_k$  is machine closed. Let  $S$  be the conjunction of the  $S_k$  and  $L$  the conjunction of the  $L_k$ , so the specification equals  $S \wedge L$ . The conjunction of safety properties is a safety property,<sup>6</sup> so  $S$  is a safety property. Hence, we can ask if the pair  $S, L$  is machine closed.

In general,  $S, L$  need not be machine closed. But, for an interleaving composition, it usually is. Liveness properties are usually expressed as the conjunction of weak and strong fairness properties of actions. As stated on page 111, a specification is machine closed if its liveness property is the conjunction of fairness properties for subactions of the next-state action. In an interleaving composition, each  $S_k$  usually has the form  $I_k \wedge \Box[N_k]_{v_k}$  where the  $v_k$  satisfy the hypothesis of the Composition Rule (page 136), and each  $N_k$  implies  $v_i' = v_i$ , for all  $i$  in  $C \setminus \{k\}$ . In this case, the Composition Rule implies that a subaction of  $N_k$  is also a subaction of the next-state relation of  $S$ . Hence, if we write an interleaving composition in the usual way, and we write machine-closed component specifications in the usual way, then the composite specification is machine closed.

Machine closure is defined in Section 8.8.2 on page 110.

<sup>5</sup>For a shared-state composition, the tuples  $v_c$  for different components need not be disjoint.

<sup>6</sup>Recall that a safety property is one that is violated by a behavior iff it is violated at some particular point in the behavior. A behavior violates a conjunction of safety properties  $S_k$  iff it violates some particular  $S_k$ , and that  $S_k$  is violated at some specific point.

It is not so easy to obtain a machine-closed noninterleaving composition—especially with a joint-action composition. We have actually seen an example of a joint-action specification in which each component is machine closed but the composition is not. In Chapter 9, we wrote a real-time specification by conjoining one or more *RTBound* formulas and an *RTnow* formula to an untimed specification. A pathological example was the following, which is formula (9.2) on page 129:

$$HC \wedge RTBound(hr' = hr - 1, hr, 0, 3600) \wedge RTNow(hr)$$

*HC* is the hour-clock specification from Chapter 2

We can view this formula as the conjunction of three component specifications:

1. *HC* specifies a clock, represented by the variable *hr*.
2. *RTBound*(*hr'* = *hr* − 1, *hr*, 0, 3600) specifies a timer, represented by the hidden (existentially quantified) timer variable.
3. *RTNow*(*hr*) specifies the behavior of time, represented by the variable *now*.

It's a joint-action composition, with two kinds of joint actions:

- Joint actions of the first and second components that change both *hr* and the timer.
- Joint actions of the second and third components that change both the timer and *now*.

The first two specifications are trivially machine closed because they assert no liveness condition, so their liveness property is TRUE. The third specification's safety property asserts that *now* is a real number that is changed only by steps that increment it and leave *hr* unchanged; its liveness property *NZ* asserts that *now* increases without bound. Any finite behavior satisfying the safety property can easily be extended to an infinite behavior satisfying the entire specification, so the third specification is also machine closed. However, as we observed in Section 9.4, the composite specification is Zeno, meaning that it's not machine closed.

## 10.6.2 Hiding

Suppose we can write a specification *S* as the composition of two component specifications *S*<sub>1</sub> and *S*<sub>2</sub>. Can we write  $\exists h : S$ , the specification *S* with variable *h* hidden, as a composition—that is, as the conjunction of two separate component specifications? If *h* represents state that is accessed by both components, then the answer is no. If the two components communicate through some part of the state, then that part of the state cannot be made internal to the separate components.

The simplest situation in which  $h$  doesn't represent shared state is when it occurs in only one of the component specifications—say,  $S_2$ . If  $h$  doesn't occur in  $S_1$ , then the equivalence

$$(\exists h : S_1 \wedge S_2) \equiv S_1 \wedge (\exists h : S_2)$$

provides the desired decomposition.

Now suppose that  $h$  occurs in both component specifications, but does not represent state accessed by both components. This can be the case only if different “parts” of  $h$  occur in the two component specifications. For example,  $h$  might be a record with components  $h.c1$  and  $h.c2$ , where  $S_1$  mentions only  $h.c1$  and  $S_2$  mentions only  $h.c2$ . In this case, we have

$$(\exists h : S_1 \wedge S_2) \equiv (\exists h1 : T_1) \wedge (\exists h2 : T_2)$$

Where  $T_1$  is obtained from  $S_1$  by substituting the variable  $h1$  for the expression  $h.c1$ , and  $T_2$  is defined similarly. Of course we can use any variables in place of  $h1$  and  $h2$ ; in particular, we can replace them both by the same variable.

We can generalize this result as follows to the composition of any finite number<sup>7</sup> of components:

**Compositional Hiding Rule** If the variable  $h$  does not occur in the formula  $T_i$ , and  $S_i$  is obtained from  $T_i$  by substituting  $h[i]$  for  $q$ , then

$$(\exists h : \forall i \in C : S_i) \equiv (\forall i \in C : \exists q : T_i)$$

for any finite set  $C$ .

The assumption that  $h$  does not occur in  $T_i$  means that the variable  $h$  occurs in formula  $S_i$  only in the expression  $h[i]$ . This in turn implies that the composition  $\forall i \in C : S_i$  does not determine the value of  $h$ , just of its components  $h[i]$  for  $i \in C$ . As noted in Section 10.2, we can make the composite specification determine the value of  $h$  by conjoining the formula  $\Box IsFcnOn(h, C)$  to it, where  $IsFcnOn$  is defined on page 138. After  $h$  is hidden, it makes no difference whether or not its value is determined. The hypotheses of the Compositional Hiding Rule imply:

$$(\exists h : \Box IsFcnOn(h, C) \wedge \forall i \in C : S_i) \equiv (\forall i \in C : \exists q : T_i)$$

Now consider the common case in which  $\forall i \in C : S_i$  is an interleaving composition, where each specification  $S_i$  describes changes to  $h[i]$  and asserts that steps of component  $i$  leave  $h[j]$  unchanged for  $j \neq i$ . We cannot apply the Compositional Hiding Rule because  $S_i$  must mention other components of  $h$  besides  $h[i]$ . For example, it probably contains an expression of the form

$$(10.6) \ h' = [h \text{ EXCEPT } ![i] = exp]$$

<sup>7</sup>The Compositional Hiding Rule is not true in general if  $C$  is an infinite set; but the examples in which it doesn't hold are pathological and don't arise in practice.

See the discussion  
in Section 10.2  
(page 136).

which mentions all of  $h$ . However, we can transform  $S_i$  into a specification  $\widehat{S}_i$  that describes only the changes to  $h[i]$  and makes no assertions about other components. For example, we can replace (10.6) with  $h'[i] = \text{exp}$ , and we can replace an assertion that  $h$  is unchanged by the assertion that  $h[i]$  is unchanged. The composition  $\forall i \in C : \widehat{S}_i$  will not be equivalent to  $\forall i \in C : S_i$ . In particular, it need not be an interleaving composition, since it might allow steps that change two different components  $h[i]$  and  $h[j]$ , while leaving all other variables unchanged. However, it can be shown that the two specifications are equivalent when  $h$  is hidden (assuming  $C$  is a finite set). So, we can apply the Composition Hiding Rule with  $S_i$  replaced by  $\widehat{S}_i$ .

## 10.7 Open-System Specifications

A specification describes the interaction between a system and its environment. For example, the FIFO buffer specification of Chapter 4 specifies the interaction between the buffer (the system) and an environment consisting of the sender and receiver. So far, all the specifications we have written have been complete-system specifications, meaning that they are satisfied by a behavior that represents the correct operation of both the system and its environment. When we write such a specification as the composition of an environment specification  $E$  and a system specification  $M$ , it has the form  $E \wedge M$ .

An open-system specification is one that can serve as a contract between a user of the system and its implementor. A behavior satisfies an open-system specification iff the system acts correctly, even if the environment misbehaves. An obvious choice of such a specification is the formula  $M$  that describes the correct behavior of the system component by itself. However, such a specification would be unimplementable. A system cannot behave as expected in the face of arbitrary behavior of its environment. It would be impossible to build a buffer that satisfies the buffer component's specification regardless of what the sender and receiver did. For example, if the sender sends a value before the previous value has been acknowledged, then the buffer could read the value while it is changing, causing unpredictable results.

A contract between a user and an implementor should require the system to act correctly only if the environment does. If  $M$  describes correct behavior of the system and  $E$  describes correct behavior of the environment, such a specification should require that  $M$  be true if  $E$  is. This suggests that we take as our open-system specification the formula  $E \Rightarrow M$ , which is true if the system behaves correctly or the environment behaves incorrectly. However,  $E \Rightarrow M$  is too weak a specification for the following reason. Consider again the example of a *FIFO* buffer, where  $M$  describes the buffer and  $E$  the sender and receiver. Suppose now that the buffer sends a new value before the receiver has acknowledged the previous one. This could cause the receiver to act incorrectly, possibly modifying

Open-system specifications are sometimes called *rely-guarantee* or *assume-guarantee* specifications.

the output channel in some way not allowed by the receiver's specification. This situation is described by a behavior in which both  $E$  and  $M$  are false—a behavior that satisfies the specification  $E \Rightarrow M$ . However, the buffer should not be considered to act correctly in this case, since it was the buffer's error that caused the receiver to act incorrectly. Hence, this behavior should not satisfy the buffer's specification.

An open-system specifications should assert that the system behaves correctly at least as long as the environment does. To express this, we introduce a new temporal operator  $\vdash$ , where  $E \vdash M$  asserts that  $M$  remains true at least one step longer than  $E$  does, remaining true forever if  $E$  does. Somewhat more precisely,  $E \vdash M$  asserts that:

- $E$  implies  $M$ , and
- If the safety property of  $E$  is not violated by the first  $n$  states of a behavior, then the safety property of  $M$  is not violated by the first  $n+1$  states, for any natural number  $n$ . (Recall that a safety property is one that, if violated, is violated at some definite point in the behavior.)

A more precise definition of  $\vdash$  appears in Section 15.2.4 (page 298). If  $E$  describes the desired behavior of the environment and  $M$  describes the desired behavior of the system, then we take as our open-system specification the formula  $E \vdash M$ .

Once we write separate specifications of the components, we can usually transform a complete-system specification into an open-system one by simply replacing a  $\wedge$  with a  $\vdash$ . This requires first deciding whether each conjunct of the complete-system specification belongs to the specification of the environment, of the system, or of neither. As an example, consider the composite specification of the FIFO buffer in module *CompositeFIFO* on page 141. We take the system to consist of just the buffer, with the sender and receiver forming the environment. The closed-system specification *Spec* has three main conjuncts. We consider them separately.

*Sender*  $\wedge$  *Buffer*  $\wedge$  *Receiver*

The conjuncts *Sender* and *Receiver* are obviously part of the environment specification and the conjunct *Buffer* is part of the system specification.

$(in.ack = in.rdy) \wedge (ack.out = out.rdy)$

These two initial conjuncts can be assigned to either, depending on which component we want to blame if they are violated. Let's assign to the component sending on a channel  $c$  the responsibility for establishing that  $c.ack = c.rdy$  holds initially. We then assign  $in.ack = in.rdy$  to the environment and  $ack.out = out.rdy$  to the system.

$\Box(IsChannel(in) \wedge IsChannel(out))$

This formula is not naturally attributed to either the system or the environment. We regard it as a property inherent in our way of modeling

the system, which assumes that *in* and *out* are records with *ack*, *val*, and *rdy* components. We therefore take the formula to be a separate conjunct of the complete specification, not belonging to either the system or the environment.

We then have the following open-system specification for the FIFO buffer:

$$\begin{aligned} & \wedge \Box(IsChannel(in) \wedge IsChannel(out)) \\ & \wedge ((in.ack = in.rdy) \wedge Sender \wedge Receiver) \stackrel{\pm}{\triangleright} ((ack.out = out.rdy) \wedge Buffer) \end{aligned}$$

As this example suggests, there is little difference between writing a composite complete-system specification and an open-system specification. Most of the specification doesn't depend on which we choose. The two differ only at the very end, when we put the pieces together.

## 10.8 Interface Refinement

An *interface refinement* is a method of obtaining a lower-level specification by refining the variables of a higher-level specification. I start with two examples and then discuss interface refinement in general.

### 10.8.1 A Binary Hour Clock

In specifying an hour clock, we described its display with a variable *hr* whose value (in a behavior satisfying the specification) is an integer from 1 to 12. Suppose we want to specify a *binary hour clock*. This is an hour clock for use in a computer, where the display consists of a four-bit register that displays the hour as one of the twelve values 0001, 0010, ..., 1100. We can easily specify such a clock from scratch. But suppose we want to describe it informally to someone who already knows what an hour clock is. We would simply say that a binary clock is the same as an hour clock, except that the value of the display is represented in binary. We now formalize that description.

We begin by describing what it means for a four-bit value to represent a number. There are several reasonable ways to represent a four-bit value mathematically. We could use a four-element sequence, which in TLA<sup>+</sup> is a function whose domain is 1 .. 4. However, a mathematician would find it more natural to represent an *n*-bit number as a function from 0 .. (*n* - 1) to {0, 1}, the function *b* representing the number  $b[0] * 2^0 + b[1] * 2^1 + \dots + b[n-1] * 2^{n-1}$ . In TLA<sup>+</sup>,

We can also write {0, 1} as 0 .. 1.

we can define  $BitArrayVal(b)$  to be the numerical value of such a function  $b$  by:

$$\begin{aligned}
 BitArrayVal(b) \triangleq & \text{ LET } n \triangleq \text{ CHOOSE } i \in Nat : \text{ DOMAIN } b = 0 \dots (i-1) \\
 & val[i \in 0 \dots (n-1)] \triangleq \boxed{\text{Defines } val[i] \text{ to equal } 2^0 + \dots + 2^i.} \\
 & \text{ IF } i = 0 \text{ THEN } 2^0 \text{ ELSE } 2^i + val[i-1] \\
 & \text{ IN } val[n-1]
 \end{aligned}$$

To specify a binary hour clock whose display is described by the variable  $bits$ , we would simply say that  $BitArrayVal(bits)$  changes the same way that the specification  $HC$  of the hour clock allows  $hr$  to change. Mathematically, this means that we obtain the specification of the binary hour clock by substituting  $BitArrayVal(bits)$  for the variable  $hr$  in  $HC$ . In  $TLA^+$ , substitution is expressed with the `INSTANCE` statement. Writing

$$B \triangleq \text{ INSTANCE } HourClock \text{ WITH } hr \leftarrow BitArrayVal(bits)$$

defines (among other things)  $B!HC$  to be the formula obtained from  $HC$  by substituting  $BitArrayVal(bits)$  for  $hr$ .

Unfortunately, this specification is wrong. The value of  $BitArrayVal(b)$  is specified only if  $b$  is a function with domain  $0 \dots n$  for some natural number  $n$ . We don't know what  $BitArrayVal(\{\text{"abc"}\})$  equals. It might equal 7. If it did, then  $B!HC$  would allow a behavior in which the initial value of  $bits$  is  $\{\text{"abc"}\}$ . We must rule out this possibility by substituting for  $hr$  not  $BitArrayVal(bits)$ , but some expression  $HourVal(bits)$  whose value is an element of  $1 \dots 12$  only if  $b$  is a function in  $[(0 \dots 3) \rightarrow \{0, 1\}]$ . For example, we can write

$$\begin{aligned}
 HourVal(b) \triangleq & \text{ IF } b \in [(0 \dots 3) \rightarrow \{0, 1\}] \text{ THEN } BitArrayVal(b) \\
 & \text{ ELSE } 99 \\
 B \triangleq & \text{ INSTANCE } HourClock \text{ WITH } hr \leftarrow HourVal(bits)
 \end{aligned}$$

This defines  $B!HC$  to be the desired specification of the binary hour clock. Because  $HC$  is not satisfied by a behavior in which  $hr$  ever assumes the value 99,  $B$  is not satisfied by any behavior in which  $bits$  ever assumes a value not in the set  $[(0 \dots 3) \rightarrow \{0, 1\}]$ .

There is another way to use the specification  $HC$  of the hour clock to specify the binary hour clock. Instead of substituting for  $hr$  in the hour-clock specification, we first specify a system consisting of both an hour clock and a binary hour clock that keep the same time, and we then hide the hour clock. This specification has the form

$$(10.7) \quad \exists hr : IR \wedge HC$$

where  $IR$  is a temporal formula that is true iff  $bits$  is always the four-bit value representing the value of  $hr$ . This formula asserts that  $bits$  is the representation of  $hr$  as a four-bit array, for some choice of values for  $hr$  that satisfies  $HC$ .

---

MODULE *BinaryHourClock*

---

EXTENDS *Naturals*

VARIABLE *bits*

$H(hr) \triangleq$  INSTANCE *HourClock*

$BitArrayVal(b) \triangleq$  LET  $n \triangleq$  CHOOSE  $i \in Nat : \text{DOMAIN } b = 0 \dots (i - 1)$

$val[i \in 0 \dots n - 1] \triangleq$   $\text{Defines } val[i] \text{ to equal } b[0] * 2^0 + \dots + b[i] * 2^i.$

IF  $i = 0$  THEN  $b[0] * 2^0$  ELSE  $(b[i] * 2^i) + val[i - 1]$

IN  $val[n - 1]$

$HourVal(b) \triangleq$  IF  $b \in [(0 \dots 3) \rightarrow \{0, 1\}]$  THEN  $BitArrayVal(b)$

ELSE 99

$IR(b, h) \triangleq \Box(h = HourVal(b))$

$BHC \triangleq \exists hr : IR(bits, hr) \wedge H(hr)!HC$

---

Figure 10.3: A specification of a binary hour clock.

Using the definition of *HourVal* given above, we can define *IR* simply to equal  $\Box(h = HourVal(b))$ .

If *HC* is defined as in module *HourClock*, then (10.7) can't appear in a  $TLA^+$  specification. For *HC* to be defined in the context of the formula, the variable *hr* must be declared in that context. If *hr* is already declared, then it can't be used as the bound variable of the quantifier  $\exists$ . As usual, this problem is solved with parametrized instantiation. The complete  $TLA^+$  specification *BHC* of the binary hour clock appears in module *BinaryHourClock* of Figure 10.3 on this page.

### 10.8.2 Refining a Channel

As our second example of interface refinement, consider a system that interacts with its environment by sending numbers from 1 through 12 over a channel. We refine it to a lower-level system that is the same, except it sends a number as a sequence of four bits. Each bit is sent separately, starting with the left-most (most significant) one. For example, to send the number 3, the lower-level system sends the sequence of bits 0, 1, 0, 1. We specify both channels with the *Channel* module of Figure 3.2 on page 30, so each value that is sent must be acknowledged before the next one can be sent.

Suppose *HSpec* is the higher-level system's specification, and its channel is represented by the variable *h*. Let *l* be the variable representing the lower-level channel. We write the lower-level system's specification as

$$(10.8) \quad \exists h : IR \wedge HSpec$$



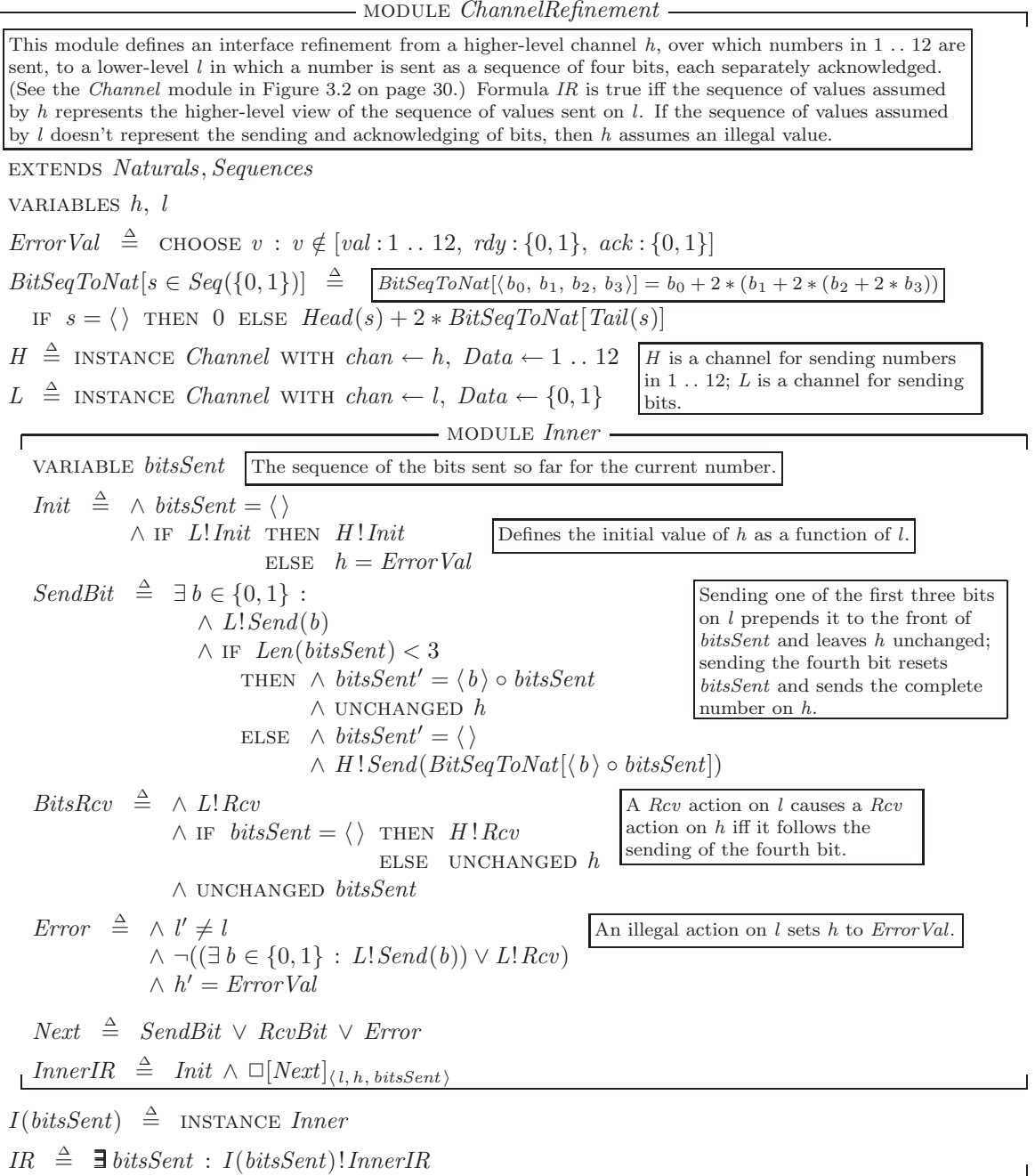


Figure 10.4: Refining a channel.

bits, then  $BitSeqToNat[s]$  is its numeric value interpreted as a binary number whose low-order bit is at the head of  $s$ .

There follows a submodule that defines the internal specification—the one with the internal variable  $bitsSent$  visible. The internal specification's initial predicate  $Init$  asserts that if  $l$  has a legal initial value, then  $h$  can have any legal initial value; otherwise  $h$  has an illegal value. Initially  $bitsSent$  is the empty sequence  $\langle \rangle$ . The internal specification's next-state action is the disjunction of three actions:

The use of a submodule to define an internal specification was introduced in the real-time hour clock specification of Section 9.1.

*SendBit* A *SendBit* step is one in which a bit is sent on  $l$ . If  $bitsSent$  has fewer than three elements, so fewer than three bits have already been sent, then the bit is prepended to the head of  $bitsSent$  and  $h$  is left unchanged. Otherwise, the value represented by the four bits sent so far, including the current bit, is sent on  $h$  and  $bitsSent$  is reset to  $\langle \rangle$ .

*BitsRcv* A *BitsRcv* step is one in which an acknowledgement is sent on  $l$ . It represents the sending of an acknowledgement on  $h$  iff this is an acknowledgement of the fourth bit, which is true iff  $bitsSent$  is the empty sequence.

*Error* An *Error* step is one in which an illegal change to  $l$  occurs. It sets  $h$  to an illegal value.

The inner specification *InnerIR* has the usual form. (There is no liveness requirement.) The outer module then instantiates the inner submodule with  $bitsSent$  as a parameter, and it defines *IR* to equal *InnerIR* with  $bitsSent$  hidden.

Now suppose we have a module *HigherSpec* that defines a specification *HSpec* of a system that communicates by sending numbers from 1 through 12 over a channel  $hchan$ . We obtain, as follows, a lower-level specification *LowerSpec* in which the numbers are sent as sequences of bits on a channel  $lchan$ . We first declare  $lchan$  and all the variables and constants of the *HigherSpec* module except  $hchan$ . We then write:

$$\begin{aligned} HS(hchan) &\triangleq \text{INSTANCE } HigherSpec \\ IR(h) &\triangleq \text{INSTANCE } ChannelRefinement \text{ WITH } l \leftarrow lchan \\ LowerSpec &\triangleq \exists h : IR(h)!Spec \wedge HS(h)!IR \end{aligned}$$

### 10.8.3 Interface Refinement in General

In the examples of the binary clock and of channel refinement, we defined a lower-level specification *LSpec* in terms of a higher-level one *HSpec* as:

$$(10.9) \quad LSpec \triangleq \exists h : IR \wedge HSpec$$

where  $h$  is a free variable of  $HSpec$  and  $IR$  is a relation between the  $h$  and the lower-level variable  $l$  of  $LSpec$ . We can view the internal specification  $IR \wedge HSpec$  as the composition of two components, as shown here:



We can regard  $IR$  as the specification of a component that transforms the lower-level behavior of  $l$  into the higher-level behavior of  $h$ . Formula  $IR$  is called an *interface refinement*.

In both examples, the interface refinement was independent of the system specification. It depended only on the representation of the interface—that is, on how the interaction between the system and its environment was represented. In general, for an interface refinement  $IR$  to be independent of the system using the interface, it should ascribe a behavior of the higher-level interface variable  $h$  to any behavior of the lower-level variable  $l$ . In other words, for any sequence of values for  $l$ , there should be some sequence of values for  $h$  that satisfy  $IR$ . This is expressed mathematically by the requirement that the formula  $\exists h : IR$  should be valid—that is, true for all behaviors.

So far, I have discussed refinement of a single interface variable  $h$  by a single variable  $l$ . This generalizes in the obvious way to the refinement of a collection of higher-level variables  $h_1, \dots, h_n$  by the variables  $l_1, \dots, l_m$ . The interface refinement  $IR$  specifies the values of the  $h_i$  in terms of the values of the  $l_j$  and perhaps of other variables as well. Formula (10.9) is replaced by

$$LSpec \triangleq \exists h_1, \dots, h_n : IR \wedge HSpec$$

A particularly simple type of interface refinement is a *data refinement*, in which  $IR$  has the form  $\Box P$ , where  $P$  is a state predicate that expresses the values of the higher-level variables  $h_1, \dots, h_n$  as functions of the values of the lower-level variables  $l_1, \dots, l_m$ . The interface refinement in our binary clock specification is a data refinement, where  $P$  is the predicate  $hr = HourVal(bits)$ . As another example, the two specifications of an asynchronous channel interface in Chapter 3 can each be obtained from the other by an interface refinement. The specification  $Spec$  of the *Channel* module (page 30) is equivalent to the specification obtained as a data refinement of the specification  $Spec$  of the *AsynchInterface* module (page 27) by letting  $P$  equal:

$$(10.10) \text{chan} = [val \mapsto val, rdy \mapsto rdy, ack \mapsto ack]$$

This formula asserts that *chan* is a record whose *val* component is the value of the variable *val*, whose *rdy* component is the value of the variable *rdy*, and whose *ack* component is the value of the variable *ack*. Conversely, specification  $Spec$  of the *AsynchInterface* module is equivalent to a data refinement of the specification

*Spec* of the *Channel* module. In this case, defining the state predicate  $P$  is a little tricky. The obvious choice is to let  $P$  be the formula *GoodVals* defined by:

$$\begin{aligned} \text{GoodVals} &\triangleq \wedge val = \text{chan.val} \\ &\quad \wedge rdy = \text{chan.rdy} \\ &\quad \wedge ack = \text{chan.ack} \end{aligned}$$

However, this can assert that *val*, *rdy*, and *ack* have good values even if *chan* has an illegal value—for example, if it is a record with more than three components. Instead, we let  $P$  equal

$$\text{IF } \text{chan} \in [val : \text{Data}, rdy : \{0, 1\}, ack : \{0, 1\}] \text{ THEN } \text{GoodVals} \\ \text{ELSE } \text{BadVals}$$

where *BadVals* asserts that *val*, *rdy*, and *ack* have some illegal values—that is, values that are impossible in a behavior satisfying formula *Spec* of module *AsynchInterface*. (We don't need such a trick when defining *chan* as a function of *val*, *rdy*, and *ack* because our definition (10.10) assures that the value of *chan* is legal iff the values of all three variables *val*, *rdy*, and *ack* are legal.)

Data refinement is the simplest form of interface refinement. In a more complicated interface refinement, the value of the higher-level variables cannot be expressed as a function of the current values of the lower-level variables. In the channel refinement example of Section 10.8.2, the number being sent on the higher-level channel depends on the values of bits that were previously sent on the lower-level channel, not just on the lower-level channel's current state.

We often refine both a system and its interface at the same time. For example, we may implement a specification  $H$  of a system that communicates by sending numbers over a channel with a lower-level specification  $LImpl$  of a system that sends individual bits. In this case,  $LImpl$  is not itself obtained from  $HSpec$  by an interface refinement. Rather,  $LImpl$  implements some specification  $LSpec$  that is obtained from  $HSpec$  by an interface refinement  $IR$ . In that case, we say that  $LImpl$  implements  $HSpec$  under the interface refinement  $IR$ .

### 10.8.4 Open-System Specifications

So far, I have been discussing interface refinement for complete-system specifications. Let's now consider what happens if the higher-level specification  $HSpec$  is the kind of open-system specification discussed in Section 10.7 above. For simplicity, we consider the refinement of a single higher-level interface variable  $h$  by a single lower-level variable  $l$ . The generalization to more variables will be obvious.

Let's suppose first that  $HSpec$  is a safety property, with no liveness condition. As explained in Section 10.7, the specification attributes each change to  $h$  either to the system or to the environment. Any change to a lower-level interface

variable  $l$  that produces a change to  $h$  is therefore attributed to the system or the environment. A bad change to  $h$  that is attributed to the environment makes  $HSpec$  true; a bad change that is attributed to the system makes  $HSpec$  false. Thus, (10.9) defines  $LSpec$  to be an open-system specification. For this to be a sensible specification, the interface refinement  $IR$  must ensure that the way changes to  $l$  are attributed to the system or environment is sensible.

If  $HSpec$  contains liveness conditions, then interface refinement can be more subtle. Suppose  $IR$  is the interface refinement defined in the *ChannelRefinement* module of Figure 10.4 on page 160, and suppose that  $HSpec$  requires that the system eventually send some number on  $h$ . Consider a behavior in which the system sends the first bit of a number on  $l$ , but the environment never acknowledges it. Under the interface refinement  $IR$ , this behavior is interpreted as one in which  $h$  never changes. Such a behavior fails to satisfy the liveness condition of  $HSpec$ . Thus, if  $LSpec$  is defined by (10.9), then failure of the environment to do something can cause  $LSpec$  to be violated, through no fault of the system.

In this example, we want the environment to be at fault if it causes the system to halt by failing to acknowledge any of the first three bits of a number sent by the system. (The acknowledgment of the fourth bit is interpreted by  $IR$  as the acknowledgment of a value sent on  $h$ , so blame for its absence is properly assigned to the environment.) Putting the environment at fault means making  $LSpec$  true. We can achieve this by modifying (10.9) to define  $LSpec$  as follows:

$$(10.11) LSpec \triangleq Liveness \Rightarrow \exists h : IR \wedge HSpec$$

where *Liveness* is a formula requiring that any bit sent on  $l$ , other than the last bit of a number, must eventually be acknowledged. However, if illegal values are sent on  $l$ , then we want the safety part of the specification to determine who is responsible. So, we want *Liveness* to be true in this case.

Here's one way to define *Liveness*. We use the definitions from the *Inner* submodule of module *ChannelRefinement*, where  $l$ ,  $h$ , and *bitsSent* are visible. The action that acknowledges receipt of one of the first three bits of the number is  $BitsRcv \wedge (bitsSent \neq \langle \rangle)$ . Weak fairness of this action asserts that the required acknowledgments must eventually be sent. For the case of bad values, recall that *InnerIR* implies that sending a bad value on  $l$  causes  $h$  to equal *ErrorVal*. We want *Liveness* to be true if this ever happens, which means if it eventually happens. We therefore add the following definition to the submodule *Inner* of the *ChannelRefinement* module:

$$\begin{aligned} InnerLiveness \triangleq & \quad \vee \wedge InnerIR \\ & \quad \wedge WF_{\langle l, h, bitsSent \rangle} (BitsRcv \wedge (bitsSent \neq \langle \rangle)) \\ & \quad \vee \Diamond (h = ErrorVal) \end{aligned}$$

To define *Liveness*, we just have to hide  $h$  and *bitsSent* in *InnerLiveness*. We

can do this, in a context in which  $l$  is declared, as follows:

$$\begin{aligned} ICR(h) &\triangleq \text{INSTANCE } ChannelRefinement \\ Liveness &\triangleq \exists h, bitsSent : ICR(h)!I(bitsSent)!InnerLiveness \end{aligned}$$

Now, suppose it is the environment that sends numbers over  $h$  and the system is supposed to acknowledge their receipt and then process them in some way. In this case, we want failure to acknowledge a bit to be a system error. So,  $LSpec$  should be false if  $Liveness$  is. The specification should then be

$$LSpec \triangleq Liveness \wedge (\exists h : IR \wedge HSpec)$$

Since  $h$  does not occur free in  $Liveness$ , this definition is equivalent to

$$LSpec \triangleq \exists h : Liveness \wedge IR \wedge HSpec$$

which has the form (10.9) if the interface refinement  $IR$  of (10.9) is taken to be  $Liveness \wedge IR$ . In other words, in this case, we make the liveness condition part of the interface refinement.

In general, if  $HSpec$  is an open-system specification that describes liveness as well as safety, then an interface refinement may have to take the form (10.11). Both  $Liveness$  and the liveness condition of  $IR$  may depend on which changes to the lower-level interface variable  $l$  are attributed to the system and which to the environment. For the channel refinement, this means that they will depend on whether the system or the environment is sending values on the channel.

## 10.9 Should You Compose?

When specifying a system, should we write a monolithic specification with a single next-state action, a closed-system composition that is the conjunction of specifications of individual components, or an open-system specification? The answer is that it usually makes little difference. For a real system, the definitions of the components' actions will take hundreds or thousands of lines. The different forms of specification differ only in the few lines where we assemble the initial predicates and next-state actions into the final formula.

If you are writing a specification from scratch, it's probably better to write a monolithic specification. It is usually easier to understand. Of course, there are exceptions. We write a real-time specification as the conjunction of an untimed specification and timing constraints; describing the changes to the system variables and the timers with a single next-state action usually makes the specification harder to understand.

Writing a composite specification may be sensible when you are starting from an existing specification. If you already have a specification of one component, you may want to write a separate specification of the other component and

compose the two specifications. If you have a higher-level specification, you may want to write a lower-level version as an interface refinement. However, these are rather rare situations. Moreover, it's likely to be just as easy to modify the original specification or re-use it in another way. For example, instead of conjoining a new component to the specification of an existing one, you can simply include the definition of the existing component's next-state action, with an EXTENDS or INSTANCE statement, as part of the new specification.

Composition provides a new way of writing a complete-system specification; it doesn't change the specification. The choice between a composite specification and a monolithic one is therefore ultimately a matter of taste. Disjoint-state compositions are generally straightforward and present no problems. Shared-state compositions can be tricky and require care.

Open-system specifications introduce a mathematically different form of specification. A closed-system specification  $E \wedge M$  and its open-system counterpart  $E \dashv M$  are not equivalent. If we really want a specification to serve as a legal contract between a user and an implementor, then we have to write an open-system specification. We also need open-system specifications if we want to specify and reason about systems built by composing off-the-shelf components with pre-existing specifications. All we can assume about such a component is that it satisfies a contract between the system builder and the supplier, and such a contract can be formalized only as an open-system specification. However, you are unlikely to encounter off-the-shelf component specifications during the early part of the twenty-first century. In the near future, open-system specifications are likely to be of theoretical interest only.

## Chapter 11

# Advanced Examples

It would be nice to provide an assortment of typical examples that cover most of the specification problems that arise in practice. However, there is no such thing as a typical specification. Every real specification seems to pose its own problems. But we can partition all specifications into two classes, depending on whether or not they contain VARIABLE declarations.

A specification with no variables defines data structures and operations on those structures. For example, the *Sequences* module defines various operations on sequences. When specifying a system, you may need some kind of data structure other than the ones provided by the standard modules like *Sequences* and *Bags*, described in Chapter 17. Section 11.1 gives some examples of data structure specifications.

A system specification contains variables that represent the system's state. We can further divide system specifications into two classes—high-level specifications that describe what it means for a system to be correct, and lower-level specifications that describe what the system actually does. In the memory example of Chapter 5, the linearizable memory specification of Section 5.3 is a high-level specification of correctness, while the write-through cache specification of Section 5.6 describes how a particular algorithm works. This distinction is not precise; whether a specification is high- or low-level is a matter of perspective. But it can be a useful way of categorizing system specifications.

Lower-level system specifications tend to be relatively straightforward. Once the level of abstraction has been chosen, writing the specification is usually just a matter of getting the details right when describing what the system does. Specifying high-level correctness can be much more subtle. Section 11.2 considers a high level specification problem—formally specifying a multiprocessor memory.

## 11.1 Specifying Data Structures

Most of the data structures required for writing specifications are mathematically very simple and are easy to define in terms of sets, functions, and records. Section 11.1.2 describes the specification of one such structure—a graph. On rare occasions, a specification might require sophisticated mathematical concepts. The only ones I know of are hybrid system specifications, discussed in Section 9.5. There, we used a module for describing the solutions to differential equations. That module is specified in Section 11.1.3 below. But, before specifying data structures, you should know how to make local definitions.

### 11.1.1 Local Definitions

In the course of specifying a system, we write lots of auxiliary definitions. A system specification may consist of a single formula *Spec*, but we define dozens of other identifiers in terms of which we define *Spec*. These other identifiers often have fairly common names—for example, the identifier *Next* is defined in many specifications. The different definitions of *Next* don't conflict with one another because, if a module that defines *Next* is used as part of another specification, it is usually instantiated with renaming. For example, the *Channel* module is used in module *InnerFIFO* on page 38 with the statement:

$$InChan \triangleq \text{INSTANCE } Channel \text{ WITH } \dots$$

The action *Next* of the *Channel* module is then instantiated as *InChan!Next*, so its definition doesn't conflict with the definition of *Next* in the *InnerFIFO* module.

A module that defines operations on a data structure is likely to be used in an *EXTENDS* statement, which does no renaming. The module might define some auxiliary operators that are used only to define the operators in which we're interested. For example, we need the *DifferentialEquations* module only to define the single operator *Integrate*. However, *Integrate* is defined in terms of other defined operators with names like *Nbhd* and *IsDeriv*. We don't want these definitions to conflict with other uses of those identifiers in a module that extends *DifferentialEquations*. So, we want the definitions of *Nbhd* and *IsDeriv* to be local to the *DifferentialEquations* module.<sup>1</sup>

TLA<sup>+</sup> provides a *LOCAL* modifier for making definitions local to a module. If a module *M* contains the definition

$$\text{LOCAL } Foo(x) \triangleq \dots$$


---

<sup>1</sup>We could use the *LET* construct to put these auxiliary definitions inside the definition of *Integrate*, but that trick wouldn't work if the *DifferentialEquations* module exported other operators besides *Integrate* that were defined in terms of *Nbhd* and *IsDeriv*.

then *Foo* can be used inside module *M* just like any ordinary defined identifier. However, a module that extends or instantiates *M* does not obtain the definition of *Foo*. That is, the statement `EXTENDS M` in another module does not define *Foo* in that module. Similarly, the statement

$$N \triangleq \text{INSTANCE } M$$

does not define *N!Foo*. The `LOCAL` modifier can also be applied to an instantiation. The statement

$$\text{LOCAL INSTANCE } \textit{Sequences}$$

in module *M* incorporates into *M* the definitions from the *Sequences* module. However, another module that extends or instantiates *M* does not obtain those definitions. Similarly, a statement like

$$\text{LOCAL } P(x) \triangleq \text{INSTANCE } N$$

makes all the instantiated definitions local to the current module.

The `LOCAL` modifier can be applied only to definitions and `INSTANCE` statements. It cannot be applied to a declaration or to an `EXTENDS` statement, so you *cannot* write either of the following:

$$\text{LOCAL CONSTANT } N$$

$$\text{LOCAL EXTENDS } \textit{Sequences}$$

These are not legal statements.

If a module has no `CONSTANT` or `VARIABLE` declarations, then extending it and instantiating it are equivalent. Thus, the two statements

$$\text{EXTENDS } \textit{Sequences} \qquad \text{INSTANCE } \textit{Sequences}$$

are equivalent.

In a module that defines general mathematical operators, I like to make all definitions local except for the ones that users of the module would expect. For example, users expect the *Sequences* module to define operators on sequences, such as *Append*. They don't expect it to define operators on numbers, such as *+*. The *Sequences* module uses *+* and other operators defined in the *Naturals* module. But instead of extending *Naturals*, it defines those operators with the statement

$$\text{LOCAL INSTANCE } \textit{Naturals}$$

The definitions of the operators from *Naturals* are therefore local to *Sequences*. A module that extends the *Sequences* module could then define *+* to mean something other than addition of numbers.

### 11.1.2 Graphs

A graph is an example of the kind of simple data structure often used in specifications. Let's now write a *Graphs* module for use in writing system specifications.

We must first decide how to represent a graph in terms of data structures that are already defined—either built-in TLA<sup>+</sup> data structures like functions, or ones defined in existing modules. Our decision depends on what kind of graphs we want to represent. Are we interested in directed graphs or undirected graphs? Finite or infinite graphs? Graphs with or without self-loops (edges from a node to itself)? If we are specifying graphs for a particular specification, the specification will tell us how to answer these questions. In the absence of such guidance, let's handle arbitrary graphs. My favorite way of representing both directed and undirected graphs is to specify arbitrary directed graphs, and to define an undirected graph as a directed graph that contains an edge iff it contains the opposite-pointing edge. Directed graphs have a pretty obvious representation: a directed graph consists of a set of nodes and a set of edges, where an edge from node  $m$  to node  $n$  is represented by the ordered pair  $\langle m, n \rangle$ .

In addition to deciding how to represent graphs, we must decide how to structure the *Graphs* module. The decision depends on how we expect the module to be used. For a specification that uses a single graph, it is most convenient to define operations on that specific graph. So, we want the *Graphs* module to have (constant) parameters *Node* and *Edge* that represent the sets of nodes and edges of a particular graph. A specification could use such a module with a statement

INSTANCE *Graphs* WITH *Node*  $\leftarrow \dots$ , *Edge*  $\leftarrow \dots$

where the “ $\dots$ ”s are the sets of nodes and edges of the particular graph appearing in the specification. On the other hand, a specification might use many different graphs. For example, it might include a formula that asserts the existence of a subgraph, satisfying certain properties, of some given graph  $G$ . Such a specification needs operators that take any graph as an argument—for example, a *Subgraph* operator defined so *Subgraph*( $G$ ) is the set of all subgraphs of a graph  $G$ . In this case, the *Graphs* module would have no parameters; specifications would incorporate it with an EXTENDS statement. Let's write this kind of module.

An operator like *Subgraph* takes a graph as an argument, so we have to decide how to represent a graph as a single value. A graph  $G$  consists of a set  $N$  of nodes and a set  $E$  of edges. A mathematician would represent  $G$  as the ordered pair  $\langle N, E \rangle$ . However,  $G.node$  is more perspicuous than  $G[1]$ , so we represent  $G$  as a record with *node* field  $N$  and *edge* field  $E$ .

Having made these decisions, it's easy to define any standard operator on graphs. We just have to decide what we should define. Here are some generally useful operators:

*IsDirectedGraph*(*G*)

True iff *G* is an arbitrary directed graph—that is, a record with *node* component *N* and *edge* component *E* such that *E* is a subset of  $N \times N$ . This operator is useful because a specification might want to assert that something is a directed graph. (To understand how to assert that *G* is a record with *node* and *edge* fields, see the definition of *IsChannel* in Section 10.3 on page 139.)

*DirectedSubgraph*(*G*)

The set of all subgraphs of a directed graph *G*. Alternatively, we could define *IsDirectedSubgraph*(*H*, *G*) to be true iff *H* is a subgraph of *G*. However, it's easy to express *IsDirectedSubgraph* in terms of *DirectedSubgraph*:

$$\text{IsDirectedSubgraph}(H, G) \equiv H \in \text{DirectedSubgraph}(G)$$

On the other hand, it's awkward to express *DirectedSubgraph* in terms of *IsDirectedSubgraph*:

$$\text{DirectedSubgraph}(G) \equiv \text{CHOOSE } S : \forall H : (H \in S) \equiv \text{IsDirectedSubgraph}(H, G)$$

Section 6.1 explains why in we can't define a set of all directed graphs, so we had to define the *IsDirectedGraph* operator.

*IsUndirectedGraph*(*G*)*UndirectedSubgraph*(*G*)

These are analogous to the operators for directed graphs. As mentioned above, an undirected graph is a directed graph *G* such that for every edge  $\langle m, n \rangle$  in *G.edge*, the inverse edge  $\langle n, m \rangle$  is also in *G.edge*. Note that, if *G* is a nontrivial undirected graph, then *DirectedSubgraph*(*G*) contains directed graphs that are not undirected graphs.

*Path*(*G*)

The set of all paths in *G*, where a path is any sequence of nodes that can be obtained by following edges in the direction they point. This definition is useful because many properties of a graph can be expressed in terms of its set of paths. It is convenient to consider the one-element sequence  $\langle n \rangle$  to be a path, for any node *n*.

*AreConnectedIn*(*m*, *n*, *G*)

True iff there is a path from node *m* to node *n* in *G*. The utility of this operator becomes evident when you try defining various common graph properties, like connectivity.

There are any number of other graph properties and classes of graphs that we might define. Let's define these two:

*IsStronglyConnected*(*G*)

True iff *G* is strongly connected, meaning that there is a path from any

node to any other node. For an undirected graph, strongly connected is equivalent to the ordinary definition of connected.

*IsTreeWithRoot*( $G, r$ )

True iff  $G$  is a tree with root  $r$ , where we represent a tree as a graph with an edge from each nonroot node to its parent. Thus, the parent of a nonroot node  $n$  equals:

CHOOSE  $m \in G.node : \langle n, m \rangle \in G.edge$

The *Graphs* module appears on the next page. By now, you should be able work out for yourself the meanings of all the definitions.

### 11.1.3 Solving Differential Equations

Section 9.5 on page 129 describes how to specify a hybrid system whose state includes a physical variable satisfying an ordinary differential equation. The specification uses an operator *Integrate* such that *Integrate*( $D, t_0, t_1, \langle x_0, \dots, x_{n-1} \rangle$ ) is the value at time  $t_1$  of the  $n$ -tuple

$$\langle x, dx/dt, \dots, d^{n-1}x/dt^{n-1} \rangle$$

where  $x$  is a solution to the differential equation

$$D[t, x, dx/dt, \dots, d^n x/dt^n] = 0$$

whose 0<sup>th</sup> through  $(n-1)$ <sup>st</sup> derivatives at time  $t_0$  are  $x_0, \dots, x_{n-1}$ . We assume that there is such a solution, and it is unique. Defining *Integrate* illustrates how sophisticated mathematics can be expressed in TLA<sup>+</sup>.

We start by defining some mathematical notation that we will use to define the derivative. As usual, we obtain from the *Reals* module the definitions of the set *Real* of real numbers and of the ordinary arithmetic operators. Let *PosReal* be the set of all positive reals:

$$PosReal \triangleq \{r \in Real : r > 0\}$$

and let *OpenInterval*( $a, b$ ) be the open interval from  $a$  to  $b$  (the set of numbers greater than  $a$  and less than  $b$ ):

$$OpenInterval(a, b) \triangleq \{s \in Real : (a < s) \wedge (s < b)\}$$

(Mathematicians usually write this set as  $(a, b)$ .) Let's also define *Nbhd*( $r, e$ ) to be the open interval of width  $2e$  centered at  $r$ , for  $e > 0$ :

$$Nbhd(r, e) \triangleq OpenInterval(r - e, r + e)$$

To explain the definitions, we need some notation for the derivative of a function. It's rather difficult to make mathematical sense of the usual notation  $df/dt$  for

| MODULE <i>Graphs</i>                                                                                                                                                                                                                   |                                                                               |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------|
| A module that defines operators on graphs. A directed graph is represented as a record whose <i>node</i> component is the set of nodes and whose <i>edge</i> component is the set of edges, where an edge is an ordered pair of nodes. |                                                                               |
| LOCAL INSTANCE <i>Naturals, Sequences</i>                                                                                                                                                                                              |                                                                               |
| $IsDirectedGraph(G) \triangleq$                                                                                                                                                                                                        | True iff $G$ is a directed graph.                                             |
| $\wedge G = [node \mapsto G.node, edge \mapsto G.edge]$<br>$\wedge G.edge \subseteq (G.node \times G.node)$                                                                                                                            |                                                                               |
| $DirectedSubgraph(G) \triangleq$                                                                                                                                                                                                       | The set of all (directed) subgraphs of a directed graph.                      |
| $\{H \in [node : SUBSET\ G.node, edge : SUBSET\ (G.node \times G.node)] :$<br>$IsDirectedGraph(H) \wedge H.edge \subseteq G.edge\}$                                                                                                    |                                                                               |
| $IsUndirectedGraph(G) \triangleq$                                                                                                                                                                                                      | An undirected graph is a directed graph in which every                        |
| $\wedge IsDirectedGraph(G)$                                                                                                                                                                                                            | edge has an oppositely-directed one.                                          |
| $\wedge \forall e \in G.edge : \langle e[2], e[1] \rangle \in G.edge$                                                                                                                                                                  |                                                                               |
| $UndirectedSubgraph(G) \triangleq$                                                                                                                                                                                                     | The set of (undirected) subgraphs of an undirected graph.                     |
| $\{H \in DirectedSubgraph(G) : IsUndirectedGraph(H)\}$                                                                                                                                                                                 |                                                                               |
| $Path(G) \triangleq$                                                                                                                                                                                                                   | The set of paths in $G$ , where a path is represented as a sequence of nodes. |
| $\{p \in Seq(G.node) : \wedge p \neq \langle \rangle$<br>$\wedge \forall i \in 1 \dots (Len(p) - 1) : \langle p[i], p[i + 1] \rangle \in G.edge\}$                                                                                     |                                                                               |
| $AreConnectedIn(m, n, G) \triangleq$                                                                                                                                                                                                   | True iff there is a path from $m$ to $n$ in graph $G$ .                       |
| $\exists p \in Path(G) : (p[1] = m) \wedge (p[Len(p)] = n)$                                                                                                                                                                            |                                                                               |
| $IsStronglyConnected(G) \triangleq$                                                                                                                                                                                                    | True iff graph $G$ is strongly connected.                                     |
| $\forall m, n \in G.node : (m \neq n) \Rightarrow AreConnectedIn(m, n, G)$                                                                                                                                                             |                                                                               |
| $IsTreeWithRoot(G, r) \triangleq$                                                                                                                                                                                                      | True if $G$ is a tree with root $r$ , where edges point                       |
| $\wedge IsDirectedGraph(G)$                                                                                                                                                                                                            | from child to parent.                                                         |
| $\wedge \forall n \in G.node : (\exists e \in G.edge : e[1] = n) \equiv (n \neq r)$                                                                                                                                                    |                                                                               |
| $\wedge \forall e, f \in G.edge : (e[1] = f[1]) \Rightarrow (e = f)$                                                                                                                                                                   |                                                                               |
| $\wedge \forall n \in G.node : AreConnectedIn(n, r, G)$                                                                                                                                                                                |                                                                               |

Figure 11.1: A module for specifying operators on graphs.

the derivative of  $f$ . (What exactly is  $t$ ?) So, let's use a mathematically simpler notation and write the  $n^{\text{th}}$  derivative of the function  $f$  as  $f^{(n)}$ . (We don't have to use  $\text{TLA}^+$  notation because differentiation does not appear explicitly in the specification.) Recall that  $f^{(0)}$ , the  $0^{\text{th}}$  derivative of  $f$ , equals  $f$ .

We can now start to define *Integrate*. If  $a$  and  $b$  are numbers, *InitVals* is an  $n$ -tuple of numbers, and  $D$  is a function from  $(n+2)$ -tuples of numbers to numbers, then

$$\text{Integrate}(D, a, b, \text{InitVals}) = \langle f^{(0)}[b], \dots, f^{(n-1)}[b] \rangle$$

where  $f$  is the function satisfying the following two conditions:

- $D[r, f^{(0)}[r], f^{(1)}[r], \dots, f^{(n)}[r]] = 0$ , for all  $r$  in some open interval containing  $a$  and  $b$ .
- $\langle f^{(0)}[a], \dots, f^{(n-1)}[a] \rangle = \text{InitVals}$

We want to define *Integrate*( $D, a, b, \text{InitVals}$ ) in terms of this function  $f$ , which we can specify using the *CHOOSE* operator. It's easiest to choose not just  $f$ , but its first  $n$  derivatives as well. So, we choose a function  $g$  such that  $g[i] = f^{(i)}$  for  $i \in 0 \dots n$ . The function  $g$  maps numbers in  $0 \dots n$  into functions. More precisely,  $g$  is an element of

$$[0 \dots n \rightarrow [\text{OpenInterval}(a - e, b + e) \rightarrow \text{Real}]]$$

for some positive  $e$ . It is the function in this set that satisfies the following conditions:

1.  $g[i]$  is the  $i^{\text{th}}$  derivative of  $g[0]$ , for all  $i \in 0 \dots n$ .
2.  $D[r, g[0][r], \dots, g[n][r]] = 0$ , for all  $r$  in  $\text{OpenInterval}(a - e, b + e)$ .
3.  $\langle g[0][a], \dots, g[n-1][a] \rangle = \text{InitVals}$

We now have to express these conditions formally.

To express the first condition, we will define *IsDeriv* so that *IsDeriv*( $i, df, f$ ) is true iff  $df$  is the  $i^{\text{th}}$  derivative of  $f$ . More precisely, this will be the case if  $f$  is a real-valued function on an open interval; we don't care what *IsDeriv*( $i, df, f$ ) equals for other values of  $f$ . Condition 1 is then:

$$\forall i \in 1 \dots n : \text{IsDeriv}(i, g[i], g[0])$$

To express the second condition formally, without the "...", we reason as follows:

$$\begin{aligned} & D[r, g[0][r], \dots, g[n][r]] \\ &= D[\langle r, g[0][r], \dots, g[n][r] \rangle] \\ &= D[\langle r \rangle \circ \langle g[0][r], \dots, g[n][r] \rangle] \\ &= D[\langle r \rangle \circ [i \in 1 \dots (n+1) \mapsto g[i-1][r]]] \end{aligned}$$

See Section 15.1.7 on page 285.

Tuples are sequences

An  $(n+1)$ -tuple is a function with domain  $1 \dots n+1$ .

The third condition is simply:

$$\forall i \in 1 \dots n : g[i-1][a] = \text{InitVals}[i]$$

We can therefore write the formula specifying  $g$  as:

$$\begin{aligned} \exists e \in \text{PosReal} : \wedge g \in [0 \dots n \rightarrow [\text{OpenInterval}(a-e, b+e) \rightarrow \text{Real}]] \\ \wedge \forall i \in 1 \dots n : \wedge \text{IsDeriv}(i, g[i], g[0]) \\ \wedge g[i-1][a] = \text{InitVals}[i] \\ \wedge \forall r \in \text{OpenInterval}(a-e, b+e) : \\ D[\langle r \rangle \circ [i \in 1 \dots (n+1) \mapsto g[i-1][r]]] = 0 \end{aligned}$$

where  $n$  is the length of  $\text{InitVals}$ . The value of  $\text{Integrate}(D, a, b, \text{InitVals})$  is the tuple  $\langle g[0][b], \dots, g[n-1][b] \rangle$ , which can be written formally as

$$[i \in 1 \dots n \mapsto g[i-1][b]]$$

To complete the definition of  $\text{Integrate}$ , we now define the operator  $\text{IsDeriv}$ . It's easy to define the  $i^{\text{th}}$  derivative inductively in terms of the first derivative. So, we define  $\text{IsFirstDeriv}(df, f)$  to be true iff  $df$  is the first derivative of  $f$ , assuming that  $f$  is a real-valued function whose domain is an open interval. Our definition actually works if the domain of  $f$  is any open set.<sup>2</sup> Elementary calculus tells us that  $df[r]$  is the derivative of  $f$  at  $r$  iff

$$df[r] = \lim_{s \rightarrow r} \frac{f[s] - f[r]}{s - r}$$

The classical “ $\delta$ - $\epsilon$ ” definition of the limit states that this is true iff, for every  $\epsilon > 0$ , there is a  $\delta > 0$  such that  $0 < |s - r| < \delta$  implies:

$$\left| df[r] - \frac{f[s] - f[r]}{s - r} \right| < \epsilon$$

Stated formally, this condition is:

$$\begin{aligned} \forall \epsilon \in \text{PosReal} : \\ \exists \delta \in \text{PosReal} : \\ \forall s \in \text{Nbhd}(r, \delta) \setminus \{r\} : \left( df[r] - \frac{f[s] - f[r]}{s - r} \right) \in \text{Nbhd}(df[r], \epsilon) \end{aligned}$$

We define  $\text{IsFirstDeriv}(df, f)$  to be true iff the domains of  $df$  and  $f$  are equal, and this condition holds for all  $r$  in their domain.

The definitions of  $\text{Integrate}$  and all the other operators introduced above appear in the *DifferentialEquations* module of Figure 11.2 on the next page. The `LOCAL` construct described in Section 11.1.1 above is used to make all these definitions local to the module, except for the definition of  $\text{Integrate}$ .

<sup>2</sup>A set  $S$  is open iff for every  $r \in S$  there exists an  $\epsilon > 0$  such that the interval from  $r - \epsilon$  to  $r + \epsilon$  is contained in  $S$ .

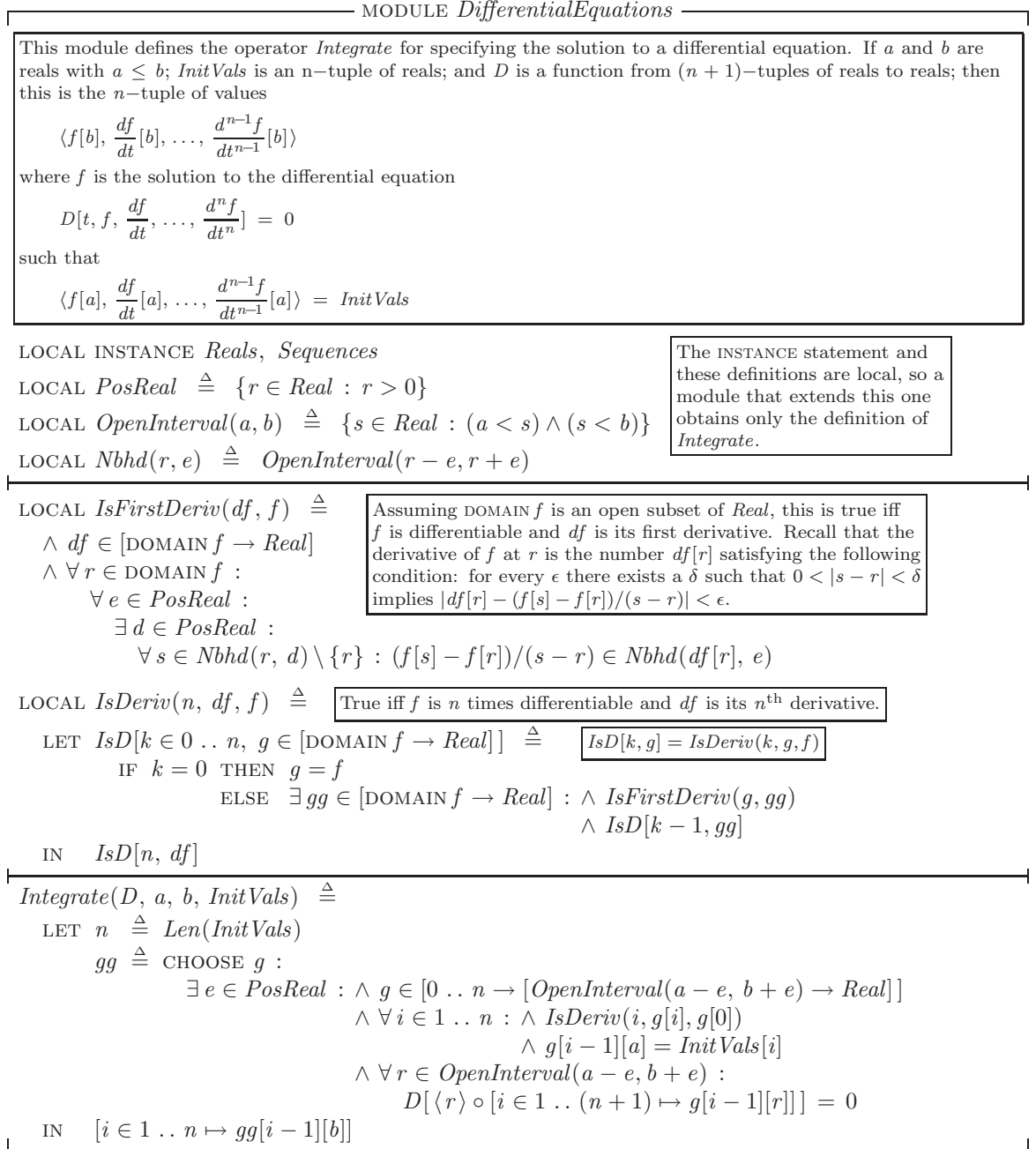


Figure 11.2: A module for specifying the solution to a differential equation.

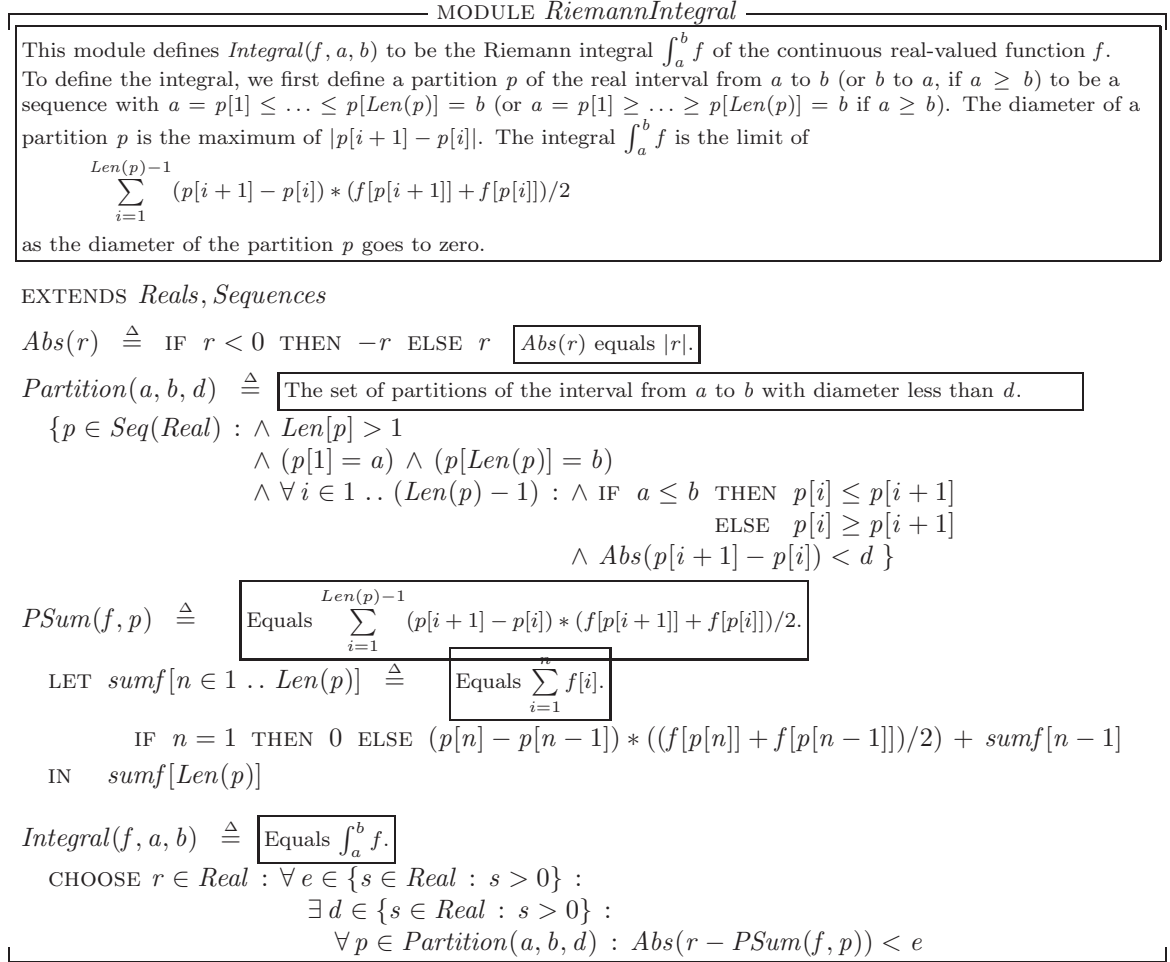


Figure 11.3: A specification of the Riemann integral.

### 11.1.4 The Riemann Integral

To demonstrate how easy it is to formalize ordinary math, we give a specification of the Riemann integral—the definite integral of elementary calculus. Though sometimes written  $\int_a^b f[x] dx$ , it's pretty hard (though by no means impossible) to make rigorous sense of the  $dx$ , so careful mathematicians usually write this integral as  $\int_a^b f$ . The specification is in Figure 11.3 on this page.

## 11.2 Other Memory Specifications

Section 5.3 specifies a multiprocessor memory. The specification is unrealistically simple for three reasons: a processor can have only one outstanding request at a time, the basic correctness condition is too restrictive, and only simple read and write operations are provided. (Real memories provide many other operations, such as partial-word writes and cache prefetches.) I will specify a memory that allows multiple outstanding requests and has a realistic, weaker correctness condition. To keep the specification short, I will still consider only the simple operations of reading and writing one word of memory.

### 11.2.1 The Interface

A modern processor performs multiple instructions concurrently. It can begin new memory operations before previous ones have been completed. The memory responds to a request as soon as it can; it need not respond to different requests in the order that they were issued.

The first thing we must do to specifying a memory is determine the interface. The interface we choose depends on the purpose of the specification. There are many different reasons why we might be specifying a multiprocessor memory. We could be specifying a computer architecture, or the semantics of a programming language. I will suppose that we are specifying the memory of an actual computer.

A processor issues a request to a memory system by setting some register. I assume that each processor has a set of registers through which it communicates with the memory. Each register has three fields: an *adr* field that holds an address, a *val* field that holds a word of memory, and an *op* field that indicates what kind of operation, if any, is in progress. The processor can issue a command using a register whose *op* field equals “Free”. It sets the *op* field to “Rd” or “Wr” to indicate the operation; it sets the *adr* field to the address of the memory word; and, for a write, it sets the *val* field to the value being written. (On a read, the processor can set the *val* field to any value.) The memory responds by setting the *op* field back to “Free” and, for a read, setting the *val* field to the value read. (The memory does not change the *val* field when responding to a write.)

Module *RegisterInterface* in Figure 11.4 on the next page contains some declarations and definitions for specifying the interface. It declares the constants *Adr*, *Val*, and *Proc*, which are the same as in the memory interface of Section 5.1, and the constant *Reg*, which is the set of registers. (More precisely, *Reg* is a set of register identifiers.) A processor has a separate register corresponding to each element of *Reg*. The variable *regFile* represents the processors’ registers, *regFile*[*p*][*r*] being register *r* of processor *p*. The module also defines the sets of requests and register values, as well as a type invariant for *regFile*.

| MODULE <i>RegisterInterface</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |                                                                                                              |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------|
| CONSTANT <i>Adr</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               | The set of memory addresses.                                                                                 |
| <i>Val</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        | The set of memory-word values.                                                                               |
| <i>Proc</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | The set of processors used by a processor.                                                                   |
| <i>Reg</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          | The set of registers.                                                                                        |
| VARIABLE <i>regFile</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             | <i>regFile</i> [ <i>p</i> ][ <i>r</i> ] represents the contents of register <i>r</i> of processor <i>p</i> . |
| $RdRequest \triangleq [adr : Adr, val : Val, op : \{ \text{“Rd”} \}]$<br>$WrRequest \triangleq [adr : Adr, val : Val, op : \{ \text{“Wr”} \}]$<br>$FreeRegValue \triangleq [adr : Adr, val : Val, op : \{ \text{“Free”} \}]$<br>$Request \triangleq RdRequest \cup WrRequest$ The set of all possible requests.<br>$RegValue \triangleq Request \cup FreeRegValue$ The set of all possible register values.<br>$RegFileTypeInvariant \triangleq$ The type correctness invariant for <i>regFile</i> .<br>$regFile \in [Proc \rightarrow [Reg \rightarrow RegValue]]$ |                                                                                                              |

Figure 11.4: A module for specifying a register interface to a memory.

### 11.2.2 The Correctness Condition

Section 5.3 specifies what is called a linearizable memory. In a linearizable memory, a processor never has more than one outstanding request. The correctness condition for the memory can be stated as:

The result of any execution is the same as if the operations of all the processors were executed in some sequential order, and each operation is executed between the request and the response.

The second clause, which requires the system to act as if each operation were executed between its request and its response, is both too weak and too strong for our specification. It’s too weak because it says nothing about the execution order of two operations from the same processor unless one is issued after the other’s response. For example, suppose a processor *p* issues a write and then a read to the same address. We want the read to obtain either the value *p* just wrote, or a value written by another processor—even if *p* issues the read before receiving the response for the write. This isn’t guaranteed by the condition. The second clause is too strong because it places unnecessary ordering constraints on operations issued by different processors. If operations *A* and *B* are issued by two different processors, then we don’t need to require that *A* precedes *B* in the execution order just because *B* was requested after *A*’s response.

We modify the second clause to require that the system act as if operations of each individual processor were executed in the order that they were issued,

obtaining the condition:

The result of any execution is the same as if the operations of all the processors were executed in some sequential order, and the operations of each individual processor appear in this sequence in the order in which the requests were issued.

In other words, we require that the values returned by the reads can be explained by some total ordering of the operation executions that is consistent with the order in which each processor issued its requests. There are a number of different ways of formalizing this condition; they differ in how bizarre the explanation may be. The differences can be described in terms of whether or not certain scenarios are permitted. In describing scenarios, I let  $Wr_p(a, v)$  represent a write operation of value  $v$  to address  $a$  by processor  $p$ , and  $Rd_p(a, v)$  represent a read of  $a$  by  $p$  that returns the value  $v$ .

The first decision we must make is whether all operations in an infinite behavior must be ordered, or if the ordering must exist only at each finite point during the behavior. Consider a scenario in which each of two processes writes its own value to the same address and then keeps reading that value forever:

Process  $p$ :  $Wr_p(a, v1), Rd_p(a, v1), Rd_p(a, v1), Rd_p(a, v1), \dots$   
 Process  $q$ :  $Wr_q(a, v2), Rd_q(a, v2), Rd_q(a, v2), Rd_q(a, v2), \dots$

In these scenarios, I assume that values and addresses with different names are different.

At each point in the execution, we can explain the values returned by the reads with a total order in which all the operations of either processor precede all the operations of the other. However, there is no way of explaining the entire infinite scenario with a single total order. In this scenario, neither processor ever sees the value written by the other. Since a multiprocessor memory is supposed to allow processors to communicate, we disallow this scenario.

The second decision we must make is whether the memory is allowed to predict the future. Consider this scenario:

Processor  $p$ :  $Wr_p(a, v1), Rd_p(a, v2)$   
 Processor  $q$ :  $Wr_q(a, v2)$

Here,  $q$  issues its write of  $v2$  after  $p$  has obtained the result of its read. The scenario is explained by the ordering  $Wr_p(a, v1), Wr_q(a, v2), Rd_p(a, v2)$ . However, this is a bizarre explanation because, to return the value  $v2$  for  $p$ 's read, the memory had to predict that another processor would write  $v2$  some time in the future. Since a real memory can't predict what requests will be issued in the future, such a behavior cannot be produced by a correct implementation. We can therefore rule out the scenario as unreasonable. Alternatively, since no correct implementation can produce it, there's no need to outlaw the scenario.

If we don't allow the memory to predict the future, then it must always be able to explain the values read in terms of the writes that have been issued so

far. In this case, we have to decide whether the explanations must be stable. For example, suppose a scenario begins as follows:

Processor  $p$ :  $Wr_p(a1, v1), Rd_p(a1, v3)$   
 Processor  $q$ :  $Wr_q(a2, v2), Wr_q(a1, v3)$

At this point, the only explanation for  $p$ 's read  $Rd_p(a1, v3)$  is that  $q$ 's write  $Wr_q(a1, v3)$  preceded it, which implies that  $q$ 's other write  $Wr_q(a2, v2)$  also preceded the read. Hence, if  $p$  now reads  $a2$ , it must obtain the value  $v2$ . But suppose the scenario continues as follows, with another processor  $r$  joining in:

Processor  $p$ :  $Wr_p(a1, v1), Rd_p(a1, v3), Rd_p(a2, v0)$   
 Processor  $q$ :  $Wr_q(a2, v2), Wr_q(a1, v3)$   
 Processor  $r$ :  $Wr_r(a1, v3)$

We can explain this scenario with the following ordering of the operations:

$Wr_p(a1, v1), Wr_r(a1, v3), Rd_p(a1, v3), Rd_p(a2, v0), Wr_q(a2, v2), Wr_q(a1, v3)$

In this explanation, processor  $r$  provided the value of  $a1$  read by  $p$ , and  $p$  read the initial value  $v0$  of memory address  $a2$ . The explanation of the value of  $a1$  read by  $p$  thus changes in mid-execution. To write a specification, we must decide whether or not to allow explanations to change in this way.

### 11.2.3 A Serial Memory

We first specify a memory that cannot predict the future and cannot change its explanations. There seems to be no standard name for such a memory; I'll call it a *serial* memory.

Our informal correctness condition is in terms of the sequence of all operations that have ever been issued. There is a general method of formalizing such a condition that works for specifying many different kinds of systems. We add an internal variable  $opQ$  that records the history of the execution. For each processor  $p$ , the value of  $opQ[p]$  is a sequence whose  $i^{\text{th}}$  element,  $opQ[p][i]$ , describes the  $i^{\text{th}}$  request issued by  $p$ , the response to that request (if it has been issued), and any other information about the operation needed to express the correctness condition. If necessary, we can also add other internal variables to record additional information not readily associated with individual requests.

For a system with the kind of register interface we are using, the next-state relation has the form

$$(11.1) \vee \exists proc \in Proc, reg \in Reg : \vee \exists req \in Request : IssueRequest(proc, req, reg) \\ \vee RespondToRequest(proc, reg) \\ \vee Internal$$

where the component actions do the following:

*IssueRequest*(*proc*, *req*, *reg*)

The action with which processor *proc* issues a request *req* in register *reg*.

*RespondToRequest*(*proc*, *reg*)

The action with which the system responds to a request in processor *proc*'s register *reg*.

*Internal*

An action that changes only the internal state.

Liveness properties are asserted by fairness conditions on the *RespondToRequest* and *Internal* actions.

A general trick for writing the specification is to choose the internal state so the safety part of the correctness condition can be expressed by the formula  $\Box P$  for some state predicate  $P$ . We guarantee that  $P$  is always true by letting  $P'$  be a conjunct of each action. I'll use this approach to specify the serial memory, taking for  $P$  a state predicate *Serializable*.

We want to require that the value returned by each read is explainable as the value written by some operation already issued, or as the initial value of the memory. Moreover, we don't want this explanation to change. We therefore add to the *opQ* entry for each completed read a *source* field that indicates where the value came from. This field is set by the *RespondToRequest* action.

We want all operations in an infinite behavior eventually to be ordered. This means that, for any two operations, the memory must eventually decide which one precedes the other—and it must stick to that decision. We introduce an internal variable *opOrder* that describes the ordering of operations to which the memory has already committed itself. An *Internal* step changes only *opOrder*, and it can only enlarge the ordering.

The predicate *Serializable* used to specify the safety part of the correctness condition describes what it means for *opOrder* to be a correct explanation. It asserts that there is some consistent total ordering of the operations that satisfies the following conditions:

- It extends *opOrder*.
- It orders all operations from the same processor in the order that they were issued.
- It orders operations so that the source of any read is the latest write to the same address that precedes the read, and is the initial value iff there is no such write.

We now translate this informal sketch of the specification into TLA<sup>+</sup>. We first choose the types of the variables *opQ* and *opOrder*. To describe the source field of a read and an order on operations, we define a set *opId* of values that identify the operations that have been issued. An operation is identified by a

pair  $\langle p, i \rangle$  where  $p$  is a processor and  $i$  is a position in the sequence  $opQ[p]$ . (The set of all such positions  $i$  is  $\text{DOMAIN } opQ[p]$ .) We let the corresponding element of  $opId$  be the record with  $proc$  field  $p$  and  $idx$  field  $i$ . Writing the set of all such records is a bit tricky because the possible values of the  $idx$  field depends on the  $proc$  field. The simplest way to define it is first to define  $OpId$  to be the larger set of records whose  $idx$  field can be any value, and then to define  $opId$  to be a subset of  $OpId$ :

$$\begin{aligned} OpId &\triangleq [proc : Proc, idx : Nat] \\ opId &\triangleq \{oiv \in OpId : oiv.idx \in \text{DOMAIN } opQ[oiv.proc]\} \end{aligned}$$

For convenience, we define  $opIdQ(oi)$  to be the value of the  $opQ$  entry identified by an element  $oi$  of  $opId$ :

$$opIdQ(oi) \triangleq opQ[oi.proc][oi.idx]$$

The source of a value need not be an operation; it can also be the initial contents of the memory. The latter possibility is represented by letting the *source* field of the  $opQ$  entry have the special value *InitWr*. We then let  $opQ$  be an element of  $[Proc \rightarrow Seq(opVal)]$ , where  $opVal$  is the union of three sets:

$[req : Request, reg : Reg]$

Represents an active request in the register of the requesting processor indicated by the *reg* field.

$[req : WrRequest, reg : \{Done\}]$

Represents a completed write request, where *Done* is a special value that is not a register.

$[req : RdRequest, reg : \{Done\}, source : opId \cup \{InitWr\}]$

Represents a completed read request whose value came from the operation indicated by the *source* field, or from the initial value of the memory location if the *source* field equals *InitWr*.

*WrRequest*, *RdRequest*, and *Request* are defined in module *RegisterInterface* on page 179.

Observe that  $opId$  and  $opVal$  are state functions whose values depend upon the value of the variable  $opQ$ .

We need to specify the initial value of memory. A program generally cannot assume anything about the memory's initial value, except that every address does contain a value in *Val*. So, the initial value of memory can be any element of  $[Adr \rightarrow Val]$ . We declare an “internal” constant *InitMem*, whose value is the initial memory value. In the final specification, *InitMem* will be hidden along with the internal variables  $opQ$  and  $opOrder$ . We hide a constant with ordinary existential quantification  $\exists$ . The requirement that *InitMem* is a function from addresses to values could be made part of the initial predicate, but it's more natural to express it in the quantifier. The final specification will therefore have the form:

$$\exists InitMem \in [Adr \rightarrow Val] : \exists opQ, opOrder : \dots$$

For later use, we define  $goodSource(oi)$  to be the set of plausible values for the source of a read operation  $oi$  in  $opId$ . By a plausible value, I mean either  $InitWr$  or a write to the same address that  $oi$  reads. It will be an invariant of the specification that the source of any completed read operation  $oi$  is an element of  $goodSource(oi)$ . Moreover, the value returned by a completed read operation must come from its source. If the source is  $InitWr$ , then the value must come from  $InitMem$ ; otherwise, it must come from the source request's  $val$  field. To express this formally, observe that the  $opQ$  entries only of completed reads have a  $source$  field. Since a record has a  $source$  field iff the string “source” is in its domain, we can write this invariant as:

Section 5.2 on page 48 explains that a record is a function whose domain is a set of strings.

$$\begin{aligned}
 (11.2) \quad \forall oi \in opId : \\
 & (“source” \in \text{DOMAIN } opIdQ(oi)) \Rightarrow \\
 & \quad \wedge opIdQ(oi).source \in goodSource(oi) \\
 & \quad \wedge opIdQ(oi).req.val = \text{IF } opIdQ(oi).source = InitWr \\
 & \quad \quad \text{THEN } InitMem[opIdQ(oi).req.adr] \\
 & \quad \quad \text{ELSE } opIdQ(opIdQ(oi).source).req.val
 \end{aligned}$$

We now choose the type of  $opOrder$ . We usually denote an ordering relation by an operator such as  $\prec$ , writing  $A \prec B$  to mean that  $A$  precedes  $B$ . However, the value of a variable cannot be an operator. So, we must represent an ordering relation as a set or a function. Mathematicians usually describe a relation  $\prec$  on a set  $S$  as a set  $R$  of ordered pairs of elements in  $S$ , with  $\langle A, B \rangle$  in  $R$  iff  $A \prec B$ . So, we let  $opOrder$  be a subset of  $opId \times opId$ , where  $\langle oi, oj \rangle \in opOrder$  means that  $oi$  precedes  $oj$ .

The difference between operators and functions is discussed in Section 6.4 on page 69.

Our internal state is redundant because, if register  $r$  of processor  $p$  contains an uncompleted operation, then there is an  $opQ$  entry that points to the register and contains the same request. This redundancy means that the following relations among the variables are invariants of the specification:

- If an  $opQ$  entry's  $reg$  field is not equal to  $Done$ , then it denotes a register whose contents is the entry's  $req$  field.
- The number of  $opQ$  entries pointing to a register equals 1 if the register contains an active operation, otherwise it equals 0.

In the specification, we combine this condition, formula (11.2), and the type invariant into a single state predicate  $DataInvariant$ .

Having chosen the types of the variables, we can now define the initial predicate  $Init$  and the predicate  $Serializable$ . The definition of  $Init$  is easy. We define  $Serializable$  in terms of  $totalOpOrder$ , the set of all total orderings of  $opId$ . A relation  $\prec$  is a total ordering of  $opId$  iff it satisfies the following three conditions, for any  $oi$ ,  $oj$ , and  $ok$  in  $opId$ .

Totality: Either  $oi = oj$ ,  $oi \prec oj$ , or  $oj \prec oi$ .

Transitivity:  $oi \prec oj$  and  $oj \prec ok$  imply  $oi \prec ok$ .

Irreflexivity:  $oi \not\prec oi$ .

The predicate *Serializable* asserts that there is a total ordering of *opId* satisfying the three conditions on page 182. We can express this formally as the assertion that there exists an *R* in *totalOpOrder* satisfying:

|                                                                                                                                                                                                                                                                                                                                                            |                                                                                                                                                                                                   |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $\wedge \text{opOrder} \subseteq R$                                                                                                                                                                                                                                                                                                                        | $R$ extends <i>opOrder</i>                                                                                                                                                                        |
| $\wedge \forall oi, oj \in \text{opId} :$<br>$(oi.\text{proc} = oj.\text{proc}) \wedge (oi.\text{idx} < oj.\text{idx}) \Rightarrow (\langle oi, oj \rangle \in R)$                                                                                                                                                                                         | $R$ correctly orders operations from the same processor.                                                                                                                                          |
| $\wedge \forall oi \in \text{opId} :$<br>$(\text{"source"} \in \text{DOMAIN } \text{opIdQ}(oi))$<br>$\Rightarrow \neg(\exists oj \in \text{goodSource}(oi) :$<br>$\quad \wedge \langle oj, oi \rangle \in R$<br>$\quad \wedge (\text{opIdQ}(oi).\text{source} \neq \text{InitWr}) \Rightarrow (\langle \text{opIdQ}(oi).\text{source}, oj \rangle \in R))$ | For every completed read <i>oi</i> in <i>opId</i> , there is no write <i>oj</i> to the same address that (i) precedes <i>oi</i> and (ii) follows the source if that source is not <i>InitWr</i> . |

We allow each step to extend *opOrder* to any relation on *opId* that satisfies *Serializable*. We do this by letting every subaction of the next-state action specify *opOrder'* with the conjunct *UpdateOpOrder*, defined by:

$$\begin{aligned}
 \text{UpdateOpOrder} &\triangleq \wedge \text{opOrder}' \subseteq (\text{opId}' \times \text{opId}') \\
 &\quad \wedge \text{opOrder} \subseteq \text{opOrder}' \\
 &\quad \wedge \text{Serializable}'
 \end{aligned}$$

The next-state action has the generic form of formula (11.1) on page 181. We split the *RespondToRequest* action into the disjunction of separate *RespondToWr* and *RespondToRd* actions that represent responding to writes and reads, respectively. *RespondToRd* is the most complicated of the next-state action's subactions, so let's examine its definition. The definition has the form:

$$\begin{aligned}
 \text{RespondToRd}(\text{proc}, \text{reg}) &\triangleq \\
 \text{LET } \text{req} &\triangleq \text{regFile}[\text{proc}][\text{reg}] \\
 \text{idx} &\triangleq \text{CHOOSE } i \in \text{DOMAIN } \text{opQ}[\text{proc}] : \text{opQ}[\text{proc}][i].\text{reg} = \text{reg} \\
 \text{IN } &\dots
 \end{aligned}$$

This defines *req* to be the request in the register and *idx* to be an element in the domain of *opQ*[*proc*] such that *opQ*[*proc*][*idx*].*reg* equals *reg*. If the register is not free, then there is exactly one such value *idx*; and *opQ*[*proc*][*idx*].*reg*, the *idx*<sup>th</sup> request issued by *proc*, equals *req*. (We don't care what *idx* equals if the register is free.) The IN expression begins with the enabling condition:

$$\wedge \text{req.op} = \text{"Rd"}$$

which asserts that the register is not free and it contains a read request. The next conjunct of the IN expression is:

$$\begin{aligned} & \wedge \exists src \in goodSource([proc \mapsto proc, idx \mapsto idx]) : \\ & \quad LET \ val \triangleq IF \ src = InitWr \ THEN \ InitMem[req.adr] \\ & \quad \quad \quad ELSE \ opIdQ(src).req.val \\ & \quad IN \ \dots \end{aligned}$$

It asserts the existence of a value  $src$ , which will be the source of the value returned by the read; and it defines  $val$  to be that value. If the source is the initial value of memory, then the value is obtained from  $InitMem$ ; otherwise, it is obtained from the source request's  $val$  field. The inner IN expression has two conjuncts that specify the values of  $regFile'$  and  $opQ'$ . The first conjunct asserts that the register's  $val$  field is set to  $val$  and its  $op$  field is set to “Free”, indicating that the register is made free.

$$\begin{aligned} & \wedge regFile' = [regFile \text{ EXCEPT } ![proc][reg].val = val, \\ & \quad \quad \quad ![proc][reg].op = \text{“Free”}] \end{aligned}$$

The second conjunct of the inner IN expression describes the new value of  $opQ$ . Only the  $idx^{\text{th}}$  element of  $opQ[proc]$  is changed. It is set to a record whose  $req$  field is the same as the original request  $req$ , except that its  $val$  field is equal to  $val$ ; whose  $reg$  field equals  $Done$ ; and whose  $source$  field equals  $src$ .

$$\begin{aligned} & \wedge opQ' = [opQ \text{ EXCEPT } ![proc][idx] = [req \mapsto [req \text{ EXCEPT } !.val = val], \\ & \quad \quad \quad reg \mapsto Done, \\ & \quad \quad \quad source \mapsto src]] \end{aligned}$$

Finally, the outer IN clause ends with the conjunct

$$\wedge UpdateOpOrder$$

that determines the value of  $opOrder'$ . It also implicitly determines the possible choices of the source of the read—that is, the value of  $opQ'[proc][idx].source$ . For some choices of this value allowed by the second outer conjunct, there will be no value of  $opOrder'$  satisfying  $UpdateOpOrder$ . The conjunct  $UpdateOpOrder$  rules out those choices for the source.

The definitions of the other subactions of the next-state action—*IssueRequest*, *RespondtoWr*, and *Internal*—are simpler and I won't explain them.

Having finished the initial predicate and the next-state action, we must determine the liveness conditions. The first condition is that the memory must eventually respond to every operation. The response to a request in register  $reg$  of processor  $proc$  is produced by a  $RespondToWr(proc, reg)$  or  $RespondToRd(proc, reg)$  action. So, the obvious way to express this condition is:

$$\begin{aligned} & \forall proc \in Proc, reg \in Reg : \\ & \quad WF_{\langle \dots \rangle}(RespondToWr(proc, reg) \vee RespondToRd(proc, reg)) \end{aligned}$$

For this fairness condition to imply that the response is eventually issued, a  $RespondToWr(proc, reg)$  or  $RespondToRd(proc, reg)$  step must be enabled whenever there is an uncompleted operation in  $proc$ 's register  $reg$ . It isn't completely obvious that a  $RespondToRd(proc, reg)$  step is enabled when there is a read operation in the register, since the step is enabled only if there exist a source for the read and a value of  $opOrder'$  that satisfy  $Serializable'$ . The required source and value do exist because  $Serializable$ , which holds in the first state of the step, implies the existence of a correct total ordering of all the operations; this ordering can be used to choose a source and a relation  $opOrder'$  that satisfy  $Serializability'$ .

The second liveness condition asserts that the memory must eventually commit to an ordering for every pair of operations. It is expressed as a fairness condition, for every pair of distinct operations  $oi$  and  $oj$  in  $opId$ , on an *Internal* action that makes  $oi$  either precede or follow  $oj$  in the ordering  $opOrder'$ . A first attempt at this condition is

$$(11.3) \quad \forall oi, oj \in opId : (oi \neq oj) \Rightarrow WF_{\langle \dots \rangle}(\wedge Internal \\ \wedge (\langle oi, oj \rangle \in opOrder') \vee (\langle oj, oi \rangle \in opOrder'))$$

However, this isn't correct. In general, a formula  $\forall x \in S : F$  is equivalent to  $\forall x : (x \in S) \Rightarrow F$ . Hence, (11.3) is equivalent to the assertion that the following formula holds, for all constant values  $oi$  and  $oj$ :

$$(oi \in opId) \wedge (oj \in opId) \Rightarrow \\ \left( (oi \neq oj) \Rightarrow WF_{\langle \dots \rangle}(\wedge Internal \\ \wedge (\langle oi, oj \rangle \in opOrder') \vee (\langle oj, oi \rangle \in opOrder')) \right)$$

In a temporal formula, a predicate with no temporal operators is an assertion about the initial state. Hence, (11.3) asserts that the fairness condition is true for all pairs of distinct values  $oi$  and  $oj$  in the initial value of  $opId$ . But  $opId$  is initially empty, so this condition is vacuously true. Hence, (11.3) is trivially implied by the initial predicate. We must instead assert fairness for the action

$$(11.4) \quad \wedge (oi \in opId) \wedge (oj \in opId) \\ \wedge Internal \\ \wedge (\langle oi, oj \rangle \in opOrder') \vee (\langle oj, oi \rangle \in opOrder'))$$

for all distinct values  $oi$  and  $oj$ . It suffices to assert this only for  $oi$  and  $oj$  of the right type. Since it's best to use bounded quantifiers whenever possible, let's write this condition as:

$$\forall oi, oj \in [proc : Proc, idx : Nat] : \boxed{\text{All operations are eventually ordered.}} \\ (oi \neq oj) \Rightarrow WF_{\langle \dots \rangle}(\wedge (oi \in opId) \wedge (oj \in opId) \\ \wedge Internal \\ \wedge (\langle oi, oj \rangle \in opOrder') \vee (\langle oj, oi \rangle \in opOrder'))$$

For this formula to imply that any two operations are eventually ordered by *opOrder*, action (11.4) must be enabled if *oi* and *oj* are unordered operations in *opId*. It is, because *Serializable* is always enabled, so it is always possible to extend *opOrder* to a total ordering of all issued operations.

The complete inner specification, with *InitMem*, *opQ*, and *opOrder* visible, is in module *InnerSerial* on pages 189–191. I have made two minor modifications to allow the specification to be checked by the TLC model checker. (Chapter 13 describes TLC and explains why these changes are needed.) Instead of the definition of *opId* given on page 183, the specification uses the equivalent definition:

$$opId \triangleq \text{UNION } \{ [proc : p, idx : \text{DOMAIN } opQ[p]] : p \in Proc \}$$

In the definition of *UpdateOpOrder*, the first conjunct is changed from

$$opOrder' \subseteq opId' \times opId'$$

to the equivalent

$$opOrder' \in \text{SUBSET } (opId' \times opId')$$

For TLC's benefit, I also ordered the conjuncts of all actions so *UpdateOpOrder* follows the “assignment of a value to” *opQ'*. This resulted in the UNCHANGED conjunct not being the last one in action *Internal*.

The complete specification is obtained by the customary use of a parametrized instantiation of *InnerSerial* to hide the constant *InitMem* and the variables *opQ* and *opOrder*:

|                                                                                                                                                                                                                                                 |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| MODULE <i>SerialMemory</i>                                                                                                                                                                                                                      |
| EXTENDS <i>RegisterInterface</i><br>$Inner(InitMem, opQ, opOrder) \triangleq \text{INSTANCE } InnerSerial$<br>$Spec \triangleq \exists InitMem \in [Adr \rightarrow Val] :$<br>$\quad \exists opQ, opOrder : Inner(InitMem, opQ, opOrder)!Spec$ |

### 11.2.4 A Sequentially Consistent Memory

The serial memory specification does not allow the memory to predict future requests. We now remove this restriction and specify what is called a *sequentially consistent* memory. The freedom to predict the future can't be used by any real implementation,<sup>3</sup> so there's little practical difference between a serial and a

<sup>3</sup>The freedom to change explanations, which a sequentially consistent memory allows, could conceivably be used to permit a more efficient implementation, but it's not easy to see how.

|                                                                                                                                                    |                                                                                                        |
|----------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------|
| MODULE <i>InnerSerial</i>                                                                                                                          |                                                                                                        |
| EXTENDS <i>RegisterInterface</i> , <i>Naturals</i> , <i>Sequences</i> , <i>FiniteSets</i>                                                          |                                                                                                        |
| CONSTANT <i>InitMem</i>                                                                                                                            | The initial contents of memory, which will be an element of $[Proc \rightarrow Adr]$ .                 |
| VARIABLE <i>opQ</i> ,                                                                                                                              | $opQ[p][i]$ is the $i^{\text{th}}$ operation issued by processor $p$ .                                 |
| <i>opOrder</i>                                                                                                                                     | The ordering of operations, which is a subset of $opId \times opId$ . ( <i>opId</i> is defined below). |
| $opId \triangleq \text{UNION } \{ [proc : \{p\}, idx : \text{DOMAIN } opQ[p]] : p \in Proc \}$<br>$opIdQ(oi) \triangleq opQ[oi.proc][oi.idx]$      |                                                                                                        |
| [ $proc \mapsto p, idx \mapsto i$ ] identifies operation $i$ of processor $p$ .                                                                    |                                                                                                        |
| <i>InitWr</i> $\triangleq$ CHOOSE $v : v \notin [proc : Proc, idx : Nat]$                                                                          | The source for an initial memory value.                                                                |
| <i>Done</i> $\triangleq$ CHOOSE $v : v \notin Reg$                                                                                                 | The <i>reg</i> field value for a completed operation.                                                  |
| <i>opVal</i> $\triangleq$                                                                                                                          | Possible values of $opQ[p][i]$ .                                                                       |
| $[req : Request, reg : Reg]$                                                                                                                       | An active request using register <i>reg</i> .                                                          |
| $\cup [req : WrRequest, reg : Done]$                                                                                                               | A completed write.                                                                                     |
| $\cup [req : RdRequest, reg : Done, source : opId \cup \{InitWr\}]$                                                                                | A completed read of value from <i>source</i> .                                                         |
| $goodSource(oi) \triangleq$<br>$\{InitWr\} \cup \{o \in opId : \wedge opIdQ(o).req.op = \text{"Wr"} \wedge opIdQ(o).req.adr = opIdQ(oi).req.adr\}$ |                                                                                                        |
| DataInvariant $\triangleq$                                                                                                                         |                                                                                                        |
| $\wedge RegFileTypeInvariant$                                                                                                                      | Simple type invariants for <i>regFile</i> ,                                                            |
| $\wedge opQ \in [Proc \rightarrow Seq(opVal)]$                                                                                                     | <i>opQ</i> , and                                                                                       |
| $\wedge opOrder \subseteq (opId \times opId)$                                                                                                      | <i>opOrder</i> .                                                                                       |
| $\wedge \forall oi \in opId :$                                                                                                                     |                                                                                                        |
| $\wedge (\text{"source"} \in \text{DOMAIN } opIdQ(oi)) \Rightarrow$                                                                                | The source of any completed read is either <i>InitWr</i> or a write operation to the same address.     |
| $\wedge opIdQ(oi).source \in goodSource(oi)$                                                                                                       |                                                                                                        |
| $\wedge opIdQ(oi).req.val = \text{IF } opIdQ(oi).source = InitWr$                                                                                  | A read's value comes from its source.                                                                  |
| $\text{THEN } InitMem[opIdQ(oi).req.adr]$                                                                                                          |                                                                                                        |
| $\text{ELSE } opIdQ(opIdQ(oi).source).req.val$                                                                                                     |                                                                                                        |
| $\wedge (opIdQ(oi).req \neq Done) \Rightarrow$                                                                                                     | <i>opQ</i> correctly describes the register contents.                                                  |
| $(opIdQ(oi).req = regFile[oi.proc][opIdQ(oi).reg])$                                                                                                |                                                                                                        |
| $\wedge \forall p \in Proc, r \in Reg :$                                                                                                           | Only nonfree registers have corresponding <i>opQ</i> entries.                                          |
| $Cardinality(\{i \in \text{DOMAIN } opQ[p] : opQ[p][i].reg = r\}) =$                                                                               |                                                                                                        |
| $\text{IF } regFile[p][r].op = \text{"Free"} \text{ THEN } 0 \text{ ELSE } 1$                                                                      |                                                                                                        |

Figure 11.5: Module *InnerSerial* (beginning).

|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |                                                                                                                                                                                                                                                                       |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $Init \triangleq$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          | The initial predicate.                                                                                                                                                                                                                                                |
| $\wedge \text{regFile} \in [\text{Proc} \rightarrow [\text{Reg} \rightarrow \text{FreeRegValue}]]$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         | Every register is free.                                                                                                                                                                                                                                               |
| $\wedge \text{opQ} = [p \in \text{Proc} \mapsto \langle \rangle]$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          | There are no operations in $\text{opQ}$ .                                                                                                                                                                                                                             |
| $\wedge \text{opOrder} = \{\}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             | The ordering relation $\text{opOrder}$ is empty.                                                                                                                                                                                                                      |
| $\text{totalOpOrder} \triangleq$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           | The set of all total orderings on the set $\text{opId}$ .                                                                                                                                                                                                             |
| $\{R \in \text{SUBSET}(\text{opId} \times \text{opId}) :$<br>$\wedge \forall oi, oj \in \text{opId} : (oi = oj) \vee (\langle oi, oj \rangle \in R) \vee (\langle oj, oi \rangle \in R)$<br>$\wedge \forall oi, oj, ok \in \text{opId} : (\langle oi, oj \rangle \in R) \wedge (\langle oj, ok \rangle \in R) \Rightarrow (\langle oi, ok \rangle \in R)$<br>$\wedge \forall oi \in \text{opId} : \langle oi, oi \rangle \notin R\}$                                                                                                                                                       |                                                                                                                                                                                                                                                                       |
| $\text{Serializable} \triangleq$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           | Asserts that there exists a total ordering $R$ of all operations that extends $\text{opOrder}$ , orders the operations of each processor correctly, and makes the source of each read the most recent write to the address.                                           |
| $\exists R \in \text{totalOpOrder} :$<br>$\wedge \text{opOrder} \subseteq R$<br>$\wedge \forall oi, oj \in \text{opId} : (oi.\text{proc} = oj.\text{proc}) \wedge (oi.\text{idx} < oj.\text{idx}) \Rightarrow (\langle oi, oj \rangle \in R)$<br>$\wedge \forall oi \in \text{opId} : (\text{"source"} \in \text{DOMAIN } \text{opIdQ}(oi)) \Rightarrow$<br>$\neg(\exists oj \in \text{goodSource}(oi) :$<br>$\wedge \langle oj, oi \rangle \in R$<br>$\wedge (\text{opIdQ}(oi).\text{source} \neq \text{InitWr}) \Rightarrow (\langle \text{opIdQ}(oi).\text{source}, oj \rangle \in R))$ |                                                                                                                                                                                                                                                                       |
| $\text{UpdateOpOrder} \triangleq$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          | An action that chooses the new value of $\text{opOrder}$ , allowing it to be any relation that equals or extends the current value of $\text{opOrder}$ and satisfies $\text{Serializable}$ . This action is used in defining the subactions of the next-state action. |
| $\wedge \text{opOrder}' \in \text{SUBSET}(\text{opId}' \times \text{opId}')$<br>$\wedge \text{opOrder} \subseteq \text{opOrder}'$<br>$\wedge \text{Serializable}'$                                                                                                                                                                                                                                                                                                                                                                                                                         |                                                                                                                                                                                                                                                                       |
| $\text{IssueRequest}(\text{proc}, \text{req}, \text{reg}) \triangleq$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | Processor $\text{proc}$ issues request $\text{req}$ in register $\text{reg}$ .                                                                                                                                                                                        |
| $\wedge \text{regFile}[\text{proc}][\text{reg}].\text{op} = \text{"Free"}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 | The register must be free.                                                                                                                                                                                                                                            |
| $\wedge \text{regFile}' = [\text{regFile} \text{ EXCEPT } ![\text{proc}][\text{reg}] = \text{req}]$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        | Put the request in the register.                                                                                                                                                                                                                                      |
| $\wedge \text{opQ}' = [\text{opQ} \text{ EXCEPT } ![\text{proc}] = \text{Append}(\text{@}, [\text{req} \mapsto \text{req}, \text{reg} \mapsto \text{reg}])]$                                                                                                                                                                                                                                                                                                                                                                                                                               | Add request to $\text{opQ}[\text{proc}]$ .                                                                                                                                                                                                                            |
| $\wedge \text{UpdateOpOrder}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |                                                                                                                                                                                                                                                                       |
| $\text{RespondToWr}(\text{proc}, \text{reg}) \triangleq$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   | The memory responds to a write request in processor $\text{proc}$ 's register $\text{reg}$ .                                                                                                                                                                          |
| $\wedge \text{regFile}[\text{proc}][\text{reg}].\text{op} = \text{"Wr"}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   | The register must contain an active write request.                                                                                                                                                                                                                    |
| $\wedge \text{regFile}' = [\text{regFile} \text{ EXCEPT } ![\text{proc}][\text{reg}].\text{op} = \text{"Free"}]$                                                                                                                                                                                                                                                                                                                                                                                                                                                                           | The register is freed.                                                                                                                                                                                                                                                |
| $\wedge \text{LET } \text{idx} \triangleq \text{CHOOSE } i \in \text{DOMAIN } \text{opQ}[\text{proc}] : \text{opQ}[\text{proc}][i].\text{reg} = \text{reg}$                                                                                                                                                                                                                                                                                                                                                                                                                                | The appropriate $\text{opQ}$ entry is updated.                                                                                                                                                                                                                        |
| $\text{IN } \text{opQ}' = [\text{opQ} \text{ EXCEPT } ![\text{proc}][\text{idx}].\text{reg} = \text{Done}]$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                       |
| $\wedge \text{UpdateOpOrder}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              | $\text{opOrder}$ is updated.                                                                                                                                                                                                                                          |

Figure 11.6: Module *InnerSerial* (middle).

|                                                                                                                 |                                                                                      |
|-----------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------|
| $RespondToRd(proc, reg) \triangleq$                                                                             | The memory responds to a read request in processor $proc$ 's register $reg$ .        |
| LET $req \triangleq regFile[proc][reg]$                                                                         | $proc$ 's register $reg$ contains the request $req$ , which is in $opQ[proc][idx]$ . |
| $idx \triangleq$ CHOOSE $i \in \text{DOMAIN } opQ[proc] : opQ[proc][i].reg = req$                               |                                                                                      |
| IN $\wedge req.op = \text{"Rd"}$                                                                                | The register must contain an active read request.                                    |
| $\wedge \exists src \in goodSource([proc \mapsto proc, idx \mapsto idx]) :$                                     | The read obtains its value from a source $src$ .                                     |
| LET $val \triangleq$ IF $src = InitWr$ THEN $InitMem[req.adr]$                                                  | The value returned by the read.                                                      |
| ELSE $opIdQ(src).req.val$                                                                                       |                                                                                      |
| IN $\wedge regFile' = [regFile \text{ EXCEPT } ![proc][reg].val = val,$                                         | The register's $val$ field is set, and it is freed.                                  |
| $![proc][reg].op = \text{"Free"}]$                                                                              |                                                                                      |
| $\wedge opQ' = [opQ \text{ EXCEPT } opQ[proc][idx] \text{ is updated appropriately.}$                           |                                                                                      |
| $![proc][idx] = [req \mapsto [req \text{ EXCEPT } !.val = val],$                                                |                                                                                      |
| $reg \mapsto Done,$                                                                                             |                                                                                      |
| $source \mapsto src]]$                                                                                          |                                                                                      |
| $\wedge UpdateOpOrder$                                                                                          | $opOrder$ is updated.                                                                |
| $Internal \triangleq$                                                                                           | $\wedge \text{UNCHANGED } \langle regFile, opQ \rangle$                              |
|                                                                                                                 | $\wedge UpdateOpOrder$                                                               |
| $Next \triangleq$                                                                                               | The next-state action.                                                               |
| $\vee \exists proc \in Proc, reg \in Reg : \vee \exists req \in Request : IssueRequest(proc, req, reg)$         |                                                                                      |
| $\vee RespondToRd(proc, reg)$                                                                                   |                                                                                      |
| $\vee RespondToWr(proc, reg)$                                                                                   |                                                                                      |
| $\vee Internal$                                                                                                 |                                                                                      |
| $Spec \triangleq$                                                                                               | The complete internal specification.                                                 |
| $\wedge Init$                                                                                                   |                                                                                      |
| $\wedge \Box [Next]_{\langle regFile, opQ, opOrder \rangle}$                                                    |                                                                                      |
| $\wedge \forall proc \in Proc, reg \in Reg :$                                                                   | The memory eventually responds to every request.                                     |
| $WF_{\langle regFile, opQ, opOrder \rangle}(RespondToWr(proc, reg) \vee RespondToRd(proc, reg))$                |                                                                                      |
| $\wedge \forall oi, oj \in [proc : Proc, idx : Nat] :$                                                          | All operations are eventually ordered.                                               |
| $(oi \neq oj) \Rightarrow WF_{\langle regFile, opQ, opOrder \rangle}(\wedge (oi \in opId) \wedge (oj \in opId)$ |                                                                                      |
| $\wedge Internal$                                                                                               |                                                                                      |
| $\wedge (\langle oi, oj \rangle \in opOrder') \vee (\langle oj, oi \rangle \in opOrder'))$                      |                                                                                      |
| THEOREM $Spec \Rightarrow \Box(DataInvariant \wedge Serializable)$                                              |                                                                                      |

Figure 11.7: Module *InnerSerial* (end).

sequentially consistent memory. However, the sequentially consistent memory has a simpler specification. This specification is surprising and instructive.

The next-state action of the sequential memory specification has the same structure as that of the serial memory specification, with actions *IssueRequest*, *RespondToRd*, *RespondToWr*, and *Internal*. Like the serial memory specification, it has an internal variable *opQ* to which the *IssueRequest* operation appends an entry with *req* (request) and *reg* (register) fields. However, an operation does not remain forever in *opQ*. Instead, an *Internal* step removes it after it has been completed.<sup>4</sup> The specification has a second internal variable *mem* that represents the contents of a memory—that is, the value of *mem* is a function from *Adr* to *Val*. The value of *mem* is changed only by an *Internal* action that removes a write from *opQ*.

Recall that the correctness condition has two requirements:

1. There is a sequential execution order of all the operations that explains the values returned by reads.
2. This execution order is consistent with the order in which operations are issued by each individual processor.

The order in which operations are removed from *opQ* is an explanatory execution order that satisfies requirement 1 if the *Internal* action satisfies these properties:

- When a write of value *val* to address *adr* is removed from *opQ*, the value of *mem[adr]* is set to *val*.
- A read of address *adr* that returned a value *val* can be removed from *opQ* only if *mem[adr] = val*.

Requirement 2 is satisfied if operations issued by processor *p* are appended by the *IssueRequest* action to the tail of *opQ[p]*, and are removed by the *Internal* action only from the head of *opQ[p]*.

We have now determined what the *IssueRequest* and *Internal* actions should do. The *RespondToWr* action is obvious; it's essentially the same as in the serial memory specification. The problem is the *RespondToRd* action. How can we define it so that the value returned by a read is one that *mem* will contain when the *Internal* action has to remove the read from *opQ*? The answer is surprisingly simple: we allow the read to return any value. If the read were to return a bad value—for example, one that is never written—then the *Internal* action would never be able to remove the read from *opQ*. We rule out that possibility with a liveness condition requiring that every operation in *opQ* eventually be removed. This makes it easy to write the *Internal* action. The only remaining problem is expressing the liveness condition.

---

<sup>4</sup>We could allow uncompleted writes to be removed, but that would complicate the specification.

To guarantee that every operation is eventually removed from  $opQ$ , it suffices to guarantee that, for every processor  $Proc$ , the operation at the head of  $opQ[proc]$  is eventually removed. The desired liveness condition can therefore be expressed as:

$$\forall proc \in PProc : WF_{\langle \dots \rangle}(RemoveOp(proc))$$

where  $RemoveOp(proc)$  is an action that unconditionally removes the operation from the head of  $opQ[proc]$ . For convenience, we let the  $RemoveOp(proc)$  action also update  $mem$ . We then define a separate action  $Internal(proc)$  for each processor  $proc$ . It conjoins to  $RemoveOp(proc)$  the following enabling condition, which asserts that if the operation being removed is a read, then it has returned the correct value.

$$\begin{aligned} (Head(opQ[proc]).req.op = \text{“Rd”}) \Rightarrow \\ (mem[Head(opQ[proc]).req.adr] = Head(opQ[proc]).req.val) \end{aligned}$$

The complete internal specification, with the variables  $opQ$  and  $mem$  visible, appears in module *InnerSequential* on the following two pages. At this point, you should have no trouble understanding it. You should also have no trouble writing a module that instantiates *InnerSequential* and hides the internal variables  $opQ$  and  $mem$  to produce the final specification, so I won’t bother doing it for you.

### 11.2.5 The Memory Specifications Considered

Almost every specification we write admits a direct implementation, based on the initial predicate and next-state action. Such an implementation may be completely impractical, but it is theoretically possible. It’s easy to implement the linearizable memory with a single central memory. A direct implementation of the serial memory would require maintaining queues of all operations issued thus far, and a computationally infeasible search for possible total orderings. But, in theory, it’s easy.

Our specification of a sequentially consistent memory cannot be directly implemented. A direct implementation would have to guess the correct value to return on a read, which is impossible. The specification is not directly implementable because it is not machine closed. As explained in Section 8.8.2 on page 110, a non-machine closed specification is one in which a direct implementation can “paint itself into a corner,” reaching a point at which it is no longer possible to satisfy the specification. Any finite scenario of memory operations can be produced by a behavior satisfying the sequentially consistent memory’s initial predicate and next-state action—namely, a behavior that contains no *Internal* steps. However, not every finite scenario can be extended to one that is explainable by a sequential execution. For example, no scenario that begins as follows is possible in a two-processor system:

This notation for describing scenarios was introduced on page 180.

MODULE *InnerSequential*

EXTENDS *RegisterInterface*, *Naturals*, *Sequences*, *FiniteSets*

VARIABLE *opQ*, *opQ*[*p*][*i*] is the *i*<sup>th</sup> operation issued by processor *p*.

*mem* An internal memory.

---

*Done*  $\triangleq$  CHOOSE *v* : *v*  $\notin$  *Reg* The *reg* field value for a completed operation.

*DataInvariant*  $\triangleq$

$\wedge$  *RegFileTypeInvariant* Simple type invariants for *regFile*,

$\wedge$  *opQ*  $\in$  [*Proc*  $\rightarrow$  Seq(*req* : *Request*, *reg* : *Reg*  $\cup$  {*Done*})] *opQ*, and

$\wedge$  *mem*  $\in$  [*Adr*  $\rightarrow$  *Val*] *opOrder*.

$\wedge$   $\forall p \in$  *Proc*, *r*  $\in$  *Reg* : Only nonfree registers have corresponding *opQ* entries.

$\text{Cardinality}(\{i \in \text{DOMAIN } opQ[p] : opQ[p][i].reg = r\}) =$   
IF *regFile*[*p*][*r*].*op* = “Free” THEN 0 ELSE 1

---

*Init*  $\triangleq$  The initial predicate.

$\wedge$  *regFile*  $\in$  [*Proc*  $\rightarrow$  [*Reg*  $\rightarrow$  *FreeRegValue*]] Every register is free.

$\wedge$  *opQ* = [*p*  $\in$  *Proc*  $\mapsto$   $\langle \rangle$ ] There are no operations in *opQ*.

$\wedge$  *mem*  $\in$  [*Adr*  $\rightarrow$  *Val*] The internal memory can have any initial contents.

---

*IssueRequest*(*proc*, *req*, *reg*)  $\triangleq$  Processor *proc* issues request *req* in register *reg*.

$\wedge$  *regFile*[*proc*][*reg*].*op* = “Free” The register must be free.

$\wedge$  *regFile*' = [*regFile* EXCEPT ![*proc*][*reg*] = *req*] Put request in register.

$\wedge$  *opQ*' = [*opQ* EXCEPT ![*proc*] = *Append*(@, [*req*  $\mapsto$  *req*, *reg*  $\mapsto$  *reg*])] Add request to *opQ*[*proc*].

$\wedge$  UNCHANGED *mem*

---

*RespondToRd*(*proc*, *reg*)  $\triangleq$  The memory responds to a read request in processor *proc*'s register *reg*.

$\wedge$  *regFile*[*proc*][*reg*].*op* = “Rd” The register must contain an active read request.

$\wedge$   $\exists val \in$  *Val* : *val* is the value returned.

$\wedge$  *regFile*' = [*regFile* EXCEPT ![*proc*][*reg*].*val* = *val*,  
![*proc*][*reg*].*op* = “Free”] Set the register's *val* field,  
and free the register.

$\wedge$  *opQ*' = LET *idx*  $\triangleq$  *opQ*[*proc*][*idx*] contains the request in register *reg*.

CHOOSE *i*  $\in$  DOMAIN *opQ*[*proc*] : *opQ*[*proc*][*i*].*reg* = *reg*

IN [*opQ* EXCEPT ![*proc*][*idx*].*req.val* = *val*,  
![*proc*][*idx*].*reg* = *Done*] Set *opQ*[*proc*][*idx*]'s *val* field to  
*val* and its *reg* field to *Done*.

$\wedge$  UNCHANGED *mem*

Figure 11.8: Module *InnerSequential* (beginning).

|                                                                                                                                                                                                             |                                                                                                                          |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------|
| $RespondToWr(proc, reg) \triangleq$                                                                                                                                                                         | The memory responds to a write request in processor $proc$ 's register $reg$ .                                           |
| $\wedge regFile[proc][reg].op = \text{"Wr"}$                                                                                                                                                                | The register must contain an active write request.                                                                       |
| $\wedge regFile' = [regFile \text{ EXCEPT } ![proc][reg].op = \text{"Free"}]$                                                                                                                               | The register is freed.                                                                                                   |
| $\wedge \text{LET } idx \triangleq \text{CHOOSE } i \in \text{DOMAIN } opQ[proc] : opQ[proc][i].reg = reg$<br>$\text{IN } opQ' = [opQ \text{ EXCEPT } ![proc][idx].reg = Done]$                             | The appropriate $opQ$ entry is updated.                                                                                  |
| $\wedge \text{UNCHANGED } mem$                                                                                                                                                                              |                                                                                                                          |
| $RemoveOp(proc) \triangleq$                                                                                                                                                                                 | Unconditionally remove the operation at the head of $opQ[proc]$ and update $mem$ .                                       |
| $\wedge opQ[proc] \neq \langle \rangle$                                                                                                                                                                     | $opQ[p]$ must be nonempty.                                                                                               |
| $\wedge Head(opQ[proc]).reg = Done$                                                                                                                                                                         | The operation must have been completed.                                                                                  |
| $\wedge mem' = \text{IF } Head(opQ[proc]).req.op = \text{"Rd"}$<br>$\text{THEN } mem$<br>$\text{ELSE } [mem \text{ EXCEPT } ![Head(opQ[proc]).req.adr] =$<br>$Head(opQ[proc]).req.val]$                     | Leave $mem$ unchanged for a read operation, update it for write operation.                                               |
| $\wedge opQ' = [opQ \text{ EXCEPT } ![proc] = Tail(@)]$                                                                                                                                                     | Remove the operation from $opQ[proc]$ .                                                                                  |
| $\wedge \text{UNCHANGED } regFile$                                                                                                                                                                          | No register is changed.                                                                                                  |
| $Internal(proc) \triangleq$                                                                                                                                                                                 | Remove the operation at the head of $opQ[proc]$ . But if it's a read, only do so if it returned the value now in $mem$ . |
| $\wedge RemoveOp(proc)$                                                                                                                                                                                     |                                                                                                                          |
| $\wedge (Head(opQ[proc]).req.op = \text{"Rd"}) \Rightarrow$<br>$(mem[Head(opQ[proc]).req.adr] = Head(opQ[proc]).req.val)$                                                                                   |                                                                                                                          |
| $Next \triangleq$                                                                                                                                                                                           | The next-state action.                                                                                                   |
| $\exists proc \in Proc : \vee \exists reg \in Reg : \vee \exists req \in Request : IssueRequest(proc, req, reg)$<br>$\vee RespondToRd(proc, reg)$<br>$\vee RespondToWr(proc, reg)$<br>$\vee Internal(proc)$ |                                                                                                                          |
| $Spec \triangleq \wedge Init$                                                                                                                                                                               |                                                                                                                          |
| $\wedge \square[Next]_{\langle regFile, opQ, mem \rangle}$                                                                                                                                                  |                                                                                                                          |
| $\wedge \forall proc \in Proc, reg \in Reg :$                                                                                                                                                               | The memory eventually responds to every request.                                                                         |
| $\text{WF}_{\langle regFile, opQ, mem \rangle}(RespondToWr(proc, reg) \vee RespondToRd(proc, reg))$                                                                                                         |                                                                                                                          |
| $\wedge \forall proc \in Proc :$                                                                                                                                                                            | Every operation is eventually removed from $opQ$ .                                                                       |
| $\text{WF}_{\langle regFile, opQ, mem \rangle}(RemoveOp(proc))$                                                                                                                                             |                                                                                                                          |
| THEOREM $Spec \Rightarrow \square(DataInvariant)$                                                                                                                                                           |                                                                                                                          |

Figure 11.9: Module *InnerSequential* (end).

Processor  $p$ :  $Wr_p(a1, v1)$ ,  $Rd_p(a1, v2)$ ,  $Wr_p(a2, v2)$

Processor  $q$ :  $Wr_q(a2, v1)$ ,  $Rd_q(a2, v2)$ ,  $Wr_q(a1, v2)$

Here's why:

|                         |                                                         |
|-------------------------|---------------------------------------------------------|
| $Wr_q(a1, v2)$          |                                                         |
| precedes $Rd_p(a1, v2)$ | This is the only explanation of the value read by $p$ . |
| precedes $Wr_p(a2, v2)$ | By the order in which operations are issued.            |
| precedes $Rd_q(a2, v2)$ | This is the only explanation of the value read by $q$ . |
| precedes $Wr_q(a1, v2)$ | By the order in which operations are issued.            |

Hence  $q$ 's write of  $a1$  must precede itself, which is impossible.

As mentioned in Section 8.8.2, a specification is machine closed if the liveness property is the conjunction of fairness properties for actions that imply the next-state action. The sequential memory specification asserts weak fairness of  $RemoveOp(proc)$ , for processors  $proc$ , and  $RemoveOp(proc)$  does not imply the next-state action. (The next-state action does not allow a  $RemoveOp(proc)$  step that removes from  $opQ[proc]$  a read that has returned the wrong value.)

Very high-level system specifications, such as our memory specifications are subtle. It's easy to get them wrong. The approach we used in the serial memory specification—namely, writing conditions on the history of all operations—is dangerous. It's easy to forget some conditions. A non-machine closed specification can occasionally be the simplest way to express what you want so say.

# **Part III**

## **The Tools**



## Chapter 12

# The Java Front End

The Java Front End is a parser for TLA<sup>+</sup> written in Java by Jean-Charles Grégoire. It provides a front end for other tools, such as TLC (see Chapter 13). It can also be run by itself to find syntax errors in a specification. You should obtain directions for running the Java Front End's parser on your particular system when you obtain the software.

### 12.1 Finding an Error

When the parser reports an error, finding what caused it can be tricky. The errors that the parser detects fall into two separate classes, which are usually called *syntactic* and *semantic* errors. A syntax error is one that makes the specification grammatically incorrect, meaning that it violates the BNF grammar, or the precedence and alignment rules, described in Chapter 14. A semantic error is one that violates the legality conditions mentioned in Chapter 16. The term *semantic error* is misleading, because it suggests an error that makes a specification have the wrong meaning. All errors found by the parser are ones that make the specification illegal—that is, not syntactically well-formed—and hence make it have no meaning at all.

The parser reads the file sequentially, starting from the beginning, and it reports a syntax error if and when it reaches a point at which it becomes impossible for any continuation to produce a grammatically correct specification. For example, if you leave out a colon and type  $\backslash A \ x \ P(x)$  instead of  $\backslash A \ x : P(x)$ , the parser will print something like:

```
Encountered "P" at line 7, column 11.
Was expecting one of:
    "," ...
    <OpSymbol> ...
```

`":" ...`

The “was expecting” list describes every possible symbol that could lead to a grammatically correct specification. Knowing what the parser was expecting can sometimes help find the error.

The parser may detect a grammatical error far from the actual mistake. For example, suppose you type `[` instead of `{` to produce  $[x \in \dots : P(x)]$ , where “ $\dots$ ” is a very long expression. The parser will discover the error only when it sees the colon, well past the erroneous `[`. If you can’t find the source of an error, try the “divide and conquer” method: keep removing different parts of the module until you isolate the source of the problem.

A typical semantic error is an undefined symbol that arises because you mistype an identifier. For example, if you define an operator *Cat* but spell it *cat* somewhere by mistake, the parser may report

`unresolved identifier cat at [line: 87, col: 6] to [87,8].`

The source of a semantic error is usually easy to find.

The parser stops when it encounters the first syntactic error. It can detect multiple semantic errors in a single run.

The current version of the parser has the following limitations, which should be corrected in future versions:

- It does not detect certain semantic errors. In particular, it does not do any level checking. (See Section 16.2 on page 305.)
- It does not properly handle strings with “escape sequences” such as `\`. (See Section 15.1.10.)
- It does not properly handle parametrized instantiation. (See Section 4.2.2 on page 39.)
- It does not handle the symbol `(\X)`. You must type `\otimes` to represent  $\otimes$ .
- It produces meaningless warnings of the form

`numbers are used but NUMERAL isn't defined`

This warning is a remnant of a minor change to  $\text{TLA}^+$ .

- It does not handle numbers written in binary, octal, or hexadecimal notation. (See Section 15.1.11.)

## Chapter 13

# The TLC Model Checker

This is a description of how we expect Version 2.0 of TLC to behave. Version 1.0 of TLC does not completely implement this description; Section 13.4 on page 231 describes its limitations. Section 13.5 on page 233 describes improvements that we may make to later versions.

### 13.1 Introduction to TLC

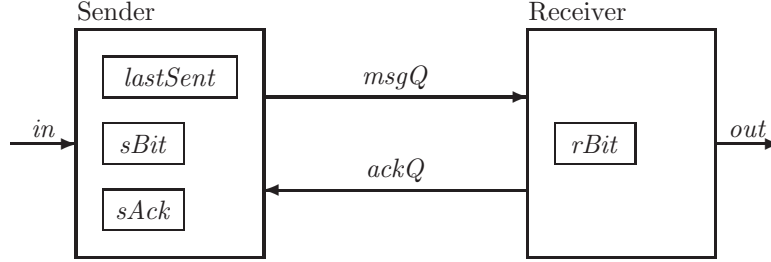
TLC is a program for finding errors in TLA<sup>+</sup> specifications. A syntactically correct specification can have two kinds of errors: it may contain “silliness” or it may not capture the intention of its author. As explained in Section 6.2, a silly expression is one whose meaning is not determined by the semantics of TLA<sup>+</sup>—for example,  $3 + \langle 1, 2 \rangle$ . A specification is incorrect if whether or not some particular behavior satisfies it depends on the meaning of a silly expression. Intention isn’t a well-defined concept, and there may be a fine line between errors and unintended features.

Experience has shown that one of the most effective ways of finding both kinds of errors is by trying to verify invariance properties of a specification. TLC tries to find errors by looking for counterexamples to invariance properties—that is, to assertions of the form

$$(13.1) \quad \textit{Init} \wedge \Box[\textit{Next}]_v \Rightarrow \Box \textit{Inv}$$

where *Inv* is a state predicate. One example of an invariance property is the absence of deadlock, in which *Inv* equals `ENABLED Next`. A counterexample to this instance of (13.1) shows the possibility of deadlock—the ability of the system to reach a state in which no further progress is possible. (Of course, for some systems, deadlock may just mean successful termination.) As explained in Sections 13.3.4 and 13.3.5 below, TLC can also be used to check other properties besides invariance.

I will illustrate the use of TLC with a simple example: a specification of the alternating bit protocol for sending data over a lossy FIFO transmission line. The protocol might be described as a system that looks like this:



The sender receives data values from the environment over communication channel *in*, and the receiver then sends the values to the environment on channel *out*. The protocol uses two lossy FIFO transmission lines: the sender sends data and control information on *msgQ*, and the receiver sends acknowledgments on *ackQ*. The variable *lastSent* records the last message the sender received from the environment. The variables *sBit*, *sAck*, and *rBit* are one-bit values used to control the sending of messages between the sender and receiver.

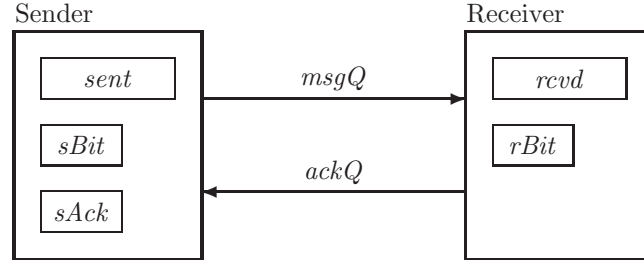
The protocol is supposed to implement a FIFO transmission line—in fact, a bounded FIFO of length 1. Correctness of the protocol means that its specification implements (implies) formula *Spec* of module *BoundedFIFO* (Figure 4.2 on page 43), with 1 substituted for *N*. TLC is most easily used to check invariance properties, so we want to express correctness as an invariance property.

Intuitively, correctness of the protocol means that every value sent by the environment on channel *in*, except possibly for the last one sent, has been received by the environment on channel *out*. To express this condition as an invariant, we add two variables, *sent* and *rcvd*, that record the sequences of messages sent over channels *in* and *out*, respectively. Correctness of the algorithm is then expressed in the form (13.1) with

$$\begin{aligned} Inv \triangleq & \wedge Len(rcvd) \in \{Len(sent) - 1, Len(sent)\} \\ & \wedge \forall i \in 1 \dots Len(rcvd) : rcvd[i] = sent[i] \end{aligned}$$

Changes to the variables *sent* and *rcvd* record the communication over the channels *in* and *out*. If our goal is to verify the correctness of the protocol, rather than to specify an entire system, there's no need to represent the actual input and output channels. So, we eliminate the variables *in* and *out* from the specification. Since the last message sent by the environment is the last element of the sequence *sent*, we don't need the variable *lastSent*. So, we can describe the protocol as a system that looks like this:

*Len(s)* is defined in the *Sequences* module to be the length of sequence *s*, and  $1 \dots n$  is the set  $\{1, \dots, n\}$  of integers.



The complete protocol specification appears in module *AlternatingBit* in Figures 13.1 and 13.2 on the following two pages.

However, for now all you need to know are the declarations:

```

CONSTANT Data = {The set of data values that can be sent.}
VARIABLES msgQ, ackQ, sBit, sAck, rBit, sent, rcvd

```

and the types of the variables:

- *msgQ* is a sequence of elements in  $\{0, 1\} \times Data$ .
- *ackQ* is a sequence of elements in  $\{0, 1\}$ .
- *sBit*, *sAck*, and *rBit* are elements of  $\{0, 1\}$ .
- *sent* and *rcvd* are sequences of elements in *Data*.

The input to TLC consists of a  $TLA^+$  module and a configuration file. The configuration file tells TLC the names of the initial predicate, the next-state relation, and the invariant to be checked. For example, the configuration file for the alternating bit protocol will contain the declaration

```
INIT ABInit
```

telling TLC to take *ABInit* as the formula *Init* in (13.1).

TLC works by generating behaviors that satisfy the specification. To do this for the alternating bit protocol, it needs to know what elements are in the set *Data* of data values. We can tell TLC to let *Data* equal the set containing two elements, named *d1* and *d2*, by putting the following declaration in the configuration file.

```
CONSTANT Data = {d1, d2}
```

(We can use any sequence of letters and digits containing at least one letter as the name of an element.)

There are two ways to use TLC. The default method is to have it try to check all reachable states—that is, all states that can occur in behaviors

The keywords  
CONSTANT and  
CONSTANTS are  
equivalent.

| MODULE <i>AlternatingBit</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       |                                                                                                  |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------|
| <p>This specification describes a protocol for using lossy FIFO transmission lines to transmit a sequence of values from a sender to a receiver. The sender sends a data value <math>d</math> by sending a sequence of <math>\langle b, d \rangle</math> messages on <math>msgQ</math>, where <math>b</math> is a control bit. It knows that the message has been received when it receives the ack <math>b</math> from the receiver on <math>ackQ</math>. It sends the next value with a different control bit. The receiver knows that a message on <math>msgQ</math> contains a new value when its control bit differs from the last one it has received. The receiver keeps sending the last control bit it received on <math>ackQ</math>.</p> |                                                                                                  |
| EXTENDS <i>Naturals, Sequences</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |                                                                                                  |
| CONSTANTS <i>Data</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              | The set of data values that can be sent.                                                         |
| VARIABLES <i>msgQ</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            | The sequence of $\langle$ control bit, data value $\rangle$ messages in transit to the receiver. |
| <i>ackQ</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | The sequence of one-bit acknowledgments in transit to the sender.                                |
| <i>sBit</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | The last control bit sent by sender; it is complemented when sending a new data value.           |
| <i>sAck</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | The last acknowledgment bit received by the sender.                                              |
| <i>rBit</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | The last control bit received by the receiver.                                                   |
| <i>sent</i> ,                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | The sequence of values sent by the sender.                                                       |
| <i>rcvd</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        | The sequence of values received by the receiver.                                                 |
| $ABInit \triangleq \wedge msgQ = \langle \rangle$<br>$\wedge ackQ = \langle \rangle$<br>$\wedge sBit \in \{0, 1\}$<br>$\wedge sAck = sBit$<br>$\wedge rBit = sBit$<br>$\wedge sent = \langle \rangle$<br>$\wedge rcvd = \langle \rangle$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |                                                                                                  |
| <p>The initial condition:<br/>Both message queues are empty.<br/>All the bits equal 0 or 1<br/>and are equal to each other.<br/><br/>No values have been sent<br/>or received.</p>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |                                                                                                  |
| $TypeInv \triangleq \wedge msgQ \in Seq(\{0, 1\} \times Data)$<br>$\wedge ackQ \in Seq(\{0, 1\})$<br>$\wedge sBit \in \{0, 1\}$<br>$\wedge sAck \in \{0, 1\}$<br>$\wedge rBit \in \{0, 1\}$<br>$\wedge sent \in Seq(Data)$<br>$\wedge rcvd \in Seq(Data)$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |                                                                                                  |
| <p>The type-correctness invariant.</p>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |                                                                                                  |
| $SndNewValue(d) \triangleq$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |                                                                                                  |
| <p>The action in which the sender sends a new data value <math>d</math>.</p>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       |                                                                                                  |
| $\wedge sAck = sBit$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               | Enabled iff $sAck$ equals $sBit$ .                                                               |
| $\wedge sent' = Append(sent, d)$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   | Append $d$ to $sent$ .                                                                           |
| $\wedge sBit' = 1 - sBit$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          | Complement control bit $sBit$                                                                    |
| $\wedge msgQ' = Append(msgQ, \langle sBit', d \rangle)$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            | Send value on $msgQ$ with new control bit.                                                       |
| $\wedge UNCHANGED \langle ackQ, sAck, rBit, rcvd \rangle$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |                                                                                                  |

Figure 13.1: The Alternating Bit Protocol (beginning)

|                                                                            |                                                                                                                                                 |
|----------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------|
| $ReSndMsg \triangleq$                                                      | The sender resends the last message it sent on $msgQ$ .                                                                                         |
| $\wedge sAck \neq sBit$                                                    | Enabled iff $sAck$ doesn't equal $sBit$ .                                                                                                       |
| $\wedge msgQ' = Append(msgQ, \langle sBit, sent[Len(sent)] \rangle)$       | Resend the last value in $sent$ .                                                                                                               |
| $\wedge UNCHANGED \langle ackQ, sBit, sAck, rBit, sent, rcvd \rangle$      |                                                                                                                                                 |
| $RcvMsg \triangleq$                                                        | The receiver receives the message at the head of $msgQ$ .                                                                                       |
| $\wedge msgQ \neq \langle \rangle$                                         | Enabled iff $msgQ$ not empty.                                                                                                                   |
| $\wedge msgQ' = Tail(msgQ)$                                                | Remove message from head of $msgQ$ .                                                                                                            |
| $\wedge rBit' = Head(msgQ)[1]$                                             | Set $rBit$ to message's control bit.                                                                                                            |
| $\wedge rcvd' = IF \ rBit' \neq rBit$                                      | If control bit $\neq$ last one received,                                                                                                        |
| THEN $Append(rcvd, Head(msgQ)[2])$                                         | then append message's value to $rcvd$ .                                                                                                         |
| ELSE $rcvd$                                                                |                                                                                                                                                 |
| $\wedge UNCHANGED \langle ackQ, sBit, sAck, sent \rangle$                  |                                                                                                                                                 |
| $SndAck \triangleq$                                                        |                                                                                                                                                 |
| $\wedge ackQ' = Append(ackQ, rBit)$                                        | The receiver sends $rBit$ on $ackQ$ at any time.                                                                                                |
| $\wedge UNCHANGED \langle msgQ, sBit, sAck, rBit, sent, rcvd \rangle$      |                                                                                                                                                 |
| $RcvAck \triangleq$                                                        |                                                                                                                                                 |
| $\wedge ackQ \neq \langle \rangle$                                         | The sender receives an ack on $ackQ$ . It removes the ack and sets $sAck$ to its value.                                                         |
| $\wedge ackQ' = Tail(ackQ)$                                                |                                                                                                                                                 |
| $\wedge sAck' = Head(ackQ)$                                                |                                                                                                                                                 |
| $\wedge UNCHANGED \langle msgQ, sBit, rBit, sent, rcvd \rangle$            |                                                                                                                                                 |
| $Lose(c) \triangleq$                                                       | The action of losing a message on transmission line $c$ .                                                                                       |
| $\wedge c \neq \langle \rangle$                                            | Enabled iff $c$ is not empty. For some $i$ , remove the $i$ th message from $c$ . Leave everything unchanged except $c$ , $msgQ$ , and $ackQ$ . |
| $\wedge \exists i \in 1 \dots Len(c) :$                                    |                                                                                                                                                 |
| $c' = [j \in 1 \dots (Len(c) - 1) \mapsto IF \ j < i \ THEN \ c[j]$        |                                                                                                                                                 |
| ELSE $c[j + 1]]$                                                           |                                                                                                                                                 |
| $\wedge UNCHANGED \langle sBit, sAck, rBit, sent, rcvd \rangle$            |                                                                                                                                                 |
| $LoseMsg \triangleq Lose(msgQ) \wedge UNCHANGED ackQ$                      | Lose a message on $msgQ$ .                                                                                                                      |
| $LoseAck \triangleq Lose(ackQ) \wedge UNCHANGED msgQ$                      | Lose a message on $ackQ$ .                                                                                                                      |
| $ABNext \triangleq$                                                        | The next-state relation.                                                                                                                        |
| $\vee \exists d \in Data : SndNewValue(d)$                                 |                                                                                                                                                 |
| $\vee ReSndMsg \vee RcvMsg \vee SndAck \vee RcvAck$                        |                                                                                                                                                 |
| $\vee LoseMsg \vee LoseAck$                                                |                                                                                                                                                 |
| $vars \triangleq \langle msgQ, ackQ, sBit, sAck, rBit, sent, rcvd \rangle$ | The tuple of all variables.                                                                                                                     |
| $Spec \triangleq ABInit \wedge \Box[ABNext]_{vars}$                        | The complete specification.                                                                                                                     |
| $Inv \triangleq$                                                           | The invariant that expresses correctness of the protocol.                                                                                       |
| $\wedge Len(rcvd) \in \{Len(sent) - 1, Len(sent)\}$                        |                                                                                                                                                 |
| $\wedge \forall i \in 1 \dots Len(rcvd) : rcvd[i] = sent[i]$               |                                                                                                                                                 |
| THEOREM $Spec \Rightarrow \Box Inv$                                        |                                                                                                                                                 |

Figure 13.2: The Alternating Bit Protocol (end)

|                                                                                                                                                                      |                                                                                                                                                              |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|
| MODULE <i>MCAAlternatingBit</i>                                                                                                                                      |                                                                                                                                                              |
| EXTENDS <i>AlternatingBit</i>                                                                                                                                        |                                                                                                                                                              |
| CONSTANTS <i>msgQLen</i> , <i>ackQLen</i> , <i>sentLen</i>                                                                                                           |                                                                                                                                                              |
| <i>SeqConstraint</i> $\triangleq$                                                                                                                                    | <div style="border: 1px solid black; padding: 5px; margin-top: 10px;">             A constraint on the lengths of sequences for use by TLC.           </div> |
| $\wedge \text{Len}(\text{msg}Q) \leq \text{msgQLen}$<br>$\wedge \text{Len}(\text{ack}Q) \leq \text{ackQLen}$<br>$\wedge \text{Len}(\text{sent}) \leq \text{sentLen}$ |                                                                                                                                                              |

Figure 13.3: Module *MCAAlternatingBit*.

satisfying the specification.<sup>1</sup> You can also use TLC in simulation mode, in which it randomly generates behaviors, without trying to check all reachable states. We now consider the default mode; simulation mode is described in Section 13.2.7 on page 215.

Exhaustively checking all possible reachable states is impossible for the alternating bit protocol because the sequences of messages and values can get arbitrarily long, so there are infinitely many reachable states. To bound the number of states, we define a state predicate called the *constraint* that asserts bounds on the lengths of the sequences. For example, the following constraint asserts that *msgQ* and *ackQ* have lengths at most 2 and *sent* has length at most 3:

$$\begin{aligned} &\wedge \text{Len}(\text{msg}Q) \leq 2 \\ &\wedge \text{Len}(\text{ack}Q) \leq 2 \\ &\wedge \text{Len}(\text{sent}) \leq 3 \end{aligned}$$

Since the sequence *rcvd* of values received can never be longer than the sequence *sent* of values sent, there is no need to constrain its length.

Instead of specifying the bounds on the lengths of sequences in this way, I prefer to make them parameters and to assign them values in the configuration file. We usually don't want to put into the specification itself declarations and definitions that are just for TLC's benefit. So, let's write a new module, called *MCAAlternatingBit*, that EXTENDS the *AlternatingBit* module and can be used as input to TLC. That module appears in Figure 13.3 on this page. A possible configuration file for this module appears in Figure 13.4 on the next page. The INVARIANT section tells TLC to check the invariant *Inv* and the type-correctness invariant *TypeInv*, both of which are defined in module *AlternatingBit*. Observe that the configuration file must specify values for all the constant parameters of the specification—in this case, the parameter *Data* from the *AlternatingBit* module and the three parameters declared in module *MCAAlternatingBit* itself.

The keywords INVARIANT and INVARIANTS are equivalent.

<sup>1</sup>As explained in Section 2.3 (page 18), a state is an assignment of values to all possible variables. However, when discussing a particular specification, we usually consider a state to be an assignment of values to that specification's variables. That's what we're doing in this chapter.

```

CONSTANTS  Data      = {d1, d2}
           msgQLen = 2
           ackQLen = 2
           sentLen = 3
INIT       AInit
NEXT       ABNext
INVARIANTS Inv
           TypeInv
CONSTRAINT SeqConstraint

```

Figure 13.4: A configuration file for module *MCAalternatingBit*.

The constraint and the assignment of values to the constant parameters define what we call a *model* of the specification. Given the specification and a model, TLC uses essentially the following algorithm to compute a set  $\mathcal{R}$  of reachable states.

- Initialize  $\mathcal{R}$  to the set of states that satisfy the initial predicate, and initialize the set  $\mathcal{U}$  of unexplored states to the subset of those states that satisfy the constraint.
- While  $\mathcal{U}$  is nonempty, do the following for each state  $s$  in  $\mathcal{U}$ . Remove  $s$  from  $\mathcal{U}$  and compute all states  $t$  such that the step  $s \rightarrow t$  satisfies the next-state action and  $t$  is not in  $\mathcal{R}$ . For each such  $t$ : add  $t$  to  $\mathcal{R}$  and, if  $t$  satisfies the constraint, add it to  $\mathcal{U}$ .

This computes  $\mathcal{R}$  to be the set of all states  $s$  for which there is some behavior  $\sigma$  satisfying the specification that contains  $s$  and such that every state preceding  $s$  in  $\sigma$  satisfies the constraint. TLC reports an error and stops if, during this computation, it adds to  $\mathcal{R}$  a state that fails to satisfy the invariant.

How long it takes TLC to check a specification depends on the specification and the size of the model. Run on a 600MHz work station, TLC finds about 1000 reachable states per second for a specification as simple as that of the alternating bit protocol. For some specifications, the rate at which TLC generates states varies (inversely) with the size of the model; it can also go down as the states it generates become more complicated. For some specifications run on larger models, TLC can find fewer than one reachable state per second.

You should always begin testing a specification with a very small model, which TLC can check quickly. A small model will catch most simple errors. For example, a typical “off-by-one error” will make the next state depend upon the value of an expression  $f[e]$  when  $e$  is not in the domain of the function  $f$ , and TLC will report this as an error. When a very small model reveals no more errors, you can then run TLC with larger models to try to catch more subtle errors.

One way to figure out how large a model TLC can handle is to estimate the approximate number of reachable states as a function of the parameters. For the alternating bit protocol, a calculation based on knowledge of how the protocol works shows that there are about

$$(13.2) \ 3 * D^{sentLen+1} * (msgQLen+1) * (msgQLen+2) * (ackQLen+1) * (ackQLen+2)$$

reachable states that satisfy the constraint, where  $D$  is the number of elements in *Data*. For the model specified by the configuration file of Figure 13.4, this is about 6900 states. The estimate is somewhat high; TLC finds 2312 distinct states, 1158 of which satisfy the constraint. (Remember that TLC examines states reachable in one step from a state that satisfies the constraint.) However, what's important is not the precise number of states, but how that number varies with the parameters. From (13.2), we see that the number of states depends only quadratically on *msgLen* and *ackLen*, but exponentially on *sentLen*. For example, letting *Data* have 3 instead of 2 elements and increasing *sentLen* from 3 to 4 can be expected to increase the number of states, and hence the running time, by a factor of 81/8, or about 10. In fact, TLC then finds 22058 states, 10050 of which satisfy the constraint.

Calculating the number of reachable states can be hard. If you can't do it, increase the model size very gradually. The number of reachable states is typically an exponential function of the model's parameters; and the value of  $a^b$  grows very fast with increasing values of  $b$ .

Many systems have errors that will show up only on models too large for TLC to check exhaustively. After TLC has exhaustively checked your specification on as large a model as it can, you can run it in simulation mode on larger models. Simulation can't catch all errors, but it's worth trying.

## 13.2 How TLC Works

A program like TLC can't handle all the specifications that can be written in a language as expressive as  $TLA^+$ . To explain what kinds of specifications TLC can and cannot handle, I will now sketch how TLC works. A more complete description is given in Section 13.6.

### 13.2.1 TLC Values

A state is an assignment of values to variables.  $TLA^+$  allows you to describe a wide variety of values—for example, the set of all sequences of prime numbers. TLC can compute only a restricted class of values. Those values are built up from the following four types of primitive values:

*Booleans* The values TRUE and FALSE.

*Integers* Values like 3 and  $-1$ .

*Strings* Values like “ab3”.

*Model Values* in the **CONSTANT** section of the configuration file. For example, the configuration file shown in Figure 13.4 on page 207 introduces the model values  $d1$  and  $d2$ . Model values with different names are assumed to be different.

A TLC value is defined inductively to be either

1. a primitive value, or
2. a finite set of comparable TLC values (*comparable* is defined below), or
3. a function  $f$  whose domain  $D$  is a TLC value such that  $f[x]$  is a TLC value, for all  $x \in D$ .

For example, the first two rules imply that

(13.3)  $\{\{\text{“a”}, \text{“b”}\}, \{\text{“b”}, \text{“c”}\}, \{\text{“c”}, \text{“d”}\}\}$

is a TLC value because rules 1 and 2 imply that  $\{\text{“a”}, \text{“b”}\}$ ,  $\{\text{“b”}, \text{“c”}\}$ , and  $\{\text{“c”}, \text{“d”}\}$  are TLC values, and the second rule then implies that (13.3) is a TLC value. Since tuples and records are functions, rule 3 implies that a record or tuple whose components are TLC values is a TLC value. For example,  $\langle 1, \text{“a”}, 2, \text{“b”} \rangle$  is a TLC value.

To complete the definition of what a TLC value is, I must explain what *comparable* means in rule 2. The basic idea is that two values should be comparable if the semantics of  $\text{TLA}^+$  determines whether or not they are equal. For example, strings and numbers are not comparable because the semantics of  $\text{TLA}^+$  doesn’t tell us whether or not “abc” equals 42. TLC considers a model value to be comparable to, and unequal to, any other value. The precise rules for comparability are given in Section 13.6.1.

### 13.2.2 How TLC Evaluates Expressions

To check whether an invariant is true in a state, TLC must evaluate the invariant, meaning that it must compute the TLC value (TRUE or FALSE) of the invariant. It does this in a straightforward way, generally evaluating subexpressions “from left to right”. In particular:

- TLC evaluates  $p \wedge q$  by first evaluating  $p$  and, if it equals TRUE, then evaluating  $q$ .
- TLC evaluates  $p \vee q$  by first evaluating  $p$  and, if it equals FALSE, then evaluating  $q$ . It evaluates  $p \Rightarrow q$  as  $\neg p \vee q$ .

- TLC evaluates  $\text{IF } p \text{ THEN } e_1 \text{ ELSE } e_2$  by first evaluating  $p$ , then evaluating either  $e_1$  or  $e_2$ .

TLC evaluates  $\exists x \in S : p$  by enumerating the elements  $s_1, \dots, s_n$  of  $S$  in some order and then evaluating  $p$  with  $s_i$  substituted for  $x$ , successively for  $i = 1, \dots, n$ . TLC enumerates the elements of a set  $S$  in a very straightforward way, and it gives up and declares an error if it isn't obvious from the form of the expression  $S$  that the set is finite. For example, it can enumerate the elements of the following three set expressions:

$$\{0, 1, 2, 3\} \quad 0 \dots 3 \quad \{i \in 0 \dots 5 : i < 4\}$$

It cannot enumerate the elements of

$$\{i \in \text{Nat} : i < 4\}$$

The rules for what sets TLC can enumerate, along with a complete specification of how TLC evaluates expressions, are given elsewhere in Section 13.6.3 below.

TLC evaluates the expressions  $\forall x \in S : p$  and  $\text{CHOOSE } x \in S : p$  by first enumerating the elements of  $S$ , much the same way as it evaluates  $\exists x \in S : p$ . The semantics of  $\text{TLA}^+$  states that  $\text{CHOOSE } x \in S : p$  is an arbitrary value if there is no  $x$  in  $S$  for which  $p$  is true. However, this case almost always arises because of a mistake, so TLC treats it as an error. Note that evaluating the expression

$$\text{IF } n > 5 \text{ THEN } \text{CHOOSE } i \in 1 \dots n : i > 5 \text{ ELSE } 42$$

will not produce an error if  $n \leq 5$  because TLC will not evaluate the  $\text{CHOOSE}$  expression in that case.

TLC cannot evaluate “unbounded” quantifiers or  $\text{CHOOSE}$  expressions—that is, expressions having one of the forms:

$$\exists x : p \quad \forall x : p \quad \text{CHOOSE } x : p$$

It cannot evaluate any expression whose value is not a TLC value; in particular, it can evaluate a set-valued expression only if it equals a finite set; it can evaluate a function-valued expression only if it equals a function with finite domain. TLC will evaluate expressions of the following forms only if it can tell syntactically that  $S$  is a finite set:

$$\begin{array}{lll} \exists x \in S : p & \forall x \in S : p & \text{CHOOSE } x \in S : p \\ \{x \in S : p\} & \{e : x \in S\} & [x \in S \mapsto e] \\ \text{SUBSET } S & \text{UNION } S & \end{array}$$

TLC can often evaluate an expression even when it can't evaluate all subexpressions. For example, it can evaluate

$$[n \in \text{Nat} \mapsto n * (n + 1)][3]$$

The rules explaining exactly what TLC can evaluate appear in Section 13.6, but you don't have to know them to use TLC.

which equals the TLC value 12, even though it can't evaluate

$$[n \in \text{Nat} \mapsto n * (n + 1)]$$

which equals a function whose domain is the set *Nat*. (A function can be a TLC value only if its domain is a finite set.)

### 13.2.3 Assignment and Replacement

As we saw in the alternating bit example, the configuration file must determine the value of each constant parameter. To assign a TLC value  $v$  to a constant parameter  $c$  of the specification, we write  $c = v$  in the **CONSTANT** section of the configuration file. The value  $v$  may be a primitive TLC value or a finite set of primitive TLC values written in the form  $\{v_1, \dots, v_n\}$ —for example,  $\{1, -3, 2\}$ . In  $v$ , any sequence of characters like **a1** or **foo** that is not a number, a quoted string, or **TRUE** or **FALSE** is taken to be a model value.

In the assignment  $c = v$ , the symbol  $c$  need not be a constant parameter; it can also be a defined symbol. This assignment causes TLC to ignore the actual definition of  $c$  and to take  $v$  to be its value. Such an assignment is often used when TLC cannot compute the value of  $c$  from its definition. For example, TLC cannot compute the value of *NotAnS* from the definition:

$$\text{NotAnS} \triangleq \text{CHOOSE } n : n \notin S$$

because it cannot evaluate the unbounded **CHOOSE** expression. You can override this definition by assigning *NotAnS* a value in the **CONSTANT** section of the configuration file. For example, the assignment

$$\text{NotAnS} = \text{NS}$$

causes TLC to assign to *NotAnS* the model value *NS*. TLC ignores the actual definition of *NotAnS*. If you used the name *NotAnS* in the specification, you'd probably want TLC's error messages to call it **NotAnS** rather than **NS**. So, you'd probably use the assignment

$$\text{NotAnS} = \text{NotAnS}$$

which assigns to the symbol *NotAnS* the model value **NotAnS**. Remember that, in the assignment  $c = v$ , the symbol  $c$  must be defined or declared in the  $\text{TLA}^+$  module, and  $v$  must be a primitive TLC value or a finite set of such values.

The **CONSTANT** section of the configuration file can also contain *replacements* of the form  $c \leftarrow d$ , where  $c$  and  $d$  are symbols defined in the  $\text{TLA}^+$  module. This causes TLC to replace  $c$  by  $d$  when performing its calculations. One use of replacement is to give a value to an operator parameter. For example, suppose we wanted to use TLC to check the write-through cache specification of Section 5.6 (page 54). The *WriteThroughCache* module extends the *MemoryInterface* module, which contains the declaration

Note that **d** is a defined symbol in the replacement  $c \leftarrow d$ , while **v** is a TLC value in the substitution  $c = v$ .

CONSTANTS  $Send(-, -, -, -)$ ,  $Reply(-, -, -, -)$ , ...

To use TLC, we have to tell it how to evaluate the operators  $Send$  and  $Reply$ . We do this by first writing a module  $MCWriteThroughCache$  that extends the  $WriteThroughCache$  module and defines two operators

$$\begin{aligned} MCSend(p, d, old, new) &\triangleq \dots \\ MCreply(p, d, old, new) &\triangleq \dots \end{aligned}$$

We then add to the CONSTANT section of the configuration file the replacements:

```
Send <- MCSend
Reply <- MCreply
```

A replacement can also replace one defined symbol by another. In a specification, we usually write the simplest possible definitions. A simple definition is not always the easiest one for TLC to use. For example, suppose our specification requires an operator  $Sort$  such that  $Sort(S)$  is a sequence containing the elements of  $S$  in increasing order, if  $S$  is a finite set of numbers. Our specification in module  $SpecMod$  might use the simple definition:

$$\begin{aligned} Sort(S) &\triangleq \text{CHOOSE } s \in [1 \dots Cardinality(S) \rightarrow S] : \\ &\quad \forall i, j \in \text{DOMAIN } s : (i < j) \Rightarrow (s[i] < s[j]) \end{aligned}$$

To evaluate  $Sort(S)$  for a set  $S$  containing  $n$  elements, TLC has to enumerate the  $n^n$  elements in the set  $[1 \dots n \rightarrow S]$  of functions. This may be unacceptably slow. We could write a module  $MCSpecMod$  that extends  $SpecMod$  and defines  $FastSort$  so it equals  $Sort$ , when applied to finite sets of numbers, but can be evaluated more efficiently by TLC. We could then run TLC with a configuration file containing the replacement

```
Sort <- FastSort
```

The following definition of  $FastSort$  requires TLC to perform only on the order of  $n^2$  operations to sort an  $n$ -element set:

$$\begin{aligned} FastSort(S) &\triangleq \\ \text{LET } Insert[s \in Seq(Nat), e \in Nat] &\triangleq \\ \text{IF } Len(s) = 0 \text{ THEN } \langle e \rangle & \\ \text{ELSE IF } e < s[1] \text{ THEN } \langle e \rangle \circ \langle s \rangle & \\ \text{ELSE } \langle Head(s) \rangle \circ Insert[Tail(s), e] & \\ MCS[s \in Seq(Nat), SS \in SUBSET S] &\triangleq \\ \text{IF } SS = \{\} \text{ THEN } s & \\ \text{ELSE LET } e \triangleq \text{CHOOSE } ee \in SS : \text{TRUE} & \\ \text{IN } MCS[Insert[s, e], SS \setminus e] & \\ \text{IN } MCS[\langle \rangle, S] & \end{aligned}$$

An even more efficient way to define  $FastSort$  is described in Section 13.3.3, on page 224 below.

### 13.2.4 Overriding Modules

TLC cannot compute  $2 + 2$  from the definition of  $+$  contained in the standard *Naturals* module. Even if we did use a definition of  $+$  from which TLC could compute sums, it would not do so very quickly. Arithmetic operators like  $+$  are implemented directly in Java, the language in which TLC is written. This is achieved by a general mechanism of TLC that allows a module to be overridden by a Java class that implements the operators defined in the module. When TLC encounters an `EXTENDS Naturals` statement, it reads in the Java class that overrides the *Naturals* module rather than reading the module itself. There are Java classes to override the following standard modules: *Naturals*, *Integers*, *Sequences*, *FiniteSets*, and *Bags*. (The *TLC* module described below in Section 13.3.3 is also overridden by a Java class.) Instructions for implementing Java classes to override other modules will appear elsewhere.

### 13.2.5 How TLC Computes States

When TLC evaluates the invariant, it is calculating the invariant's value, which is either `TRUE` or `FALSE`. When TLC evaluates the initial predicate or the next-state action, it is computing a set of states—the set of initial states or the set of possible *successor states* (primed states) for a given (unprimed) state. I will describe how TLC does this for the next-state relation; the evaluation of the initial predicate is analogous.

Recall that a state is an assignment of values to variables. TLC begins computing the successors to a given state  $s$  by assigning to all unprimed variables their values in state  $s$ , and assigning no values to the primed variables. It then starts computing the next-state action.

TLC computes a set of successor states by simultaneously performing a set of computations, each trying to compute a single successor state. A computation splits into multiple separate computations when multiple possibilities are encountered. TLC begins a single computation to evaluate the next-state relation. The evaluation proceeds as described in Section 13.2.2 (page 209), except that when it evaluates a subformula  $A \vee B$ , it splits the computation into two separate computations—in one taking the subformula to be  $A$  and in the other taking the subformula to be  $B$ . Similarly, when it evaluates  $\exists x \in S : p$ , it splits the computation into multiple subcomputations, one for each element of  $S$ .

TLC reports an error if, in its evaluation, it encounters a primed variable that is not assigned a value—except that:

- It evaluates a conjunct of the form  $x' = e$  when  $x'$  has no value by evaluating  $e$  and assigning to  $x'$  the value it obtains.
- It evaluates a conjunct of the form  $x' \in S$  as if it were  $\exists v \in S : x' = v$ .

A computation that obtains the value FALSE from evaluating the next-state action finds no state. A computation that obtains the value TRUE finds the state determined by the values assigned to the primed variables. In the latter case, TLC reports an error if some primed variable has not been assigned a value.

Since TLC evaluates expressions from left to right, the order in which conjuncts appear can affect whether or not TLC can evaluate the next-state action. For example, it can evaluate:

$$\begin{array}{ll} (x' = x + 1) \wedge (y' = x' + 1) & \text{but not } (y' = x' + 1) \wedge (x' = x + 1) \\ (x' \in \{1, 2, 3\}) \wedge (x' \neq x) & \text{but not } (x' \neq x) \wedge (x' \in \{1, 2, 3\}) \end{array}$$

### 13.2.6 When TLC Computes What

When trying to figure out what caused an error, it helps to understand the exact order in which TLC performs its computations. TLC executes the following algorithm. This algorithm depends on whether or not TLC is checking for deadlock, which is determined by the switches with which TLC is run. (See Section 13.3.1.)

1. Precompute all constant definitions.
2. Compute all initial states.
3. For each initial state, evaluate the invariant on the state. If the invariant is satisfied, put the state in  $\mathcal{R}$ ; otherwise, report an error and stop.
4. Set the queue  $\mathcal{U}$  equal to all the initial states that satisfy the constraint, arranged in some arbitrary order.
5. If the queue  $\mathcal{U}$  is empty, stop.
6. Remove the state  $s$  from the head of  $\mathcal{U}$  and do the following:
  - (a) If there is some state  $t$  such that  $s \rightarrow t$  satisfies the next-state relation, then go to step 6b. If not, then if TLC is checking for deadlock, stop and report an error; otherwise go to step 5.
  - (b) Compute some states  $t$  such that  $s \rightarrow t$  satisfies the next-state relation. For each of these states  $t$  that is not in  $\mathcal{R}$ , do the following:
    - i. If  $t$  does not satisfy the invariant, report an error and stop.
    - ii. If  $t$  satisfies the constraint, add it to the end of the queue  $\mathcal{U}$ .
  - (c) If there is any other state  $t$  such that  $s \rightarrow t$  satisfies the next-state relation, then go to step 6b. Otherwise, go to step 5.

TLC can use multiple threads, so step 6 may be performed concurrently by different threads for different states  $s$ .

### 13.2.7 Random Simulation

I have described how TLC tries to find all reachable states that satisfy a model specified by the configuration file. TLC can also be used in *simulation mode* to generate randomly chosen behaviors that satisfy the specification and check that they satisfy the invariant. TLC generates a behavior by randomly choosing a state  $s_0$  that satisfies the initial predicate, and then randomly choosing a sequence of states  $s_1, s_2, \dots$ , such that each transition  $s_i \rightarrow s_{i+1}$  satisfies the next-state action. You can specify a maximum behavior length. When TLC has generated a behavior with that many states, it starts the process again to generate another behavior. TLC continues until it either finds an error or you stop it.

In simulation mode, there is no need to specify a constraint. (TLC ignores it if you do.) You will probably first run TLC on models that are small enough so it can generate all reachable states within a reasonable length of time. When you have found all the errors you can in this way, you can then search for more errors by letting TLC generate random behaviors on larger models.

For some specifications, it can take TLC a long time to generate a transition when the state gets large. For example, suppose the next-state relation contains a disjunct of the form

$$\begin{aligned} &\wedge \dots \\ &\wedge T' \in \text{SUBSET } S \times S \\ &\wedge \text{Pred}(T') \end{aligned}$$

where  $S$  and  $T$  are set-valued variables and  $\text{Pred}$  is some Boolean-valued operator. To choose a possible next state, TLC may have to examine all the subsets of  $S \times S$  to find a value of  $T'$  satisfying  $\text{Pred}(T')$ . If  $S$  has  $n$  elements, then there are  $2^{n^2}$  such subsets. You will probably want to limit the length of behaviors generated by TLC to be small enough so that  $S$  cannot become too large.

TLC makes its random choices using a pseudorandom number generator. Pseudorandom number generation is controlled with a seed. Running a random simulation twice with the same seed produces identical results. You can either let TLC choose a random seed, or specify the seed with a switch, as described on the next page in Section 13.3.1. (TLC always prints the seed it is using.) If TLC finds an error, you can also get it to rerun just the error trace; see the description of the *aril* switch on page 217.

## 13.3 How to Use TLC

### 13.3.1 Running TLC

Exactly how you run TLC depends upon what operating system you are using and how it is configured. You will probably type a command of the following form

*program\_name switches spec\_file*

where:

*program\_name* is specific to your system. It might be something like `java TLC`.

*spec\_file* is the name of the file containing the TLA<sup>+</sup> specification. Each TLA<sup>+</sup> module that appears in the specification must be in a separate file named *M.tla*, where *M* is the name of the module. The extension *.tla* may be omitted from *spec\_file*.

*switches* is a possibly empty sequence of switches. TLC accepts the following switches:

**-config** *config\_file*

Specifies that the configuration file is named *config\_file*, which must be a file with extension *.cfg*. The extension *.cfg* may be omitted from *config\_file*. If this switch is omitted, the configuration file is assumed to have the same name as *spec\_file*, except with the extension *.cfg*.

**-deadlock**

Tells TLC not to check for deadlock. TLC checks for deadlock unless this switch is present.

**-simulate**

Tells TLC to generate randomly chosen behaviors, instead of generating all reachable states. (See Section 13.2.7 above.)

**-depth** *num*

This switch tells TLC that, in simulation mode, it should generate behaviors of length at most *num*. If the switch is absent, then TLC will use a default value of 100. This switch is meaningful only when the **-simulate** switch is present.

**-seed** *num*

In simulation mode, the behaviors generated by TLC are determined by the initial “seed” given to a pseudorandom number generator. This switch tells TLC to let the seed be *num*, which must be an integer from  $-2^{63}$  to  $2^{63} - 1$ . Running TLC twice in simulation mode

with the same seed and `aril` (see the `aril` switch below) will produce identical results. If this switch is omitted, TLC chooses a random seed. This switch is meaningful only when the `-simulate` switch is present.

**-aril *num***

The switch tells TLC that, in simulation mode, it should use *num* as the *aril*. The *aril* is a modifier of the initial seed. When TLC finds an error in simulation mode, it prints out both the initial seed and an *aril* number. Using this initial seed and *aril* will cause the first trace generated to be that error trace. Adding *Print* expressions will usually not change the order in which TLC generates traces. So, if the trace doesn't tell you what went wrong, you can try running TLC again on just that trace to print out additional information.

**-recover *run\_id***

This switch tells TLC to start executing the specification not from the beginning, but where it left off at the last checkpoint. When TLC takes a checkpoint, it prints the run identifier. (That identifier is the same throughout an execution of TLC.) The value of *run\_id* should be that run identifier.

**-cleanup**

TLC creates a number of files when it runs. When it completes, it erases all of them. If TLC finds an error, or if you stop it before it finishes, TLC can leave some large files around. The `-cleanup` option tells TLC to delete all files created by previous runs. Do not use this option if you are currently running another copy of TLC in the same directory; if you do, it can cause the other copy to fail.

**-workers *num***

Step 6 of the TLC execution algorithm described on page 214 can be speeded up on a multiprocessor computer by the use of multiple threads. This switch tells TLC to use *num* threads. There is never any point to using more threads than there are actual processors on your computer. If the switch is omitted, TLC uses a single thread.

### 13.3.2 Debugging a Specification

When you write a specification, it usually contains errors. The purpose of running TLC is to find as many of those errors as possible. Hopefully, an error in the specification causes TLC to report an error. The challenge of debugging is to find the error in the specification that leads to the error that TLC reports. Before addressing that problem, let's first understand TLC's output when it finds no error.

## TLC's Normal Output

When you run TLC, the first thing it prints is the version number and creation date—something like:

```
TLC Version 1.0 of 26 May 1999
```

Always include this information when reporting any problems with TLC. Next, TLC describes the mode in which it's being run. The possibilities are

```
Model-checking
```

in which it is exhaustively checking all reachable states, or

```
Running Random Simulation with seed 1901803014088851111.
```

in which it is doing random simulation using the indicated seed. (See section 13.2.7.) Let's suppose it's doing model checking. TLC next types something like:

```
Finished computing initial states:
4 states generated, with 2 of them distinct.
```

This indicates that, when evaluating the initial predicate, TLC generated 4 states, among which there were 2 distinct ones. TLC then types one or more messages like

```
Progress: 2846 states generated, 984 distinct states found.
856 states left on queue.
```

This message indicates that TLC has thus far generated and examined 2846 states, it has found 984 distinct ones, and the queue  $\mathcal{U}$  of unexplored states contains 856 states. (See Section 13.2.6 on page 214.) After running for a while, TLC generates these progress reports about once every five minutes. For most specifications, the number of states on the queue increases monotonically at the beginning of the execution and decreases monotonically at the end. The progress reports therefore provide a useful guide to how much longer the execution is likely to take.

When TLC successfully completes, it prints

```
Model checking completed. No error has been found.
```

It then prints something like:

```
Estimates of the probability that TLC did not check all reachable
states because two distinct states had the same fingerprint:
  calculated (optimistic): .000003
  based on the actual fingerprints: .00007
```

As explained in Section 13.6.4 on page 244, there is a chance that TLC did not examine the complete set of reachable states. TLC prints two different estimates of that probability. The first estimate is generally lower and more optimistic; the second is perhaps a more realistic one.

Finally, TLC prints something like

```
2846 transitions taken. 984 states discovered. 0 states left
on queue.
```

with the grand totals.

While TLC is running, it may also print something like

```
-- Checkpointing run states/99-05-20-15-47-55 completed
```

This indicates that it has written a checkpoint that you can use to restart TLC in the event of a computer failure. (As explained in Section 13.3.6 on page 230, checkpoints have other uses as well.) The run identifier

```
states/99-05-20-15-47-55
```

is used with the `-recover` switch to restart TLC from where the checkpoint was taken. (If only part of this message was printed—for example, because your computer crashed while TLC was taking the checkpoint—there is an extremely small chance that all the checkpoints are corrupted and you must start TLC again from the beginning.)

## Error Reports

The first problems you find in your specification will probably be syntax errors. TLC reports them with

```
ParseException in parseSpec:
```

followed by the error message generated by the parser. Chapter 12 explains how to interpret the parser’s error messages. (Note: TLC Version 1.0 does not use the parser’s full error detection mechanism; you should check your specification with the parser before running TLC on it.)

After parsing, TLC executes two basic phases: in the first, it computes the initial states and in the second it generates the successor states of states on the queue of unexplored states. You can tell which phase TLC is in by whether or not it has printed the “initial states computed” message.

TLC evaluates the invariant in both phases—on the initial states in the first phase, and on newly generated successor states in the second. TLC’s most straightforward error report occurs when the invariant is violated. Suppose we introduce an error into our alternating bit specification (Figures 13.1 and 13.2

on pages 204 and 205) by replacing the first conjunct of the invariant *TypeInv* with

$$\wedge \text{msgQ} \in \text{Seq}(\text{Data})$$

TLC quickly finds the error and types

**Invariant TypeInv is violated**

It next prints a minimal-length<sup>2</sup> behavior that leads to the state not satisfying the invariant:

The behavior up to this point is:

STATE 1:

```

/\ rBit = 0
/\ ackQ = << >>
/\ rcvd = << >>
/\ sent = << >>
/\ sAck = 0
/\ sBit = 0
/\ msgQ = << >>

```

STATE 2:

```

/\ rBit = 0
/\ ackQ = << >>
/\ rcvd = << >>
/\ sent = << d1 >>
/\ sAck = 0
/\ sBit = 1
/\ msgQ = << << 1, d1 >> >>

```

Observe that TLC prints each state as a TLA<sup>+</sup> predicate that determines the state. When printing a state, TLC describes functions using the operators  $:>$  and  $@@$ , where

$$(d_1 :> e_1 @@ \dots d_n :> e_n)$$

is the function  $f$  with domain  $\{d_1, \dots, d_n\}$  such that  $f[d_i] = e_i$ , for  $i = 1, \dots, n$ . These operators are defined by the *TLC* module, described in Section 13.3.3 (page 222). For example, the sequence  $\langle \text{"ab"}, \text{"cd"} \rangle$ , which is a function with domain  $\{1, 2\}$ , can be written as

$$(1 :> \text{"ab"} @@ 2 :> \text{"cd"})$$

---

<sup>2</sup>When using multiple threads, it there is a slight chance that there exists a shorter behavior that also violates the invariant.

TLC generally prints values the way they appear in the specification, so a sequence will be printed as a sequence, rather than with this function notation.

The hardest errors to locate are usually the ones that TLC detects when evaluating an expression. They may occur when evaluating the initial predicate, the next-state action, or an invariant. These errors are detected when TLC is forced to evaluate an expression that it can't handle, or one that is "silly" because its value is not specified by the semantics of TLA<sup>+</sup>. As an example of a silliness error, let's return again to the alternating bit protocol and replace the THEN clause in the fourth conjunct of the definition of *RcvMsg* with

```
THEN Append(rcvd, Head(msgQ)[3])
```

TLC discovers the error because the elements of *msgQ* are pairs, so 3 is not an element in the domain of *Head(msgQ)*. TLC reports the error by printing:

```
Error: Applying tuple to an integer out of domain.
```

It then prints a behavior leading to the error. TLC finds the error when evaluating the next-state action to compute the successor states for some state *s*, and *s* is the last state in that behavior. Had it found the error when evaluating the invariant, the behavior would end with the state in which TLC was evaluating the invariant. Finally, TLC prints the location of the error:

```
The error occurred when TLC was evaluating the nested
expressions at the following positions:
0. Line 44, column 6 to line 45, column 61 in AlternatingBit
```

This position identifies the entire conjunct

```
rcvd' = IF rBit' ≠ rBit THEN Append(rcvd, Head(msgQ)[3])
        ELSE rcvd
```

TLC is not very precise when indicating the position of an error; usually it just narrows it down to a conjunct or disjunct of a formula. In general, it prints a tree of nested expressions—higher-level ones first. This can be helpful if the error occurs when evaluating the definition of an operator that is used in several places.

## Debugging

Tracking down an error can be difficult—especially in TLC Version 1, which does not provide you with very much information. (Sometimes, it doesn't even print the location of the error.) The *TLC* module provides some operators that can help in debugging.<sup>3</sup>

<sup>3</sup>Actually, it is the Java class that overrides the *TLC* module that provides the useful functionality of these operators.

The *TLC* module defines the operator *Print* so that *Print(out, val)* equals *val*. But, when TLC evaluates this expression, it prints the values of *out* and *val*. You can add *Print* expressions to a specification to help locate an error. For example, if your specification contains

$$\begin{aligned} &\wedge \textit{Print}(\text{"a"}, \text{TRUE}) \\ &\wedge P \\ &\wedge \textit{Print}(\text{"b"}, \text{TRUE}) \end{aligned}$$

and TLC prints the "a" but not the "b" before reporting an error, then the error occurs while TLC is evaluating *P*. If you know where the error is but don't know why it's occurring, you can add *Print* expressions to give you more information about what values TLC has computed.

To understand what will be printed when, you must know how TLC evaluates expressions, which is explained in Section 13.2.2 on page 209. An expression is typically evaluated many times by TLC, so inserting a *Print* expression in the specification can produce a lot of output. You can limit the amount of output by putting the *Print* expression inside an IF/THEN expression, so it is executed only in interesting cases.

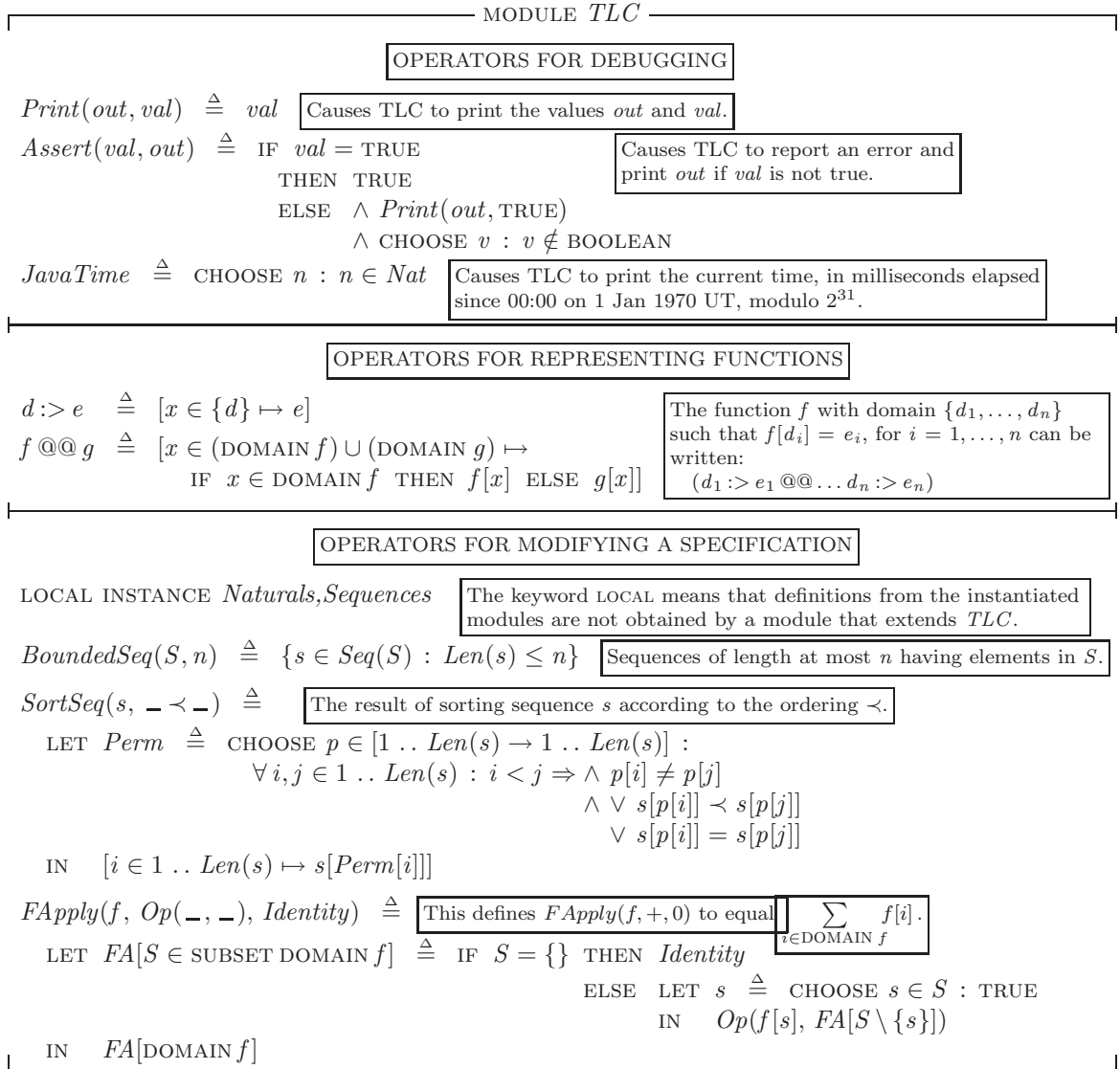
Two other debugging operators defined in the *TLC* module are *Assert* and *JavaTime*. The expression *Assert(val, out)* equals TRUE if *val* equals TRUE. However, if *val* does not equal TRUE, then *Assert(val, out)* equals an illegal value, and evaluating it causes TLC to print the value of *out* and to halt. The expression *JavaTime* equals the current time at which TLC evaluates the expression. That time is given as the number of milliseconds elapsed since 00:00 Universal Time on 1 January 1970, modulo  $2^{31}$ . If TLC is generating states slowly, using the *JavaTime* operator in conjunction with *Print* expressions can help you understand why. If TLC is spending too much time evaluating an operator, you may be able to replace the operator with an equivalent one that TLC can evaluate more efficiently. (See Section 13.2.3 on page 211.)

### 13.3.3 The *TLC* Module

The standard *TLC* module, in Figure 13.5 on the next page, contains operators that are handy when using TLC. The module on which you run TLC will normally extend the *TLC* module. The TLC module is overridden by its Java implementation.

Module *TLC* first defines three operators *Print*, *Assert*, and *JavaTime* that are of no use except when running TLC. They are explained in Section 13.3.2 on the preceding page. That section also describes the operators *:>* and *@@*, which are used for explicitly writing functions. These operators could be useful when writing specifications, even if you're not using TLC.

The module next defines *BoundedSeq(S, n)* to be the set of sequences of elements of *S* of length at most *n*. TLC cannot evaluate *BoundedSeq(S, n)* from

Figure 13.5: The standard module *TLC*

the definition in the module, since it would first have to evaluate the infinite set  $Seq(S)$ . TLC could evaluate it from the equivalent definition

$$BoundedSeq(S, n) \triangleq \text{UNION } \{[1 \dots m \rightarrow S] : m \in 0 \dots n\}$$

but evaluation is faster with the overriding Java implementation.

The TLC module next defines the operator  $SortSeq$ . If  $s$  is a finite sequence and  $<$  is a total ordering relation on its elements, then  $SortSeq(s, <)$  is the sequence obtained from  $s$  by sorting its elements according to  $<$ . For example,  $SortSeq(\langle 3, 1, 3, 8 \rangle, >)$  equals  $\langle 8, 3, 3, 1 \rangle$ . The Java implementation of  $SortSeq$  allows TLC to evaluate it more efficiently than a user-defined sorting operator. For example, because the following definition of the operator  $FastSort$ , which makes use of  $SortSeq$  to do the actual sorting, TLC can evaluate it faster than the definition given above on page 212.

$$\begin{aligned} FastSort(S) &\triangleq \\ &\text{LET } MakeSeq[SS \in \text{SUBSET } S] \triangleq \\ &\quad \text{IF } SS = \{\} \text{ THEN } \langle \rangle \\ &\quad \quad \text{ELSE LET } ss \triangleq \text{CHOOSE } ss \in SS : \text{TRUE} \\ &\quad \quad \text{IN } Append(MakeSeq[SS \setminus \{ss\}], ss) \\ &\text{IN } SortSeq(MakeSeq[S], <) \end{aligned}$$

The *TLC* module ends by defining the operator  $FApply$ . If  $f$  is a function with finite domain, and  $Op$  is an operator that takes two arguments, then  $FApply(f, Op, Id)$  equals

$$Op(f[d_1], Op(f[d_2], Op(\dots Op(f[d_n], Id) \dots)))$$

where  $d_1, \dots, d_n$  are the elements in the domain of  $f$ , listed in some arbitrary order. Using  $FApply$  to avoid a recursive definition can speed up TLC. For example, the factorial function  $fact$  can be defined by

$$fact(n) \triangleq FApply([i \in 1 \dots n \mapsto i], *, 1)$$

TLC can compute factorials from this definition faster than from the standard recursive definition on page 54.

### 13.3.4 Checking Action Invariance

Section 5.7 on page 60 discusses two different kinds of invariants. A state predicate  $Inv$  that satisfies  $Spec \Rightarrow \Box Inv$  is called an invariant of the specification  $Spec$ . This is the kind of invariance property that TLC checks. If  $Inv$  satisfies  $Inv \wedge [Next]_v \Rightarrow Inv'$ , then  $Inv$  is called an invariant of the action  $[Next]_v$ . It's not hard to see that  $Inv$  is an invariant of action  $[Next]_v$  iff it satisfies  $Inv \wedge \Box [Next]_v \Rightarrow \Box Inv$ . We can therefore check if  $Inv$  is an invariant of an

When  $v$  is the tuple of all relevant variables,  $Inv$  is an invariant of  $[Next]_v$  iff it is an invariant of  $Next$ .

$$\begin{aligned}
ABInv &\triangleq \\
&\wedge TypeInv \\
&\wedge Inv \\
&\wedge \text{IF } rBit = sBit \\
&\quad \text{THEN } \wedge Len(sent) = Len(rcvd) \\
&\quad \quad \wedge (sent = \langle \rangle) \Rightarrow (msgQ = \langle \rangle) \wedge (sAck = sBit) \\
&\quad \quad \wedge \forall i \in 1 \dots Len(msgQ) : msgQ[i] = \langle sBit, sent[Len(sent)] \rangle \\
&\quad \text{ELSE } \wedge Len(sent) \neq Len(rcvd) \\
&\quad \quad \wedge (msgQ \neq \langle \rangle) \Rightarrow \exists i \in 0 \dots Len(msgQ) : \\
&\quad \quad \quad \wedge (i \neq 0) \Rightarrow (rcvd \neq \langle \rangle) \\
&\quad \quad \quad \wedge \forall j \in 1 \dots i : msgQ[j] = \langle rBit, rcvd[Len(rcvd)] \rangle \\
&\quad \quad \quad \wedge \forall j \in (i + 1) \dots Len(msgQ) : \\
&\quad \quad \quad \quad msgQ[j] = \langle sBit, sent[Len(sent)] \rangle \\
&\quad \quad \wedge sAck \neq sBit \\
&\wedge \text{IF } rBit = sAck \\
&\quad \text{THEN } \forall i \in 1 \dots Len(ackQ) : ackQ[i] = rBit \\
&\quad \text{ELSE } (ackQ \neq \langle \rangle) \Rightarrow \\
&\quad \quad \exists i \in 0 \dots Len(ackQ) : \wedge \forall j \in 1 \dots i : ackQ[j] = sAck \\
&\quad \quad \wedge \forall j \in (i + 1) \dots Len(ackQ) : ackQ[j] = rBit
\end{aligned}$$

Figure 13.6: An invariant of the alternating bit protocol's next-state action.

action *Next* by running TLC with *Inv* as both the initial condition and the invariant.

As an example, let's return to the protocol of module *AlternatingBit* (pages 204–205). The predicate *Inv* is an invariant of the specification *Spec*, but not of the next-state action *ABNext*. To prove that *Inv* is an invariant of *Spec*, we must find an invariant of *ABNext* that is true in an initial state and implies *Inv*. Such an invariant *ABInv* is defined in Figure 13.6 on this page. Don't worry about the details of this invariant; just observe that it is the conjunction of the type invariant *TypeInv*, the invariant *Inv* of the specification, and another formula.

Before checking that *ABInv* is an invariant of the next-state relation, we should first make sure that it's an invariant of the specification. We add the definition of *ABInv* to module *MCAAlternatingBit*, modify the configuration file to use *ABInv* as the invariant, and run TLC. After correcting the inevitable typing mistakes, we find that *ABInv* does appear to be an invariant of the specification. (Remember that TLC only checks *ABInv* on a finite model, it doesn't prove invariance.) We now try running TLC again letting *ABInv* also be the initial predicate. TLC reports the error:

```

While computing initial states, TLC was trying to compute
the values of a variable v from an expression
v \in S, but S was not enumerable.

```

at line 20, column 15 to line 20, column 41 in: *AlternatingBit*

Recall that TLC computes the initial states by computing the conjuncts of the initial predicate in order. (See Section 13.2.5 on page 213). From the definition of *TypeInv*, we see that it first tries to compute the possible values of *msgQ* by evaluating

$$msgQ \in Seq(\{0, 1\} \times Data)$$

Since  $Seq(\{0, 1\} \times Data)$  is an infinite set, containing sequences of any length, this yields an infinite number of possible values of *msgQ*, so TLC gives up. (Remember that TLC first computes all states satisfying the initial condition, then it throws away those that don't satisfy the constraint.) So, we must modify *ABInv* to specify only a finite set of initial values. We add to the *MCAAlternatingBit* module a predicate *BTypeInv* that is a bounded version of *TypeInv*. For example, it specifies the initial value of *msgQ* by

$$msgQ \in BoundedSeq(\{0, 1\} \times Data, msgQLen)$$

where the operator *BoundedSeq*, defined in the *TLC* module, is explained on pages 222–224, and *msgQLen* is the parameter of *MCAAlternatingBit* used in the constraint to bound the length of *msgQ*. We then define *BdedABInv* to be the same as *ABInv*, except with *TypeInv* replaced by *BTypeInv*, and run TLC with *BdedABInv* as the initial condition and *ABInv* as the invariant, using the same constraint as before. (Do you see why we can't use *BdedABInv* as the invariant? If not, review Section 13.2.6.)

When computing the initial states, TLC examines every state satisfying *BTypeInv*, throwing away those that don't satisfy the rest of *BdedABInv*. Since *BoundedSeq(S, n)* has more than  $Cardinality(S)^n$  elements, the number of initial states that TLC examines is greater than

$$8 * 2^{msgQLen+ackQLen} * D^{msgQLen+2*sentLen}$$

where *D* is the cardinality of *Data*. (I am taking *sentLen* to be the bound on the length of the queues *sent* and *rcvd*.) This number grows much faster than the number of reachable states, estimated by formula (13.2) on page 208. Typically, TLC can check an invariant of an action only on a very small model—much smaller than the models on which it can check an invariant of a specification.

We can use *BdedABInv* as the initial predicate only because its first conjunct is the type invariant *BTypeInv*. If we interchange its first two conjuncts, then TLC will try to evaluate *Inv* before evaluating *BTypeInv*. Evaluating the first conjunct

$$Len(rcvd) \in \{Len(sent) - 1, Len(sent)\}$$

of *Inv* causes TLC to produce the error message

The identifier `rcvd` is not assigned a value in the current context...

When using TLC to check that a predicate  $P$  is an invariant of an action, evaluating  $P$  must specify the value of each variable  $v$  with a conjunct of the form  $v \in S$  before that value is used.

### 13.3.5 Checking Implementation

In industrial applications, a TLA<sup>+</sup> specification is likely to be the only formal description of a system. Such a specification is correct if it means what we intended, which is not a formalizable concept. By formalizing some of our intentions as invariants, we can use TLC to help catch errors in the specification. Sometimes, a specification  $S_L$  is a lower-level view of a system for which we have a higher-level specification  $S_H$ . We usually regard  $S_H$  as the system's specification and  $S_L$  as an implementation, and take correctness of  $S_L$  to mean that it implies  $S_H$ . This situation occurred in Chapter 5, with the implementation  $S_L$  being the specification of the write-through cache (formula *Spec* of module *WriteThroughCache*) and the specification  $S_H$  being the specification of a linearizable memory (formula *Spec* of the *Memory* module).

We saw in Section 5.8 that this requires proving a formula of the form:

$$(13.4) \quad Init \wedge \Box[Next]_v \Rightarrow HInit \wedge \Box[HNext]_{hv}$$

This formula is true iff the following two conditions hold:

1. *Init* implies *HInit*
2. In every behavior satisfying  $Init \wedge \Box[Next]_v$ , every step is an  $[HNext]_{hv}$  step.

TLC allows you to check these conditions as follows. First, specify as usual in the configuration file that *Init* and *Next* are the initial predicate and next-state action. To check condition 1, specify *HInit* as the IMPLIED-INIT predicate by adding the following to the configuration file:<sup>4</sup>

IMPLIED-INIT *HInit*

TLC will report an error if it finds an initial state (one satisfying *Init*) that does not satisfy *HInit*. To check condition 2, first define a new symbol, say *SqHNext*, in the TLA module:

$$SqHNext \triangleq [HNext]_{hv}$$

Then, specify *SqHNext* as the IMPLIED-ACTION formula by adding the following to the configuration file:

---

<sup>4</sup>TLC does not yet support the IMPLIED-INIT field. Until it does, you can have it check condition 1 by running it with the initial predicate  $Init \wedge Assert(HInit, "HInit \text{ false}')$ .

|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| MODULE <i>ABSpec</i>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       |
| EXTENDS <i>Sequences</i><br>CONSTANT <i>Data</i><br>VARIABLES <i>sent, rcvd</i><br>$Init \triangleq \wedge sent = \langle \rangle$<br>$\wedge rcvd = \langle \rangle$<br>$Send(d) \triangleq \wedge sent = rcvd$<br>$\wedge sent' = Append(sent, d)$<br>$\wedge UNCHANGED rcvd$<br>$Rcv \triangleq \wedge sent \neq rcvd$<br>$\wedge rcvd' = Append(rcvd, sent[Len(sent)])$<br>$\wedge UNCHANGED sent$<br>$Next \triangleq Rcv \vee (\exists d \in Data : Send(d))$<br>$HSpec \triangleq Init \wedge \Box[Next]_{\langle sent, rcvd \rangle} \wedge WF_{\langle sent, rcvd \rangle}(Next)$ |

Figure 13.7: A higher-level specification of the alternating bit protocol.

#### IMPLIED-ACTION *SqHNext*

TLC will report an error, and provide an error trace, if it encounters a legal step (a *Next* step from a reachable state) that is not a *SqHNext* step. Don't forget the  $[\dots]_{hv}$ ; otherwise, TLC will report an error if a lower-level step leaves the higher-level variables unchanged.

As an example, let's return again to the alternating bit protocol. Module *ABSpec* of Figure 13.7 on this page contains a high-level specification *HSpec* of the transmission protocol that is implemented by the alternating bit protocol. (In that specification, values are simply transferred directly from *sent* to *rcvd*.)

We check that the alternating bit protocol implies specification *HSpec* by running TLC on module *MCABSpec* in Figure 13.8 on the next page, using the initial predicate *ABInit*, the next-state relation *ABNext*, the IMPLIED-INIT *Init*, and the IMPLIED-ACTION *SqNext*. Note that module *MCABSpec* can extend both the *ABSpec* and *AlternatingBit* modules because there happen to be no name conflicts between the two modules; usually we would have to instantiate one of the modules with renaming.

### 13.3.6 Some Hints

Here are some suggestions for using TLC effectively.

```

┌────────────────────────── MODULE MCABSpec ───────────────────────────┐
EXTENDS AlternatingBit, ABSpec
CONSTANTS msgQLen, ackQLen, sentLen

SqNext  $\triangleq$  [Next]{sent,rcvd}
SeqConstraint  $\triangleq$   $\wedge$  Len(msgQ)  $\leq$  msgQLen
                   $\wedge$  Len(ackQ)  $\leq$  ackQLen
                   $\wedge$  Len(sent)  $\leq$  sentLen
└──────────────────────────────────────────────────────────────────────────┘

```

Figure 13.8: A module for checking the alternating bit protocol.

**Start small**

The specification of a real system probably has more reachable states than you expect. When you start testing it with TLC, use the smallest model you possibly can. Let every set parameter, such as a set of processes, have only one element. Let queues be of length one—or even of length zero, if possible. A specification that has not been tested probably has lots of trivial errors that can be found with any model. You will find them faster with a tiny model.

**Be suspicious of success**

The easiest way to build a system that satisfies an invariant is to have it do nothing. If TLC finds no states in which the invariant is false, that may be because it isn't generating many states. TLA<sup>+</sup> specifications use liveness properties to rule out that possibility. TLC doesn't check liveness properties, so you'll have to use some tricks to look for errors that affect liveness. One trick is to check that progress is possible—namely, that TLC reaches the states that you expect it to. This can be done by using an invariant asserting that the expected states are *not* reached, and making sure that TLC reports that the invariant is violated. For example, we can check that the alternating bit protocol allows progress by using an invariant asserting that the length of the *rcvd* queue is less than *sentLen*. You can also check that states are reached by the judicious use of *Print* or *Assert* expressions. (See Section 13.3.2.)

**Make model values parameters**

When debugging a specification, you may want to refer to model values within TLA<sup>+</sup> expressions. You can do this by making the model values constant parameters. For example, we might add to module *MCAlternatingBit* the declaration

```
CONSTANTS d1, d2
```

and add to the **CONSTANT** section of the configuration file the assignments

`d1 = d1      d2 = d2`

which assign to the constant parameters  $d1$  and  $d2$  the model values `d1` and `d2`, respectively.

### Don't start over after every error

After you've eliminated the errors that are easy to find, TLC may have to run for a long time before finding an error. It often takes more than one try to correctly correct an error, and it can be frustrating to run TLC for an hour only to find that you made a silly mistake in the correction. If the error was discovered when taking a step from a correct state, then it's a good idea to check your correction by starting TLC from that state. You can do this by using the state printed by TLC to define the initial predicate to be used by TLC.

Another way to avoid starting from scratch after an error is by using checkpoints. A checkpoint saves the set of reached states and the queue of unexplored states.<sup>5</sup> It does not save any other information about the specification. You can restart a specification from a checkpoint even if you have changed the specification, as long as the specification's variables and the values that they can assume haven't changed. If TLC finds an error after running for a long time, you may want to continue it from the last checkpoint instead of having it recheck all the states it had already checked.

### Check everything you can

A single invariant seldom expresses correctness of a specification. Write and let TLC check as many different invariance properties as you can. If you think that some predicate should be an invariant, let TLC test if it is. Discovering that the predicate is not an invariant may not signal an error in the specification, but it will probably teach you something about your specification.

### Try to check liveness

Although TLC can't directly check liveness properties, you can sometimes check them indirectly by modifying the specification. A liveness property asserts that something must eventually happen. Sometimes, you can check such a property by checking that the property must hold after the system has taken a certain number of steps. For example, suppose you want to check that the specification

$$Spec \triangleq Init \wedge \Box[Next]_v \wedge WF_v(Next)$$

satisfies the property  $P \leadsto Q$  for some predicates  $P$  and  $Q$ . In most cases,  $Spec$  implies  $P \leadsto Q$  iff it implies that, whenever  $P$  is true,  $Q$  must become true

<sup>5</sup>Some of the reached states may not be saved, so TLC may explore again some states that had been reached before the checkpoint.

The operator  $\leadsto$  (leads-to) is explained on page 91.

within  $N$  steps, where  $N$  is some function of the constant parameters. You can check this by adding a variable  $ctr$  that counts the number of steps taken since  $P$  became true, and is reset when  $Q$  becomes true. You can check that  $P$  always leads to  $Q$  within  $N$  steps by checking that  $ctr$  is never greater than  $N$ . More precisely, you can define a new specification whose initial predicate is

$$Init \wedge (ctr = \text{IF } P \wedge \neg Q \text{ THEN } 0 \text{ ELSE } -1)$$

and whose next-state action is

$$\begin{aligned} &\wedge Next \\ &\wedge ctr' = \text{IF } \neg Q' \wedge (ctr \geq 0 \vee P') \text{ THEN } ctr + 1 \text{ ELSE } -1 \end{aligned}$$

You can then use TLC to check that  $ctr \leq N$  is an invariant of this specification.

#### Use symmetry to save time

For many specifications, TLC will take a long time to check all the reachable states for a nontrivial model. One way to have it use less time is to reduce the number of reachable states it must check. There is often some symmetry in the specification that makes certain states equivalent in terms of finding errors. Two states  $s$  and  $t$  are equivalent if a behavior starting from state  $s$  can produce an error iff one starting from state  $t$  can. If  $s$  and  $t$  are equivalent, there is no need to explore states reachable from both of them. For example, the alternating bit protocol does not depend on the actual values being sent. So, permuting all the data values in any state produces an equivalent state.

The constraint can often be used to keep TLC from exploring the successors of different equivalent states. For example, we can exploit the symmetry of the alternating bit protocol under permutation of data values as follows. We use a model in which *Data* is a set of numbers and conjoin to the constraint the requirement that the elements of *sent* are sorted:

$$\forall i, j \in 1 \dots Len(sent) : (i \leq j) \Rightarrow (sent[i] \leq sent[j])$$

For the model size specified by the configuration file of Figure 13.4, this extra constraint reduces the number of reachable states from 2312 to 1482. When the number of elements in *Data* is increased to 3 and *sentLen* is increased to 4, it reduces the number of reachable states from 22058 to 6234.

## 13.4 What TLC Doesn't Do

We would like TLC to generate all the behaviors that satisfy a specification. But no program can do this for an arbitrary specification. I have already mentioned various limitations of TLC. There are some other limitations that you may

stumble on. One of them is that the Java classes that override the *Naturals* and *Integers* modules will handle only numbers in the interval  $-2^{31} \dots (2^{31} - 1)$ .

An easier goal would be for every behavior that TLC does generate to satisfy the (safety part) of the specification. However, for reasons of efficiency, TLC doesn't always meet this goal. In particular, it doesn't preserve the precise semantics of **CHOOSE**. As explained in Section 15.1, if  $S$  equals  $T$ , then  $\text{CHOOSE } x \in S : P$  should equal  $\text{CHOOSE } x \in T : P$ . However, if the expressions don't have a uniquely-defined value, then TLC guarantees this only if  $S$  and  $T$  are syntactically the same. For example, TLC might compute different values for the two expressions

$$\text{CHOOSE } x \in \{1, 2, 3\} : x < 3 \quad \text{CHOOSE } x \in \{3, 2, 1\} : x < 3$$

A similar violation of the semantics of  $\text{TLA}^+$  exists with **CASE** expressions, whose semantics are defined (in Section 15.1.4) in terms of **CHOOSE**.

TLC Version 1 has the following additional limitations, most of which should be fixed in version 2. Telling us which ones you find to be a nuisance will help ensure that they really are fixed in Version 2.

- TLC doesn't handle specifications that use the **INSTANCE** statement.
- TLC doesn't properly handle Cartesian products of more than two sets. Thus, instead of writing  $S \times T \times U$ , you have to write

$$\{\langle s, t, u \rangle : s \in S, t \in T, u \in U\}$$

- You cannot put infix, prefix, or postfix operators in the **CONSTANT** section of the configuration file. Hence, neither of the following replacements may appear in the configuration file:

`++ <- Foo      Bar <- &`

- You cannot use an infix, prefix, or postfix operator as the argument of a higher-order operator. For example, if *IsPartialOrder* is the operator described on pages 70–71, you cannot write *IsPartialOrder*( $<$ , *Nat*). To get around this problem, you can define

$$\text{LessThan}(a, b) \triangleq a < b$$

and then write *IsPartialOrder*(*LessThan*, *Nat*).

- You cannot define an infix, prefix, or postfix operator in a **LET** expression.
- The Java class that overrides the *Sequences* module does not properly handle the *BoundedSeq* operator. You can define a replacement operator for it, using the definition on pages 222–224, that works properly.
- The Java class that overrides the standard *Bags* module does not implement the *SubBag* operator.

- TLC does not check that a symbol is defined before it is used. This means that TLC will accept illegal recursive operator definitions like

$$\text{Silly}(\text{foo}) \triangleq \text{Silly}(\text{foo} + 1)$$

- The Java classes that override the *Naturals* and *Integers* modules produce incorrect results, rather than an error message, if a computation yields a number outside the interval  $-2^{31} \dots (2^{31} - 1)$ .
- TLC doesn't implement the  $\cdot$  (action composition) operator.
- TLC treats strings as primitive values, and not as functions. It thus considers the legal TLA<sup>+</sup> expression “**abc**”[2] to be an error. It also handles only strings containing letters, numbers, spaces, and the following ASCII characters:

< > ? , . / : ; [ ] { } | ~ ! @ # \$ % ^ & \* ( ) \_ - + =

- TLC allows constants to be replaced by nonconstant operators. Doing so may cause TLC to produce incorrect results. (Version 1 also allows nonconstants to be replaced; this might not be allowed in Version 2.)

Replacement is explained in Section 13.2.3.

## 13.5 Future Plans

The following additions and improvements to future versions of TLC are being considered. It is unlikely that they will all be implemented; recommendations of which ones are most important would be useful.

- Check liveness properties.
- Handle the type of compositional specifications, described in Section 10.2, that have conjuncts of the form  $\Box \text{IsFcnOn}(f, S)$ , and the next-state action specifies only the values of  $f[s]$  for  $s \in S$ .
- Have the progress reports contain separate counts of generated states that do and don't satisfy the constraint.
- Introduce subclasses of model values, where a model value is comparable only with model values in its own subclass. This requires also adding some way of handling the construct  $\text{CHOOSE } x : x \notin S$ .
- Improve debugging information as follows:
  - Print out the context at the point of the error—including variable values, and the values of bound variables.
  - Identify the disjunct of the next-state action responsible for each step when printing a behavior.

- More precisely identify the place within the specification where an error occurs.
- When an invariant is found to be false, have TLC print out why it is false—that is, which conjunct or conjuncts are false.
- Add a debugging mode in which TLC won’t stop for an error, but will just keep going.
- Add some coverage analysis, indicating which parts of the next-state relation never evaluated to TRUE.
- Add some support for using symmetry to reduce the state space searched.
- Have TLC check all assumptions.

## 13.6 The Fine Print: What TLC Really Does

We now describe in more detail what TLC does. This section is for the intellectually curious and mathematically sophisticated. Most users will not want to read it.

This section specifies the results produced by TLC, not the actual algorithms it uses to compute those results. Moreover, it specifies only what TLC guarantees to do. TLC may actually do better, producing a correct result when the specification allows it to report an error.

### 13.6.1 TLC Values

Section 13.2.1 (page 208) describes TLC values. That description was incomplete in two ways. First, it did not precisely define when values are comparable. The precise definition is that two TLC values are comparable iff the following rules imply that they are:

1. Two primitive values are comparable if they have the same value type.  
This rule implies that “abc” and “123” are comparable. But “abc” and 123 are not comparable.
2. A model value is comparable with any value. (It is equal only to itself.)
3. Two sets are comparable if they have different numbers of elements, or if they have the same numbers of elements and all the elements in one set are comparable with all the elements in the other.  
This rule implies that  $\{1\}$  and  $\{“a”, “b”\}$  are comparable and that  $\{1, 2\}$  and  $\{2, 3\}$  are comparable. However,  $\{1, 2\}$  and  $\{“a”, “b”\}$  are not comparable.

4. Two functions  $f$  and  $g$  are comparable if (i) their domains are comparable and (ii) if their domains are equal, then  $f[x]$  and  $g[x]$  are comparable for every element  $x$  in their domains.

This rule implies that  $\langle 1, 2 \rangle$  and  $\langle \text{"a"}, \text{"b"}, \text{"c"} \rangle$  are comparable, and that  $\langle 1, \text{"a"} \rangle$  and  $\langle 2, \text{"bc"} \rangle$  are comparable. However,  $\langle 1, 2 \rangle$  and  $\langle \text{"a"}, \text{"b"} \rangle$  are not comparable.

Section 13.2.1 also failed to mention that additional primitive value types can be introduced by Java classes that override modules. The primitive TLC value types are thus Booleans, integers, strings, model values, and any additional value types introduced by overriding.

TLC Version 1 does not support the introduction of new primitive value types.

### 13.6.2 Overridden Values

As explained in Section 13.2.4 above, an extended module may be overridden by a Java class. All the values and operators defined by the module are replaced by special overridden values. For example, the Java class that overrides the *Naturals* module replaces *Nat* and  $+$  by overridden values. An overridden module may not have parameters.

### 13.6.3 Expression Evaluation

TLC evaluates expressions when performing three different tasks:

- It evaluates the initial predicate to compute the set of initial states.
- It evaluates the next-state action to compute the set of states reachable from a given state by a step of that action.
- It evaluates an invariant when checking that a reachable state is correct.

The result of successfully evaluating an expression is a TLC value. Evaluation may fail, in which case TLC reports an error and halts.

TLC evaluates an expression in a context, which is an assignment of meanings to user-defined symbols, parameters, variables, and primed variables. A meaning is a  $\text{TLA}^+$  expression formed from the built-in operators of  $\text{TLA}^+$ , TLC values, and overridden values. Variables and primed variables can be assigned only TLC values. For example, consider a module that has two variables  $x$  and  $y$ . To compute the next states starting in a state with  $x = 2$  and  $y = \{\text{"a"}\}$ , TLC evaluates the next-state relation in the context  $\mathcal{C}$  in which the defined symbols have the meanings assigned to them by the module, the module's parameters have the values assigned to them by the configuration file,  $x$  is assigned the value 2,  $y$  is assigned the value  $\{\text{"a"}\}$ , and  $x'$  and  $y'$  have no meanings assigned to them. Suppose the next-state action has the form  $(x' = e) \wedge A$ . TLC computes

the value  $e$  in context  $\mathcal{C}$ , and then evaluates  $A$  in the context  $\mathcal{C}$  augmented with the assignment of the value of  $e$  to  $x'$ .

We now define the result  $\mathcal{V}(e)$  of evaluating a  $\text{TLA}^+$  expression  $e$ , which is a TLC value. If the expression  $e$  does not denote a TLC value, then its evaluation fails. We use an inductive definition that defines  $\mathcal{V}(e)$  in terms of two other operations: *one-step evaluation*  $\mathcal{O}(e)$  and *set enumeration*  $\mathcal{E}(e)$  of an expression  $e$ . While  $\mathcal{V}(e)$  is a TLC value,  $\mathcal{O}(e)$  is an arbitrary  $\text{TLA}^+$  expression and  $\mathcal{E}(e)$  is a finite sequence of arbitrary  $\text{TLA}^+$  expressions. One-step evaluation and enumeration are explained below.

In  $\text{TLA}^+$ , the  $.h$  in an expression  $e.h$  or in the  $!$  clause of an **EXCEPT** construct just means  $[\text{“h”}]$ . To evaluate an expression, TLC first replaces all such instances of  $.h$  by  $[\text{“h”}]$ .

The inductive rules for computing  $\mathcal{V}(e)$  are given below. The rules are stated somewhat informally, using  $\text{TLA}^+$  notation and the operators from the standard *Naturals* and *Sequences* modules (see Chapter 17). Evaluation fails if the rules do not imply that  $\mathcal{V}(e)$  is a TLC value. For example, evaluation of the nonsensical expression  $\{1, 2\}[3]$  fails because no rule applies to an expression of this form. Evaluation of  $e$  also fails if a rule implies that computing  $\mathcal{V}(e)$  requires evaluating an expression whose evaluation fails.

## Quantification

The rules for quantifiers use the operator  $\text{Perm}$ , where  $\text{Perm}(s)$  is an arbitrary permutation of a sequence  $s$ . It can be defined by<sup>6</sup>

$$\begin{aligned} \text{Perm}(s) \triangleq & \text{LET } Pi \triangleq \text{CHOOSE } f \in [1 \dots \text{Len}(s) \rightarrow 1 \dots \text{Len}(s)] : \\ & \forall n \in 1 \dots \text{Len}(s) : \exists m \in 1 \dots \text{Len}(s) : f[m] = n \\ \text{IN } & [i \in 1 \dots \text{Len}(s) \mapsto s[Pi[i]]] \end{aligned}$$

For example,  $\text{Perm}(\langle 1, 2, 3 \rangle)$  might equal  $\langle 2, 1, 3 \rangle$ .

$$\begin{aligned} \mathcal{V}(\forall x \in S : p) = & \\ \text{LET } s \triangleq & \text{Perm}(\mathcal{E}(S)) \\ AV[n \in \text{Nat}] \triangleq & \text{IF } n = 0 \text{ THEN TRUE} \\ & \text{ELSE IF } \mathcal{V}(\text{LET } x \triangleq s[n] \text{ IN } p) \\ & \text{THEN } AV[n - 1] \\ & \text{ELSE FALSE} \\ \text{IN } & AV[\text{Len}(s)] \end{aligned}$$

$$\mathcal{V}(\exists x \in S : p) =$$

---

<sup>6</sup>This defines  $\text{Perm}(s)$  in terms of a permutation of the elements that depends only on the length of  $s$ , not on  $s$  itself. A more precise definition would use the *Choice* operator described on page 279 to make  $\text{Perm}(s)$  depend on  $s$ .

$$\begin{aligned}
& \text{LET } s \triangleq \text{Perm}(\mathcal{E}(S)) \\
& \quad EV[n \in \text{Nat}] \triangleq \text{IF } n = 0 \text{ THEN FALSE} \\
& \qquad \qquad \qquad \text{ELSE IF } \mathcal{V}(\text{LET } x \triangleq s[n] \text{ IN } p) \\
& \qquad \qquad \qquad \quad \text{THEN TRUE} \\
& \qquad \qquad \qquad \quad \text{ELSE } EV[n - 1] \\
& \text{IN } EV[\text{Len}(s)]
\end{aligned}$$

### One-Step Evaluable Expressions

The result  $\mathcal{O}(e)$  of one-step evaluation is a  $\text{TLA}^+$  expression that is obtained by performing just the first step in the evaluation of  $e$ . For example, suppose  $e$  is the expression

$$\text{IF } x > 2 \text{ THEN } x + 2 \text{ ELSE } x - 2$$

In a context in which  $x$  is assigned the value 3,  $\mathcal{O}(e)$  is the expression  $x + 2$ . This differs from the result  $\mathcal{V}(e)$  of evaluating  $e$ , which is the TLC value 5. TLC may be able to perform one-step evaluation even if it can't evaluate the expression. For example, one-step evaluation of

$$\text{IF } x > 2 \text{ THEN } \{j \in \text{Nat} : j > i\} \text{ ELSE } \{\}$$

in a context in which  $x$  equals 3 is the expression  $\{j \in \text{Nat} : j > 5\}$ . TLC will fail if it tries to evaluate this expression because it denotes infinite set, which is not a TLC value.

An expression  $e$  is *one-step evaluable* if  $\mathcal{O}(e)$  is defined. The following kinds of expressions are one-step evaluable: LET, IF/THEN/ELSE, CASE, CHOOSE, and function application (expressions of the form  $f[e]$ ). The rule for evaluating any one-step evaluable expression  $e$  is  $\mathcal{V}(e) = \mathcal{V}(\mathcal{O}(e))$ . In other words, a one-step evaluable expression is evaluated by evaluating the result of one-step evaluation.

Below are the rules for one-step evaluation for all of these except function application, whose rules are given later. (The operator  $\text{Perm}$ , appearing in the rule for CHOOSE, is defined above.)

$\mathcal{O}(\text{LET } x \triangleq v \text{ IN } e)$  evaluated in the current context is equal to  $e$ , with the evaluation proceeding in the context augmented by assigning  $v$  to  $x$ .

$$\begin{aligned}
\mathcal{O}(\text{IF } p \text{ THEN } e_1 \text{ ELSE } e_2) &= \\
\text{CASE } \mathcal{V}(p) = \text{TRUE} &\rightarrow e_1 \\
\quad \square \mathcal{V}(p) = \text{FALSE} &\rightarrow e_2 \\
\quad \square \text{OTHER} &\rightarrow \text{evaluation fails}
\end{aligned}$$

$$\begin{aligned}
\mathcal{O}(\text{CASE } p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n) &= \text{LET } j \triangleq \mathcal{V}(\text{CHOOSE } i \in 1 \dots n : p_i) \\
&\text{IN } e_j
\end{aligned}$$

$$\begin{aligned}
\mathcal{O}(\text{CASE } p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n \square \text{OTHER} \rightarrow e) &= \\
&\text{IF } \mathcal{V}(\exists i \in 1 \dots n : p_i) \text{ THEN } \mathcal{O}(\text{CASE } p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n) \\
&\quad \text{ELSE } e \\
\mathcal{O}(\text{CHOOSE } x \in S : p) &= \\
&\text{LET } s \triangleq \text{Perm}(\mathcal{E}(S)) \\
&\quad CV[n \in \text{Nat}] \triangleq \text{IF } n = 0 \\
&\quad \quad \text{THEN evaluation fails} \\
&\quad \quad \text{ELSE IF } \mathcal{V}(\text{LET } x \triangleq s[n] \text{ IN } p) \\
&\quad \quad \quad \text{THEN } s[n] \\
&\quad \quad \quad \text{ELSE } CV[n - 1] \\
&\text{IN } CV[\text{Len}(s)]
\end{aligned}$$

### Enumerating Set-Valued Expressions

If  $e$  is an expression whose value is a set, then the result  $\mathcal{E}(e)$  of enumerating  $e$  is a finite sequence of  $\text{TLC}^+$  expressions obtained by enumerating the elements of  $e$  without evaluating them. The order of elements in the sequence doesn't matter.<sup>7</sup> For example, if  $e$  is the expression  $\{1 + 1, 2, 3 - 1\}$ , then  $\mathcal{E}(e)$  is the sequence  $\langle 1 + 1, 2, 3 - 1 \rangle$  of expressions. Note that the sets  $\{1, \text{"a"}\}$  and  $\{\text{Nat}, \text{Int}\}$  can be enumerated even though they cannot be evaluated. (The set  $\{1, \text{"a"}\}$  isn't a TLC value because its elements are not comparable; the set  $\{\text{Nat}, \text{Int}\}$  isn't a TLC value because its elements are not TLC values.)

A set-valued expression is evaluated by enumerating it, then evaluating each expression in the enumeration, and finally eliminating duplicate values. More precisely, for any set-valued expression  $S$ :

$$\begin{aligned}
\mathcal{V}(S) &= \\
&\text{LET } n \triangleq \text{Len}(\mathcal{E}(S)) \\
&\quad s[i \in 0 \dots n] \triangleq \text{IF } i = 0 \\
&\quad \quad \text{THEN } \langle \rangle \\
&\quad \quad \text{ELSE IF } \exists j \in 1 \dots \text{Len}(s[i - 1]) : s[i - 1][j] = \mathcal{V}(\mathcal{E}(S)[i]) \\
&\quad \quad \quad \text{THEN } s[i - 1] \\
&\quad \quad \quad \text{ELSE } \text{Append}(s[i - 1], \mathcal{V}(\mathcal{E}(S)[i])) \\
&\text{IN } \{s[n][1], \dots, s[n][\text{Len}(s[n])]\}
\end{aligned}$$

Below are the rules for enumerating set-valued expressions. They use the *SelectSeq* operator, which is defined in the *Sequences* module so that *SelectSeq*( $s$ ,  $\text{Test}$ ) is the subsequence of the sequence  $s$  consisting of all elements  $e$  such that  $\text{Test}(e)$  is true.

$$\mathcal{E}(S \cup T) = \mathcal{E}(S) \circ \mathcal{E}(T)$$

<sup>7</sup>It would perhaps be better to define an enumeration to be a bag, but I won't bother introducing notation for bags here.

$$\begin{aligned} \mathcal{E}(S \cap T) &= \text{LET } InT(s) \triangleq \mathcal{V}(s \in T) \\ &\quad \text{IN } SelectSeq(\mathcal{E}(S), InT) \end{aligned}$$

$$\begin{aligned} \mathcal{E}(S \setminus T) &= \text{LET } NotInT(s) \triangleq \mathcal{V}(s \notin T) \\ &\quad \text{IN } SelectSeq(\mathcal{E}(S), NotInT) \end{aligned}$$

$$\mathcal{E}(\{e_1, \dots, e_n\}) = \langle e_1, \dots, e_n \rangle$$

$$\begin{aligned} \mathcal{E}(\{x \in S : p\}) &= \text{LET } PTrue(s) \triangleq \mathcal{V}(\text{LET } x \triangleq s \text{ IN } p) \\ &\quad \text{IN } SelectSeq(\mathcal{E}(S), PTrue) \end{aligned}$$

$$\mathcal{E}(\{e : x \in S\}) = [i \in 1 \dots Len(\mathcal{E}(S)) \mapsto \text{LET } x \triangleq \mathcal{E}(S)[i] \text{ IN } e]$$

$\mathcal{E}(\text{SUBSET } S)$  is a sequence of length  $2^{Len(\mathcal{E}(S))}$  whose elements consist of all expressions of the form  $\{\mathcal{E}(S)[i_1], \dots, \mathcal{E}(S)[i_k]\}$  for all subsets  $\{i_1, \dots, i_k\}$  of  $1 \dots Len(\mathcal{E}(S))$ .

$$\mathcal{E}(\text{UNION } S) = \mathcal{E}(\mathcal{E}(S)[1]) \circ \dots \circ \mathcal{E}(\mathcal{E}(S)[Len(\mathcal{E}(S))])$$

$\mathcal{E}([S \rightarrow T])$  is a sequence of length  $Len(\mathcal{E}(T))^{Cardinality(S)}$  whose elements consist of all expressions of the form

$$(s_1 :> \mathcal{E}(T)[i_1] \ @\@ \dots \ @\@ s_n :> \mathcal{E}(T)[i_n])$$

for all possible choices of  $i_j$  in  $1 \dots Len(\mathcal{E}(T))$ , where  $\mathcal{V}(S) = \{s_1, \dots, s_n\}$ . The enumeration fails if  $S$  cannot be evaluated.

$\mathcal{E}([h_1 : S_1, \dots, h_n : S_n])$  is a sequence of length  $Len(\mathcal{E}(S_1)) * \dots * Len(\mathcal{E}(S_n))$  whose elements consist of all expressions of the form

$$[h_1 \mapsto \mathcal{E}(S_1)[i_1], \dots, h_n \mapsto \mathcal{E}(S_n)[i_n]]$$

for all possible choices of  $i_j$  in  $1 \dots Len(\mathcal{E}(S_j))$ .

$\mathcal{E}(S_1 \times \dots \times S_n)$  is a sequence of length  $Len(\mathcal{E}(S_1)) * \dots * Len(\mathcal{E}(S_n))$  whose elements consist of all expressions of the form

$$\langle \mathcal{E}(S_1)[i_1], \dots, \mathcal{E}(S_n)[i_n] \rangle$$

for all possible choices of  $i_j$  in  $1 \dots Len(\mathcal{E}(S_j))$ .

$\mathcal{E}(\text{DOMAIN } f)$  depends as follows on the form of the expression  $f$ :

$$\mathcal{E}(\text{DOMAIN } [x \in S \mapsto e]) = \mathcal{E}(S)$$

$$\mathcal{E}(\text{DOMAIN } [g \text{ EXCEPT } \dots]) = \mathcal{E}(\text{DOMAIN } g)$$

$$\mathcal{E}(\text{DOMAIN } (s_1 :> t_1 \ @\@ \dots \ @\@ s_n :> t_n)) = \{s_1, \dots, s_n\}$$

$$\mathcal{E}(\text{DOMAIN } \langle s_1, \dots, s_n \rangle) = \langle 1, \dots, n \rangle$$

$$\mathcal{E}(\text{DOMAIN } [h_1 \mapsto e_1, \dots, h_n \mapsto e_n]) = \{\text{"h}_1\text{"}, \dots, \text{"h}_n\text{"}\}.$$

In addition, if  $f$  is a one-step evaluable expression, then  $\mathcal{E}(\text{DOMAIN } f) = \mathcal{E}(\text{DOMAIN } \mathcal{O}(f))$ .

In addition, if  $e$  is a one-step evaluable expression, then  $\mathcal{E}(e) = \mathcal{E}(\mathcal{O}(e))$ .

## Boolean-Valued Expressions

Evaluation of Boolean expressions is defined by the following rules.

$$\begin{aligned}
 \mathcal{V}(p \vee q) &= \text{IF } \mathcal{V}(p) \text{ THEN TRUE} \\
 &\quad \text{ELSE IF } \mathcal{V}(q) \text{ THEN TRUE ELSE FALSE} \\
 \mathcal{V}(p \wedge q) &= \text{IF } \mathcal{V}(p) \text{ THEN IF } \mathcal{V}(q) \text{ THEN TRUE ELSE FALSE} \\
 &\quad \text{ELSE FALSE} \\
 \mathcal{V}(\neg p) &= \text{IF } \mathcal{V}(p) \text{ THEN FALSE ELSE TRUE} \\
 \mathcal{V}(p \Rightarrow q) &= \text{IF } \mathcal{V}(p) \text{ THEN IF } \mathcal{V}(q) \text{ THEN TRUE ELSE FALSE} \\
 &\quad \text{ELSE TRUE} \\
 \mathcal{V}(p \equiv q) &= \text{IF } \mathcal{V}(p) \text{ THEN IF } \mathcal{V}(q) \text{ THEN TRUE ELSE FALSE} \\
 &\quad \text{ELSE IF } \mathcal{V}(q) \text{ THEN FALSE ELSE TRUE}
 \end{aligned}$$

## Equality

An expression of the form  $v = w$  is evaluated by first evaluating  $v$  and  $w$ . Evaluation of  $v = w$  fails if the evaluation of  $v$  or  $w$  fails, or if the resulting values are not comparable, according to the rules of Section 13.6.1 above. If they are comparable, then equality of TLC values is computed in the obvious way from the equality relation on primitive values—namely, by using the following rules. In these rules, the  $s_i$  are assumed to be all distinct, as are the  $t_j$ .

$$\begin{aligned}
 \mathcal{V}(\{s_1, \dots, s_n\} = \{t_1, \dots, t_m\}) &= \\
 &\quad \text{IF } n = m \text{ THEN } \mathcal{V}(\forall i \in 1 \dots n : \exists j \in 1 \dots m : s_i = t_j) \\
 &\quad \text{ELSE FALSE} \\
 \mathcal{V}((s_1 :> d_1 @@@ \dots @@@ s_n :> d_n) = (t_1 :> e_1 @@@ \dots @@@ t_m :> e_m)) &= \\
 &\quad \text{IF } \mathcal{V}(\{s_1, \dots, s_n\} = \{t_1, \dots, t_m\}) \\
 &\quad \text{THEN } \forall i \in 1 \dots n : \text{LET } j \triangleq \mathcal{V}(\text{CHOOSE } k \in 1 \dots m : s_i = t_k) \\
 &\quad \quad \text{IN } \mathcal{V}(e_i = d_j) \\
 &\quad \text{ELSE FALSE}
 \end{aligned}$$

## Set Membership

The value of an expression of the form  $v \in w$  depends as follows on the form of  $w$ .

$$\begin{aligned}
\mathcal{V}(v \in S \cup T) &= \text{IF } \mathcal{V}(v \in S) \text{ THEN TRUE ELSE } \mathcal{V}(v \in T) \\
\mathcal{V}(v \in S \cap T) &= \text{IF } \mathcal{V}(v \in S) \text{ THEN } \mathcal{V}(v \in T) \text{ ELSE FALSE} \\
\mathcal{V}(v \in S \setminus T) &= \text{IF } \mathcal{V}(v \in S) \text{ THEN IF } \mathcal{V}(v \in T) \text{ THEN FALSE ELSE TRUE} \\
&\quad \text{ELSE FALSE} \\
\mathcal{V}(v \in \{e_1, \dots, e_n\}) &= \mathcal{V}(\exists s \in \{e_1, \dots, e_n\} : v = s) \\
\mathcal{V}(v \in \{x \in S : p\}) &= \text{IF } \mathcal{V}(v \in S) \text{ THEN } \mathcal{V}(\text{LET } x \triangleq v \text{ IN } p) \\
&\quad \text{ELSE FALSE} \\
\mathcal{V}(v \in \{e : x \in S\}) &= \mathcal{V}(\exists s \in S : v = \text{LET } x \triangleq s \text{ IN } e) \\
\mathcal{V}(v \in \text{SUBSET } S) &= \mathcal{V}(\forall s \in v : s \in S) \\
\mathcal{V}(v \in \text{UNION } S) &= \mathcal{V}(\exists s \in S : v \in s) \\
\mathcal{V}(v \in \text{DOMAIN } f) &= \mathcal{V}(v \in \mathcal{V}(\text{DOMAIN } f)), \text{ except in the following two cases:} \\
&\quad \mathcal{V}(v \in \text{DOMAIN } [x \in S \mapsto e]) = \mathcal{V}(v \in S) \\
&\quad \mathcal{V}(v \in \text{DOMAIN } [g \text{ EXCEPT } \dots]) = \mathcal{V}(v \in \text{DOMAIN } g) \\
\mathcal{V}(v \in [S \rightarrow T]) &= \text{IF } \mathcal{V}(\text{DOMAIN } v = S) \text{ THEN } \mathcal{V}(\forall s \in S : v[s] \in T) \\
&\quad \text{ELSE FALSE} \\
\mathcal{V}(v \in [h_1 : S_1, \dots, h_n : S_n]) &= \text{IF } \mathcal{V}(\text{DOMAIN } v = \{\text{"h}_1\text{", } \dots, \text{"h}_n\text{"}\}) \\
&\quad \text{THEN } \mathcal{V}(\forall i \in 1 \dots n : v[\text{"h}_i\text{"}] \in S_i) \\
&\quad \text{ELSE FALSE} \\
\mathcal{V}(v \in S_1 \times \dots \times S_n) &= \text{IF } \mathcal{V}(\text{DOMAIN } v = 1 \dots n) \\
&\quad \text{THEN } \mathcal{V}(\forall i \in 1 \dots n : v[i] \in S_i) \\
&\quad \text{ELSE FALSE}
\end{aligned}$$

If  $w$  is one of a class of *special set constants*, then TLC evaluates  $v \in w$  by determining if  $\mathcal{V}(v)$  is an element of the set represented by  $w$ . For example,  $\mathcal{V}(v \in \text{STRING})$  equals TRUE if  $\mathcal{V}(v)$  is a string and it equals FALSE if  $\mathcal{V}(v)$  is a model value; otherwise, evaluation fails. The built-in special set constants are STRING, BOOLEAN. An overridden symbol can also be a special set constant. In particular, the standard Java classes that override the *Naturals* and *Integers* modules define *Nat* and *Int* to be special set constants.

## Function Application

As stated above, an expression of the form  $v[w]$  is one-step evaluable. Here are the rules that define  $\mathcal{O}(v[w])$  for the possible function-valued expressions  $w$ .

$$\mathcal{O}([x \in S \mapsto e][w]) = \text{IF } w \in S \text{ THEN LET } x \stackrel{\Delta}{=} w \text{ IN } e \\ \text{ELSE evaluation fails}$$

$$\mathcal{O}([f \text{ EXCEPT } ![e_1] \dots [e_n] = d][w]) = \\ \text{IF } w \in \text{DOMAIN } f \\ \text{THEN IF } w = e_1 \\ \text{THEN IF } n = 1 \text{ THEN } d \\ \text{ELSE } [f \text{ EXCEPT } ![e_2] \dots [e_n] = d] \\ \text{ELSE } f[w] \\ \text{ELSE evaluation fails}$$

$$\mathcal{O}([h_1 \mapsto e_1, \dots, h_n \mapsto e_n][w]) = \text{LET } j \stackrel{\Delta}{=} \text{CHOOSE } i \in 1 \dots n : w = h_i \\ \text{IN } e_j$$

$$\mathcal{O}(\langle e_1, \dots, e_n \rangle[w]) = \text{IF } u \in 1 \dots n \text{ THEN } e_u \\ \text{ELSE evaluation fails ,}$$

where  $u$  is the value obtained by evaluating  $w$ .

$$\mathcal{O}((s_1 :> e_1 @@@ \dots @@@ s_n :> e_n)[w]) = \\ \text{IF } \mathcal{V}(\exists i \in 1 \dots n : w = s_i) \\ \text{THEN LET } j \stackrel{\Delta}{=} \mathcal{V}(\text{CHOOSE } i \in 1 \dots n : w = s_i) \\ \text{IN } e_j \\ \text{ELSE evaluation fails}$$

If  $v$  is one-step evaluable, then  $\mathcal{O}(v[w]) = \mathcal{O}(\mathcal{O}(v)[w])$ .

## Functions and Records

$$\mathcal{V}([x \in S \mapsto e]) = (s_1 :> (\text{LET } x \stackrel{\Delta}{=} s_1 \text{ IN } e) @@@ \\ \dots @@@ s_n :> (\text{LET } x \stackrel{\Delta}{=} s_n \text{ IN } e)) ,$$

where  $\mathcal{V}(S)$  equals  $\{s_1, \dots, s_n\}$ .

$$\mathcal{V}([f \text{ EXCEPT } ![e_1] \dots [e_n] = d]) = \\ (\dots @@@ s_i :> (\text{IF } \mathcal{V}(s_i = e_1) \\ \text{THEN IF } n = 1 \\ \text{THEN } \mathcal{V}(\text{LET } @ \stackrel{\Delta}{=} t_i \text{ IN } d) \\ \text{ELSE } \mathcal{V}([t_i \text{ EXCEPT } ![e_2] \dots [e_n] = d]) \\ \text{ELSE } t_i \\ @@@ \dots)),$$

where  $\mathcal{V}(f)$  equals  $(\dots @@@ s_i :> t_i @@@ \dots)$ .

$$\mathcal{V}([h_1 \mapsto e_1, \dots, h_n \mapsto e_n]) = ("h_1" :> \mathcal{V}(e_1) @@@ \dots @@@ "h_n" :> \mathcal{V}(e_n))$$

$$\mathcal{V}(\langle e_1, \dots, e_n \rangle) = (1 :> \mathcal{V}(e_1) @@@ \dots @@@ n :> \mathcal{V}(e_n))$$

### Other Constant Operators

$$\mathcal{V}(v \neq w) = \mathcal{V}(\neg(v = w))$$

$$\mathcal{V}(e \notin S) = \mathcal{V}(\neg(e \in S))$$

$$\mathcal{V}(S \subseteq T) = \mathcal{V}(\forall s \in S : s \in T)$$

### Primitive Values and Overridden Operators

If  $v$  is a primitive or overridden value, then  $\mathcal{V}(v) = v$ .

If  $Op$  is an overridden operator, then  $\mathcal{V}(Op(e_1, \dots, e_n))$  equals

$$Op(\mathcal{V}(e_1), \dots, \mathcal{V}(e_n))$$

if this is a TLC value. If it isn't, then evaluation fails. For example, if the concatenation operator  $\circ$  of the *Sequences* module is overridden, then  $\mathcal{V}(s \circ t)$  equals  $\mathcal{V}(s) \circ \mathcal{V}(t)$ , which is the function

$$\begin{aligned} &(1 :> \mathcal{V}(s[1]) \text{ @@ } \dots \text{ @@ } \mathcal{V}(Len(s)) :> \mathcal{V}(s[Len(s)]) \text{ @@} \\ &\quad (1 + \mathcal{V}(Len(s))) :> \mathcal{V}(t[1]) \text{ @@ } \dots \text{ @@} \\ &\quad \mathcal{V}(Len(s)) + \mathcal{V}(Len(t)) :> \mathcal{V}(t[Len(t)]) ) \end{aligned}$$

if  $\mathcal{V}(s)$  and  $\mathcal{V}(t)$  are sequences.

### Action Expressions

$\mathcal{V}(e')$  in a context  $\mathcal{C}$  equals  $\mathcal{V}(e)$  in the context obtained from  $\mathcal{C}$  by assigning to each unprimed variable  $x$  the value assigned by  $\mathcal{C}$  to  $x'$ , and assigning no value to any primed variable.

$$\mathcal{V}([A]_e) = \mathcal{V}(A \vee e' = e)$$

$$\mathcal{V}(\langle A \rangle_e) = \mathcal{V}(A \wedge e' \neq e)$$

$\mathcal{V}(\text{ENABLED } A)$  in a context  $\mathcal{C}$  is computed by attempting to find a state  $t$  such that  $s \rightarrow t$  is an  $A$  step, where  $s$  is the state that assigns to each variable the value assigned to it by  $\mathcal{C}$ . The value of  $\mathcal{V}(\text{ENABLED } A)$  is TRUE if such a  $t$  exists, otherwise it is FALSE.

$\mathcal{V}(A \cdot B)$  in a context  $\mathcal{C}$  is computed by finding all states  $t$  such that  $s \rightarrow t$  is an  $A$  step, where  $s$  is the state that assigns to each variable the value assigned to it by  $\mathcal{C}$ , then finding all states  $u$  such that  $t \rightarrow u$  is a  $B$  step. Note that TLC evaluates an action only when either computing the next-state action to find successor states to a state  $s$ , or when evaluating an ENABLED expression.

TLC Version 1 cannot evaluate an expression of the form  $A \cdot B$ .

### 13.6.4 Fingerprinting

The description of TLC's computation in Section 13.2.6 is inaccurate in one important respect. TLC does not actually keep the set of states  $\mathcal{R}$ . Instead, it keeps the set of fingerprints of those states. A fingerprint of a state is a 64-bit value generated by a "hashing" function. Ideally, the probability that two different states have the same fingerprint is  $2^{-64}$ , which is a very small number.

Rather than checking if a state  $s$  is in  $\mathcal{R}$ , TLC checks if its fingerprint equals the fingerprint of a state already in  $\mathcal{R}$ . With very high probability, that will be true iff  $s$  is in  $\mathcal{R}$ . However, it is possible for a *collision* to occur, meaning that the fingerprint of  $s$  equals the fingerprint of some other state in  $\mathcal{R}$ . If this happens, TLC will not explore the successor states of  $s$ , and thus may not find all reachable states.

It is essentially impossible to do an accurate *a posteriori* calculation of the probability that TLC did not check all reachable states because of a fingerprint collision. If the probability that any two states have the same fingerprint is  $2^{-64}$ , then a simple calculation shows that if TLC generated  $n$  states with  $m$  distinct fingerprints, then the probability of a collision is about  $m * (n - m) * 2^{-64}$ . However, the process of generating states is highly nonrandom, and no known fingerprinting scheme can guarantee that the probability of any two distinct states generated by TLC having the same fingerprint is actually  $2^{-64}$ . So, this estimate is an optimistic one.

Another way to estimate the probability of collision is empirical. If there was a collision, then it is likely that there was also a "near miss". One estimate of the probability that there was a collision is the maximum value of  $1/|f_1 - f_2|$  over all pairs  $\langle f_1, f_2 \rangle$  of distinct fingerprints generated by TLC for the states that it found.

TLC prints out both probability estimates when it finishes. You can use these values to decide how much confidence to place in the completeness of TLC's exploration of the reachable states. You may wish to supplement TLC's model checking calculation with random simulation, which does not maintain the set  $\mathcal{R}$  and thus uses no fingerprinting.

Part IV

The TLA<sup>+</sup> Language



This part of the book describes TLA<sup>+</sup> in detail. Chapter 14 explains the syntax; Chapters 15 and 16 explain the semantics; and Chapter 17 contains the standard modules. Almost all of the TLA<sup>+</sup> language has already been described—mainly through examples. In fact, most of the language was described in Chapters 1–6. Here, we give a complete specification of the language.

A completely formal specification of TLA<sup>+</sup> would consist of a formal definition of the set of legal (syntactically well-formed) modules, and a precisely-defined meaning operator that assigns to every legal module  $M$  its mathematical meaning  $\llbracket M \rrbracket$ . Such a specification would be quite long and of limited interest. Instead, I have tried to provide a fairly informal specification that is detailed enough to show mathematically sophisticated readers how they could write a completely formal one.

These chapters are heavy going, and only a few sophisticated readers will want to read them completely. However, I hope they can serve as a reference manual for anyone who reads or writes TLA<sup>+</sup> specifications. If you have a question about the finer details of the syntax or the meaning of some part of the language, you should be able to find the answer here.

Figures 13.9–13.16 on the next page through page 253 provide a tiny reference manual. Figures 13.9–13.12 very briefly describe all the built-in operators of TLA<sup>+</sup>. Figure 13.13 lists all the user-definable operator symbols, and indicates which ones are already used by the standard modules. Figure 13.14 gives the precedence of the operators; it is explained in Section 14.2.1. Figure 13.15 lists all operators defined by the standard modules. Finally, Figure 13.16 shows how to type any symbol that doesn't have an obvious ASCII equivalent.

**Logic**
 $\wedge \vee \neg \Rightarrow \equiv$ 

 TRUE FALSE BOOLEAN [the set  $\{\text{TRUE}, \text{FALSE}\}$ ]

 $\forall x : p \quad \exists x : p \quad \forall x \in S : p \quad (1) \quad \exists x \in S : p \quad (1)$ 

 CHOOSE  $x : p$  [An  $x$  satisfying  $p$ ] CHOOSE  $x \in S : p$  [An  $x$  in  $S$  satisfying  $p$ ]
**Sets**
 $= \neq \in \notin \cup \cap \subseteq \setminus$  [set difference]

 $\{e_1, \dots, e_n\}$  [Set consisting of elements  $e_i$ ]

 $\{x \in S : p\} \quad (2)$  [Set of elements  $x$  in  $S$  satisfying  $p$ ]

 $\{e : x \in S\} \quad (1)$  [Set of elements  $e$  such that  $x$  in  $S$ ]

 SUBSET  $S$  [Set of subsets of  $S$ ]

 UNION  $S$  [Union of all elements of  $S$ ]
**Functions**
 $f[e]$  [Function application]

 DOMAIN  $f$  [Domain of function  $f$ ]

 $[x \in S \mapsto e] \quad (1)$  [Function  $f$  such that  $f[x] = e$  for  $x \in S$ ]

 $[S \rightarrow T]$  [Set of functions  $f$  with  $f[x] \in T$  for  $x \in S$ ]

 $[f \text{ EXCEPT } ![e_1] = e_2] \quad (3)$  [Function  $\hat{f}$  equal to  $f$  except  $\hat{f}[e_1] = e_2$ ]
**Records**
 $e.h$  [The  $h$ -component of record  $e$ ]

 $[h_1 \mapsto e_1, \dots, h_n \mapsto e_n]$  [The record whose  $h_i$  component is  $e_i$ ]

 $[h_1 : S_1, \dots, h_n : S_n]$  [Set of all records with  $h_i$  component in  $S_i$ ]

 $[r \text{ EXCEPT } !.h = e] \quad (3)$  [Record  $\hat{r}$  equal to  $r$  except  $\hat{r}.h = e$ ]
**Tuples**
 $e[i]$  [The  $i^{\text{th}}$  component of tuple  $e$ ]

 $\langle e_1, \dots, e_n \rangle$  [The  $n$ -tuple whose  $i^{\text{th}}$  component is  $e_i$ ]

 $S_1 \times \dots \times S_n$  [The set of all  $n$ -tuples with  $i^{\text{th}}$  component in  $S_i$ ]
**Strings and Numbers**
 $\text{"c}_1 \dots \text{c}_n\text{"}$  [A literal string of  $n$  characters]

STRING [The set of all strings]

 $d_1 \dots d_n \quad d_1 \dots d_n . d_{n+1} \dots d_m$  [Numbers (where the  $d_i$  are digits)]

(1)  $x \in S$  may be replaced by a comma-separated list of items  $v \in S$ , where  $v$  is either a comma-separated list or a tuple of identifiers.

(2)  $x$  may be an identifier or tuple of identifiers.

(3)  $![e_1]$  or  $!.h$  may be replaced by a comma separated list of items  $!a_1 \dots a_n$ , where each  $a_i$  is  $[e_i]$  or  $.h_i$ .

Figure 13.9: The constant operators.

---

|                                                                                                         |                                                                                                       |
|---------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------|
| IF $p$ THEN $e_1$ ELSE $e_2$                                                                            | $[e_1 \text{ if } p \text{ true, else } e_2]$                                                         |
| CASE $p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n$                                    | $[\text{Some } e_i \text{ such that } p_i \text{ true}]$                                              |
| CASE $p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n \square \text{OTHER} \rightarrow e$ | $[\text{Some } e_i \text{ such that } p_i \text{ true, or } e \text{ if all } p_i \text{ are false}]$ |
| LET $d_1 \triangleq e_1 \dots d_n \triangleq e_n$ IN $e$                                                | $[e \text{ in the context of the definitions}]$                                                       |
| $\wedge p_1$                                                                                            | $[\text{the conjunction } p_1 \wedge \dots \wedge p_n]$                                               |
| $\vdots$                                                                                                | $\vee p_1$                                                                                            |
| $\wedge p_n$                                                                                            | $\vee p_n$                                                                                            |

Figure 13.10: Miscellaneous constructs.

|                       |                                                                 |
|-----------------------|-----------------------------------------------------------------|
| $e'$                  | $[\text{The value of } e \text{ in the final state of a step}]$ |
| $[A]_e$               | $[A \vee (e' = e)]$                                             |
| $\langle A \rangle_e$ | $[A \wedge (e' \neq e)]$                                        |
| ENABLED $A$           | $[\text{An } A \text{ step is possible}]$                       |
| UNCHANGED $e$         | $[e' = e]$                                                      |
| $A \cdot B$           | $[\text{Composition of actions}]$                               |

Figure 13.11: Action operators.

|                                   |                                                          |
|-----------------------------------|----------------------------------------------------------|
| $\Box F$                          | $[F \text{ is always true}]$                             |
| $\Diamond F$                      | $[F \text{ is eventually true}]$                         |
| $\text{WF}_e(A)$                  | $[\text{Weak fairness for action } A]$                   |
| $\text{SF}_e(A)$                  | $[\text{Strong fairness for action } A]$                 |
| $F \leadsto G$                    | $[F \text{ leads to } G]$                                |
| $F \triangleleft\triangleright G$ | $[F \text{ guarantees } G]$                              |
| $\exists x : F$                   | $[\text{Temporal existential quantification (hiding).}]$ |
| $\forall x : F$                   | $[\text{Temporal universal quantification.}]$            |

Figure 13.12: Temporal operators.

### Infix Operators

|                         |                          |                              |                          |                        |                     |
|-------------------------|--------------------------|------------------------------|--------------------------|------------------------|---------------------|
| $+$ <sup>(1)</sup>      | $-$ <sup>(1)</sup>       | $*$ <sup>(1)</sup>           | $/$ <sup>(2)</sup>       | $\circ$ <sup>(3)</sup> | $++$                |
| $\div$ <sup>(1)</sup>   | $\%$ <sup>(1)</sup>      | $\cdot$ <sup>(1,4)</sup>     | $\ddots$ <sup>(1)</sup>  | $\dots$                | $--$                |
| $\oplus$ <sup>(5)</sup> | $\ominus$ <sup>(5)</sup> | $\otimes$                    | $\oslash$                | $\odot$                | $**$                |
| $<$ <sup>(1)</sup>      | $>$ <sup>(1)</sup>       | $\leq$ <sup>(1)</sup>        | $\geq$ <sup>(1)</sup>    | $\square$              | $//$                |
| $\prec$                 | $\succ$                  | $\preceq$                    | $\succeq$                | $\sqsubset$            | $\sim$              |
| $\ll$                   | $\gg$                    | $\prec :$                    | $: \succ$ <sup>(6)</sup> | $\&$                   | $\&\&$              |
| $\sqcap$                | $\sqcup$                 | $\sqsubseteq$ <sup>(5)</sup> | $\sqsupset$              | $ $                    | $  $                |
| $\cap$                  | $\cup$                   |                              | $\cup$                   | $\star$                | $\%\%$              |
| $\top$                  | $\bot$                   | $\Vdash$                     | $\Vdash$                 | $\bullet$              | $\#\#$              |
| $\sim$                  | $\approx$                | $\approx$                    | $\approx$                | $\$$                   | $\$\$$              |
| $:=$                    | $::=$                    | $\rangle$                    | $\dot{=}$                | $?$                    | $??$                |
| $\propto$               | $\wr$                    | $\oplus$                     | $\bigcirc$               | $!!$                   | $@@$ <sup>(6)</sup> |

### Postfix Operators <sup>(7)</sup>

$\hat{+}$      $\hat{*}$      $\hat{\#}$

### Prefix Operator

$-$  <sup>(8)</sup>

(1) Defined by the *Naturals*, *Integers*, and *Reals* modules.

(2) Defined by the *Reals* module.

(3) Defined by the *Sequences* module.

(4)  $x \hat{\cdot} y$  is printed as  $x^y$ .

(5) Defined by the *Bags* module.

(6) Defined by the *TLC* module.

(7)  $e \hat{+}$  is printed as  $e^+$ , and similarly for  $\hat{*}$  and  $\hat{\#}$ .

(8) Defined by the *Integers* and *Reals* modules.

Figure 13.13: User-definable operator symbols.

| Prefix Operators     |           |               |          |              |          |
|----------------------|-----------|---------------|----------|--------------|----------|
| ENABLED              | 3-14      | SUBSET        | 9-12     | UNION        | 9-12     |
| UNCHANGED            | 3-14      | $\square$     | 4-14     | DOMAIN       | 9-12     |
| $\neg$               | 4-4       | $\diamond$    | 4-14     | $-$          | 11-11    |
| Infix Operators      |           |               |          |              |          |
| $\Rightarrow$        | 1-1       | $\gg$         | 7-7      | $\dots$      | 7-7      |
| $\pm \triangleright$ | 2-2       | $\in$         | 7-7      | $\dots$      | 7-7      |
| $\equiv$             | 2-2       | $\leq$        | 7-7      | $!!$         | 7-12     |
| $\leadsto$           | 2-2       | $\ll$         | 7-7      | $##$         | 7-12 (a) |
| $\wedge$             | 3-3 (a)   | $\prec$       | 7-7      | $\$$         | 7-12 (a) |
| $\vee$               | 3-3 (a)   | $\preceq$     | 7-7      | $\$ \$$      | 7-12 (a) |
| $@ @$                | 5-5 (a)   | $\propto$     | 7-7      | $?$          | 7-12 (a) |
| $:>$                 | 6-6       | $\sim$        | 7-7      | $??$         | 7-12 (a) |
| $<:$                 | 6-6       | $\simeq$      | 7-7      | $\square$    | 7-12 (a) |
| $\neq$               | 7-7       | $\sqsubset$   | 7-7      | $\sqcup$     | 7-12 (a) |
| $\vdash$             | 7-7       | $\sqsubseteq$ | 7-7      | $\oplus$     | 7-12 (a) |
| $::=$                | 7-7       | $\sqsupset$   | 7-7      | $\wr$        | 7-13     |
| $: =$                | 7-7       | $\sqsupseteq$ | 7-7      | $\oplus$     | 9-9 (a)  |
| $<$                  | 7-7       | $\subset$     | 7-7      | $+$          | 9-9 (a)  |
| $=$                  | 7-7       | $\subseteq$   | 7-7      | $++$         | 9-9 (a)  |
| $\pm$                | 7-7       | $\supset$     | 7-7      | $\%$         | 9-10     |
| $>$                  | 7-7       | $\supseteq$   | 7-7      | $\% \%$      | 9-10 (a) |
| $\approx$            | 7-7       | $\cup$        | 7-7      | $ $          | 9-10 (a) |
| $\asymp$             | 7-7       | $\supseteq$   | 7-7      | $  $         | 9-10 (a) |
| $\cong$              | 7-7       | $\vdash$      | 7-7      | $\backslash$ | 9-12     |
| $\doteq$             | 7-7       | $\Vdash$      | 7-7      | $\cap$       | 9-12 (a) |
| $\geq$               | 7-7       | $\cdot^{(1)}$ | 7-13 (a) | $\cup$       | 9-12 (a) |
| $\ominus$            | 10-10 (a) |               |          |              |          |
| $-$                  | 10-10 (a) |               |          |              |          |
| $--$                 | 10-10 (a) |               |          |              |          |
| $\&$                 | 12-12 (a) |               |          |              |          |
| $\&\&$               | 12-12 (a) |               |          |              |          |
| $\odot$              | 12-12 (a) |               |          |              |          |
| $\oslash$            | 12-12     |               |          |              |          |
| $\otimes$            | 12-12 (a) |               |          |              |          |
| $*$                  | 12-12 (a) |               |          |              |          |
| $**$                 | 12-12 (a) |               |          |              |          |
| $/$                  | 12-12     |               |          |              |          |
| $//$                 | 12-12     |               |          |              |          |
| $\bigcirc$           | 12-12 (a) |               |          |              |          |
| $\bullet$            | 12-12 (a) |               |          |              |          |
| $\div$               | 12-12     |               |          |              |          |
| $\circ$              | 12-12 (a) |               |          |              |          |
| $\star$              | 12-12 (a) |               |          |              |          |
| $\wedge$             | 13-13     |               |          |              |          |
| $\wedge \wedge$      | 13-13     |               |          |              |          |
| $\cdot^{(2)}$        | 16-16 (a) |               |          |              |          |
| Postfix Operators    |           |               |          |              |          |
| $\hat{+}$            | 14-14     | $\hat{*}$     | 14-14    | $\hat{\#}$   | 14-14    |
|                      |           |               |          | $'$          | 14-14    |

(1) Action composition ( $\cdot$ ).

(2) Record component (period).

Figure 13.14: The precedence ranges of operators. The relative precedence of two operators is unspecified if their ranges overlap. Left-associative operators are indicated by (a).

**Modules** *Naturals, Integers, Reals*

|        |           |        |           |                |         |                  |
|--------|-----------|--------|-----------|----------------|---------|------------------|
| $+$    | $-^{(1)}$ | $*$    | $/^{(2)}$ | $\wedge^{(3)}$ | $\dots$ | $Infinity^{(2)}$ |
| $\div$ | $\%$      | $\leq$ | $\geq$    | $<$            | $>$     |                  |

(1) Only infix  $-$  is defined in *Naturals*.

(3) Exponentiation.

(2) Defined only in *Reals* module.**Module** *Sequences*

|               |             |                  |               |
|---------------|-------------|------------------|---------------|
| $\circ$       | <i>Head</i> | <i>SelectSeq</i> | <i>SubSeq</i> |
| <i>Append</i> | <i>Len</i>  | <i>Seq</i>       | <i>Tail</i>   |

**Module** *FiniteSets*

|                    |                    |
|--------------------|--------------------|
| <i>IsFiniteSet</i> | <i>Cardinality</i> |
|--------------------|--------------------|

**Module** *Bags*

|                       |                 |                 |               |
|-----------------------|-----------------|-----------------|---------------|
| $\oplus$              | <i>BagIn</i>    | <i>CopiesIn</i> | <i>SubBag</i> |
| $\ominus$             | <i>BagOfAll</i> | <i>EmptyBag</i> |               |
| $\sqsubseteq$         | <i>BagToSet</i> | <i>IsABag</i>   |               |
| <i>BagCardinality</i> | <i>BagUnion</i> | <i>SetToBag</i> |               |

**Module** *RealTime*

|                |              |                                        |
|----------------|--------------|----------------------------------------|
| <i>RTBound</i> | <i>RTnow</i> | <i>now</i> (declared to be a variable) |
|----------------|--------------|----------------------------------------|

**Module** *TLC*

|      |                   |                 |                |
|------|-------------------|-----------------|----------------|
| $:>$ | <i>Assert</i>     | <i>FApply</i>   | <i>Print</i>   |
| $@@$ | <i>BoundedSeq</i> | <i>JavaTime</i> | <i>SortSeq</i> |

Figure 13.15: Operators defined in the standard modules.

|                                                                                  |                                                          |                                                                                  |                                               |                                                                                  |                                  |
|----------------------------------------------------------------------------------|----------------------------------------------------------|----------------------------------------------------------------------------------|-----------------------------------------------|----------------------------------------------------------------------------------|----------------------------------|
| $\wedge$                                                                         | <code>\</code> or <code>\land</code>                     | $\vee$                                                                           | <code>\</code> or <code>\lor</code>           | $\Rightarrow$                                                                    | <code>=&gt;</code>               |
| $\neg$                                                                           | <code>~</code> or <code>\not</code> or <code>\neg</code> | $\equiv$                                                                         | <code>&lt;=&gt;</code> or <code>\equiv</code> | $\triangle$                                                                      | <code>==</code>                  |
| $\in$                                                                            | <code>\in</code>                                         | $\neq$                                                                           | <code>#</code> or <code>/=</code>             | $'$                                                                              | <code>,</code>                   |
| $\langle$                                                                        | <code>&lt;&lt;</code>                                    | $\rangle$                                                                        | <code>&gt;&gt;</code>                         | $\square$                                                                        | <code>[]</code>                  |
| $<$                                                                              | <code>&lt;</code>                                        | $>$                                                                              | <code>&gt;</code>                             | $\diamond$                                                                       | <code>&lt;&gt;</code>            |
| $\leq$                                                                           | <code>\leq</code> or <code>=&lt;</code>                  | $\geq$                                                                           | <code>\geq</code> or <code>&gt;=</code>       | $\gtrsim$                                                                        | <code>~&gt;</code>               |
| $\ll$                                                                            | <code>\ll</code>                                         | $\gg$                                                                            | <code>\gg</code>                              | $\pm$                                                                            | <code>-+&gt;</code>              |
| $\prec$                                                                          | <code>\prec</code>                                       | $\succ$                                                                          | <code>\succ</code>                            | $\mapsto$                                                                        | <code> -&gt;</code>              |
| $\preceq$                                                                        | <code>\preceq</code>                                     | $\succeq$                                                                        | <code>\succeq</code>                          | $\div$                                                                           | <code>\div</code>                |
| $\subseteq$                                                                      | <code>\subseteq</code>                                   | $\supseteq$                                                                      | <code>\supseteq</code>                        | $\wr$                                                                            | <code>\wr</code>                 |
| $\subset$                                                                        | <code>\subset</code>                                     | $\supset$                                                                        | <code>\supset</code>                          | $\bullet$                                                                        | <code>\bullet</code>             |
| $\sqsubset$                                                                      | <code>\sqsubset</code>                                   | $\sqsupset$                                                                      | <code>\sqsupset</code>                        | $\star$                                                                          | <code>\star</code>               |
| $\sqsubseteq$                                                                    | <code>\sqsubseteq</code>                                 | $\sqsupseteq$                                                                    | <code>\sqsupseteq</code>                      | $\sim$                                                                           | <code>\sim</code>                |
| $\vdash$                                                                         | <code> -</code>                                          | $\dashv$                                                                         | <code>- </code>                               | $\simeq$                                                                         | <code>\simeq</code>              |
| $\models$                                                                        | <code> =</code>                                          | $\models$                                                                        | <code>= </code>                               | $\asymp$                                                                         | <code>\asymp</code>              |
| $\rightarrow$                                                                    | <code>-&gt;</code>                                       | $\leftarrow$                                                                     | <code>&lt;-</code>                            | $\approx$                                                                        | <code>\approx</code>             |
| $\cap$                                                                           | <code>\cap</code> or <code>\intersect</code>             | $\cup$                                                                           | <code>\cup</code> or <code>\union</code>      | $\cong$                                                                          | <code>\cong</code>               |
| $\sqcap$                                                                         | <code>\sqcap</code>                                      | $\sqcup$                                                                         | <code>\sqcup</code>                           | $\doteq$                                                                         | <code>\doteq</code>              |
| $\oplus$                                                                         | <code>(+)</code> or <code>\oplus</code>                  | $\uplus$                                                                         | <code>\uplus</code>                           | $\propto$                                                                        | <code>\propto</code>             |
| $\ominus$                                                                        | <code>(-)</code> or <code>\ominus</code>                 | $\times$                                                                         | <code>\X</code> or <code>\times</code>        | $x^y$                                                                            | <code>x^y</code> <sup>(2)</sup>  |
| $\odot$                                                                          | <code>(.)</code> or <code>\odot</code>                   | $\cdot$                                                                          | <code>\cdot</code>                            | $x^+$                                                                            | <code>x^+</code> <sup>(2)</sup>  |
| $\otimes$                                                                        | <code>(\X)</code> or <code>\otimes</code>                | $\circ$                                                                          | <code>\o</code> or <code>\circ</code>         | $x^*$                                                                            | <code>x^*</code> <sup>(2)</sup>  |
| $\oslash$                                                                        | <code>(/)</code> or <code>\oslash</code>                 | $\text{"s"}$                                                                     | <code>"s"</code> <sup>(1)</sup>               | $x^\#$                                                                           | <code>x^\#</code> <sup>(2)</sup> |
| $\exists$                                                                        | <code>\E</code>                                          | $\forall$                                                                        | <code>\A</code>                               |                                                                                  |                                  |
| $\exists$                                                                        | <code>\EE</code>                                         | $\forall$                                                                        | <code>\AA</code>                              |                                                                                  |                                  |
| $\overline{\hspace{1cm}}$                                                        |                                                          | $\overline{\hspace{1cm}}$                                                        |                                               | $\overline{\hspace{1cm}}$                                                        |                                  |
| $\underline{\hspace{1cm}}$                                                       |                                                          | $\underline{\hspace{1cm}}$                                                       |                                               | $\underline{\hspace{1cm}}$                                                       |                                  |
| $\overline{\hspace{1cm}} \hspace{0.5cm} \text{-----} \hspace{0.5cm} \text{(3)}$  |                                                          | $\overline{\hspace{1cm}} \hspace{0.5cm} \text{-----} \hspace{0.5cm} \text{(3)}$  |                                               | $\overline{\hspace{1cm}} \hspace{0.5cm} \text{-----} \hspace{0.5cm} \text{(3)}$  |                                  |
| $\underline{\hspace{1cm}} \hspace{0.5cm} \text{-----} \hspace{0.5cm} \text{(3)}$ |                                                          | $\underline{\hspace{1cm}} \hspace{0.5cm} \text{-----} \hspace{0.5cm} \text{(3)}$ |                                               | $\underline{\hspace{1cm}} \hspace{0.5cm} \text{=====} \hspace{0.5cm} \text{(3)}$ |                                  |

(1)  $s$  is a sequence of characters. See Section 15.1.10 on page 291.

(2)  $x$  and  $y$  are any expressions.

(3) a sequence of four or more  $-$  or  $=$  characters.

Figure 13.16: The ASCII representations of typeset symbols.



## Chapter 14

# The Syntax of TLA<sup>+</sup>

The term *syntax* has two different usages, which I will somewhat arbitrarily attribute to mathematicians and computer scientists. A computer scientist would say that  $\langle a, a \rangle$  is a syntactically correct TLA<sup>+</sup> expression. A mathematician would say that the expression is syntactically correct iff it appears in a context in which  $a$  is defined or declared. A computer scientist would call this requirement a *semantic* rather than a syntactic condition. A mathematician would say that  $\langle a, a \rangle$  is meaningless if  $a$  isn't defined or declared, and one can't talk about the semantics of a meaningless expression.

This chapter describes the syntax of TLA<sup>+</sup>, in the computer scientist's sense of syntax. (The "semantic" part of the syntax is specified in Chapters 15 and 16.) TLA<sup>+</sup> is designed to be easy for humans to read and write. In particular, its syntax for expressions tries to capture some of the richness of ordinary mathematical notation. This makes a precise specification of the syntax rather complicated. Such a specification has been written in TLA<sup>+</sup>, but it is quite detailed and you probably don't want to look at it unless you are writing a parser for the language. This chapter gives a less formal description of the syntax that should answer any questions likely to arise in practice. Section 14.1 specifies precisely a simple grammar that ignores some aspects of the syntax such as operator precedence, indentation rules for  $\wedge$  and  $\vee$  lists, and comments. These other aspects are explained informally in Section 14.2.1. Finally, Section 14.3 lists the correspondence between the typed and typeset versions of symbols, as well as all user-definable operator symbols.

## 14.1 The Simple Grammar

The simple grammar of TLA<sup>+</sup> is described in BNF. Just for fun, this BNF grammar is specified in TLA<sup>+</sup>. Specifying a BNF grammar in TLA<sup>+</sup> provides a nice little exercise in “mathematicizing” a simple concept. Moreover, it demonstrates the flexibility of ordinary mathematics, as formalized in TLA<sup>+</sup>. The syntax of TLA<sup>+</sup> doesn’t permit us to write BNF grammars exactly the way you’re used to seeing them, but we can come reasonably close.

BNF, which stands for Backus-Naur Form, is a standard way of describing computer languages.

Section 14.1.1 explains how to specify a BNF grammar. The BNF grammar for TLA<sup>+</sup> is specified in Section 14.1.2 as a TLA<sup>+</sup> module, with comments that describe how to read the TLA<sup>+</sup> specification as an ordinary BNF grammar. If you are familiar with BNF grammars and just want to learn the syntax of TLA<sup>+</sup>, you can skip directly to Section 14.1.2 (page 260).

### 14.1.1 BNF Grammars

Let’s start by reviewing BNF grammars. Consider the little language SE of simple expressions described by this BNF grammar:

$$\begin{aligned} \text{expr} &::= \text{ident} \mid \text{expr op expr} \mid (\text{expr}) \mid \text{LET } \text{def} \text{ IN } \text{expr} \\ \text{def} &::= \text{ident} == \text{expr} \end{aligned}$$

where **op** is some class of infix operators like  $+$ , and **ident** is some class of identifiers such as *abc* and *x*. The language SE contains expressions like

$$abc + (\text{LET } x == y + abc \text{ IN } x * x)$$

I will represent this expression as the sequence

$$\langle \text{“abc”, “+”, “(”, “LET”, “x”, “==”, “y”, “+”, “abc”, “IN”, “x”, “*”, “x”, “)”} \rangle$$

of strings. The strings such as “abc” and “+” appearing in this sequence are usually called *lexemes*. In general, a sequence of lexemes is called a *sentence*; and a set of sentences is called a *language*. So, we want to define the language SE to consist of the set of all such sentences described by the BNF grammar.<sup>1</sup>

To represent a BNF grammar in TLA<sup>+</sup>, we must assign a mathematical meaning to nonterminal symbols like *def*, to terminal symbols like **op**, and to the grammar’s two productions. The method that I find simplest is to let the meaning of a nonterminal symbol be the language that it generates. Thus, the meaning of *expr* is the language SE itself. I define a *grammar* to be a function *G* such that, for any string “str”, the value of *G*[“str”] is the language generated by

<sup>1</sup>BNF grammars are also used to specify how an expression is parsed—for example that  $a + b * c$  is parsed as  $a + (b * c)$  rather than  $(a + b) * c$ . By considering the grammar to specify only a set of sentences, we are deliberately not capturing that use in our TLA<sup>+</sup> representation of BNF grammars.

the nonterminal *str*. Thus, if  $G$  is the BNF grammar above, then  $G[\text{“expr”}]$  is the complete language SE, and  $G[\text{“def”}]$  is the language defined by the production for *def*, which contains sentences like

$\langle \text{“y”}, \text{“==”}, \text{“qq”}, \text{“+”}, \text{“abc”} \rangle$

Instead of letting the domain of  $G$  consist of just the two strings “expr” and “def”, it turns out to be more convenient to let its domain be the entire set STRING of strings, and to let  $G[s]$  be the empty language (the empty set) for all strings  $s$  other than “expr” and “def”. So, a grammar is a function from the set of all strings to the set of sequences of strings. We can therefore define the set *Grammar* of all grammars by

$$\text{Grammar} \triangleq [\text{STRING} \rightarrow \text{SUBSET Seq}(\text{STRING})]$$

In describing the mathematical meaning of records, Section 5.2 explained that *r.ack* is an abbreviation for  $r[\text{“ack”}]$ . This is the case even if  $r$  isn’t a record. So, we can write  $G.op$  instead of  $G[\text{“op”}]$ . (A grammar isn’t a record because its domain is the set of all strings rather than a finite set of strings.)

A terminal like **ident** can appear anywhere to the right of a “::=” that a nonterminal like *expr* can, so a terminal should also be a set of sentences. A terminal is a set of sentences each containing a single lexeme. I will call such a sentence a *token*. Thus, the terminal **ident** is a set containing tokens such as  $\langle \text{“abc”} \rangle$ ,  $\langle \text{“x”} \rangle$ , and  $\langle \text{“qq”} \rangle$ . Any terminal appearing in the BNF grammar must be represented by a set of tokens, so the == in the grammar for SE is the set  $\{\langle \text{“==”} \rangle\}$ . Let’s define the operator *tok* by

$$\text{tok}(s) \triangleq \{\langle s \rangle\}$$

*tok* is short for *token*.

so we can write this set of tokens as  $\text{tok}(\text{“==”})$ .

A production expresses a relation between the values of  $G.str$  for some grammar  $G$  and some strings “str”. For example, the production

$$\text{def} ::= \text{ident} == \text{expr}$$

asserts that a sentence  $s$  is in  $G.\text{def}$  iff it has the form  $i \circ \langle \text{“==”} \rangle \circ e$  for some token  $i$  in **ident** and some sentence  $e$  in  $G.\text{expr}$ . In mathematics, a formula about  $G$  must mention  $G$  (perhaps indirectly by using a symbol defined in terms of  $G$ ). So, we can try writing this production in TLA<sup>+</sup> as

$$G.\text{def} ::= \text{ident } \text{tok}(\text{“==”}) G.\text{expr}$$

In the expression to the right of the “::=”, adjacency is expressing some operation. Just as we have to make multiplication explicit by writing  $2 * x$  instead of  $2x$ , we must express this operation by an explicit operator. Let’s use  $\&$ , so we can write the production as

$$(14.1) \quad G.\text{def} ::= \text{ident} \& \text{tok}(\text{“==”}) \& G.\text{expr}$$

This expresses the desired relation between the sets  $G.def$  and  $G.expr$  of sentences if  $::=$  is defined to be equality and  $\&$  is defined so that  $L \& M$  is the set of all sentences obtained by concatenating a sentence in  $L$  with a sentence in  $M$ :

$$L \& M \triangleq \{s \circ t : s \in L, t \in M\}$$

The production

$$expr ::= \mathbf{ident} \mid expr \mathbf{op} expr \mid (expr) \mid \mathbf{LET} \mathit{def} \mathbf{IN} expr$$

can similarly be expressed as

$$(14.2) \quad G.expr ::= \begin{array}{l} ident \\ \mid G.expr \& op \& G.expr \\ \mid tok("(") \& G.expr \& tok(")") \\ \mid tok("LET") \& G.def \& tok("IN") \& G.expr \end{array}$$

The precedence rules of  $TLA^+$  imply that  $a \mid b \& c$  is interpreted as  $a \mid (b \& c)$ .

This expresses the desired relation if  $\mid$  (which means *or* in the BNF grammar) is defined to be set union ( $\cup$ ).

We can also define the following operators that are sometimes used in BNF grammars:

- $Nil$  is defined so that  $Nil \& S$  equals  $S$  for any set  $S$  of sentences:

$$Nil \triangleq \{\langle \rangle\}$$

- $L^+$  equals  $L \mid L \& L \mid L \& L \& L \mid \dots$ :

$$L^+ \triangleq \mathbf{LET} \quad LL[n \in Nat] \triangleq \begin{array}{l} \text{IF } n = 0 \text{ THEN } L \\ \text{ELSE } LL[n-1] \mid LL[n-1] \& L \end{array} \quad \text{IN } \mathbf{UNION} \{LL[n] : n \in Nat\}$$

$LL[n] = L \mid \dots \mid \underbrace{L \& \dots L}_{n+1 \text{ copies}}$

$L^+$  is typed  $L^{*+}$  and  $L^*$  is typed  $L^{*}$ .

- $L^*$  equals  $Nil \mid L \mid L \& L \mid L \& L \& L \mid \dots$ :

$$L^* \triangleq Nil \mid L^+$$

The BNF grammar for SE consists of two productions, expressed by the  $TLA^+$  formulas (14.1) and (14.2). The entire grammar is the single formula that is the conjunction of these two formulas. We must turn this formula into a mathematical definition of a grammar  $GSE$ , which is a function from strings to languages. The formula is an assertion about a grammar  $G$ . We define  $GSE$  to be the smallest grammar  $G$  satisfying the conjunction of (14.1) and (14.2), where grammar  $G_1$  smaller than  $G_2$  means that  $G_1[s] \subseteq G_2[s]$  for every string  $s$ . To express this in  $TLA^+$ , we define an operator  $LeastGrammar$  so that  $LeastGrammar(P)$  is the smallest grammar  $G$  satisfying  $P(G)$ :

$$\begin{aligned} LeastGrammar(P(-)) &\triangleq \\ &\mathbf{CHOOSE} \ G \in Grammar : \wedge P(G) \\ &\quad \wedge \forall H \in Grammar : P(H) \Rightarrow (\forall s \in \mathbf{STRING} : G[s] \subseteq H[s]) \end{aligned}$$

Letting  $P(G)$  be the conjunction of (14.1) and (14.2), we can define the grammar  $GSE$  to be  $LeastGrammar(P)$ . We can then define the language  $SE$  to equal  $GSE.expr$ . The smallest grammar  $G$  satisfying a formula  $P$  must have  $G[s]$  equal to the empty language for any string  $s$  that doesn't appear in  $P$ . Thus,  $GSE[s]$  equals the empty language  $\{\}$  for any string  $s$  other than “expr” and “def”.

To complete our specification of  $GSE$ , we must define the sets *ident* and *op* of tokens. We can define the set *op* of operators by enumerating them—for example:

$$op \triangleq tok(“+”) \mid tok(“-”) \mid tok(“*”) \mid tok(“/”)$$

To express this a little more compactly, let's define  $Tok(S)$  to be the set of all tokens formed from elements in the set  $S$  of lexemes:

$$Tok(S) \triangleq \{\langle s \rangle : s \in S\}$$

We can then write

$$op \triangleq Tok(\{“+”, “-”, “*”, “/”\})$$

Let's define *ident* to be the set of tokens whose lexemes are words made entirely of lower-case letters, such as “abc”, “qq”, and “x”. To learn how to do that, we must first understand what strings in  $TLA^+$  really are. In  $TLA^+$ , a string is a sequence of characters. (We don't care, and the semantics of  $TLA^+$  doesn't specify, what a character is.) We can therefore apply the usual string operators on them. For example,  $Tail(“abc”)$  equals “bc”, and “abc”  $\circ$  “de” equals “abcde”.

The operators like  $\&$  that we just defined for expressing BNF were applied to sets of sentences, where a sentence is a sequence of lexemes. These operators can be applied just as well to sets of sequences of any kind—including sets of strings. For example,  $\{\text{“one”}, \text{“two”}\} \& \{\text{“s”}\}$  equals  $\{\text{“ones”}, \text{“twos”}\}$ , and  $\{\text{“ab”}\}^+$  is the set consisting of all the strings “ab”, “abab”, “ababab”, etc. So, we can define *ident* to equal  $Tok(Letter^+)$ , where *Letter* is the set of all lexemes consisting of a single lower-case letter:

$$Letter \triangleq \{\text{“a”}, \text{“b”}, \dots, \text{“z”}\}$$

Writing this definition out in full (without the “...”) is tedious. We can make this a little easier as follows. We first define the operator  $OneOf(s)$  to be the set of all one-character strings made from the characters of the string  $s$ :

$$OneOf(s) \triangleq \{[j \in \{1\} \mapsto s[i]] : i \in \text{DOMAIN } s\}$$

We can then define

$$Letter \triangleq OneOf(“abcdefghijklmnopqrstuvwxyz”)$$

See Section 15.1.10 on page 291 for more about strings. Remember that we take *sequence* and *tuple* to be synonymous.

Figure 14.1: The definition of the grammar  $GSE$  for the language SE.

All the operators we've defined here for specifying grammars are grouped into module *BNFGrammars*, which appears in Figure 14.2 on the next page.

The following TLA<sup>+</sup> module specifies the simple BNF grammar of TLA<sup>+</sup>. It makes use of the operators explained in Section 14.1.1 above, which are in the *BNFGrammars* module. However, if you are already familiar with BNF grammars, you should be able to read it without having read the previous section.

$$AtLeast4(s) \triangleq Tok(\{s \circ s \circ s\} \ \& \ \{s\}^+)$$



We now define some sets of lexemes. First is *ReservedWord*, the set of words that can't be used as identifiers. (Note that `BOOLEAN`, `TRUE`, `FALSE`, and `STRING` are identifiers that are predefined.)

$$\begin{aligned} \textit{ReservedWord} &\triangleq \\ &\{ \text{"ASSUME"}, \quad \text{"DOMAIN"}, \quad \text{"INSTANCE"}, \quad \text{"THEOREM"}, \\ &\quad \text{"ASSUMPTION"}, \quad \text{"ELSE"}, \quad \text{"LET"}, \quad \text{"UNCHANGED"}, \\ &\quad \text{"AXIOM"}, \quad \text{"ENABLED"}, \quad \text{"MODULE"}, \quad \text{"UNION"}, \\ &\quad \text{"CASE"}, \quad \text{"EXCEPT"}, \quad \text{"OTHER"}, \quad \text{"VARIABLE"}, \\ &\quad \text{"CHOOSE"}, \quad \text{"EXTENDS"}, \quad \text{"SF\_"}, \quad \text{"VARIABLES"}, \\ &\quad \text{"CONSTANT"}, \quad \text{"IF"}, \quad \text{"SUBSET"}, \quad \text{"WF\_"}, \\ &\quad \text{"CONSTANTS"}, \quad \text{"IN"}, \quad \text{"THEN"}, \quad \text{"WITH"} \} \end{aligned}$$

Here are three sets of characters—more precisely, sets of 1-character lexemes. They are the sets of letters, numbers, and characters that can appear in an identifier.

$$\begin{aligned} \textit{Letter} &\triangleq \\ &\textit{OneOf}(\text{"abcdefghijklmnopqrstuvwxyzABCDEFGHIJKLMNOPQRSTUVWXYZ"}) \\ \textit{Numeral} &\triangleq \textit{OneOf}(\text{"0123456789"}) \\ \textit{NameChar} &\triangleq \textit{Letter} \cup \textit{Numeral} \cup \{ \text{"\_"} \} \end{aligned}$$

We now define some sets of tokens. A *Name* is a token composed of letters, numbers, and `_` characters that contains at least one letter. It can be used as the name of a record component or a module. An *Identifier* is a *Name* that isn't a reserved word.

$$\begin{aligned} \textit{Name} &\triangleq \textit{Tok}(\textit{NameChar}^* \ \& \ \textit{Letter} \ \& \ \textit{NameChar}^*) \\ \textit{Identifier} &\triangleq \textit{Name} \setminus \textit{Tok}(\textit{ReservedWord}) \end{aligned}$$

An *IdentifierOrTuple* is either an identifier or a tuple of identifiers. Note that `<` is typed as `<<` and `>` as `>>`.

$$\begin{aligned} \textit{IdentifierOrTuple} &\triangleq \\ &\textit{Identifier} \mid \textit{tok}(\text{"<<"}) \ \& \ \textit{CommaList}(\textit{Identifier}) \ \& \ \textit{tok}(\text{">>"}) \end{aligned}$$

A *Number* is a token representing a number. You can write the integer 63 in the following ways: `63`, `63.00`, `\b111111` or `\B111111` (binary), `\o77` or `\O77` (octal), or `\h3f`, `\H3f`, `\h3F`, or `\H3F` (hexadecimal).

$$\begin{aligned} \textit{NumberLexeme} &\triangleq \quad \textit{Numeral}^+ \\ &\quad \mid (\textit{Numeral}^* \ \& \ \{ \text{"."} \} \ \& \ \textit{Numeral}^+) \\ &\quad \mid \{ \text{"\b"}, \text{"\B"} \} \ \& \ \textit{OneOf}(\text{"01"})^+ \\ &\quad \mid \{ \text{"\o"}, \text{"\O"} \} \ \& \ \textit{OneOf}(\text{"01234567"})^+ \\ &\quad \mid \{ \text{"\h"}, \text{"\H"} \} \ \& \ \textit{OneOf}(\text{"0123456789abcdefABCDEF"})^+ \end{aligned}$$

$$\textit{Number} \triangleq \textit{Tok}(\textit{NumberLexeme})$$

A *String* token represents a literal string. See Section 15.1.10 on page 291 to find out how special characters are typed in a string.

$$\textit{String} \triangleq \textit{Tok}(\{ \text{"\"} \} \ \& \ \textit{STRING} \ \& \ \{ \text{"\"} \})$$

We next define the sets of tokens that represent prefix operators (like  $\square$ ), infix operators (like  $+$ ), and postfix operators (like prime ( $'$ )). See Figure 13.16 on page 253 to find out what symbols these ASCII strings represent.

$$\text{PrefixOp} \triangleq \text{Tok}(\{ \text{"-"}, \text{"~"}, \text{"["}, \text{"<>"}, \text{"DOMAIN"}, \text{"ENABLED"}, \text{"SUBSET"}, \text{"UNCHANGED"}, \text{"UNION"} \})$$

$$\begin{aligned} \text{InfixOp} \triangleq & \\ & \text{Tok}(\{ \text{"!!"}, \text{"\#"}, \text{"\#\#"}, \text{"\$"}, \text{"\$\$"}, \text{"\%"}, \text{"\%\%"}, \\ & \text{"\&"}, \text{"\&\&"}, \text{"(+)"}, \text{"(-)"}, \text{"(.)"}, \text{"(/)"}, \text{"(\X)"}, \\ & \text{"*"}, \text{"**"}, \text{"+"}, \text{"++"}, \text{"-"}, \text{"-->"}, \text{"--"}, \\ & \text{"-|"}, \text{".."}, \text{"..."}, \text{"/"}, \text{"//"}, \text{"="}, \text{"\/"}, \\ & \text{"::="}, \text{":="}, \text{">"}, \text{"<"}, \text{"<:"}, \text{"<=>"}, \text{"="}, \\ & \text{"=<"}, \text{"=>"}, \text{"=|"}, \text{">"}, \text{">="}, \text{"?"}, \text{"??"}, \\ & \text{"@@"}, \text{"\"}, \text{"\"}, \text{"^"}, \text{"^"}, \text{"|"}, \text{"|"}, \\ & \text{"|="}, \text{"||"}, \text{"~>"}, \text{"."}, \\ & \text{"\approx"}, \text{"\asymp"}, \text{"\bigcirc"}, \text{"\bullet"}, \\ & \text{"\cap"}, \text{"\cdot"}, \text{"\circ"}, \text{"\cong"}, \\ & \text{"\cup"}, \text{"\div"}, \text{"\doteq"}, \text{"\equiv"}, \\ & \text{"\geq"}, \text{"\gg"}, \text{"\in"}, \text{"\intersect"}, \\ & \text{"\land"}, \text{"\leq"}, \text{"\ll"}, \text{"\lor"}, \\ & \text{"\mod"}, \text{"\o"}, \text{"\odot"}, \text{"\ominus"}, \\ & \text{"\oplus"}, \text{"\oslash"}, \text{"\otimes"}, \text{"\prec"}, \\ & \text{"\preceq"}, \text{"\propto"}, \text{"\sim"}, \text{"\simeq"}, \\ & \text{"\sqcap"}, \text{"\sqcup"}, \text{"\sqsubset"}, \text{"\sqsubseteq"}, \\ & \text{"\sqsupset"}, \text{"\sqsupseteq"}, \text{"\star"}, \text{"\subset"}, \\ & \text{"\subseteq"}, \text{"\succ"}, \text{"\succeq"}, \text{"\supset"}, \\ & \text{"\supseteq"}, \text{"\union"}, \text{"\uplus"}, \text{"\wr"} \}) \end{aligned}$$

$$\text{PostfixOp} \triangleq \text{Tok}(\{ \text{"^+"}, \text{"^*"}, \text{"^#"}, \text{"^,"} \})$$

Formally, the grammar *TLAPlusGrammar* of  $\text{TLA}^+$  is the smallest grammar satisfying the BNF productions below.

$$\begin{aligned} \text{TLAPlusGrammar} & \triangleq \\ \text{LET } P(G) & \triangleq \end{aligned}$$

Here is the BNF grammar. Terms that begin with “ $G$ .”, like  $G.\text{Module}$ , represent nonterminals. The terminals are sets of tokens, either defined above or described with the operators *tok* and *Tok*. The operators *AtLeast4* and *CommaList* are defined above.

$$\begin{aligned} \wedge G.\text{Module} ::= & \text{AtLeast4}(\text{"-"} ) \ \& \ \text{tok}(\text{"MODULE"} ) \ \& \ \text{Name} \ \& \ \text{AtLeast4}(\text{"-"} ) \\ & \ \& \ (G.\text{Unit})^* \\ & \ \& \ \text{AtLeast4}(\text{"="} ) \end{aligned}$$

$$\begin{aligned}
\wedge G.Unit &::= && G.VariableDeclaration \\
&| && G.ConstantDeclaration \\
&| && (Nil \mid tok("LOCAL")) \& G.OperatorDefinition \\
&| && (Nil \mid tok("LOCAL")) \& G.FunctionDefinition \\
&| && (Nil \mid tok("LOCAL")) \& G.Instance \\
&| && (Nil \mid tok("LOCAL")) \& G.ModuleDefinition \\
&| && G.Assumption \\
&| && G.Theorem \\
&| && G.Module \\
&| && AtLeast4("-") \\
\\
\wedge G.VariableDeclaration &::= && Tok(\{"VARIABLE", "VARIABLES"\}) \& CommaList(Identifier) \\
\\
\wedge G.ConstantDeclaration &::= && Tok(\{"CONSTANT", "CONSTANTS"\}) \& CommaList(G.OpDecl) \\
\\
\wedge G.OpDecl &::= && Identifier \\
&| && Identifier \& tok("(") \& CommaList(tok("_")) \& tok(")") \\
&| && PrefixOp \& tok("_") \\
&| && tok("_") \& InfixOp \& tok("_") \\
&| && tok("_") \& PostfixOp \\
\\
\wedge G.OperatorDefinition &::= && ( \quad G.NonFixLHS \\
&| && PrefixOp \& Identifier \\
&| && Identifier \& InfixOp \& Identifier \\
&| && Identifier \& PostfixOp ) \\
&&& \& tok("==") \\
&&& \& G.Expression \\
\\
\wedge G.NonFixLHS &::= && Identifier \\
&&& \& ( \quad Nil \\
&&& \quad | \quad tok("(") \& CommaList(Identifier \mid G.OpDecl) \& tok(")") ) \\
\\
\wedge G.FunctionDefinition &::= && Identifier \\
&&& \& tok("[") \& CommaList(G.QuantifierBound) \& tok("]") \\
&&& \& tok("==") \\
&&& \& G.Expression
\end{aligned}$$

$$\wedge G.QuantifierBound ::= \quad ( IdentifierOrTuple \mid CommaList(Identifier) ) \\ \quad \quad \quad \& \ tok(\backslash in) \\ \quad \quad \quad \& \ G.Expression$$

$$\wedge G.Instance ::= \quad tok(INSTANCE) \\ \quad \quad \quad \& \ Name \\ \quad \quad \quad \& \ ( Nil \mid tok(WITH) \& \ CommaList(G.Substitution) )$$

$$\wedge G.Substitution ::= \quad ( Identifier \mid PrefixOp \mid InfixOp \mid PostfixOp ) \\ \quad \quad \quad \& \ tok(<-) \\ \quad \quad \quad \& \ G.Argument$$

$$\wedge G.Argument ::= \quad G.Expression \\ \quad \quad \quad \mid G.GeneralPrefixOp \\ \quad \quad \quad \mid G.GeneralInfixOp \\ \quad \quad \quad \mid G.GeneralPostfixOp$$

$$\wedge G.InstancePrefix ::= \\ \quad ( Identifier \\ \quad \quad \& \ ( Nil \\ \quad \quad \quad \mid tok(") \& \ CommaList(G.Expression) \& \ tok(") ) \\ \quad \quad \& \ tok(!) )^*$$

$$\wedge G.GeneralIdentifier ::= G.InstancePrefix \& \ Identifier$$

$$\wedge G.GeneralPrefixOp ::= G.InstancePrefix \& \ PrefixOp$$

$$\wedge G.GeneralInfixOp ::= G.InstancePrefix \& \ InfixOp$$

$$\wedge G.GeneralPostfixOp ::= G.InstancePrefix \& \ PostfixOp$$

$$\wedge G.ModuleDefinition ::= G.NonFixLHS \& \ tok(==) \& \ G.Instance$$

$$\wedge G.Assumption ::= \\ \quad Tok(\{ ASSUME, ASSUMPTION, AXIOM \}) \& \ G.Expression$$

$$\wedge G.Theorem ::= tok(THEOREM) \& \ G.Expression$$

The comments give examples of each of the different types of expression.

$$\wedge G.Expression ::=$$

$$G.GeneralIdentifier \quad \boxed{A(x+7)!B!Id}$$

$$\mid G.GeneralIdentifier \& \ tok("(") \quad \boxed{A!Op(x+1,y)} \\ \quad \quad \& \ CommaList(G.Argument) \& \ tok(",")$$

$$\mid G.GeneralPrefixOp \& \ G.Expression \quad \boxed{SUBSET S.foo}$$

|  |                                                                                                                                                              |                                                                  |
|--|--------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------|
|  | $G.Expression \ \& \ G.GeneralInfixOp \ \& \ G.Expression$                                                                                                   | $a + b$                                                          |
|  | $G.Expression \ \& \ G.GeneralPostfixOp$                                                                                                                     | $x[1]'$                                                          |
|  | $tok("(") \ \& \ G.Expression \ \& \ tok(")")$                                                                                                               | $(x + 1)$                                                        |
|  | $Tok(\{"\backslash A", "\backslash E"\}) \ \& \ CommaList(G.QuantifierBound)$<br>$\ \& \ tok(":".)$ $G.Expression$                                           | $\forall x \in S, \langle y, z \rangle \in T : F(x, y, z)$       |
|  | $Tok(\{"\backslash A", "\backslash E", "\backslash AA", "\backslash EE"\}) \ \& \ CommaList(Identifier)$<br>$\ \& \ tok(":".)$ $G.Expression$                | $\exists x, y : x + y > 0$                                       |
|  | $tok("CHOOSE")$<br>$\ \& \ IdentifierOrTuple$<br>$\ \& \ (Nil \mid tok("\backslash in") \ \& \ G.Expression)$<br>$\ \& \ tok(":".)$<br>$\ \& \ G.Expression$ | $CHOOSE \ \langle x, y \rangle \in S : F(x, y)$                  |
|  | $tok("{") \ \& \ (Nil \mid CommaList(G.Expression)) \ \& \ tok("}")$                                                                                         | $\{1, 2, 2 + 2\}$                                                |
|  | $tok("{")$<br>$\ \& \ IdentifierOrTuple \ \& \ tok("\backslash in") \ \& \ G.Expression$<br>$\ \& \ tok(":".)$<br>$\ \& \ G.Expression$<br>$\ \& \ tok("}")$ | $\{x \in Nat : x > 0\}$                                          |
|  | $tok("{")$<br>$\ \& \ G.Expression$<br>$\ \& \ tok(":".)$<br>$\ \& \ CommaList(G.QuantifierBound)$<br>$\ \& \ tok("}")$                                      | $\{F(x, y, z) : x, y \in S, z \in T\}$                           |
|  | $G.Expression \ \& \ tok("[") \ \& \ CommaList(G.Expression) \ \& \ tok("]")$                                                                                | $f[i + 1, j]$                                                    |
|  | $tok("[")$<br>$\ \& \ CommaList(G.QuantifierBound)$<br>$\ \& \ tok(" ->")$<br>$\ \& \ G.Expression$<br>$\ \& \ tok("]")$                                     | $[i, j \in S, \langle p, q \rangle \in T \mapsto F(i, j, p, q)]$ |
|  | $tok("[") \ \& \ G.Expression \ \& \ tok(" ->") \ \& \ G.Expression \ \& \ tok("]")$                                                                         | $[(S \cup T) \rightarrow U]$                                     |
|  | $tok("[") \ \& \ CommaList(Name \ \& \ tok(" ->") \ \& \ G.Expression)$<br>$\ \& \ tok("]")$                                                                 | $[a \mapsto x + 1, b \mapsto y]$                                 |

|                                                                      |                                                                  |
|----------------------------------------------------------------------|------------------------------------------------------------------|
| tok("[") & CommaList( Name & tok(":") & G.Expression )               | $[a : Nat, b : S]$                                               |
| & tok("]")                                                           |                                                                  |
| tok("[")                                                             | $[f \text{ EXCEPT } ![1, x].r = 4, ![\langle 2, y \rangle] = e]$ |
| & G.Expression                                                       |                                                                  |
| & tok("EXCEPT")                                                      |                                                                  |
| & CommaList( tok("!=")                                               |                                                                  |
| & ( tok(".") & Name                                                  |                                                                  |
| tok("[") & CommaList(G.Expression) & tok("]") ) <sup>+</sup>         |                                                                  |
| & tok("=") & G.Expression )                                          |                                                                  |
| & tok("]")                                                           |                                                                  |
| tok("<<") & CommaList(G.Expression) & tok(">>")                      | $\langle 1, 2, 3 \rangle$                                        |
| G.Expression & (Tok({ "\X", "\times" }) & G.Expression) <sup>+</sup> | $Nat \times Nat \times Real$                                     |
| tok("[") & G.Expression & tok("_]") & G.Expression                   | $[Next]_{\langle x, y \rangle}$                                  |
| tok("<<") & G.Expression & tok(">>_") & G.Expression                 | $\langle Send \rangle_{vars}$                                    |
| Tok({ "WF_", "SF_" }) & G.Expression                                 | $WF_{vars}(Next)$                                                |
| & tok("(") & G.Expression & tok(")")                                 |                                                                  |
| tok("IF") & G.Expression & tok("THEN")                               | $IF \ p \ THEN \ A \ ELSE \ B$                                   |
| & G.Expression & tok("ELSE") & G.Expression                          |                                                                  |
| tok("CASE")                                                          | $CASE \ p1 \ \rightarrow \ e1$                                   |
| & ( LET CaseArm $\triangleq$                                         | $\square \ p2 \ \rightarrow \ e2$                                |
| G.Expression & tok(">") & G.Expression                               | $OTHER \ p3 \ \rightarrow \ e3$                                  |
| IN CaseArm & (tok("[") & CaseArm) <sup>*</sup> )                     |                                                                  |
| & ( Nil                                                              |                                                                  |
| ( tok("[") & tok("OTHER") & tok(">") & G.Expression) )               |                                                                  |
| tok("LET")                                                           | $LET \ x \triangleq y + 1$                                       |
| & ( L.OperatorDefinition                                             | $f[t \in Nat] \triangleq t^2$                                    |
| L.FunctionDefinition                                                 | $IN \ x + f[y]$                                                  |
| L.ModuleDefinition ) <sup>+</sup>                                    |                                                                  |
| & tok("IN")                                                          |                                                                  |
| & G.Expression                                                       |                                                                  |
| (tok("/") & G.Expression) <sup>+</sup>                               | $\wedge \ x = 1$                                                 |
|                                                                      | $\wedge \ y = 2$                                                 |
| (tok("\") & G.Expression) <sup>+</sup>                               | $\vee \ x = 1$                                                   |
|                                                                      | $\vee \ y = 2$                                                   |

|  |                  |                                                                                                             |
|--|------------------|-------------------------------------------------------------------------------------------------------------|
|  | <i>Number</i>    | <span style="border: 1px solid black; padding: 0 5px;">09001</span>                                         |
|  | <i>String</i>    | <span style="border: 1px solid black; padding: 0 5px;">"foo"</span>                                         |
|  | <i>tok</i> ("Ⓒ") | <span style="border: 1px solid black; padding: 0 5px;">@ (Can be used only in an EXCEPT expression.)</span> |

IN *LeastGrammar*(*P*)

## 14.2 The Complete Grammar

We now complete our explanation of the syntax of TLA<sup>+</sup> by giving the details that are not described by the BNF grammar in the previous section. Section 14.2.1 gives the precedence rules, Section 14.2.2 gives the alignment rules for conjunction and disjunction lists, and Section 14.2.3 describes comments. Section 14.2.4 briefly discusses the syntax of temporal formulas. Finally, for completeness, Section 14.2.5 explains the handling of two anomalous cases that you're unlikely ever to encounter.

### 14.2.1 Precedence and Associativity

The expression  $a + b * c$  is interpreted as  $a + (b * c)$  rather than  $(a + b) * c$ . This convention is described by saying that the operator  $*$  has *higher precedence* than the operator  $+$ . In general, operators with higher precedence are applied before operators of lower precedence. This applies to prefix operators (like SUBSET) and postfix operators (like ') as well as to infix operators like  $+$  and  $*$ . Thus,  $a + b'$  is interpreted as  $a + (b')$ , rather than as  $(a + b)'$ , because  $'$  has higher precedence than  $+$ . Application order can also be determined by associativity. The expression  $a - b - c$  is interpreted as  $(a - b) - c$  because  $-$  is a left-associative infix operator.

In TLA<sup>+</sup>, the precedence of an operator is a range of numbers, like 7–13. The operator  $\$$  has higher precedence than the operator  $:>$  because the precedence of  $\$$  is 7–12, and this entire range is greater than the precedence range of  $:>$ , which is 6–6. An expression is illegal (not syntactically well-formed) if the order of application of two operators is not determined because their precedence ranges overlap and they are not two instances of an associative infix operator. For example, the expression  $a + b * c' \% d$  is illegal for the following reason. The precedence range of  $'$  is higher than that of  $*$ , and the precedence range of  $*$  is higher than that of  $+$  and  $\%$ , so this expression is interpreted as  $a + (b * (c')) \% d$ . However, the precedences of  $+$  (9–9) and  $\%$  (9–10) overlap, so the expression is illegal.

TLA<sup>+</sup> embodies the philosophy that it's better to require parentheses than

to allow expressions that could be easily be misinterpreted. Thus,  $*$  and  $/$  have overlapping precedence, making an expression like  $a/b * c$  illegal. (This also makes  $a * b/c$  illegal, even though  $(a * b)/c$  and  $a * (b/c)$  happen to be equal when  $*$  and  $/$  have their usual definitions.) Unconventional operators like  $\$$  have wide precedence ranges for safety. But, even when the precedence rules imply that parentheses aren't needed, it's often a good idea to use them anyway if you think there's any chance that a reader might not understand how an expression is parsed.

Figure 13.14 on page 251 gives the precedence ranges of all operators and tells which infix operators are left associative. (No  $\text{TLA}^+$  operators are right associative.) Below are some additional precedence rules not covered by the operator precedence ranges.

### Function Application

Function application has higher precedence than any operator. Thus,  $a + b[c]'$  is interpreted as  $a + (b[c])'$ . (This in turn is interpreted as  $a + ((b[c])')$ , since the  $'$  operator has higher precedence than  $+$ .)

### Cartesian Products

In the Cartesian product construct,  $\times$  (typed as `\X` or `\times`) acts somewhat like an associative infix operator with precedence range 9–12. Thus,  $A \times B \subseteq C$  is interpreted as  $(A \times B) \subseteq C$ , rather than as  $A \times (B \subseteq C)$ . However,  $\times$  is part of a special construct, not an infix operator. For example, the three sets  $A \times B \times C$ ,  $(A \times B) \times C$ , and  $A \times (B \times C)$  are all different:

$$\begin{aligned} A \times B \times C &= \{\langle a, b, c \rangle : a \in A, b \in B, c \in C\} \\ (A \times B) \times C &= \{\langle \langle a, b \rangle, c \rangle : a \in A, b \in B, c \in C\} \\ A \times (B \times C) &= \{\langle a, \langle b, c \rangle \rangle : a \in A, b \in B, c \in C\} \end{aligned}$$

The first is a set of triples; the last two are sets of pairs.

### Undelimited Constructs

$\text{TLA}^+$  has several expression-making constructs with no explicit right-hand terminator. They are: CHOOSE, IF/THEN/ELSE, CASE, LET/IN, and quantifier constructs. These constructs are treated as prefix operators with the lowest possible precedence, so an expression made with one of them extends as far as possible. More precisely, the expression is ended only by one of the following:

- The beginning of the next module unit. (Module units are produced by the *Unit* nonterminal in the BNF grammar of Section 14.1.2; they include definition and declaration statements.)

- A right delimiter whose matching left delimiter occurs before the beginning of the construct. Delimiter pairs are ( ), [ ], { }, and ⟨ ⟩.
- Any of the following lexemes: THEN, ELSE, IN, comma (,), colon (:), →.
- The CASE separator □ (not the prefix temporal operator that is typed the same) ends all of these constructs except a CASE statement without an OTHER clause. That is, the □ acts as a delimiter except when it can be part of a CASE statement.
- Any symbol not to the right of the ∧ or ∨ prefixing a conjunction or disjunction list element containing the construct. (See Section 14.2.2 below.)

Here is how some expressions are interpreted under this rule:

$$\begin{array}{lcl}
 \text{IF } x > 0 \text{ THEN } y + 1 & & \text{IF } x > 0 \text{ THEN } y + 1 \\
 \text{ELSE } y - 1 & \text{means} & \text{ELSE } (y - 1 + 2) \\
 + 2 & & \\
 \\ 
 \forall x \in S : P(x) & \text{means} & \forall x \in S : (P(x) \vee Q) \\
 \vee Q & & 
 \end{array}$$

As these examples show, indentation is ignored—except in conjunction and disjunction lists, discussed below. The absence of a terminating lexeme (an END) for an IF/THEN/ELSE or CASE construct usually makes an expression less cluttered, but sometimes it does require you to add parentheses.

### Subscripts

TLA uses subscript notation in the following constructs:  $[A]_e$ ,  $\langle A \rangle_e$ ,  $\text{WF}_e(A)$ , and  $\text{SF}_e(A)$ . In TLA<sup>+</sup>, these are written with a “\_” character, as in  $\ll A \gg\_e$ . This notation is, in principle, problematic. The expression  $\ll A \gg\_x \wedge B$ , which we expect to mean  $(\langle A \rangle_x) \wedge B$ , could conceivably be interpreted as  $\langle A \rangle_{(x \wedge B)}$ . The precise rule for parsing these constructs isn’t important; you should put parentheses around the subscript except in the following two cases.

- The subscript is a *GeneralIdentifier* in the BNF grammar.
- The subscript is an expression enclosed by one of the following matching delimiter pairs: ( ), [ ], ⟨ ⟩, or { }—for example,  $\langle x, y \rangle$  or  $(x + y)$ .

Although  $[A]_f[x]$  is interpreted correctly as  $[A]_{f[x]}$ , it will be easier to read in the ASCII text if you write it as  $[A]_-(f[x])$ .

### 14.2.2 Alignment

The most novel aspect of TLA<sup>+</sup> syntax is the aligned conjunction and disjunction lists. If you write such a list in a straightforward manner, then it will mean what you expect it to. However, you might wind up doing something weird through a typing error. So, it's a good idea to know what the exact syntax rules are for these lists. I give the rules here for conjunction lists; the rules for disjunction lists are the same.

A conjunction list is an expression that begins with  $\wedge$ , which is typed as  $\wedge$ . Let  $c$  be the column in which the  $/$  occurs. (The first character in a line is in column 1.) The conjunction list consists of a sequence of conjuncts, each beginning with a  $\wedge$ . A conjunct is ended by any one of the following that occurs after the  $\wedge$ :

1. Another  $\wedge$  whose  $/$  character is in column  $c$  and is the first nonspace character on the line.
2. Any nonspace character in column  $c$  or a column to the left of column  $c$ .
3. A right delimiter whose matching left delimiter occurs before the beginning of the conjunction list. Delimiter pairs are  $()$ ,  $[]$ ,  $\{\}$ , and  $\langle \rangle$ .
4. The beginning of the next module unit. (Module units are produced by the *Unit* nonterminal in the BNF grammar; they include definition and declaration statements.)

In case 1, the  $\wedge$  begins the next conjunct in the same conjunction list. In the other three cases, the end of the conjunct is the end of the entire conjunction list. In all cases, the character ending the conjunct does not belong to the conjunct. With these rules, indentation properly delimits expressions in a conjunction list—for example:

|                                      |       |                                         |
|--------------------------------------|-------|-----------------------------------------|
| $\wedge$ IF $e$ THEN $P$<br>ELSE $Q$ | means | $\wedge$ (IF $e$ THEN $P$<br>ELSE $Q$ ) |
| $\wedge$ $R$                         |       | $\wedge R$                              |

It's best to indent each conjunction completely to the right of its  $\wedge$  symbol. These examples illustrate precisely what happens if you don't:

|                        |       |                                    |       |                                                                                 |
|------------------------|-------|------------------------------------|-------|---------------------------------------------------------------------------------|
| $\wedge$ $x'$<br>$= y$ | means | $\wedge x' = y$<br>$\wedge y' = x$ | means | $\wedge x'$ $(\wedge x')$<br>$= y$ $= (y$<br>$\wedge y' = x$ $\wedge (y' = x))$ |
|------------------------|-------|------------------------------------|-------|---------------------------------------------------------------------------------|

In the second example, the  $\wedge x'$  is interpreted as a conjunction list containing only one conjunct, and the second  $\wedge$  is interpreted as an infix operator.

You can't use parentheses to circumvent the indentation rules. For example, this is illegal:

```

/\ (x'
= y)
/\ y'=x

```

The rules imply that the first  $\wedge$  begins a conjunction list that is ended before the  $=$ . That conjunction list is therefore  $\wedge(x'$ , which has an unmatched left parenthesis.

The conjunction/disjunction list notation is quite robust. Even if you mess up the alignment by typing one space too few or too many—something that's easy to do when the conjuncts are long—the formula is still likely to mean what you intended. Here's an example of what happens if you misalign a conjunct:

```

/\ A      ((\ A)
/\ B  means  \wedge B)
/\ C      \wedge C

```

(The last two  $\wedge$  symbols are interpreted as infix operators.) While not interpreted the way you expected, this formula is equivalent to  $A \wedge B \wedge C$ , which is what you meant in the first place.

Most keyboards contain one key that is the source of a lot of trouble: the tab key (sometimes marked on the keyboard with a right arrow). On my computer screen, I can produce

```

A ==
      /\ x' = 1
      /\ y' = 2

```

by beginning the second line with eight space characters and the third with one tab character. In this case, it is unspecified whether or not the two  $/$  characters occur in the same column. Tab characters are an anachronism left over from the days of typewriters and of computers with memory capacity measured in kilobytes. I strongly advise you never to use them. But, if you insist on using them, here are the rules:

- A tab character is considered to be equivalent to one or more space characters, so it occupies one or more columns.
- Identical sequences of space and tab characters that occur at the beginning of a line occupy the same number of columns.

There are no other guarantees if you use tab characters.

### 14.2.3 Comments

Comments are described in Section 3.5 on page 32. A comment may appear between any two lexemes in a specification. There are two types of comments:

- A delimited comment is a string of the form “ $( * \circ s \circ * )$ ”, where  $s$  is any string not containing the substring “ $*$ ”.
- An end-of-line comment is a string of the form “ $\backslash * \circ s \circ \langle \text{LF} \rangle$ ”, where  $s$  is any string not containing an end-of-line character  $\langle \text{LF} \rangle$ .

I like to write comments as shown here:

```
BufRcv == /\ InChan!Rcv          (*****)
          /\ q' = Append(q, in.val) (* Receive message from channel *)
          /\ out                  (* in and append to tail of q. *)
                                   (*****)
```

Grammatically, this piece of specification has four distinct comments, the first and last consisting of the same string  $(**\dots**)$ . But a person reading it would regard them as a single comment, spread over four lines. This kind of commenting convention is not part of the  $\text{TLA}^+$  language, but it might be supported by tools such as a pretty-printer.

## 14.2.4 Temporal Formulas

The BNF grammar treats  $\Box$  and  $\Diamond$  simply as prefix operators. However, as explained in Section 8.1 (page 88), the syntax of temporal formulas places restrictions on their use. For example,  $\Box(x' = x + 1)$  is not a legal formula. It's not hard to write a BNF grammar that specifies legal temporal formulas made from the temporal operators and ordinary Boolean operators like  $\neg$  and  $\wedge$ . However, such a BNF grammar won't tell you which of these two expressions is legal:

|                                          |                                          |
|------------------------------------------|------------------------------------------|
| LET $F(P, Q) \triangleq P \wedge \Box Q$ | LET $F(P, Q) \triangleq P \wedge \Box Q$ |
| IN $F(x = 1, x = y + 1)$                 | IN $F(x = 1, x' = y + 1)$                |

The first is legal; the second isn't because it represents the illegal formula

$(x = 1) \wedge \Box(x' = y + 1)$  This formula is illegal.

The precise rules for determining if a temporal formula is syntactically well-formed involve first replacing all defined operators by their definitions, using the procedure described in Section 16.4 below. I won't bother specifying these rules.

In practice, temporal operators are not used very much in  $\text{TLA}^+$  specifications, and one rarely writes definitions of new ones such as

$$F(P, Q) \triangleq P \wedge \Box Q$$

The syntactic rules for expressions involving such operators are of academic interest only.

### 14.2.5 Two Anomalies

There are two sources of ambiguity in the grammar of  $TLA^+$  that you are unlikely to encounter and that have *ad hoc* resolutions. The least unlikely of these arises from the use of  $-$  as both an infix operator (as in  $2 - 2$ ) and a prefix operator (as in  $2 + -2$ ). This poses no problem when  $-$  is used in an ordinary expression. However, there are two places in which an operator can appear by itself:

- As the argument of a higher-order operator, as in  $HOp(+, -)$ .
- In an INSTANCE substitution, such as

INSTANCE  $M$  WITH  $Plus \leftarrow +$ ,  $Minus \leftarrow -$

In both these cases, the symbol  $-$  is interpreted as the infix operator. You must type  $-.$  to denote the prefix operator. You also have to type  $-.$  if you should ever want to define the prefix  $-$  operator, as in:

$$- . a \triangleq UMinus(a)$$

Remember that, in ordinary expressions, you just type  $-$  as usual for both operators.

The second source of ambiguity in the  $TLA^+$  syntax is an unlikely expression of the form  $\{x \in S : y \in T\}$ , which might be taken to mean either of the following:

$$\begin{aligned} \text{LET } p &\triangleq y \in T \text{ IN } \{x \in S : p\} \\ \text{LET } p &\triangleq x \in S \text{ IN } \{p : y \in T\} \end{aligned}$$

It is interpreted as the first formula.

## 14.3 What You Type

The grammar in Section 14.1.2 describes the ASCII syntax for  $TLA^+$  specifications. Typeset versions of specifications appear in this book. For example, the grammar lists the infix operator `\prec`, but that operator is printed in specifications as  $\prec$ . Figure 13.16 on page 253 gives the correspondence between the ASCII and typeset versions of all  $TLA^+$  symbols for which the correspondence may not be obvious.

Finally, Figure 13.13 on page 250 lists all the user-definable infix, postfix, and prefix operator symbols of  $TLA^+$ . It also indicates which of them are defined by the standard operator modules. This is a good place to look when choosing notation for your specification.

## Chapter 15

# The Operators of $\text{TLA}^+$

This chapter describes the built-in operators of  $\text{TLA}^+$ . Most of these operators have been described in Part I. Here, you can find brief explanations of the operators, along with references to the longer descriptions in Part I. The explanations cover some subtle points that are not mentioned elsewhere. The chapter can serve as a reference manual for readers who have finished Part I or who are already familiar enough with the mathematical concepts that the brief explanations are all they need.

The chapter includes a formal semantics of the operators. The rigorous description of  $\text{TLA}^+$  that a formal semantics provides is usually needed only by people building  $\text{TLA}^+$  tools. If you're not building a tool and don't have a special fondness for formalism, you will probably want to skip all the subsections titled *Formal Semantics*. However, you may some day encounter an obscure question about the meaning of a  $\text{TLA}^+$  operator that is answered only by the formal semantics.

This chapter also defines some of the “semantic” conditions on the syntax of  $\text{TLA}^+$  that are omitted from the grammar of Chapter 14. For example, it tells you that  $[a : \text{Nat}, a : \text{BOOLEAN}]$  is an illegal expression. Other semantic conditions on expressions arise from a combination of the definitions in this chapter and the conditions stated in Chapter 16. For example, this chapter defines  $\exists x, x : p$  to equal  $\exists x : (\exists x : p)$ , and Chapter 16 tells you that the latter expression is illegal.

### 15.1 Constant Operators

We first define the constant operators of  $\text{TLA}^+$ . These are the operators of ordinary mathematics, having nothing to do with TLA or temporal logic. All the constant operators of  $\text{TLA}^+$  are listed in Figure 13.9 on page 248 and Fig-

ure 13.10 on page 249.

An operator combines one or more expressions into a “larger” expression. For example, the set union operator  $\cup$  combines two expressions  $e_1$  and  $e_2$  into the expression  $e_1 \cup e_2$ . Some operators don’t have such simple names as  $\cup$ . For example, there’s no simple name for the operator that combines the  $n$  expressions  $e_1, \dots, e_n$  to form the expression  $\{e_1, \dots, e_n\}$ . We could name it  $\{\dots\}$  or  $\{-, \dots, -\}$ , but that would be awkward. Instead of explicitly mentioning the operator, I’ll refer to the *construct*  $\{e_1, \dots, e_n\}$ . The distinction between an operator like  $\cup$  and the nameless one used in the construct  $\{e_1, \dots, e_n\}$  is purely syntactic, with no mathematical significance. In Chapter 16, I will abstract away from this syntactic difference and treat all operators uniformly. For now, I’ll stay closer to the syntax.

## Formal Semantics

A formal semantics for a language is a translation from that language into some form of mathematics. We assign a mathematical expression  $\llbracket e \rrbracket$ , called the *meaning* of  $e$ , to certain terms  $e$  in the language. Since we presumably understand the mathematics, we know what  $\llbracket e \rrbracket$  means, and that tells us what  $e$  means.

Meaning is generally defined inductively. For example, the meaning  $\llbracket e_1 \cup e_2 \rrbracket$  of the expression  $e_1 \cup e_2$  would be defined in terms of the meanings  $\llbracket e_1 \rrbracket$  and  $\llbracket e_2 \rrbracket$  of its subexpressions. This definition is said to define the semantics of the operator  $\cup$ .

Because much of  $TLA^+$  is a language for expressing ordinary mathematics, much of its semantics is trivial. For example, the semantics of  $\cup$  can be defined by

$$\llbracket e_1 \cup e_2 \rrbracket \triangleq \llbracket e_1 \rrbracket \cup \llbracket e_2 \rrbracket$$

In this definition, the  $\cup$  to the left of the  $\triangleq$  is the  $TLA^+$  symbol, while the one to the right is the set-union operator of ordinary mathematics. I could make the distinction between the two uses of the symbol  $\cup$  more obvious by writing

$$\llbracket e_1 \setminus\text{cup} e_2 \rrbracket \triangleq \llbracket e_1 \rrbracket \cup \llbracket e_2 \rrbracket$$

But, that wouldn’t make the definition any less trivial.

Instead of trying to maintain a distinction between the  $TLA^+$  operator  $\cup$  and the operator of set theory that’s written the same, I simply use  $TLA^+$  as the language of mathematics in which to define the semantics of  $TLA^+$ . That is, I take as primitive certain  $TLA^+$  operators that, like  $\cup$ , correspond to well-known mathematical operators. I describe the formal semantics of the constant operators of  $TLA^+$  by defining them in terms of these primitive operators. I also describe the semantics of some of the primitive operators by stating the axioms that they satisfy.

### 15.1.1 Boolean Operators

The truth values of logic are written in  $\text{TLA}^+$  as `TRUE` and `FALSE`. The built-in constant `BOOLEAN` is the set consisting of those two values:

$$\text{BOOLEAN} \triangleq \{\text{TRUE}, \text{FALSE}\}$$

$\text{TLA}^+$  provides the usual operators of propositional logic:

$$\wedge \quad \vee \quad \neg \quad \Rightarrow \text{ (implication)} \quad \equiv \quad \text{TRUE} \quad \text{FALSE}$$

They are explained in Section 1.1. Conjunctions and disjunctions can be written as aligned lists:

$$\begin{array}{ccc} \wedge & p_1 & \vee & p_1 \\ & \vdots & & \vdots \\ & \triangleq & p_1 \wedge \dots \wedge p_n & \triangleq & p_1 \vee \dots \vee p_n \\ & \wedge & p_n & & \vee & p_n \end{array}$$

The standard quantified formulas of predicate logic are written in  $\text{TLA}^+$  as:

$$\forall x : P \quad \exists x : P$$

I call these the *unbounded* quantifier constructions. The *bounded* versions are written as:

$$\forall x \in S : p \quad \exists x \in S : p$$

The meanings of these expressions are described in Section 1.3.  $\text{TLA}^+$  allows some common abbreviations—for example:

$$\begin{aligned} \forall x, y : p &\triangleq \forall x : (\forall y : p) \\ \exists x, y \in S, z \in T : p &\triangleq \exists x \in S : (\exists y \in S : (\exists z \in T : p)) \end{aligned}$$

$\text{TLA}^+$  also allows bounded quantification over tuples, such as

$$\forall \langle x, y \rangle \in S : p$$

This formula is true iff, for any pair  $\langle a, b \rangle$  in  $S$ , the formula obtained from  $p$  by substituting  $a$  for  $x$  and  $b$  for  $y$  is true.

### Formal Semantics

Propositional and predicate logic, along with set theory, form the foundation of ordinary mathematics. In defining the semantics of  $\text{TLA}^+$ , we therefore take as primitives the operators of propositional logic and the simple unbounded quantifier constructs  $\exists x : p$  and  $\forall x : p$ , where  $x$  is an identifier. Among the Boolean operators described above, this leaves only the general forms of the quantifiers, given by the BNF grammar of Chapter 14, whose meanings must

be defined. This is done by defining those general forms in terms of the simple forms.

The unbounded operators have the general forms:

$$\forall x_1, \dots, x_n : p \quad \exists x_1, \dots, x_n : p$$

where each  $x_i$  is an identifier. They are defined in terms of quantification over a single variable by:

$$\forall x_1, \dots, x_n : p \triangleq \forall x_1 : (\forall x_2 : (\dots \forall x_n : p) \dots)$$

and similarly for  $\exists$ . The bounded operators have the general forms:

$$\forall \mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n : p \quad \exists \mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n : p$$

where each  $\mathbf{y}_i$  has the form  $x_1, \dots, x_k$  or  $\langle x_1, \dots, x_k \rangle$ , and each  $x_j$  is an identifier. The general forms of  $\forall$  are defined inductively by

$$\forall x_1, \dots, x_k \in S : p \triangleq \forall x_1, \dots, x_k : (x_1 \in S) \wedge \dots \wedge (x_k \in S) \Rightarrow p$$

$$\forall \mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n : p \triangleq \forall \mathbf{y}_1 \in S_1 : \dots \forall \mathbf{y}_n \in S_n : p$$

$$\forall \langle x_1, \dots, x_k \rangle \in S : p \triangleq \forall x_1, \dots, x_k : (\langle x_1, \dots, x_k \rangle \in S) \Rightarrow p$$

where the  $\mathbf{y}_i$  are as before. In these expressions,  $S$  and the  $S_i$  lie outside the scope of the quantifier's bound identifiers. The definitions for  $\exists$  are similar. In particular:

$$\exists \langle x_1, \dots, x_k \rangle \in S : p \triangleq \exists x_1, \dots, x_k : (\langle x_1, \dots, x_k \rangle \in S) \wedge p$$

See Section 15.1.9 for further details about tuples.

### 15.1.2 The Choose Operator

A simple unbounded CHOOSE expression has the form

$$\text{CHOOSE } x : p$$

As explained in Section 6.6, the value of this expression is some arbitrary value  $v$  such that  $p$  is true if  $v$  is substituted for  $x$ , if such a  $v$  exists. If no such  $v$  exists, then the expression has a completely arbitrary value.

The bounded form of the CHOOSE expression is:

$$\text{CHOOSE } x \in S : p$$

It is defined in terms of the unbounded form by

$$(15.1) \quad \text{CHOOSE } x \in S : p \triangleq \text{CHOOSE } x : (x \in S) \wedge p$$

It is equal to some arbitrary value  $v$  in  $S$  such that  $p$ , with  $v$  substituted for  $x$ , is true—if such a  $v$  exists. If no such  $v$  exists, the CHOOSE expression has a completely arbitrary value.

A CHOOSE expression can also be used to choose a tuple. For example

$$\text{CHOOSE } \langle x, y \rangle \in S : p$$

equals some pair  $\langle v, w \rangle$  in  $S$  such that  $p$ , with  $v$  substituted for  $x$  and  $w$  substituted for  $y$ , is true—if such a pair exists. If no such pair exists, it has an arbitrary value, which need not be a pair.

The unbounded CHOOSE operator satisfies the following two rules:

$$(15.2) \quad (\exists x : P(x)) \equiv P(\text{CHOOSE } x : P(x))$$

$$(\forall x : P(x) = Q(x)) \Rightarrow ((\text{CHOOSE } x : P(x)) = (\text{CHOOSE } x : Q(x)))$$

for any operators  $P$  and  $Q$ . We know nothing about the value chosen by CHOOSE except what we can deduce from these rules.

The second rule allows us to deduce the equality of certain CHOOSE expressions that we might expect to be different. In particular, for any operator  $P$ , if there exists no  $x$  satisfying  $P(x)$ , then  $\text{CHOOSE } x : P(x)$  equals the unique value  $\text{CHOOSE } x : \text{FALSE}$ . For example, the *Reals* module defines division by

$$a/b \triangleq \text{CHOOSE } c \in \text{Real} : a = b * c$$

For any nonzero number  $a$ , there exists no number  $c$  such that  $a = 0 * c$ . Hence,  $a/0$  equals  $\text{CHOOSE } c : \text{FALSE}$ , for any nonzero  $a$ . We can therefore deduce that  $1/0$  equals  $2/0$ .

We would expect to be unable to deduce anything about the nonsensical expression  $1/0$ . It's a bit disquieting to prove that it equals  $2/0$ . If this upsets you, here's a way to define division that will make you happier. First define an operator *Choice* so that  $\text{Choice}(v, P)$  equals  $\text{CHOOSE } x : P(x)$  if there exists an  $x$  satisfying  $P(x)$ , and otherwise equals some arbitrary value that depends on  $v$ . There are many ways to define *Choice*; here's one:

$$\text{Choice}(v, P(-)) \triangleq \text{IF } \exists x : P(x) \text{ THEN } \text{CHOOSE } x : P(x) \\ \text{ELSE } (\text{CHOOSE } x : x.a = v).b$$

You can then define division by

$$a/b \triangleq \text{LET } P(c) \triangleq (c \in \text{Real}) \wedge (a = b * c) \\ \text{IN } \text{Choice}(a, P)$$

This definition makes it impossible to deduce any relation between  $1/0$  and  $2/0$ . You can use *Choice* instead of CHOOSE whenever this kind of problem arises—if you consider  $1/0$  equaling  $2/0$  to be a problem. But there is seldom any practical reason for worrying about it.

### Formal Semantics

We take the construct  $\text{CHOOSE } x : p$ , where  $x$  is an identifier, to be primitive. This form of the  $\text{CHOOSE}$  operator is known to mathematicians as *Hilbert's  $\varepsilon$* . Its meaning is defined mathematically by the rules (15.2). Leisenring [3] presents a detailed mathematical exposition of Hilbert's  $\varepsilon$ .

An unbounded  $\text{CHOOSE}$  of a tuple is defined in terms of the simple unbounded  $\text{CHOOSE}$  construct by

$$\text{CHOOSE } \langle x_1, \dots, x_n \rangle : p \triangleq \text{CHOOSE } y : (\exists x_1, \dots, x_n : (y = \langle x_1, \dots, x_n \rangle) \wedge p)$$

where  $y$  is an identifier that is different from the  $x_i$  and does not occur in  $p$ . The bounded  $\text{CHOOSE}$  construct is defined in terms of unbounded  $\text{CHOOSE}$  by (15.1), where  $x$  can be either an identifier or a tuple.

### 15.1.3 The Three Interpretations of Boolean Operators

The meaning of a Boolean operator when applied to Boolean values is a standard part of traditional mathematics. Everyone agrees that  $\text{TRUE} \wedge \text{FALSE}$  equals  $\text{FALSE}$ . However, because  $TLA^+$  is untyped, an expression like  $2 \wedge \langle 5 \rangle$  is legal. We must therefore decide what it means. There are three ways of doing this, which I call the *conservative*, *moderate*, and *liberal* interpretations.

In the conservative interpretation, the value of an expression like  $2 \wedge \langle 5 \rangle$  is completely unspecified. It could equal  $\sqrt{2}$ . It need not equal  $\langle 5 \rangle \wedge 2$ . Hence, the ordinary laws of logic, such as the commutativity of  $\wedge$ , are valid only for Boolean values.

In the liberal interpretation, the value of  $2 \wedge \langle 5 \rangle$  is specified to be a Boolean. It is not specified whether it equals  $\text{TRUE}$  or  $\text{FALSE}$ . However, all the ordinary laws of logic, such as the commutativity of  $\wedge$ , are valid. Hence,  $2 \wedge \langle 5 \rangle$  equals  $\langle 5 \rangle \wedge 2$ . More precisely, any tautology of propositional or predicate logic, such as

$$(\forall x : p) \equiv \neg(\exists x : \neg p)$$

is valid, even if  $p$  is not necessarily a Boolean for all values of  $x$ .<sup>1</sup> It is easy to show that the liberal approach is sound.<sup>2</sup> For example, one way of defining operators that satisfy the liberal interpretation is to consider any nonBoolean value to be equivalent to  $\text{FALSE}$ .

The conservative and liberal interpretations are equivalent for most specifications, except for ones that use Boolean-valued functions. In practice, the

<sup>1</sup>Equality ( $=$ ) is not an operator of propositional or predicate logic; this tautology need not be valid for nonBoolean values if  $\equiv$  is replaced by  $=$ .

<sup>2</sup>A sound logic in one in which  $\text{FALSE}$  is not provable.

conservative interpretation doesn't permit you to use  $f[x]$  as a Boolean expression even if  $f$  is defined to be a Boolean-valued function. For example, suppose we define the function  $tnat$  by

$$tnat \triangleq [n \in Nat \mapsto \text{TRUE}]$$

so  $tnat[n]$  equals TRUE for all  $n$  in  $Nat$ . The formula

$$(15.3) \quad \forall n \in Nat : tnat[n]$$

equals TRUE in the liberal interpretation, but not in the conservative interpretation. Formula (15.3) is equivalent to

$$\forall n : (n \in Nat) \Rightarrow tnat[n]$$

which asserts that  $(n \in Nat) \Rightarrow tnat[n]$  is true for all  $n$ , including, for example,  $n = 1/2$ . For (15.3) to equal TRUE, the formula  $(1/2 \in Nat) \Rightarrow tnat[1/2]$ , which equals  $\text{FALSE} \Rightarrow tnat[1/2]$ , must equal TRUE. But the value of  $tnat[1/2]$  is not specified; it might equal  $\sqrt{2}$ . The formula  $\text{FALSE} \Rightarrow \sqrt{2}$  equals TRUE in the liberal interpretation; its value is unspecified in the conservative interpretation. Hence, the value of (15.3) is unspecified in the conservative interpretation. If we are using the conservative interpretation, instead of (15.3), we should write

$$\forall n \in Nat : (tnat[n] = \text{TRUE})$$

This formula equals TRUE in both interpretations.

The conservative interpretation is philosophically more satisfying, since it makes no assumptions about a silly expression like  $2 \wedge \langle 5 \rangle$ . However, as we have just seen, it would be nice if the not-so-silly formula  $\text{FALSE} \Rightarrow \sqrt{2}$  equaled TRUE. We therefore introduce the moderate interpretation, which lies between the conservative and liberal interpretations. It assumes only that expressions involving FALSE and TRUE have their expected values—for example,  $\text{FALSE} \Rightarrow \sqrt{2}$  equals TRUE, and  $\text{FALSE} \wedge 2$  equals FALSE. In the moderate interpretation, (15.3) equals TRUE, but the value of  $\langle 5 \rangle \wedge 2$  is still completely unspecified.

The laws of logic still do not hold unconditionally in the moderate interpretation. The formulas  $p \wedge q$  and  $q \wedge p$  are equivalent only if  $p$  and  $q$  are both Booleans, or if one of them equals FALSE. When using the moderate interpretation, we still have to check that all the relevant values are Booleans before applying any of the ordinary rules of logic in a proof. This can be burdensome in practice.

The semantics of  $\text{TLA}^+$  asserts that the rules of the moderate interpretation are valid. The liberal interpretation is neither required nor forbidden. You should write specifications that make sense under the moderate interpretation. However, you (and the implementer of a tool) are free to use the liberal interpretation if you wish.

### 15.1.4 Conditional Constructs

$TLA^+$  provides two conditional constructs for forming expressions that are inspired by constructs from programming languages: IF/THEN/ELSE and CASE.

The IF/THEN/ELSE construct was introduced on page 16 of Section 2.2. Its general form is:

IF  $p$  THEN  $e_1$  ELSE  $e_2$

It equals  $e_1$  if  $p$  is true, and  $e_2$  if  $p$  is false.

An expression can sometimes be simplified by using a CASE construct instead of nested IF/THEN/ELSE constructs. The CASE construct has two general forms:

(15.4) CASE  $p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n$   
CASE  $p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n \square \text{OTHER} \rightarrow e$

If some  $p_i$  is true, then the value of these expressions is some  $e_i$  such that  $p_i$  is true. For example, the expression

CASE  $n \geq 0 \rightarrow e_1 \square n \leq 0 \rightarrow e_2$

equals  $e_1$  if  $n > 0$  is true, equals  $e_2$  if  $n < 0$  is true, and equals either  $e_1$  or  $e_2$  if  $n = 0$  is true. In the latter case, the semantics of  $TLA^+$  does not specify whether the expression equals  $e_1$  or  $e_2$ . The CASE expressions (15.4) are generally used when the  $p_i$  are mutually disjoint, so at most one  $p_i$  can be true.

The two expressions (15.4) differ when  $p_i$  is false for all  $i$ . In that case, the value of the first is unspecified, while the value of the second is  $e$ , the OTHER expression. If you use a CASE expression without an OTHER clause, the value of the expression should matter only when  $\exists i \in 1 \dots n : p_i$  is true.

### Formal Semantics

The IF/THEN/ELSE and CASE constructs are defined as follows in terms of CHOOSE:

IF  $p$  THEN  $e_1$  ELSE  $e_2 \triangleq \text{CHOOSE } v : (p \Rightarrow v = e_1) \wedge (\neg p \Rightarrow v = e_2)$   
CASE  $p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n \triangleq$   
CHOOSE  $v : (p_1 \wedge (v = e_1)) \vee \dots \vee (p_n \wedge (v = e_n))$   
CASE  $p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n \square \text{OTHER} \rightarrow e \triangleq$   
CASE  $p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n \square \neg(p_1 \vee \dots \vee p_n) \rightarrow e$

### 15.1.5 LET

The LET/IN construct was introduced on page 59 of Section 5.6. The expression

LET  $d \triangleq f$  IN  $e$

equals  $e$  in the context of the definition  $d \triangleq f$ . For example

$$\text{LET } sq(i) \triangleq i * i \text{ IN } sq(1) + sq(2) + sq(3)$$

equals  $1 * 1 + 2 * 2 + 3 * 3$ , which equals 14. The general form of the construct is:

$$\text{LET } \Delta_1 \dots \Delta_n \text{ IN } e$$

where each  $\Delta_i$  has the syntactic form of any  $\text{TLA}^+$  definition. Its value is  $e$  in the context of the definitions  $\Delta_i$ . More precisely, it equals

$$\text{LET } \Delta_1 \text{ IN } (\text{LET } \Delta_2 \text{ IN } (\dots \Delta_n \text{ IN } e) \dots)$$

Hence, the symbol defined in  $\Delta_1$  can be used in the definitions  $\Delta_2, \dots, \Delta_n$ .

## Formal Semantics

The formal semantics of the LET construct is defined in Section 16.4 (page 309) below.

### 15.1.6 The Operators of Set Theory

$\text{TLA}^+$  provides the following operators on sets:

$$\in \quad \notin \quad \cup \quad \cap \quad \subseteq \quad \setminus \quad \text{UNION} \quad \text{SUBSET}$$

and the following set constructors:

$$\{e_1, \dots, e_n\} \quad \{x \in S : p\} \quad \{e : x \in S\}$$

They are all described in Section 1.2 (page 11) and Section 6.1 (page 65). Equality is also an operator of set theory, since it formally mean equality of sets.  $\text{TLA}^+$  provides the usual operators  $=$  and  $\neq$ .

The set construct  $\{x \in S : p\}$  can also be used with  $x$  a tuple of identifiers. For example,

$$\{\langle a, b \rangle \in \text{Nat} \times \text{Nat} : a > b\}$$

is the set of all pairs of natural numbers whose first component is greater than its second—pairs such as  $\langle 3, 1 \rangle$ . In the set construct  $\{e : x \in S\}$ , the clause  $x \in S$  can be generalized in exactly the same way as in a bounded quantifier such as  $\forall x \in S : p$ . For example,

$$\{\langle a, b, c \rangle : a, b \in \text{Nat}, c \in \text{Real}\}$$

is the set of all triples whose first two components are natural numbers and whose third component is a real number.

## Formal Semantics

$TLA^+$  is based on Zermelo-Fr ankel set theory, in which every value is a set. In set theory,  $\in$  is taken as a primitive, undefined operator. We could define all the other operators of set theory in terms of  $\in$ , using predicate logic and the CHOOSE operator. For example, set union could be defined by

$$S \cup T \triangleq \text{CHOOSE } U : \forall x : (x \in U) \equiv (x \in S) \vee (x \in T)$$

(To reason about  $\cup$ , we would need axioms from which we can deduce the existence of the chosen set  $U$ .) Another approach we could take is to let certain of the operators be primitives, and define the rest in terms of them. For example,  $\cup$  can be defined in terms of UNION and the construct  $\{e_1, \dots, e_n\}$  by:

$$S \cup T \triangleq \text{UNION } \{S, T\}$$

I do not try to distinguish a small set of primitive operators, and I treat  $\cup$  and UNION as equally primitive. Operators that I take to be primitive are defined them mathematically in terms of the rules that they satisfy. For example,  $S \cup T$  is defined by:

$$\forall x : (x \in (S \cup T)) \equiv (x \in S) \vee (x \in T)$$

However, there is no such defining rule for the primitive operator  $\in$ . I take only the simple forms of the constructs  $\{x \in S : p\}$  and  $\{e : x \in S\}$  as primitive, and I define the more general forms in terms of them.

$$e_1 = e_2 \triangleq \forall x : (x \in S) \equiv (x \in T).$$

$$e_1 \neq e_2 \triangleq \neg(e_1 = e_2).$$

$$e \notin S \triangleq \neg(e \in S).$$

$$S \cup T \text{ is defined by } \forall x : (x \in (S \cup T)) \equiv (x \in S) \vee (x \in T).$$

$$S \cap T \text{ is defined by } \forall x : (x \in (S \cap T)) \equiv (x \in S) \wedge (x \in T).$$

$$S \subseteq T \triangleq \forall x : (x \in S) \Rightarrow (x \in T).$$

$$S \setminus T \text{ is defined by } \forall x : (x \in (S \setminus T)) \equiv (x \in S) \wedge (x \notin T).$$

$$\text{SUBSET } S \text{ is defined by } \forall T : (T \in \text{SUBSET } S) \equiv (T \subseteq S)$$

$$\text{UNION } S \text{ is defined by } \forall x : (x \in \text{UNION } S) \equiv (\exists T \in S : x \in T).$$

$$\{e_1, \dots, e_n\} \triangleq \{e_1\} \cup \dots \cup \{e_n\},$$

where  $\{e\}$  is defined by:

$$\forall x : (x \in \{e\}) \equiv (x = e)$$

For  $n = 0$ , this construct is the empty set  $\{\}$ , defined by:

$$\forall x : x \notin \{\}$$

$\{x \in S : p\}$

where  $x$  is a bound identifier or a tuple of bound identifiers. The expression  $S$  is outside the scope of the bound identifier(s). For  $x$  an identifier, this is a primitive expression that is defined mathematically by

$$\forall y : (y \in \{x \in S : p\}) \equiv (y \in S) \wedge \hat{p}$$

where the identifier  $y$  does not occur in  $S$  or  $p$ , and  $\hat{p}$  is  $p$  with  $y$  substituted for  $x$ . For  $x$  a tuple, the expression is defined by

$$\begin{aligned} \{\langle x_1, \dots, x_n \rangle \in S : p\} &\triangleq \\ \{y \in S : (\exists x_1, \dots, x_n : (y = \langle x_1, \dots, x_n \rangle) \wedge p)\} \end{aligned}$$

where  $y$  is an identifier different from the  $x_i$  that does not occur in  $S$  or  $p$ . See Section 15.1.9 for further details about tuples.

$\{e : \mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n\}$

where each  $\mathbf{y}_i$  has the form  $x_1, \dots, x_k$  or  $\langle x_1, \dots, x_k \rangle$ , and each  $x_j$  is an identifier that is bound in the expression. The expressions  $S_i$  lie outside the scope of the bound identifiers. The simple form  $\{e : x \in S\}$ , for  $x$  an identifier, is taken to be primitive and is defined by:

$$\forall y : (y \in \{e : x \in S\}) \equiv (\exists x \in S : e = y)$$

The general form is defined inductively in terms of the simple form by:

$$\begin{aligned} \{e : \mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n\} &\triangleq \\ \text{UNION } \{ \{e : \mathbf{y}_1 \in S_1, \dots, \mathbf{y}_{n-1} \in S_{n-1}\} : \mathbf{y}_n \in S_n \} \\ \{e : x_1, \dots, x_n \in S\} &\triangleq \{e : x_1 \in S, \dots, x_n \in S\} \\ \{e : \langle x_1, \dots, x_n \rangle \in S\} &\triangleq \\ \{(\text{LET } z &\triangleq \text{CHOOSE } \langle x_1, \dots, x_n \rangle : y = \langle x_1, \dots, x_n \rangle \\ x_1 &\triangleq z[1] \\ &\vdots \\ x_n &\triangleq z[n] \text{ IN } e) : y \in S\} \end{aligned}$$

where the  $x_i$  are identifiers, and  $y$  and  $z$  are identifiers distinct from the  $x_i$  that do not occur in  $e$  or  $S$ . See section 15.1.9 for further details about tuples.

## 15.1.7 Functions

Functions are described in Section 5.2 (page 48); the difference between functions and operators is discussed in Section 6.4 (page 69). In  $\text{TLA}^+$ , we write  $f[v]$  for

the value of the function  $f$  applied to  $v$ . A function  $f$  has a domain  $\text{DOMAIN } f$ , and the value of  $f[v]$  is specified only if  $v$  is an element of  $\text{DOMAIN } f$ . We let  $[S \rightarrow T]$  denote the set of all functions  $f$  such that  $\text{DOMAIN } f = S$  and  $f[v] \in T$  for all  $v \in S$ .

Functions can be described explicitly with the construct

$$(15.5) \quad [x \in S \mapsto e]$$

This is the function  $f$  with domain  $S$  such that  $f[v]$  equals the value obtained by substituting  $v$  for  $x$  in  $e$ , for any  $v \in S$ . For example,

$$[n \in \text{Nat} \mapsto 1/n + 1]$$

is the function  $f$  with domain  $\text{Nat}$  such that  $f[0] = 1$ ,  $f[1] = 1/2$ ,  $f[2] = 1/3$ , etc. We can define an identifier  $fcn$  to equal the function (15.5) by writing

$$(15.6) \quad fcn[x \in S] \triangleq e$$

The identifier  $fcn$  can appear in the expression  $e$ , in which case this is a recursive function definition. Recursive function definitions were introduced in Section 5.5 (page 53) and discussed in Section 6.3 (page 67).

The **EXCEPT** construct describes a function that is “almost the same as” another function. For example,

$$(15.7) \quad [f \text{ EXCEPT } ![u] = a, ![v] = b]$$

is the function  $\hat{f}$  that is the same as  $f$ , except that  $\hat{f}[u] = a$  and  $\hat{f}[v] = b$ . More precisely, (15.7) equals

$$\begin{aligned} [x \in \text{DOMAIN } f \mapsto \text{IF } x = v \text{ THEN } b \\ \text{ELSE IF } x = u \text{ THEN } a \text{ ELSE } f[x]] \end{aligned}$$

Hence, if neither  $u$  nor  $v$  is in the domain of  $f$ , then (15.7) equals  $f$ . If  $u = v$ , then (15.7) equals  $[f \text{ EXCEPT } ![v] = b]$ .

An exception clause can have the general form  $![v_1] \cdots [v_n] = e$ . For example,

$$(15.8) \quad [f \text{ EXCEPT } ![u][v] = a]$$

is the function  $\tilde{f}$  that is the same as  $f$ , except that  $\tilde{f}[u][v]$  equals  $a$ . That is,  $\tilde{f}$  is the same as  $f$ , except that  $\tilde{f}[u]$  is the function that is the same as  $f[u]$ , except that  $\tilde{f}[u][v] = a$ . The symbol  $@$  occurring in an exception clause stands for the “original value”. For example, an  $@$  in the expression  $a$  of (15.8) denotes  $f[u][v]$ .

In  $TLA^+$ , a function of multiple arguments is one whose domain is a set of tuples; and  $f[v_1, \dots, v_n]$  is an abbreviation for  $f[\langle v_1, \dots, v_n \rangle]$ . The  $x \in S$  clause (15.5) and (15.6) can be generalized in the same way as in a bounded quantifier—for example here are two different ways of writing the same function:

$$[m, n \in \text{Nat}, r \in \text{Real} \mapsto e] \quad [\langle m, n, r \rangle \in \text{Nat} \times \text{Nat} \times \text{Real} \mapsto e]$$

This is a function whose domain is a set of triples. It is not the same as the function

$$[\langle m, n \rangle \in Nat, r \in Real \mapsto e]$$

whose domain is the set  $(Nat \times Nat) \times Real$  of pairs like  $\langle \langle 1, 3 \rangle, 1/3 \rangle$ , whose first element is a pair.

### Formal Semantics

Mathematicians traditionally define a function to be a set of pairs. In  $TLA^+$ , pairs (and all tuples) are functions. We take as primitives the constructs:

$$f[e] \quad \text{DOMAIN } f \quad [S \rightarrow T] \quad [x \in S \mapsto e]$$

where  $x$  is an identifier. These constructs are defined mathematically by the rules they satisfy. The other constructs, and the general forms of the construct  $[x \in S \mapsto e]$ , are defined in terms of them. These definitions use the operator  $IsAFcn$ , which is defined as follows so that  $IsAFcn(f)$  is true iff  $f$  is a function:

$$IsAFcn(f) \triangleq f = [x \in \text{DOMAIN } f \mapsto f[x]]$$

The first rule, which is not naturally associated with any one construct, is that two functions are equal iff they have the same domain and assign the same value to each element in their domain:

$$\begin{aligned} \forall f, g : IsAFcn(f) \wedge IsAFcn(g) \Rightarrow \\ ((f = g) \equiv \wedge \text{DOMAIN } f = \text{DOMAIN } g \\ \wedge \forall x \in \text{DOMAIN } f : f[x] = g[x]) \end{aligned}$$

The rest of the semantics of functions is given below. There is no separate defining rule for the  $\text{DOMAIN}$  operator.

$$f[e_1, \dots, e_n]$$

where the  $e_i$  are expressions. For  $n = 1$ , this is a primitive expression. For  $n > 1$ , it is defined by

$$f[e_1, \dots, e_n] = f[\langle e_1, \dots, e_n \rangle]$$

The tuple  $\langle e_1, \dots, e_n \rangle$  is defined in Section 15.1.9.

$$[\mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n \mapsto e]$$

where each  $\mathbf{y}_i$  has the form  $x_1, \dots, x_k$  or  $\langle x_1, \dots, x_k \rangle$ , and each  $x_j$  is an identifier that is bound in the expression. The expressions  $S_i$  lie outside the scope of the bound identifiers. The simple form  $[x \in S \mapsto e]$ , for  $x$  an identifier, is primitive and is defined by two rules:

$$(\text{DOMAIN } [x \in S \mapsto e]) = S$$

$$\forall y \in S : [x \in S \mapsto e][y] = \text{LET } x \triangleq y \text{ IN } e$$

where  $y$  is an identifier different from  $x$  that does not occur in  $S$  or  $e$ . The general form of the construct is defined inductively in terms of the simple form by:

$$\begin{aligned}
[x_1 \in S_1, \dots, x_n \in S_n \mapsto e] &\triangleq [\langle x_1, \dots, x_n \rangle \in S_1 \times \dots \times S_n \mapsto e] \\
[\dots, x_1, \dots, x_k \in S_i, \dots \mapsto e] &\triangleq [\dots, x_1 \in S_i, \dots, x_k \in S_i, \dots \mapsto e] \\
[\dots, \langle x_1, \dots, x_k \rangle \in S_i, \dots \mapsto e] &\triangleq \\
[\dots, y \in S_i, \dots \mapsto \text{LET } z &\triangleq \text{CHOOSE } \langle x_1, \dots, x_k \rangle : y = \langle x_1, \dots, x_k \rangle \\
x_1 &\triangleq z[1] \\
&\vdots \\
x_k &\triangleq z[k] \text{ IN } e]
\end{aligned}$$

where  $y$  and  $z$  are identifiers that do not appear anywhere in the original expression. See Section 15.1.9 for details about tuples.

$[S \rightarrow T]$  is defined by

$$\begin{aligned}
\forall f : f \in [S \rightarrow T] &\equiv \\
&IsAFcn(f) \wedge (S = \text{DOMAIN } f) \wedge (\forall x \in S : f[x] \in T)
\end{aligned}$$

where  $x$  and  $f$  do not occur in  $S$  or  $T$ , and  $IsAFcn$  is defined above.

$[f \text{ EXCEPT } !\mathbf{a}_1 = e_1, \dots, !\mathbf{a}_n = e_n]$

where each  $\mathbf{a}_i$  has the form  $[d_1] \dots [d_k]$  and each  $d_j$  is an expression. For the simple case when  $n = 1$  and  $\mathbf{a}_1$  is  $[d]$ , this is defined by<sup>3</sup>

$$\begin{aligned}
[f \text{ EXCEPT } ![d] = e] &\triangleq \\
[y \in \text{DOMAIN } f \mapsto \text{IF } y = d &\text{ THEN LET } @ \triangleq f[d] \text{ IN } e \\
&\text{ELSE } f[y]]
\end{aligned}$$

where  $y$  does not occur in  $f$ ,  $d$ , or  $e$ . The general form is defined inductively in terms of this simple case by:

$$\begin{aligned}
[f \text{ EXCEPT } !\mathbf{a}_1 = e_1, \dots, !\mathbf{a}_n = e_n] &\triangleq \\
[[f \text{ EXCEPT } !\mathbf{a}_1 = e_1, \dots, !\mathbf{a}_{n-1} = e_{n-1}] &\text{ EXCEPT } !\mathbf{a}_n = e_n] \\
[f \text{ EXCEPT } ![d_1] \dots [d_k] = e] &\triangleq \\
[f \text{ EXCEPT } ![d_1] = [ @ \text{ EXCEPT } ![d_2] \dots [d_k] = e]] &
\end{aligned}$$

$f[\mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n] \triangleq e$  is defined to be an abbreviation for:

$$f \triangleq \text{CHOOSE } f : f = [\mathbf{y}_1 \in S_1, \dots, \mathbf{y}_n \in S_n \mapsto e]$$

---

<sup>3</sup>Since  $@$  is not actually an identifier,  $\text{LET } @ \triangleq \dots$  isn't legal TLA<sup>+</sup> syntax. However, its meaning should be clear.

### 15.1.8 Records

TLA<sup>+</sup> borrows from programming languages the concept of a record. Records were introduced in Section 3.2 (page 28) and further explained in Section 5.2 (page 48). As in programming languages,  $r.h$  is the  $h$  component of record  $r$ . Records can be written explicitly as

$$[h_1 \mapsto e_1, \dots, h_n \mapsto e_n]$$

which equals the record with  $n$  components, whose  $h_i$  component equals  $e_i$ , for  $i = 1, \dots, n$ . The expression

$$[h_1 : S_1, \dots, h_n : S_n]$$

is the set of all such records with  $e_i \in S_i$ , for  $i = 1, \dots, n$ . These expressions are legal only if the  $h_i$  are all different. For example,  $[a : S, a : T]$  is illegal.

The EXCEPT construct, explained in Section 15.1.7 above, can be used for records as well as functions. For example,

$$[r \text{ EXCEPT } !.a = e]$$

is the record  $\hat{r}$  that is the same as  $r$ , except that  $\hat{r}.a = e$ . An exception clause can mix function application and record components. For example,

$$[f \text{ EXCEPT } ![v].a = e]$$

is the function  $\hat{f}$  that is the same as  $f$ , except that  $\hat{f}[v].a = e$ .

In TLA<sup>+</sup>, a record is a function whose domain is a finite set of strings, where  $r.h$  means  $r["h"]$ , for a record component  $h$ . Thus, the following two expressions describe the same record:

$$[fo \mapsto 7, ba \mapsto 8] \quad [x \in \{“fo”, “ba”\} \mapsto \text{IF } x = “fo” \text{ THEN } 7 \text{ ELSE } 8]$$

The name of a record component is syntactically an identifier. In the ASCII version, it is a string of letters, digits, and the underscore character ( $\_$ ) that contains at least one letter. Strings are described below in Section 15.1.10.

### Formal Semantics

The record constructs are defined in terms of function constructs.

$$e.h \triangleq e[“h”]$$

$$[h_1 \mapsto e_1, \dots, h_n \mapsto e_n] \triangleq [y \in \{“h_1”, \dots, “h_n”\} \mapsto \text{CASE } (y = “h_1”) \rightarrow e_1 \square \dots \square (y = “h_n”) \rightarrow e_n]$$

where  $y$  does not occur in any of the expressions  $e_i$ . The  $h_i$  must all be distinct.

$[h_1 : S_1, \dots, h_n : S_n] \triangleq \{[h_1 \mapsto y_1, \dots, h_n \mapsto y_n] : y_1 \in S_1, \dots, y_n \in S_n\}$   
 where the  $y_i$  do not occur in any of the expressions  $S_j$ . The  $h_i$  must all be distinct.

$[r \text{ EXCEPT } !\mathbf{a}_1 = e_1, \dots, !\mathbf{a}_n = e_n]$   
 where  $\mathbf{a}_i$  has the form  $b_1 \dots b_k$  and each  $b_j$  is either (i)  $[d]$ , where  $d$  is an expression, or (ii)  $.h$ , where  $h$  is a record component. It is defined to equal the corresponding function EXCEPT construct in which each  $.h$  is replaced by  $[“h”]$ .

### 15.1.9 Tuples

An  $n$ -tuple is written in TLA<sup>+</sup> as  $\langle e_1, \dots, e_n \rangle$ . As explained in Section 5.4. An  $n$ -tuple is defined to be a function whose domain is the set  $\{1, \dots, n\}$ , where  $\langle e_1, \dots, e_n \rangle[i] = e_i$ , for  $1 \leq i \leq n$ . The Cartesian product  $S_1 \times \dots \times S_n$  is the set of all  $n$ -tuples  $\langle e_1, \dots, e_n \rangle$  such that  $e_i \in S_i$ , for  $1 \leq i \leq n$ .

In TLA<sup>+</sup>,  $\times$  is not an associative operator. For example,

$$\begin{aligned} \langle 1, 2, 3 \rangle &\in Nat \times Nat \times Nat \\ \langle \langle 1, 2 \rangle, 3 \rangle &\in (Nat \times Nat) \times Nat \\ \langle 1, \langle 2, 3 \rangle \rangle &\in Nat \times (Nat \times Nat) \end{aligned}$$

and the tuples  $\langle 1, 2, 3 \rangle$ ,  $\langle \langle 1, 2 \rangle, 3 \rangle$ , and  $\langle 1, \langle 2, 3 \rangle \rangle$  are not equal. More precisely, the triple  $\langle 1, 2, 3 \rangle$  is unequal to either of the pairs  $\langle \langle 1, 2 \rangle, 3 \rangle$  or  $\langle 1, \langle 2, 3 \rangle \rangle$  because a triple and a pair have unequal domains. The semantics of TLA<sup>+</sup> does not specify if  $\langle 1, 2 \rangle$  equals 1 or if 3 equals  $\langle 2, 3 \rangle$ , so we don't know whether or not  $\langle \langle 1, 2 \rangle, 3 \rangle$  and  $\langle 1, \langle 2, 3 \rangle \rangle$  are equal.

The 1-tuple  $\langle e \rangle$  is different from  $e$ . That is, the semantics does not specify whether or not they are equal. There is no special notation for writing a set of 1-tuples. The easiest way to denote the set of all 1-tuples  $\langle e \rangle$  with  $e \in S$  is  $\{\langle e \rangle : e \in S\}$ .

In the standard *Sequences* module, described in Section 17.1 (page 323), an  $n$ -element sequence is represented as an  $n$ -tuple. The module defines several useful operators on sequences/tuples.

### Formal Semantics

Tuples and Cartesian products are defined in terms of functions (defined in Section 15.1.7) and the set  $Nat$  of natural numbers (defined in Section 15.1.11).

$$\langle e_1, \dots, e_n \rangle \triangleq [i \in \{j \in Nat : (0 < j) \wedge (j \leq n)\} \mapsto e_i]$$

where  $i$  does not occur in any of the expressions  $e_j$ .

$$S_1 \times \dots \times S_n \triangleq \{\langle y_1, \dots, y_n \rangle : y_1 \in S_1, \dots, y_n \in S_n\}$$

where the identifiers  $y_i$  do not occur in any of the expressions  $S_j$ .

### 15.1.10 Strings

$\text{TLA}^+$  defines a string to be a tuple of characters. (Tuples are defined in Section 15.1.9 above.) Thus, “abc” equals

$$\langle \text{“abc”}[1], \text{“abc”}[2], \text{“abc”}[3] \rangle$$

The semantics of  $\text{TLA}^+$  does not specify what a character is. However, it does specify that different characters (those having different computer representations) are different. Thus “a”[1], “b”[1], and “A”[1] (the characters *a*, *b*, and *A*) are all different. The built-in operator `STRING` is defined to be the set of all strings.

Although  $\text{TLA}^+$  doesn’t specify what a character is, it’s easy to define operators that assign values to characters. For example, here’s the definition of an operator *Ascii* that assigns to every lower-case letter its ASCII representation.<sup>4</sup>

$$\begin{aligned} \text{Ascii}(\text{char}) &\triangleq 96 + \text{CHOOSE } i \in 1 \dots 26 : \\ &\quad \text{“abcdefghijklmnopqrstuvwxyz”}[i] = \text{char} \end{aligned}$$

This defines *Ascii*(“a”[1]) to equal 97, the ASCII code for the letter *a*, and *Ascii*(“z”[1]) to equal 122, the ASCII code for *z*. Section 14.1.2 on page 260 illustrates how a specification can make use of the fact that strings are tuples.

Exactly what characters may appear in a string is system-dependent. A Japanese version of  $\text{TLA}^+$  might not allow the character *a*. The standard ASCII version contains the following characters:

```
a b c d e f g h i j k l m n o p q r s t u v w x y z
A B C D E F G H I J K L M N O P Q R S T U V W X Y Z
0 1 2 3 4 5 6 7 8 9
~ @ # $ % ^ & * _ - + = ( ) { } [ ] < > | / \ , . ? ; : ' "
<HT> (tab) <LF> (line feed) <FF> (form feed) <CR> (carriage return)
```

plus the space character. Since strings are delimited by a double-quote ("), some convention is needed for typing a string that contains a double-quote. Conventions are also needed to type characters like <LF> within a string. In the ASCII version of  $\text{TLA}^+$ , the following pairs of characters, beginning with a \ character, are used to represent these special characters:

```
\ "  "      \t  <HT>      \f  <FF>
\\  \       \n  <LF>      \r  <CR>
```

With this convention, "a\\\"b\\\"" represents the string consisting of the following five characters: *a* \ " *b* ". In the ASCII version of  $\text{TLA}^+$ , a \ character can appear in a string expression only as the first character of one of these six two-character sequences.

<sup>4</sup>This clever way of using `CHOOSE` to map from characters to numbers was pointed out to me by Georges Gonthier.

## Formal Semantics

We assume a set *Char* of characters, which may depend on the version of  $TLA^+$ . (The identifier *Char* is not a pre-defined symbol of  $TLA^+$ .)

$STRING \triangleq Seq(Char)$   
 where *Seq* is the operator defined in the *Sequences* module of Section 17.1  
 so that *Seq*(*S*) is the set of all finite sequences of elements of *S*.

" $c_1 \dots c_n$ "  $\triangleq \langle c_1, \dots, c_n \rangle$   
 where each  $c_i$  is the (system-dependent) representation of a character in *Char*.

### 15.1.11 Numbers

$TLA^+$  defines a sequence of digits like 63 to be the usual natural number—that is, 63 equals  $6 * 10^2 + 3$ .  $TLA^+$  also allows the binary representation `\b111111`, the octal representation `\o77`, and the hexadecimal representation `\h3F` of that number. (Case is ignored in the prefixes and in the hexadecimal representation, so `\H3F` and `\h3f` are equivalent to `\h3F`.) Decimal numbers are also pre-defined in  $TLA^+$ ; for example, 3.14159 equals  $314159/10^5$ .

Numbers are pre-defined in  $TLA^+$ , so 63 is defined even in a module that does not extend or instantiate one of the standard numbers modules. However, sets of numbers like *Nat* and arithmetic operators like + are not. You can write a module that defines + any way you want, in which case  $40 + 23$  need not equal 63. Of course,  $40 + 23$  does equal 63 for + defined by the standard numbers modules *Naturals*, *Integers*, and *Reals*, which are described in Section 17.4.

## Formal Semantics

The set *Nat* of natural numbers, along with its zero element *Zero* and successor function *Succ*, is defined in module *Peano* on page 329. The meaning of a representation of a natural number is defined in the usual manner:

$$0 \triangleq Zero \quad 1 \triangleq Succ(Zero) \quad 2 \triangleq Succ(Succ(Zero)) \quad \dots$$

The module *ProtoReals* module on pages 330–331 defines the set *Real* of real numbers to be a superset of the set *Nat*, and defines the usual arithmetic operators on real numbers. The meaning of a decimal number is defined in terms of these operators by:

$$c_1 \dots c_m . d_1 \dots d_n \triangleq c_1 \dots c_m d_1 \dots d_n / 10^n$$

## 15.2 Nonconstant Operators

The nonconstant operators are what distinguish  $\text{TLA}^+$  from ordinary mathematics. There are two classes of nonconstant operators: action operators, listed in Figure 13.11 on page 249, and temporal operators, listed in Figure 13.12 on page 249.

Section 15.1 above talks about the meanings of the built-in constant operators of  $\text{TLA}^+$ , without considering their arguments. We can do that for constant operators, since the meaning of  $\subseteq$  in the expression  $e_1 \subseteq e_2$  doesn't depend on whether or not the expressions  $e_1$  and  $e_2$  contain variables or primes. To understand the nonconstant operators, we need to consider their arguments. Thus, we can no longer talk about the meanings of the operators in isolation; we must describe the meanings of expressions built from those operators.

A *basic* expression is one that contains built-in  $\text{TLA}^+$  operators, declared constants, and declared variables. We now describe the meaning of all basic  $\text{TLA}^+$  expressions, including ones that contain nonconstant built-in operators. We start by considering basic constant expressions, ones containing only the constant operators we have already studied and declared constants.

### 15.2.1 The Meaning of a Basic Constant Expression

Section 15.1 above defines the meanings of the constant operators. This in turn defines the meaning of any expression built from these operators and declared constants. For example, if  $S$  and  $T$  are declared by

CONSTANTS  $S, T(-)$

then  $\exists x : S \subseteq T(x)$  is a formula that equals TRUE if there is some value  $v$  such that every element of  $S$  is an element of  $T(v)$ , and otherwise equals FALSE. Whether  $\exists x : S \subseteq T(x)$  equals TRUE or FALSE depends on what actual values we assign to  $S$  and to  $T(v)$ , for all  $v$ ; so that's as far as we can go in assigning a meaning to the expression.

There are some basic constant expressions that are true regardless of the values we assign to their declared constants—for example, the expression

$$(S \subseteq T) \equiv (S \cap T = S)$$

Such an expression is said to be *valid*.

### Formal Semantics

Section 15.1 defines all the built-in constant operators in terms of a subset of them called the primitive operators. These definitions can be formulated as an

inductive set of rules that define the meaning  $\llbracket c \rrbracket$  of any basic constant expression  $c$ . For example, from the definition

$$e \notin S \triangleq \neg(e \in S)$$

we get the rule

$$\llbracket e \notin S \rrbracket = \neg(\llbracket e \rrbracket \in \llbracket S \rrbracket)$$

These rules define the meaning of a basic constant expression to be an expression containing only primitive constant operators and declared constants.

If  $S$  and  $T$  are constants declared as above, then the meaning  $\llbracket \exists x : S \subseteq T(x) \rrbracket$  of the expression  $\exists x : S \subseteq T(x)$  is the expression itself. Logicians usually carry things further, assigning some meanings  $\llbracket S \rrbracket$  and  $\llbracket T \rrbracket$  to declared constants and defining  $\llbracket \exists x : S \subseteq T(x) \rrbracket$  to equal  $\exists x : \llbracket S \rrbracket \subseteq \llbracket T \rrbracket(x)$ . For simplicity, I have short-circuited that extra level of meaning.

I am taking as given the meaning of an expression containing only primitive constant operators and declared constants. In particular, I take as primitive the notion of validity for such expressions. Section 15.1 defines the meaning of any basic constant expression in terms of these expressions, so it defines what it means for a basic constant expression to be valid.

### 15.2.2 The Meaning of a State Function

A *state* is an assignment of values to variables. (In ZF set theory, on which the semantics of  $TLA^+$  is based, *value* is just another term for *set*.) States were discussed in Sections 2.1 and 2.3.

A *state function* is an expression that is built from declared variables, declared constants, and constant operators. (State functions can also contain ENABLED expressions, which are described below.) State functions are discussed on page 25 of Section 3.1. A state function assigns a constant expression to every state. If state function  $e$  assigns to state  $s$  the constant  $v$ , then we say that  $v$  is the value of  $e$  in state  $s$ . For example, if  $x$  is a declared variable,  $T$  is a declared constant, and  $s$  is a state that assigns to  $x$  the value 42; then the value of  $x \in T$  in state  $s$  is the constant expression  $42 \in T$ . A state function is *valid* iff it has the value TRUE in every state.

#### Formal Semantics

A state is an assignment of values to variables. Formally, a state  $s$  is a function whose domain is the set of all variable names, where  $s[“x”]$  is the value that  $s$  assigns to variable  $x$ . We write  $s\llbracket x \rrbracket$  instead of  $s[“x”]$ .

A *basic state function* is an expression that is built from declared variables, declared constants, constant operators, and ENABLED expressions, which are

expressions of the form `ENABLED  $e$` . An `ENABLED`-free basic state function is, obviously, one with no `ENABLED` expressions. The meaning of a basic state function is a mapping from states to values. We let  $s[e]$  be the value that state function  $e$  assigns to a state  $s$ . Since a variable is a state function, we thus say both that state  $s$  assigns  $s[x]$  to variable  $x$ , and that the state function  $x$  assigns  $s[x]$  to state  $s$ . A basic state function  $e$  is defined to be valid iff  $s[e]$  is valid for all states  $s$ .

Using the meanings assigned to the constant operators in Section 15.1 above, we inductively define  $s[e]$  for any `ENABLED`-free state function  $e$  to be an expression built from the primitive  $\text{TLA}^+$  constant operators, declared constants, and the values assigned by  $s$  to each variable. For example, if  $x$  is a variable and  $S$  is a constant, then

$$s[x \notin S] = \neg(s[x] \in S)$$

It is easy to see that  $s[c]$  to equal  $\llbracket c \rrbracket$ , for any constant expression  $c$ . (This expresses formally that a constant has the same value in all states.)

To define the meaning of all basic state function, not just `ENABLED`-free ones, we must define the meaning of an `ENABLED` expression. This is done below.

I described the meaning of a state function as a “mapping” on states. This mapping cannot be a function, because there is no set of all states. Since for any set  $S$  there is a state that assigns the value  $S$  to each variable, there are too many states to form a set. (See the discussion of Russell’s paradox on page 66.) To be strictly formal, I should explain what I’m doing as follows. I define an operator  $M$  such that, if  $s$  is a state and  $e$  is a syntactically correct basic state function, then  $M(s, e)$ , which I write  $s[e]$ , is the basic constant expression that is the meaning of  $e$  in state  $s$ .

Actually, this description of the semantics isn’t right either. A state is a mapping from variables to values (sets), not to constant expressions. Since there are an uncountable number of sets and only a countable number of finite sequences of strings, there are values that can’t be described by any expression. Suppose  $\xi$  is such a value, and let  $s$  be a state that assigns the value  $\xi$  to the variable  $x$ . Then  $s[x = \{\}]$  equals  $\xi = \{\}$ , which isn’t a constant expression because  $\xi$  isn’t an expression. So, to be really formal, I would have to define a semantic constant expression to be one made from primitive constant operators, declared constants, and arbitrary values. The meaning of a basic state function is a mapping from states to semantic constant expressions.

I won’t bother with these details, which are needed for a really formal definition of the semantics. I will instead define a semi-formal semantics for basic expression that I hope is easier to understand. Mathematically sophisticated readers who understand the less formal exposition should be able to fill in the missing formal details.

### 15.2.3 Action Operators

A *transition function* is an expression built from state functions using the priming operator ( $'$ ) and the other action operators of  $TLA^+$  listed in Figure 13.11 on page 249. A transition function assigns a value to every step, where a step is a pair of states. In a transition function, an unprimed occurrence of a variable  $x$  represents the value of  $x$  in the first (old) state, and a primed occurrence of  $x$  represents its value in the second (new) state. For example, if state  $s$  assigns the value 4 to  $x$  and state  $t$  assigns the value 5 to  $x$ , then transition function  $x' - x$  assigns to the step  $s \rightarrow t$  the value  $5 - 4$ , which of course equals 1.

An action is a Boolean-valued transition function, such as  $x' > x$ . We say that action  $A$  is true on step  $s \rightarrow t$ , or that  $s \rightarrow t$  is an  $A$  step, iff  $A$  assigns the value TRUE to  $s \rightarrow t$ . An action is said to be *valid* iff it is true on any step.

The action operators of  $TLA^+$  other than  $'$  have the following meanings, where  $A$  and  $B$  are actions.

$[A]_e$  equals  $A \vee (e' = e)$ .

$\langle A \rangle_e$  equals  $A \wedge (e' \neq e)$ .

ENABLED  $A$  is the state function that is true in state  $s$  iff there is some state  $t$  such that  $s \rightarrow t$  is an  $A$  step. TRUE.

UNCHANGED  $e$  equals  $e' = e$ , for any state function  $e$ .

$A \cdot B$  is the action that is true on step  $s \rightarrow t$  iff there is a state  $u$  such that  $s \rightarrow u$  is an  $A$  step and  $u \rightarrow t$  is a  $B$  step.

Priming and the construct  $[A]_v$  are introduced in Section 2.2 (page 15); the UNCHANGED operator is introduced on page 26 of Section 3.1 (24); ENABLED is introduced on page 96 of Section 8.3; the construct  $\langle A \rangle_v$  is defined on 91 of Section 8.1; and the action-composition operator “ $\cdot$ ” is introduced in Section 7.3 (page 76).

### Formal Semantics

A basic transition function is a basic expression that does not contain any temporal operators. The meaning of a basic transition function  $e$  is an assignment of a basic constant expression  $\langle s, t \rangle \llbracket e \rrbracket$  to any pair of states  $\langle s, t \rangle$ . (I use here the more conventional notation  $\langle s, t \rangle$  instead of  $s \rightarrow t$ .) A transition function is valid iff  $\langle s, t \rangle \llbracket e \rrbracket$  is valid, for all states  $s$  and  $t$ .

If  $e$  is a basic state function, then we interpret  $e$  as a basic transition function by defining  $\langle s, t \rangle \llbracket e \rrbracket$  to equal  $s \llbracket e \rrbracket$ . As indicated above, UNCHANGED and the constructs  $[A]_e$  and  $\langle A \rangle_e$  are defined in terms of priming. To define the meanings of the remaining action operators, we first define existential quantification over all states. Let  $IsAState$  be an operator such that  $IsAState(s)$  is true iff  $s$  is a

state—that is, a function whose domain is the set of all variable names. (It’s easy to define *IsAState* using the operator *IsAFcn*, defined on page 287.) Existential quantification over all states is then defined by

$$\exists_{\text{state}} s : p \triangleq \exists s : \text{IsAState}(s) \wedge p$$

for any formula  $p$ . The meanings of all transition functions and all state functions (including *ENABLED* expressions) is then defined inductively by the definitions already given and the following definitions of the remaining action operators:

$e'$  is the transition function defined by  $\langle s, t \rangle \llbracket e' \rrbracket \triangleq t \llbracket e \rrbracket$  for any state function  $e$ .

*ENABLED A* is the state function defined by

$$s \llbracket \text{ENABLED } A \rrbracket \triangleq \exists_{\text{state}} t : \langle s, t \rangle \llbracket A \rrbracket$$

for any transition function  $A$ .

$A \cdot B$  is the transition function defined by

$$\langle s, t \rangle \llbracket A \cdot B \rrbracket \triangleq \exists_{\text{state}} u : \langle s, u \rangle \llbracket A \rrbracket \wedge \langle u, t \rangle \llbracket B \rrbracket$$

for any transition functions  $A$  and  $B$ .

The formal semantics talks about transition functions, not actions. Since  $\text{TLA}^+$  is typeless, there is no formal distinction between an action and an arbitrary transition function. We could define an action  $A$  to be a transition function such that  $\langle s, t \rangle \llbracket A \rrbracket$  is a Boolean for all states  $s$  and  $t$ . However, what we usually mean by an action is a transition function  $A$  such that  $\langle s, t \rangle \llbracket A \rrbracket$  is a Boolean whenever  $s$  and  $t$  are reachable states of some specification. For example, consider a specification with a variable  $b$  of type *BOOLEAN* might contain an “action”  $b \wedge (y' = y)$ . We can calculate the meaning of *ENABLED* ( $b \wedge (y' = y)$ ) as follows:

Types are explained on page 25.

$$\begin{aligned} s \llbracket \text{ENABLED } (b \wedge (y' = y)) \rrbracket &= \exists_{\text{state}} t : \langle s, t \rangle \llbracket b \wedge (y' = y) \rrbracket \\ &= \exists_{\text{state}} t : \langle s, t \rangle \llbracket b \rrbracket \wedge (\langle s, t \rangle \llbracket y' \rrbracket = \langle s, t \rangle \llbracket y \rrbracket) \\ &= \exists_{\text{state}} t : s \llbracket b \rrbracket \wedge (t \llbracket y \rrbracket = s \llbracket y \rrbracket) \end{aligned}$$

By definition of *ENABLED*.

By definition of  $\wedge$  and  $=$ .

By definition of  $'$ , since  $\langle s, t \rangle \llbracket e \rrbracket = e$  for any state function  $e$ .

If  $s \llbracket b \rrbracket$  is a Boolean, we can use the rules of ordinary logic to simplify the last expression and obtain:

$$s \llbracket \text{ENABLED } (b \wedge (y' = y)) \rrbracket = s \llbracket b \rrbracket$$

However, if  $s$  is a state that assigns the value 2 to the variable  $b$  and the value  $-7$  to the variable  $y$ , then

$$s \llbracket \text{ENABLED } (b \wedge (y' = y)) \rrbracket = \exists_{\text{state}} t : 2 \wedge (t \llbracket y \rrbracket = -7)$$

The last expression may or may not equal 2. (See Section 15.1.3 on page 280.) If the specification we are writing makes sense, it can depend on the meaning of  $ENABLED(b \wedge (y' = y))$  only for states in which the value of  $b$  is a Boolean. We don't care about its value on a state that assigns to  $b$  the value 2, just as we don't care about the value of  $3/x$  in a state that assigns the value “abc” to  $x$ . See the discussion of silly expressions in Section 6.2 (page 67).

### 15.2.4 Temporal Operators

As explained in Section 8.1, a temporal formula  $F$  is true or false on a behavior, where a behavior is a sequence of states. The syntax of temporal formulas is defined inductively to be all formulas having one of the forms shown in Figure 13.12 on page 249, where  $e$  is a state function,  $A$  is an action, and  $F$  and  $G$  are temporal formulas. All these temporal operators are explained in Chapter 8—except for  $\dot{\vdash}$ , which is explained in Chapter 10.7.

The formula  $\Box F$  is true for a behavior  $\sigma$  iff the temporal formula  $F$  is true for  $\sigma$  and all suffixes of  $\sigma$ . To define the constructs  $\Box[A]_e$  and  $\Diamond\langle A \rangle_e$ , we regard an action  $B$  to be a temporal formula that is true of a behavior  $\sigma$  iff the first two states of  $\sigma$  form a  $B$  step. Thus,  $\Box[A]_e$  is true of  $\sigma$  iff every successive pair of states of  $\sigma$  is a  $[A]_e$  step. All the other temporal operators of  $TLA^+$ , except  $\exists$ ,  $\forall$ , and  $\dot{\vdash}$ , are defined as follows in terms of  $\Box$ :

$$\begin{aligned} \Diamond F &\triangleq \neg \Box \neg F \\ WF_e(A) &\triangleq \Box \Diamond \neg(ENABLED \langle A \rangle_e) \vee \Box \Diamond \langle A \rangle_e \\ SF_e(A) &\triangleq \Diamond \Box \neg(ENABLED \langle A \rangle_e) \vee \Box \Diamond \langle A \rangle_e \\ F \leadsto G &\triangleq \Box(F \Rightarrow \Diamond G) \neg F \end{aligned}$$

The temporal existential quantifier  $\exists$  is a hiding operator,  $\exists x : F$  meaning formula  $F$  with the variable  $x$  hidden. To define this more precisely, we first define  $\natural\sigma$  to be the (possibly finite) sequence of states obtained by removing from  $\sigma$  all stuttering steps—that is, by removing any state that is the same as the previous one. We then define  $\sigma \sim_x \tau$  to be true iff  $\natural\sigma$  and  $\natural\tau$  are the same except for the values that their states assign to the variable  $x$ . Thus,  $\sigma \sim_x \tau$  is true iff  $\sigma$  can be obtained from  $\tau$  (or vice-versa) by adding and/or removing stuttering steps and changing the values assigned to  $x$  by its states. Finally,  $\exists x : F$  is defined to be true for a behavior  $\sigma$  iff  $F$  is true for some behavior  $\tau$  such that  $\sigma \sim_x \tau$ .

The temporal universal quantifier  $\forall$  is defined in terms of  $\exists$  by

$$\forall x : F \triangleq \neg(\exists x : \neg F)$$

The formula  $F \dot{\vdash} G$  asserts that  $G$  does not become false before  $F$  does. More precisely, we define a formula  $H$  to be true for a finite prefix  $\rho$  of a behavior  $\sigma$  iff  $H$  is true for some (infinite) behavior that extends  $\rho$ . (In particular,  $H$  is

true of the empty prefix iff  $H$  satisfies some behavior.) Then  $F \stackrel{\pm}{\triangleright} G$  is defined to be true for a behavior  $\sigma$  iff (i)  $F \Rightarrow G$  is true for  $\sigma$  and (ii) for every finite prefix  $\rho$  of  $\sigma$ , if  $F$  is true for  $\rho$  then  $G$  is true for the prefix of  $\sigma$  that is one state longer than  $\rho$ .

## Formal Semantics

Formally, a behavior is a function from the set  $Nat$  of natural numbers to states. (We think of a behavior  $\sigma$  as the sequence  $\sigma[0], \sigma[1], \dots$  of states.) The meaning of a temporal formula is a predicate on behaviors—that is, a mapping from behaviors to Booleans. We write  $\sigma \models F$  for the value that the meaning of  $F$  assigns to the behavior  $\sigma$ . The temporal formula  $F$  is valid iff  $\sigma \models F$  is valid, for all behaviors  $\sigma$ .

Instead of writing  $\sigma_i$  as in Chapter 8, I use the standard functional notation  $\sigma[i]$ .

Above, we have defined all the other temporal operators in terms of  $\Box$ ,  $\exists$ , and  $\stackrel{\pm}{\triangleright}$ . Formally, since an action is not a temporal formula, the construct  $\Box[A]_e$  is not an instance of the temporal operator  $\Box$ , so its meaning should be defined separately. The construct  $\Diamond\langle A \rangle_e$ , which is similarly not an instance of  $\Diamond$ , is then defined to equal  $\neg\Box[\neg A]_e$ .

To define the meaning of  $\Box$ , we first define  $\sigma^{+n}$  to be the behavior obtained by deleting the first  $n$  states of  $\sigma$ :

$$\sigma^{+n} \triangleq [i \in Nat \mapsto \sigma[i + n]]$$

We then define the meaning of  $\Box$  as follows, for any temporal formula  $F$ , transition function  $A$  and state function  $e$ :

$$\begin{aligned} \sigma \models \Box F &\triangleq \forall n \in Nat : \sigma^{+n} \models F \\ \sigma \models \Box[A]_e &\triangleq \forall n \in Nat : \langle \sigma[n], \sigma[n+1] \rangle \llbracket [A]_e \rrbracket \end{aligned}$$

To formalize the definition of  $\exists$  given above, we first define  $\natural$  by:

$$\begin{aligned} \natural\sigma &\triangleq \text{LET } f[n \in Nat] \triangleq \begin{array}{l} \text{IF } n = 0 \\ \text{THEN } 0 \\ \text{ELSE IF } \sigma[n] = \sigma[n-1] \text{ THEN } f[n-1] \\ \text{ELSE } f[n-1] + 1 \end{array} \\ S &\triangleq \{f[n] : n \in Nat\} \\ \text{IN } [n \in S \mapsto \sigma[\text{CHOOSE } i \in Nat : f[i] = n]] \end{aligned}$$

Next, let  $s_{x \leftarrow v}$  be the state that is the same as state  $s$  except that it assigns to the variable  $x$  the value  $v$ . We then define  $\sim_x$  by:

$$\sigma \sim_x \tau \triangleq \natural\sigma = [n \in \text{DOMAIN } \natural\tau \mapsto \tau_{x \leftarrow \sigma[n]}[x]]$$

We next define existential quantification over behaviors. This is done much as we defined quantification over states on pages 296–297 above; we first define

*IsABehavior* so that *IsABehavior*( $\sigma$ ) is true iff  $\sigma$  is a behavior, and we then define:

$$\exists_{\text{behavior}} \sigma : F \triangleq \exists \sigma : \text{IsABehavior}(\sigma) \wedge F$$

We can now define the meaning of  $\exists$  by:

$$\sigma \models \exists x : F \triangleq \exists_{\text{behavior}} \tau : (\sigma \sim_x \tau) \wedge (\tau \models F)$$

Finally, we define the meaning of  $\pm\rhd$  as follows:

$$\begin{aligned} \sigma \models F \pm\rhd G &\triangleq \\ \text{LET } \text{PrefixSat}(n, H) &\triangleq \exists_{\text{behavior}} \tau : \wedge \forall i \in 0 \dots (n-1) : \tau[i] = \sigma[i] \\ &\quad \wedge \tau \models H \\ \text{IN } \wedge \sigma \models F &\Rightarrow G \\ \wedge \forall n \in \text{Nat} : &\text{PrefixSat}(n, F) \Rightarrow \text{PrefixSat}(n+1, G) \end{aligned}$$

## Chapter 16

# The Meaning of a Module

Chapter 15 defines the meaning of the built-in  $\text{TLA}^+$  operators. In doing so, it defines the meaning of a basic expression—that is, of an expression containing only built-in operators, declared constants, and declared variables. We now define the meaning of a module in terms of basic expressions. Since a  $\text{TLA}^+$  specification consists of a collection of modules, this defines the semantics of  $\text{TLA}^+$ .

We also complete the definition of the syntax of  $\text{TLA}^+$  by giving the remaining context-dependent syntactic conditions not described in Chapter 14. Here's a list of some illegal expressions that satisfy the grammar of Chapter 14, and where in this chapter you can find the conditions that make them illegal.

- $F(x)$ , if  $F$  is defined by  $F(x, y) \triangleq x + y$  (Section 16.1)
- $(x' + 1)'$  (Section 16.2)
- $x + 1$ , if  $x$  is not defined or declared (Section 16.3)
- $F \triangleq 0$ , if  $F$  is already defined (Section 16.5)

This chapter is meant to be read in its entirety. To try to make it as readable as possible, I have made the exposition somewhat informal. Wherever I could, I have used examples in place of formal definitions. The examples assume that you understand the approximate meanings of the  $\text{TLA}^+$  constructs, as explained in Part I. I hope that mathematically sophisticated readers will see how to fill in the missing formalism.

## 16.1 Operators and Expressions

Because it uses conventional mathematical notation,  $\text{TLA}^+$  has a rather rich syntax, with several different ways of expressing the same basic type of math-

ematical operation. For example, the following expressions are all formed by applying an operator to a single argument  $e$ :

$$\text{Len}(e) \quad - e \quad \{e\} \quad e'$$

This section develops a uniform way of writing all these expressions, as well as more general kinds of expressions.

### 16.1.1 The Order and Arity of an Operator

An operator has an *arity* and an *order*. An operator's arity of an operator describes the number and order of its arguments. It's the arity of the *Len* operator that tells us that  $\text{Len}(s)$  is a legal expression, while  $\text{Len}(s, t)$  and  $\text{Len}(+)$  are not. All the operators of  $\text{TLA}^+$ , whether built-in or defined, fall into three classes: 0th-, 1st-, and 2nd-order operators.<sup>1</sup> Here is how these classes, and their arities, are defined:

*Len* is defined in the *Sequences* module on page 325.

0.  $E \triangleq x' + y$  defines  $E$  to be the 0th-order operator  $x' + y$ . A 0th-order operator takes no arguments, so it is an ordinary expression. We represent the arity of such an operator by the symbol  $\_$  (underscore).
1.  $F(x, y) \triangleq x \cup \{z, y\}$  defines  $F$  to be a 1st-order operator. For any expressions  $e_1$  and  $e_2$ , it defines  $F(e_1, e_2)$  to be an expression. We represent the arity of  $F$  by  $\langle \_, \_ \rangle$ .  
In general, a 1st-order operator takes expressions as arguments. Its arity is the tuple  $\langle \_, \dots, \_ \rangle$ , with one  $\_$  for each argument.
2.  $G(f(\_, \_), x, y) \triangleq f(x, \{x, y\})$  defines  $G$  to be a 2nd-order operator. The operator  $G$  takes three arguments: its first argument is a 1st-order operator that takes two arguments; its last two arguments are expressions. For any operator  $Op$  of arity  $\langle \_, \_ \rangle$ , and any expressions  $e_1$  and  $e_2$ , this defines  $G(Op, e_1, e_2)$  to be an expression. We say that  $G$  has arity  $\langle \langle \_, \_ \rangle, \_, \_ \rangle$ .

In general, the arguments of a 2nd-order operator may be expressions or 1st-order operators. A 2nd-order operator has an arity of the form  $\langle a_1, \dots, a_n \rangle$ , where each  $a_i$  is either  $\_$  or  $\langle \_, \dots, \_ \rangle$ . (We can consider a 1st-order operator to be a degenerate case of a 2nd-order operator.)

It would be easy enough to define 3rd- and higher-order operators.  $\text{TLA}^+$  does not permit them because they are of little use and would make it harder to check level-correctness, which is discussed in Section 16.2 below.

<sup>1</sup>Even though it allows 2nd-order operators,  $\text{TLA}^+$  is still what logicians call a first-order logic because it permits quantification only over 0th-order operators. A higher-order logic would allow us to write the formula  $\exists x(\_) : \text{exp}$ .

### 16.1.2 $\lambda$ Expressions

When we define a 0th-order operator  $E$  by  $E \triangleq exp$ , we can write what the operator  $E$  equals—it equals the expression  $exp$ . We can explain the meaning of this definition by saying that it assigns the value  $exp$  to the symbol  $E$ . To explain the meaning of an arbitrary  $TLA^+$  definition, we need to be able to write what a 1st- or 2nd-order operator equals—for example, the operator  $F$  defined by:

$$F(x, y) \triangleq x \cup \{z, y\}$$

$F$  equals an operator, and  $TLA^+$  provides no way to write anything but 0th-order operators (ordinary expressions). We therefore generalize expressions to  $\lambda$  expressions, and we write the operator that  $F$  equals as the  $\lambda$  expression:

$$\lambda x, y : x \cup \{z, y\}$$

The symbols  $x$  and  $y$  in this  $\lambda$  expression are called  $\lambda$  parameters. We use  $\lambda$  expressions only to explain the meaning of  $TLA^+$  specifications; we can't write a  $\lambda$  expression in  $TLA^+$ .

We also allow 2nd-order  $\lambda$  expressions, where the operator  $G$  defined by

$$G(f(-, -), x, y) \triangleq f(y, \{x, z\})$$

is equal to the  $\lambda$  expression

$$(16.1) \lambda f(-, -), x, y : f(y, \{x, z\})$$

The general form of a  $\lambda$  expression is  $\lambda p_1, \dots, p_n : exp$ , where  $exp$  is a  $\lambda$  expression, each parameter  $p_i$  is either an identifier  $id_i$  or has the form  $id_i(- \dots, -)$ , and the  $id_i$  are all distinct. We call  $id_i$  the *identifier* of the  $\lambda$  parameter  $p_i$ . We consider the  $n = 0$  case, the  $\lambda$  expression  $\lambda : exp$  with no parameters, to be the expression  $exp$ . This makes a  $\lambda$  expression a generalization of an ordinary expressions.

A  $\lambda$  parameter identifier is a bound identifier, just like the identifier  $x$  in  $\forall x : F$ . As with any bound identifiers, renaming the  $\lambda$  parameter identifiers in a  $\lambda$  expression doesn't change the meaning of the expression. For example, (16.1) is equivalent to

$$\lambda abc(-, -), qq, m : abc(m, \{qq, z\})$$

For obscure historical reasons, this kind of renaming is called  $\alpha$  conversion.

If  $Op$  is the  $\lambda$  expression  $\lambda p_1, \dots, p_n : exp$ , then  $Op(e_1, \dots, e_n)$  equals the result of replacing the identifier of the  $\lambda$  parameter  $p_i$  in  $exp$  with  $e_i$ , for all  $i$  in  $1 \dots n$ . For example,

$$(\lambda x, y : x \cup \{z, y\})(TT, w + z) = TT \cup \{z, (w + z)\}$$

This procedure for evaluating the application of a  $\lambda$  expression is called  $\beta$  reduction.

### 16.1.3 Simplifying Operator Application

To simplify the exposition, I assume that every operator application is written in the form  $Op(e_1, \dots, e_n)$ .  $TLA^+$  provides a number of different syntactic forms for operator application, so I have to explain how they are translated into this simple form. Here are all the different forms of operator application and their translations.

- Simple constructs with a fixed number of arguments, including infix operators like  $+$ , prefix operators like `ENABLED`, and constructs like `WF`, function application, and `IF/THEN/ELSE`. These operators and constructs pose no problem. We can write  $+(a, b)$  instead of  $a + b$ , *IfThenElse*( $p, e_1, e_2$ ) instead of

$$\text{IF } p \text{ THEN } e_1 \text{ ELSE } e_2$$

and *Apply*( $f, e$ ) instead of  $f[e]$ . An expression like  $a + b + c$  is an abbreviation for  $(a + b) + c$ , so it can be written  $+(+(a, b), c)$ .

- Simple constructs with a variable number of arguments—for example,  $\{e_1, \dots, e_n\}$  and  $[h_1 \mapsto e_1, \dots, h_n \mapsto e_n]$ . We can consider each of these constructs to be repeated application of simpler operators with a fixed number of arguments. For example,

$$\begin{aligned} \{e_1, \dots, e_n\} &= \{e_1\} \cup \dots \cup \{e_n\} \\ [h_1 \mapsto e_1, \dots, h_n \mapsto e_n] &= [h_1 \mapsto e_1] @@ \dots @@ [h_n \mapsto e_n] \end{aligned}$$

where `@@` is defined in the *TLC* module, on page 223. Of course,  $\{e\}$  can be written *Singleton*( $e$ ) and  $[h \mapsto e]$  can be written *Record*(“ $h$ ”,  $e$ ). Note that an arbitrary CASE expression can be written in terms of CASE expressions of the form

$$\text{CASE } p \rightarrow e \square q \rightarrow f$$

using the relation:

$$\begin{aligned} \text{CASE } p_1 \rightarrow e_1 \square \dots \square p_n \rightarrow e_n &= \\ \text{CASE } p_1 \rightarrow e_1 \square (p_2 \vee \dots \vee p_n) \rightarrow (\text{CASE } p_2 \rightarrow e_2 \square \dots \square p_n \rightarrow e_n) \end{aligned}$$

- Constructs that introduce bound variables—for example,

$$\exists x \in S : x + z > y$$

We can rewrite this expression as

$$\text{ExistsIn}(S, \lambda x : x + z > y)$$

where *ExistsIn* is a 2nd-order operator of arity  $\langle \_, \langle \_ \rangle \rangle$ . As explained in Section 15.1.1, all the variants of the  $\exists$  construct can be expressed using either  $\exists x \in S : e$  or  $\exists x : S$ . All the other constructs that introduce bound variables, such as  $\{x \in S : \text{exp}\}$ , can similarly be expressed in the form

$Op(e_1, \dots, e_n)$  using  $\lambda$  expressions and 2nd-order operators  $Op$ . (Chapter 15 explains how to express constructs like  $\{\langle x, y \rangle \in S : exp\}$ , that have a tuple of bound identifiers, in terms of constructs with ordinary bound identifiers.)

- Operator applications such as  $M(x)!Op(y, z)$  that arise from instantiation. We write this as  $M!Op(x, y, z)$ .
- LET expressions. The meaning of a LET expression is explained in Section 16.4 below. For now, we consider only LET-free  $\lambda$  expressions—ones that contain no LET expressions.

For uniformity, I will call an operator symbol an *identifier*, even if it is a symbol like  $+$  that isn't an identifier according to the syntax of Chapter 14.

### 16.1.4 Expressions

We can now inductively define an expression to be either a 0th-order operator, or to have the form  $Op(e_1, \dots, e_n)$  where  $Op$  is an operator and each  $e_i$  is either an expression or a 1st-order operator. The expression must be *arity-correct*, meaning that  $Op$  must have arity  $\langle a_1, \dots, a_n \rangle$ , where each  $a_i$  is the arity of  $e_i$ . In other words,  $e_i$  must be an expression if  $a_i$  equals  $-$ , otherwise it must be a 1st-order operator with arity  $a_i$ . We require that  $Op$  not be a  $\lambda$  expression. (If it is, we can use  $\beta$  reduction to evaluate  $Op(e_1, \dots, e_n)$  and eliminate the  $\lambda$  expression  $Op$ .) Hence, a  $\lambda$  expression can appear in an expression only as an argument of a 2nd-order operator. This implies that only 1st-order  $\lambda$  expressions can appear in an expression.

We have eliminated all bound identifiers except the ones in  $\lambda$  expressions. We maintain the  $TLA^+$  requirement that an identifier that already has a meaning cannot be used as a bound identifier. Thus, in any  $\lambda$  expression  $\lambda p_1, \dots, p_n : exp$ , the identifiers of the parameters  $p_i$  cannot appear as parameter identifiers in any  $\lambda$  expression that occurs in  $exp$ .

Remember that  $\lambda$  expressions are used only to explain the semantics of  $TLA^+$ . They are not part of the language, and they can't be used in a  $TLA^+$  specification.

## 16.2 Levels

$TLA^+$  has a class of syntactic restrictions that come from the underlying logic  $TLA$  and have no counterpart in ordinary mathematics. The simplest of these is that “double-priming” is prohibited. For example,  $(x' + y)'$  is not syntactically well-formed, and is therefore meaningless, because the operator  $'$  (priming) can

be applied only to a state function, not to a transition function like  $x' + y$ . This class of restriction is expressed in terms of *levels*.

In TLA, an expression has one of four basic levels, which are numbered 0, 1, 2, and 3. These levels are described below, using examples that assume  $x$ ,  $y$ , and  $c$  are declared by

VARIABLES  $x, y$       CONSTANT  $c$

and symbols like  $+$  have their usual meanings.

0. A *constant*-level expression is a constant; it contains only constants and constant operators. Example:  $c + 3$ .
1. A *state*-level expression is a state function; it may contain constants, constant operators, and unprimed variables. Example:  $x + 2 * c$ .
2. A *transition*-level expression is a transition function; it may contain anything except temporal operators. Example:  $x' + y > c$ .
3. A *temporal*-level expression is a temporal formula; it may contain any TLA operator. Example:  $\Box[x' > y + c]_{\langle x, y \rangle}$ .

Chapter 15 assigns meanings to all basic expressions—ones containing only the built-in operators of  $\text{TLA}^+$  and declared constants and variables. The meaning assigned to an expression depends as follows on its level.

0. The meaning of a constant-level basic expression is a constant-level basic expression containing only primitive operators.
1. The meaning of a state-level basic expression is an assignment of a constant expression  $s[e]$  to any state  $s$ .
2. The meaning of a transition-level basic expression  $e$  is an assignment of a constant expression  $\langle s, t \rangle[e]$  to any transition  $s \rightarrow t$ .
3. The meaning of a temporal-level basic expression  $F$  is an assignment of a constant expression  $\sigma \models F$  to any behavior  $\sigma$ .

An expression of any level can be considered to be an expression of a higher level, except that a transition-level expression is not a temporal-level expression.<sup>2</sup> For example, if  $x$  is a declared variable, then the state-level expression  $x > 2$  is the temporal-level formula such that  $\sigma \models x$  is the value of  $x > 2$  in the first state of  $\sigma$ , for any behavior  $\sigma$ .<sup>3</sup>

A set of simple rules inductively defines whether a basic expression is *level-correct* and, if so, what its level is. Here are some of the rules:

<sup>2</sup>More precisely, a transition-level expression that is not a state-level expression is not a temporal-level expression.

<sup>3</sup>The expression  $x + 2$  can be considered to be a temporal-level expression that, like the temporal-level expression  $\Box(x + 2)$ , is silly. (See the discussion of silliness in Section 6.2 on page 67.)

- A declared constant is a level-correct expression of level 0.
- A declared variable is a level-correct expression of level 1.
- If  $Op$  is declared to be a 1st-order constant operator, then the expression  $Op(e_1, \dots, e_n)$  is level-correct iff each  $e_i$  is level correct, in which case its level is the maximum of the levels of the  $e_i$ .
- $e_1 \in e_2$  is level correct iff  $e_1$  and  $e_2$  are, in which case its level is the maximum of the levels of  $e_1$  and  $e_2$ .
- $e'$  is level-correct, and has level 2, iff  $e$  is level-correct and has level at most 1.<sup>4</sup>
- $\text{ENABLED } e$  is level-correct, and has level 1, iff  $e$  is level-correct and has level at most 2.
- $\exists x : e$  is level-correct, and has level  $l$ , iff  $e$  is level-correct and has level  $l$ , when  $x$  is considered to be a declared constant.
- $\exists x : e$  is level-correct, and has level 3, iff  $e$  is level-correct and has any level other than 2, when  $x$  is considered to be a declared variable.

There are additional rules for the other  $\text{TLA}^+$  operators. They should be obvious.

A useful consequence of the rules is that level-correctness of a basic expression does not depend on the levels of the declared identifiers. In other words, an expression  $e$  is level-correct when  $c$  is declared to be a constant iff it is level-correct when  $c$  is declared to be a variable. Of course, the level of  $e$  may depend on the level of  $c$ .

We can abstract these rules by generalizing the concept of a level. So far, I have defined the level only of an expression. I now define the level of a 1st- or 2nd-order operator  $Op$  to be a rule for determining the level-correctness and level of an expression  $Op(e_1, \dots, e_n)$  as a function of the levels of the arguments  $e_i$ . The level of a 1st-order operator is a rule, so the level of a 2nd-order operator  $Op$  is a rule that depends in part on rules—namely, on the levels of the arguments that are operators. This makes a rigorous general definition of levels for 2nd-order operators rather complicated. Fortunately, there's a simpler, less general definition that handles all the operators of  $\text{TLA}^+$ . Even more fortunately, you don't have to know it, so I won't bother writing it down. All you need to know is that there exists a way of assigning a level to every built-in operator of  $\text{TLA}^+$ . The level-correctness and level of any basic expression is then determined by those levels and the levels of the declared identifiers that occur in the expression.

<sup>4</sup>If  $e$  is a constant expression, then  $e'$  equals  $e$ , so we could consider  $e'$  to have level 0. For simplicity, we consider  $e'$  to have level 2 even if  $e$  is a constant.

One important class of operator levels are the *constant* levels. Any expression built from constant-level operators and declared constants has constant level. The built-in constant operators of TLA<sup>+</sup>, listed in Figures 13.10 and 13.10 (pages 249 and 249) all have constant level. Any operator defined solely in terms of constant-level operators and declared constants has constant level.

We now extend the definition of level-correctness from expressions to  $\lambda$  expressions. We define the  $\lambda$  expression  $\lambda p_1, \dots, p_n : exp$  to be level-correct iff  $exp$  is level-correct when the  $\lambda$  parameter identifiers are declared to be constants of the appropriate arity. For example,  $\lambda p, q(-) : exp$  is level-correct iff  $exp$  is level-correct with the additional declaration:

CONSTANTS  $p, q(-)$

This inductively defines level-correctness for  $\lambda$  expressions. The definition is reasonable because, as observed a couple of paragraphs ago, the level-correctness of  $exp$  doesn't depend on whether we assign level 0 or 1 to the  $\lambda$  parameters. One can also define the level of an arbitrary  $\lambda$  expression, but that would require the general definition of the level of an operator, whose complexity we want to avoid.

## 16.3 Contexts

Syntactic correctness of a basic expression depends on the arities of the declared identifiers. The expression  $Foo = \{ \}$  is syntactically correct if  $Foo$  is declared to be variable, and hence of arity  $\_$ , but not if it's declared to be a (1st-order) constant of arity  $\langle \_ \rangle$ . The meaning of a basic expression also depends on the levels of the declared identifiers. We can't determine those arities and levels just by looking at the expression itself; they are implied by the context in which the expression appears. A nonbasic expression contains defined as well as declared operators. Its syntactic correctness and meaning depend on the definitions of those operators, which also depends on the context. This section defines a precise notion of a context.

For uniformity, we can treat built-in operators the same as defined and declared operators. Just as the context can tell us that the identifier  $x$  is a declared variable, it can tell us that  $\in$  is declared to be a constant-level operator of arity  $\langle \_, \_ \rangle$  and that  $\notin$  is defined to equal  $\lambda a, b : \neg(\in(a, b))$ . Assuming a standard context that specifies all the built-in operators simplifies the description of the semantics.

To define contexts, let's first define declarations and definitions. A *declaration* assigns an arity and level to an operator name. A *definition* assigns a LET-free  $\lambda$  expression to an operator name. A *module definition* assigns the meaning of a module to a module name, where the meaning of a module is defined in

Section 16.5 below.<sup>5</sup> A *context* consists of a set of declarations, definitions, and module definitions such that:

- C1. An operator name is declared or defined at most once by the context.  
(This means that it can't be both declared and defined.)
- C2. No operator defined or declared by the context appears as the identifier of a  $\lambda$  parameter in any definition's expression.
- C3. Every operator name that appears in a definition's expression is either  $\lambda$  parameter's identifier or is declared (not defined) by the context.
- C4. No module name is assigned meanings by two different module definitions.

Module and operator names are handled separately. The same string may be both a module name that is defined by a module definition and an operator name that is either declared or defined by an ordinary definition.

Here is an example of a context that declares the symbols  $\cup$ ,  $a$ ,  $b$ , and  $\in$ , defines the symbols  $c$  and  $foo$ , and defines the module *Naturals*:

$$(16.2) \{ \cup : \langle \_, \_ \rangle, \quad a : \_, \quad b : \_, \quad \in : \langle \_, \_ \rangle, \quad c \triangleq \cup(a, b), \\ foo \triangleq \lambda p, q(-) : \in(p, \cup(q(b), a)), \quad \textit{Naturals} \stackrel{m}{=} \dots \}$$

I have not shown the levels assigned to the operators  $\cup$ ,  $a$ ,  $b$ , and  $\in$ , or the meaning assigned to *Naturals*.

If  $\mathcal{C}$  is a context, a  *$\mathcal{C}$ -basic  $\lambda$  expression* is defined to be a  $\lambda$  expression that contains only symbols declared in  $\mathcal{C}$  (in addition to  $\lambda$  parameters). For example,  $\lambda x : \in(x, \cup(a, b))$  is a  $\mathcal{C}$ -basic  $\lambda$  expression if  $\mathcal{C}$  is the context (16.2). However neither  $\cap(a, b)$  nor  $\lambda x : c(x, b)$  is a  $\mathcal{C}$ -basic  $\lambda$  expressions because neither  $\cap$  nor  $c$  is declared in  $\mathcal{C}$ . (The symbol  $c$  is defined, not declared, in  $\mathcal{C}$ .) A  $\mathcal{C}$ -basic  $\lambda$  expression is *syntactically correct* if it is arity- and level-correct with the arities and levels assigned by  $\mathcal{C}$  to the expression's operators. Condition C3 states that if  $Op \triangleq exp$  is a definition in context  $\mathcal{C}$ , then  $exp$  is a  $\mathcal{C}$ -basic  $\lambda$  expression. We add to C3 the requirement that it be syntactically correct.

We also allow a context to contain a special definition of the form  $Op \triangleq ?$  that assigns to the name  $Op$  an “illegal” value  $?$  that is not a  $\lambda$  expression. This definition indicates that, in the context, it is illegal to use the operator name  $Op$ .

## 16.4 The Meaning of a $\lambda$ Expression

I now define the meaning  $\mathcal{C}[[e]]$  of a  $\lambda$  expression  $e$  in a context  $\mathcal{C}$  to be a  $\mathcal{C}$ -basic  $\lambda$  expression. If  $e$  is an ordinary (nonbasic) expression, and  $\mathcal{C}$  is the context that

<sup>5</sup>The meaning of a module is defined in terms of contexts, so these definitions appear to be circular. In fact, the definitions of context and of the meaning of a module together form a single inductive definition.

specifies the built-in  $\text{TLA}^+$  operators and declares the constants and variables that occur in  $e$ , then  $\mathcal{C}[e]$  is a basic expression. Since Chapter 15 defines the meaning of basic expressions, this defines the meaning of an arbitrary expression. The expression  $e$  may contain LET constructs, so this defines the meaning of LET, the one operator whose meaning is not defined in Chapter 15.

Basically,  $\mathcal{C}[e]$  is obtained from  $e$  by replacing all defined operator names with their definitions, and then applying  $\beta$  reduction whenever possible. Recall that  $\beta$  reduction replaces

$$(\lambda p_1, \dots, p_n : exp)(e_1, \dots, e_n)$$

with the expression obtained from  $exp$  by replacing the identifier of  $p_i$  with  $e_i$ , for each  $i$ . The definition of  $\mathcal{C}[e]$  does not depend on the levels assigned by the declarations of  $\mathcal{C}$ . So, we ignore levels in the definition. The inductive definition of  $\mathcal{C}[e]$  consists of the following rules:

- If  $e$  is an operator symbol, then  $\mathcal{C}[e]$  equals (i)  $e$  if  $e$  is declared in  $\mathcal{C}$ , or (ii) the  $\lambda$  expression of  $e$ 's definition in  $\mathcal{C}$  if  $e$  is defined in  $\mathcal{C}$ .
- If  $e$  is  $Op(e_1, \dots, e_n)$ , where  $Op$  is declared in  $\mathcal{C}$ , then  $\mathcal{C}[e]$  equals the expression  $Op(\mathcal{C}[e_1], \dots, \mathcal{C}[e_n])$ .
- If  $e$  is  $Op(e_1, \dots, e_n)$ , where  $Op$  is defined in  $\mathcal{C}$  to equal the  $\lambda$  expression  $d$ , then  $\mathcal{C}[e]$  equals the  $\beta$  reduction of  $\overline{d}(\mathcal{C}[e_1], \dots, \mathcal{C}[e_n])$ , where  $\overline{d}$  is obtained from  $d$  by  $\alpha$  reduction (replacement of  $\lambda$  parameters) so that no  $\lambda$  parameter's identifier appears in both  $\overline{d}$  and some  $\mathcal{C}[e_i]$ .
- If  $e$  is  $\lambda p_1, \dots, p_n : exp$ , then  $\mathcal{C}[e]$  equals  $\lambda p_1, \dots, p_n : \mathcal{D}[exp]$ , where  $\mathcal{D}$  is the context obtained by adding to  $\mathcal{C}$  the declarations that, for each  $i$  in  $1 \dots n$ , assigns to the  $i^{\text{th}}$   $\lambda$  parameter's identifier the arity determined by  $p_i$ .
- If  $e$  is  $\text{LET } Op \triangleq d \text{ IN } exp$ , where  $d$  is a  $\lambda$  expression and  $exp$  an expression, then  $\mathcal{C}[e]$  equals  $\mathcal{D}[exp]$ , where  $\mathcal{D}$  is the context obtained by adding to  $\mathcal{C}$  the definition that assigns  $\mathcal{C}[d]$  to  $Op$ .

The last condition defines the meaning of any LET construct, because:

- The operator definition  $Op(p_1, \dots, p_n) \triangleq d$  in a LET means:

$$Op \triangleq \lambda p_1, \dots, p_n : d$$

- A function definition  $Op[x \in S] \triangleq d$  in a LET means:

$$Op \triangleq \text{CHOOSE } Op : Op = [x \in S \mapsto d]$$

- The expression  $\text{LET } Op_1 \triangleq d_1 \dots Op_n \triangleq d_n \text{ IN } exp$  is defined to equal

$$\text{LET } Op_1 \triangleq d_1 \text{ IN } (\text{LET } \dots \text{ IN } (\text{LET } Op_n \triangleq d_n \text{ IN } exp) \dots)$$

The  $\lambda$  expression  $e$  is defined to be legal (syntactically well-formed) in the context  $\mathcal{C}$  iff these rules define  $\mathcal{C}[e]$  to be a legal  $\mathcal{C}$ -basic expression.

## 16.5 The Meaning of a Module

The meaning of a module depends on a context. For an external module, which is not a submodule of another module, the context consists of declarations and definitions of all the built-in operators of  $\text{TLA}^+$ , together with definitions of some set of modules. This set of modules is discussed in Section 16.7 below.

The meaning of a module in a context  $\mathcal{C}$  consists of six sets:

- Dcl* A set of declarations. They come from `CONSTANT` and `VARIABLE` declarations and declarations in extended modules (modules appearing in an `EXTENDS` statement).
- GDef* A set of global definitions. They come from ordinary (non-`LOCAL`) definitions and global definitions in extended and instantiated modules.
- LDef* A set of local definitions. They come from `LOCAL` definitions and `LOCAL` instantiations of modules. (Local definitions are not obtained by other modules that extend or instantiate the module.)
- MDef* A set of module definitions. They come from submodules of the module and of extended modules.
- Ass* A set of assumptions. They come from `ASSUME` statements and from extended modules.
- Thm* A set of theorems. They come from `THEOREM` statements, from theorems in extended modules, and from the assumptions and theorems of instantiated modules, as explained below.

The  $\lambda$  expressions of definitions in *GDef* and *LDef*, as well as the expressions in *Ass* and *Thm*, are  $(\mathcal{C} \cup \text{Dcl})$ -basic  $\lambda$  expressions. In other words, the only operator symbols they contain (other than  $\lambda$  parameter identifiers) are ones declared in  $\mathcal{C}$  or in *Dcl*.

The meaning of a module in a context  $\mathcal{C}$  is defined by an algorithm for computing these six sets. The algorithm processes each statement in the module in turn, from beginning to end. The meaning of the module is the value of those sets when the end of the module is reached.

Initially, all six sets are empty. The rules for handling each possible type of statement are given below. In these rules, the *current context*  $\mathcal{CC}$  is defined to be the union of  $\mathcal{C}$ , *Dcl*, *GDef*, *LDef*, and *MDef*.

When the algorithm adds elements to the context  $\mathcal{CC}$ , it uses  $\alpha$  conversion to ensure that no defined or declared operator name appears as a  $\lambda$  parameter's identifier in any  $\lambda$  expression in  $\mathcal{CC}$ . For example, if the definition  $\text{foo} \triangleq \lambda x : x + 1$  is in *LDef*, then adding a declaration of  $x$  to *Dcl* requires  $\alpha$  conversion of this definition to rename the  $\lambda$  parameter identifier  $x$ . This  $\alpha$  conversion is not explicitly mentioned.

### 16.5.1 Extends

An EXTENDS statement has the form

$$\text{EXTENDS } M_1, \dots, M_n$$

where each  $M_i$  is a module name. This statement must be the first one in the module. The statement sets the values of  $Dcl$ ,  $GDef$ ,  $MDef$ ,  $Ass$ , and  $Thm$  equal to the union of the corresponding values for the module meanings assigned by  $\mathcal{C}$  to the module names  $M_i$ .

This statement is legal iff the module names  $M_i$  are all defined in  $\mathcal{C}$ , and the resulting current context  $\mathcal{CC}$  does not assign more than one meaning to any symbol. It is illegal for one of the modules  $M_i$  to define a symbol and another to declare the same symbol. It is also illegal for the resulting value of  $\mathcal{CC}$  to have two different declarations or two different definitions of the same operator name.

Since  $\mathcal{CC}$  is a set, by definition it cannot have multiple copies of the same definition or declaration. This means that an operator name can be declared or defined in two different modules  $M_i$ , as long as those declarations or definitions are the same. Two declarations are the same iff they assign the same level and arity. The semantics of  $\text{TLA}^+$  does not specify precisely what it means for two definitions to be the same. A tool may or may not consider the two definitions

$$A \triangleq \lambda X : \cup(X, Z) \quad A \triangleq \lambda Y : \cup(Y, Z)$$

to be the same. It does guarantee that two definitions are the same if they come from the same definition statement in the same module. For example, suppose  $M1$  and  $M2$  both extend the *Naturals* module. Then modules *Naturals*,  $M1$ , and  $M2$  all define  $+$ . However, all three definitions are the same, because they all come from the statement that provides the definition for *Naturals*. Hence, the multiple definitions of  $+$  and other operators on natural numbers obtained by the statement

$$\text{EXTENDS } \textit{Naturals}, M1, M2$$

are legal.

### 16.5.2 Declarations

A declaration statement has one of the forms

$$\text{CONSTANT } c_1, \dots, c_n \quad \text{VARIABLE } v_1, \dots, v_n$$

each  $v_i$  is an identifier and each  $c_i$  is either an identifier or has the form  $Op(\_, \dots, \_)$  for some identifier  $Op$ . This statement adds to the set  $Dcl$  the obvious declarations. It is legal iff none of the declared identifiers is defined in

$\mathcal{CC}$  or is declared in  $\mathcal{CC}$  with different arity or level. It is legal to redeclare an identifier that is declared in  $\mathcal{CC}$ , if the new declaration is the same as the one in  $\mathcal{CC}$ . For example, if module  $M$  extends module  $N$ , then both  $M$  and  $N$  can declare  $x$  to be a variable. However, it is illegal for  $M$  to declare  $x$  to be a variable and  $N$  to declare it to be a constant. It is also legal for  $x$  to be declared by two different `VARIABLE` declarations in the same module. However, redundant declarations might cause a tool to issue a warning.

### 16.5.3 Operator Definitions

A global operator definition<sup>6</sup> has one of the two forms

$$Op \triangleq exp \quad Op(p_1, \dots, p_n) \triangleq exp$$

where  $Op$  is an identifier,  $exp$  is an expression, and each  $p_i$  is either an identifier or has the form  $P(\_, \dots, \_)$ , where  $P$  is an identifier. We consider the first form an instance of the second with  $n = 0$ .

This statement is legal if the  $\lambda$  expression  $\lambda p_1, \dots, p_n : exp$  is legal in the context  $\mathcal{CC}$ . In particular, no  $\lambda$  parameter in this  $\lambda$  expression can be defined or declared in  $\mathcal{CC}$ . The statement adds to  $GDef$  the definition that assigns to  $Op$  the  $\lambda$  expression  $\mathcal{CC}[\lambda p_1, \dots, p_n : exp]$ .

A local operator definition has one of the two forms

$$\text{LOCAL } Op \triangleq exp \quad \text{LOCAL } Op(p_1, \dots, p_n) \triangleq exp$$

It is the same as a global definition, except that it adds the definition to  $LDef$  instead of  $GDef$ .

### 16.5.4 Function Definitions

A global function definition has the form

$$Op[fnargs] \triangleq exp$$

where  $fnargs$  is a comma-separated list of elements, each having the form  $Id_1, \dots, Id_n \in S$  or  $\langle Id_1, \dots, Id_n \rangle \in S$ . It is equivalent to the global operator definition

$$Op \triangleq \text{CHOOSE } Op : Op = [fnargs \mapsto exp]$$

A local function definition, which has the form

$$\text{LOCAL } Op[fnargs] \triangleq exp$$

is equivalent to the analogous local operator definition.

---

<sup>6</sup>An operator definition statement should not be confused with a definition clause in a `LET` expression. The meaning of a `LET` expression is described in Section 16.4.

### 16.5.5 Instantiation

We consider first a global instantiation of the form:

$$(16.3) \quad I(p_1, \dots, p_m) \triangleq \text{INSTANCE } N \text{ WITH } q_1 \leftarrow e_1, \dots, q_n \leftarrow e_n$$

For this to be legal,  $N$  must be a module name defined in  $\mathcal{CC}$ . Let  $NDcl$ ,  $NDef$ ,  $NAss$ , and  $NThm$  be the sets  $Dcl$ ,  $GDef$ ,  $Ass$ , and  $Thm$  in the meaning assigned to  $N$  by  $\mathcal{CC}$ . The  $q_i$  must be distinct identifiers declared by  $NDcl$ . We add a WITH clause of the form  $Op \leftarrow Op$  for any identifier  $Op$  that is declared in  $NDcl$  but is not one of the  $q_i$ , so the  $q_i$  constitute all the identifiers declared in  $NDcl$ .

Neither  $I$  nor any of the identifiers of the definition parameters  $p_i$  may be defined or declared in  $\mathcal{CC}$ . Let  $\mathcal{D}$  be the context obtained by adding to  $\mathcal{CC}$  the obvious constant-level declaration for each  $p_i$ . Then  $e_i$  must be syntactically well-formed in the context  $\mathcal{D}$ , and  $\mathcal{D}[[e_i]]$  must have the same arity as  $q_i$ , for each  $i \in 1 \dots n$ .

The instantiation must also satisfy the following level-correctness condition. Define module  $N$  to be a *constant* module iff every declaration in  $NDcl$  has constant level, and every operator appearing in every definition in  $NDef$  has constant level. If  $N$  is *not* a constant module, then for each  $i$  in  $1 \dots n$ :

- If  $q_i$  is declared in  $NDcl$  to be a constant operator, then  $\mathcal{D}[[e_i]]$  has constant level.
- If  $q_i$  is declared in  $NDcl$  to be a variable (a 0th-order operator of level 1), then  $\mathcal{D}[[e_i]]$  has level 0 or 1.

The reason for this condition is explained in Section 16.8 below.

For each definition  $Op \triangleq \lambda r_1, \dots, r_p : e$  in  $NDef$ , the definition

$$(16.4) \quad I!Op \triangleq \lambda p_1, \dots, p_m, r_1, \dots, r_p : \bar{e}$$

is added to  $GDef$ , where  $\bar{e}$  is the expression obtained from  $e$  by substituting  $e_i$  for  $q_i$ , for all  $i \in 1 \dots n$ . Before doing this substitution,  $\alpha$  conversion must be applied to ensure that  $\mathcal{CC}$  is a correct context after the definition of  $I!Op$  is added to  $GDef$ . The precise definition of  $\bar{e}$  is a bit subtle; it is given in Section 16.8 below. We require that the  $\lambda$  expression in (16.4) must be level-correct. (If  $N$  is a nonconstant module, then level-correctness of this  $\lambda$  expression is implied by the level condition on parameter instantiation described in the preceding paragraph.) Legality of the definition of  $Op$  in module  $N$  and of the WITH substitutions imply that the  $\lambda$  expression is arity-correct in the current context. Remember that  $I!Op(c_1, \dots, c_m, d_1, \dots, d_n)$  is actually written in  $\text{TLA}^+$  as  $I(c_1, \dots, c_m)!Op(d_1, \dots, d_n)$ .

Also added to  $GDef$  is the special definition  $I \triangleq ?$ . This prevents  $I$  from later being defined or declared as an operator name.

A global INSTANCE statement can also have the two forms:

$$\begin{aligned} I &\triangleq \text{INSTANCE } N \text{ WITH } q_1 \leftarrow e_1, \dots, q_n \leftarrow e_n \\ &\text{INSTANCE } N \text{ WITH } q_1 \leftarrow e_1, \dots, q_n \leftarrow e_n \end{aligned}$$

The first is just the  $m = 0$  case of (16.3); the second is similar to the first, except the definitions added to  $GDef$  do not have  $I!$  prepended to the operator names. In the second form, there is an obvious legality condition that adding the definitions to the current context does not introduce a name conflict. In all three forms of the statement, omitting the WITH clause is equivalent to the case  $n = 0$  of these statements. (Remember that all the declared identifiers of module  $N$  are either explicitly or implicitly instantiated.)

A local INSTANCE statement consists of the keyword LOCAL followed by an INSTANCE statement of the form described above. It is handled in a similar fashion to a global INSTANCE statement, except that all definitions are added to  $LDef$  instead of  $GDef$ .

### 16.5.6 Theorems and Assumptions

A theorem has one of the forms

$$\text{THEOREM } exp \qquad \text{THEOREM } Op \triangleq exp$$

where  $exp$  is an expression, which must be legal in the current context  $\mathcal{CC}$ . The first form adds the theorem  $\mathcal{CC}[[exp]]$  to the set  $Thm$ . The second form is equivalent to the pair of statements:

$$\begin{aligned} Op &\triangleq exp \\ \text{THEOREM } exp \end{aligned}$$

An assumption has one of the forms

$$\text{ASSUME } exp \qquad \text{ASSUME } Op \triangleq exp$$

The expression  $exp$  must have constant level. An assumption is similar to a theorem except that  $\mathcal{CC}[[exp]]$  is added to the set  $Ass$ .

### 16.5.7 Submodules

A module can contain a submodule, which is a complete module that begins with

$$\boxed{\text{MODULE } N}$$

for some module name  $N$ , and ends with

---

This is legal iff the module name  $N$  is not defined in  $\mathcal{CC}$  and the module is legal in the context  $\mathcal{CC}$ . In this case, the module definition that assigns to  $N$  the meaning of the submodule in context  $\mathcal{CC}$  is added to  $MDef$ .

A submodule can be used in an `INSTANCE` statement that appears later in the current module, or within a module that extends the current module.

## 16.6 Correctness of a Module

Section 16.5 above defines the meaning of a module to consist of the six sets  $Dcl$ ,  $GDef$ ,  $LDef$ ,  $MDef$ ,  $Ass$ , and  $Thm$ . Mathematically, we can view the meaning of a module to be the assertion that all the theorems in  $Thm$  are consequences of the assumptions in  $Ass$ . More precisely, let  $A$  be the conjunction of all the assumptions in  $Ass$ . The module asserts that, for every theorem  $T$  in  $Thm$ , the formula  $A \Rightarrow T$  is valid.<sup>7</sup>

An assumption or theorem of the module is a  $(\mathcal{C} \cup Dcl)$ -basic expression. For an outermost module (not a submodule),  $\mathcal{C}$  declares only the built-in operators of  $TLA^+$ , and  $Dcl$  declares the declared constants and variables of the module. Therefore, each formula  $A \Rightarrow T$  asserted by the module is a basic expression. We say that the module is *semantically correct* if each of these formulas  $A \Rightarrow T$  is valid in the context  $Dcl$ . Chapter 15, defines what it means for a basic expression to be valid.

By defining the meaning of a theorem, we have defined the meaning of a  $TLA^+$  specification. Any mathematically meaningful question we can ask about a specification can be framed as the question of whether a certain formula is a valid theorem.

## 16.7 Finding Modules

For a module  $M$  to have a meaning in a context  $\mathcal{C}$ , every module  $N$  extended or instantiated by  $M$  must have its meaning defined in  $\mathcal{C}$ —unless  $N$  is a submodule of  $M$  or of a module extended by  $M$ . In practice, a tool (or a person) begins interpreting a module  $M$  in a context  $\mathcal{C}_0$  containing only declarations and definitions of the built-in  $TLA^+$  operator names. When the tool encounters an `EXTENDS` or `INstantiate` statement that mentions a module named  $N$  not defined in the current context  $\mathcal{CC}$  of  $M$ , the tool finds the module named  $N$ , interprets it in the context  $\mathcal{C}_0$ , and then adds the module definition for  $N$  to  $\mathcal{C}_0$  and to  $\mathcal{CC}$ .

---

<sup>7</sup>In a temporal logic like  $TLA$ , the formula  $F \Rightarrow G$  is not in general equivalent to the assertion that  $G$  is a consequence of assumption  $F$ . However, the two are equivalent if  $F$  is a constant formula, and  $TLA^+$  allows only constant assumptions.

The definition of the  $TLA^+$  language does not specify how a tool finds a module named  $N$ . The tool will most likely look for the module in a file whose name is derived in some standard way from  $N$ .

## 16.8 The Semantics of Instantiation

Section 16.5.5 above defines the meaning of an `INSTANCE` statement in terms of substitution. I now define precisely how that substitution is performed and explain the level-correctness rule for instantiating nonconstant modules.

Suppose module  $M$  contains the statement

$$I \triangleq \text{INSTANCE } N \text{ WITH } q_1 \leftarrow e_1, \dots, q_n \leftarrow e_n$$

where the  $q_i$  are all the declared identifiers of module  $N$ , and that  $N$  contains the definition

$$F \triangleq e$$

where no  $\lambda$  parameter identifier in  $e$  is defined or declared in the current context of  $M$ . The `INSTANCE` statement then adds to the current context of  $M$  the definition

$$I!F \triangleq \bar{e}$$

where  $\bar{e}$  is obtained from  $e$  by substituting  $e_i$  for  $q_i$ , for all  $i$  in  $1 \dots n$ .

A fundamental principle of mathematics is that substitution preserves validity; substituting in a valid formula yields a valid formula. So, we want to define  $\bar{e}$  so that, if  $F$  is valid formula in  $N$ , then  $I!F$  is a valid formula in  $M$ .

A simple example shows that the level rule for instantiating nonconstant modules is necessary to preserve the validity of  $F$ . Suppose  $F$  is defined to equal  $\Box[c' = c]_c$ , where  $c$  is declared in  $N$  to be a constant. Then  $F$  is a temporal formula asserting that no step changes  $c$ . It is valid because a constant has the same value in every state of a behavior. If we allowed an instantiation that substitutes a variable  $x$  for the constant  $c$ , then  $I!F$  would be the formula  $\Box[x' = x]_x$ . This is not a valid formula because it is false for any behavior in which the value of  $x$  changes. Since  $x$  is a variable, such a behavior obviously exists. Preserving validity requires that we not allow substitution of a nonconstant for a declared constant when instantiating a nonconstant module. (Since  $\Box$  and  $'$  are nonconstant operators, this definition of  $F$  can appear only in a nonconstant module.)

In ordinary mathematics, there is one tricky problem in making substitution preserve validity. Consider the formula

$$(16.5) \quad (n \in \text{Nat}) \Rightarrow (\exists m \in \text{Nat} : m \geq n)$$

This formula is valid because it is true for any value of  $n$ . Now, suppose we substitute  $m + 1$  for  $n$ . A naive substitution that simply replaces  $n$  by  $m + 1$  would yield the formula

$$(16.6) \quad (m + 1 \in Nat) \Rightarrow (\exists m \in Nat : m \geq m + 1)$$

Since the formula  $\exists m \in Nat : m \geq m + 1$  is equivalent to FALSE,  $I!F$  is obviously not valid. Mathematicians call this problem *variable capture*;  $m$  is “captured” by the quantifier  $\exists m$ . Mathematicians avoid it by the rule that, when substituting for an identifier in a formula, one does not substitute for bound occurrences of the identifier. This rule requires that  $m$  be removed from (16.5) by  $\alpha$  conversion before  $m + 1$  is substituted for  $n$ .

Section 16.5.5 defines the meaning of the INSTANCE statement in a way that avoids variable capture. Indeed, formula (16.6) is illegal in  $TLA^+$  because the subexpression  $m + 1 \in Nat$  is allowed only in a context in which  $m$  is defined or declared, in which case  $m$  cannot be used as a bound identifier, so the subexpression  $\exists m \dots$  is illegal. The  $\alpha$  conversion necessary to produce a syntactically well-formed expression makes this kind of variable capture impossible.

The problem of variable capture occurs in a more subtle form in certain nonconstant operators of  $TLA^+$ , where it is not prevented by the syntactic rules. Most notable of these operators is ENABLED. Suppose  $x$  and  $y$  are declared variables of module  $N$ , and  $F$  is defined by

$$F \triangleq \text{ENABLED}(x' = 0 \wedge y' = 1)$$

Then  $F$  is equivalent to TRUE, so it is valid in module  $N$ . (For any state  $s$ , there exists a state  $t$  in which  $x = 0$  and  $y = 1$ .) Now suppose  $z$  is a declared variable of module  $M$ , and let the instantiation be

$$I \triangleq \text{INSTANCE } N \text{ WITH } x \leftarrow z, y \leftarrow z$$

With naive substitution,  $I!F$  would equal

$$\text{ENABLED}(z' = 0 \wedge z' = 1)$$

which is equivalent to FALSE. (For any state  $s$ , there is no state  $t$  in which  $z = 0$  and  $z = 1$  are both true.) Hence,  $I!F$  would not be a theorem, so instantiation would not preserve validity.

Naive substitution in a formula of the form  $\text{ENABLED } A$  does not preserve validity because the primed variables in  $A$  are really bound identifiers. The formula  $\text{ENABLED } A$  asserts that *there exist* values of the primed variables such that  $A$  is true. Substituting  $z'$  for  $x'$  and  $y'$  in the  $\text{ENABLED}$  formula is really substitution for a bound identifier. It isn't ruled out by the syntactic rules of  $TLA^+$  because the quantification is implicit.

To preserve validity, we must define  $\bar{e}$  so it avoids capture of identifiers implicitly bound in  $\text{ENABLED}$  expressions. Before performing the substitution,

we first replace the primed occurrences of variables in ENABLED expressions with new variable symbols. That is, for each subexpression of  $e$  of the form ENABLED  $A$  and each declared variable  $q$  of module  $N$ , we replace every primed occurrence of  $q$  in  $A$  with a new symbol, which I will write  $\$q$ , that does not appear in  $A$ . This new symbol is considered to be bound by the ENABLED operator. For example, the module

$$\boxed{\begin{array}{l} \text{MODULE } N \\ \text{VARIABLE } u \\ G(v, A) \triangleq \text{ENABLED } (A \vee (\{u, v\}' = \{u, v\})) \\ H \triangleq (u' = u) \wedge G(u, u' \neq u) \end{array}}$$

has as its global definitions the set:

$$\begin{array}{l} \{ G \triangleq \lambda v, A : \text{ENABLED } (A \vee (\{u, v\}' = \{u, v\})) \\ H \triangleq (u' = u) \wedge \text{ENABLED } ((u' \neq u) \vee (\{u, u\}' = \{u, u\})) \} \end{array}$$

The statement

$$I \triangleq \text{INSTANCE } N \text{ WITH } u \leftarrow x$$

adds the following definitions to the current module:

$$\begin{array}{l} I!G \triangleq \lambda v, A : \text{ENABLED } (A \vee (\{\$u, v\}' = \{u, v\})) \\ I!H \triangleq (x' = x) \wedge \text{ENABLED } ((\$u' \neq x) \vee (\{\$u, \$u\}' = \{x, x\})) \end{array}$$

Observe that even though  $H$  equals  $G(u, u' \neq u)$  in module  $N$ , and the instantiation substitutes  $x$  for  $u$ , the instantiated formula  $I!H$  does not equal  $I!G(x, x' \neq x)$ .

As another example, consider the module

$$\boxed{\begin{array}{l} \text{MODULE } N \\ \text{VARIABLES } u, v \\ A \triangleq (u' = u) \wedge (v' \neq v) \\ B(d) \triangleq \text{ENABLED } d \\ C \triangleq B(A) \end{array}}$$

The instantiation

$$I \triangleq \text{INSTANCE } N \text{ WITH } u \leftarrow x, v \leftarrow x$$

adds the following definitions to the current module

$$\begin{aligned} I!A &\triangleq (x' = x) \wedge (x' \neq x) \\ I!B &\triangleq \lambda d : \text{ENABLED } d \\ I!C &\triangleq \text{ENABLED } ((\$u' = x) \wedge (\$v' \neq x)) \end{aligned}$$

Observe that  $I!C$  is not equivalent to  $I!B(I!A)$ . In fact,  $I!C \equiv \text{TRUE}$  and  $I!B(I!A) \equiv \text{FALSE}$ .

We say that instantiation *distributes* over an operator  $Op$  if

$$\overline{Op(e_1, \dots, e_n)} = Op(\overline{e_1}, \dots, \overline{e_n})$$

for any expressions  $e_i$ , where the overlining operator ( $\overline{\phantom{x}}$ ) denotes some arbitrary instantiation. Instantiation distributes over all constant operators—for example,  $+$ ,  $\subseteq$ , and  $\exists$ .<sup>8</sup> Instantiation also distributes over most of the nonconstant operators of  $\text{TLA}^+$ , like priming ( $'$ ) and  $\square$ .

If an operator  $Op$  implicitly binds some identifiers in its arguments, then instantiation would not preserve validity if it distributed over  $Op$ . Our rules for instantiating in an  $\text{ENABLED}$  expression imply that instantiation does not distribute over  $\text{ENABLED}$ . It also does not distribute over any operator defined in terms of  $\text{ENABLED}$ —in particular, the built-in operators  $\text{WF}$  and  $\text{SF}$ .

There are two other  $\text{TLA}^+$  operators that implicitly bind identifiers: the action composition operator “ $\cdot$ ”, defined in Section 15.2.3, and the temporal operator  $\triangleleft$ , introduced in Section 10.7. The rule for instantiating an expression  $A \cdot B$  is similar to that for  $\text{ENABLED } A$ —namely, bound occurrences of variables are replaced by a new symbol. In the expression  $A \cdot B$ , primed occurrences of variables in  $A$  and unprimed occurrences in  $B$  are bound. We handle a formula of the form  $F \triangleleft G$  by replacing it with an equivalent formula in which the quantification is made explicit.<sup>9</sup> Most readers won’t care, but here’s how that equivalent formula is constructed. Let  $\mathbf{x}$  be the tuple  $\langle x_1, \dots, x_n \rangle$  of all declared variables; let  $b, \widehat{x}_1, \dots, \widehat{x}_n$  be symbols distinct from the  $x_i$  and from any bound identifiers in  $F$  or  $G$ ; and let  $\widehat{e}$  be the expression obtained from an expression  $e$  by substituting the variables  $\widehat{x}_i$  for the corresponding variables  $x_i$ . Then  $F \triangleleft G$  is equivalent to

$$\begin{aligned} (16.7) \quad \forall b : ( \wedge (b = \text{TRUE}) \wedge \square [b' = \text{FALSE}]_b \\ \wedge \exists \widehat{x}_1, \dots, \widehat{x}_n : \widehat{F} \wedge \square (b \Rightarrow (\mathbf{x} = \widehat{\mathbf{x}})) ) \\ \Rightarrow \exists \widehat{x}_1, \dots, \widehat{x}_n : \widehat{G} \wedge (\mathbf{x} = \widehat{\mathbf{x}}) \wedge \square [b \Rightarrow (\mathbf{x}' = \widehat{\mathbf{x}}')]_{\langle b, \mathbf{x}, \widehat{\mathbf{x}} \rangle} \end{aligned}$$

<sup>8</sup>Recall the explanation on pages 304–305 of how we consider  $\exists$  to be a second-order operator. Instantiation distributes over  $\exists$  because  $\text{TLA}^+$  does not permit variable capture when substituting in  $\lambda$  expressions.

<sup>9</sup>Replacing  $\text{ENABLED}$  and “ $\cdot$ ” expressions by equivalent formulas with explicit quantifiers before substituting would result in some surprising instantiations. For example, if  $N$  contains the definition  $E(A) \triangleq \text{ENABLED } A$ , then  $\text{INSTANCE } I \triangleq N$  effectively obtains the definition  $I!E(A) \triangleq A$ .

Here's a complete statement of the rules for computing  $\bar{e}$ :

1. Remove all  $\pm\triangleright$  operators by replacing each subformula of the form  $F \pm\triangleright G$  with the equivalent formula (16.7).
2. Recursively perform the following replacements, starting from the innermost subexpressions of  $e$ , for each declared variable  $x$  of  $N$ .
  - For each subexpression of the form `ENABLED  $A$` , replace each primed occurrence of  $x$  in  $A$  by a new symbol  $\$x$  that is different from any identifier and from any other symbol that occurs in  $A$ .
  - For each subexpression of the form  $B \cdot C$ , replace each primed occurrence of  $x$  in  $B$  and each unprimed occurrence of  $x$  in  $C$  by a new symbol  $\$x$  that is different from any identifier and from any other symbol that occurs in  $B$  or  $C$ .

For example, applying these rules to the inner `ENABLED` expression and to the “.” expression converts

$$\text{ENABLED } (\text{ENABLED } (x' = x)' \wedge ((y' = x) \cdot (x' = y)))$$

to

$$\text{ENABLED } (\text{ENABLED } (\$x' = x)' \wedge ((\$y' = x) \cdot (x' = \$y)))$$

and applying them again to the outer `ENABLED` expression yields

$$\text{ENABLED } (\text{ENABLED } (\$x' = \$xx)' \wedge ((\$y' = x) \cdot (\$xx' = \$y)))$$

where  $\$xx$  is some new symbol different from  $x$ ,  $\$x$ , and  $\$y$ .

3. Replace each occurrence of  $q_i$  with  $e_i$ , for all  $i$  in  $1 \dots n$ .



## Chapter 17

# The Standard Modules

We provide several standard modules for use in  $\text{TLA}^+$  specifications. Some of the definitions they contain are subtle—for example, the definitions of the set of real numbers and its operators. Others, such as the definition of  $1 \dots n$ , are obvious. There are two reasons to use standard modules. First, specifications are easier to read when they use basic operators that we’re already familiar with. Second, tools can have built-in knowledge of standard operators. For example, the TLC model checker (Chapter 13) has efficient implementations of some standard modules; and a theorem-prover might implement special decision procedures for some standard operators. The standard modules of  $\text{TLA}^+$  are described here, except for the *RealTime* module, which appears in Chapter 9.

### 17.1 Module *Sequences*

The *Sequences* module was introduced in Section 4.1 on page 35. Most of the operators it defines have already been explained. The exceptions are:

*SubSeq*( $s, m, n$ ) The subsequence  $\langle s[m], s[m + 1], \dots, s[n] \rangle$  consisting of the  $m^{\text{th}}$  through  $n^{\text{th}}$  elements of  $s$ . It is undefined if  $m < 1$  or  $n > \text{Len}(s)$ , except that it equals the empty sequence if  $m > n$ .

*SelectSeq*( $s, \text{test}$ ) The subsequence of  $s$  consisting of the elements  $s[i]$  such that  $\text{test}[s[i]]$  equals TRUE. For example:

$$\begin{aligned} \text{PosSubSeq}(s) &\triangleq \text{LET } \text{IsPos}(n) \triangleq n > 0 \\ &\quad \text{IN } \text{SelectSeq}(s, \text{IsPos}) \end{aligned}$$

defines  $\text{PosSubSeq}(\langle 0, 3, -2, 5 \rangle)$  to equal  $\langle 3, 5 \rangle$ .

The *Sequences* module uses operators on natural numbers, so we might expect it to extend the *Naturals* module. However, this would mean that any module that extends *Sequences* would then also extend *Naturals*. Just in case someone wants to use sequences without extending the *Naturals* module, the *Sequences* module contains the statement:

```
LOCAL INSTANCE Naturals
```

This statement introduces the definitions from the *Naturals* module, just as an ordinary `INSTANCE` statement would, but it does not export those definitions to another module that extends or instantiates the *Sequences* module. The `LOCAL` modifier can also precede an ordinary definition; it has the effect of making that definition usable within the current module, but not in a module that extends or instantiates it. (The `LOCAL` modifier cannot be used with parameter declarations.)

Everything else that appears in the *Sequences* module should be familiar. The module is in Figure 17.1 on the next page.

## 17.2 Module *FiniteSets*

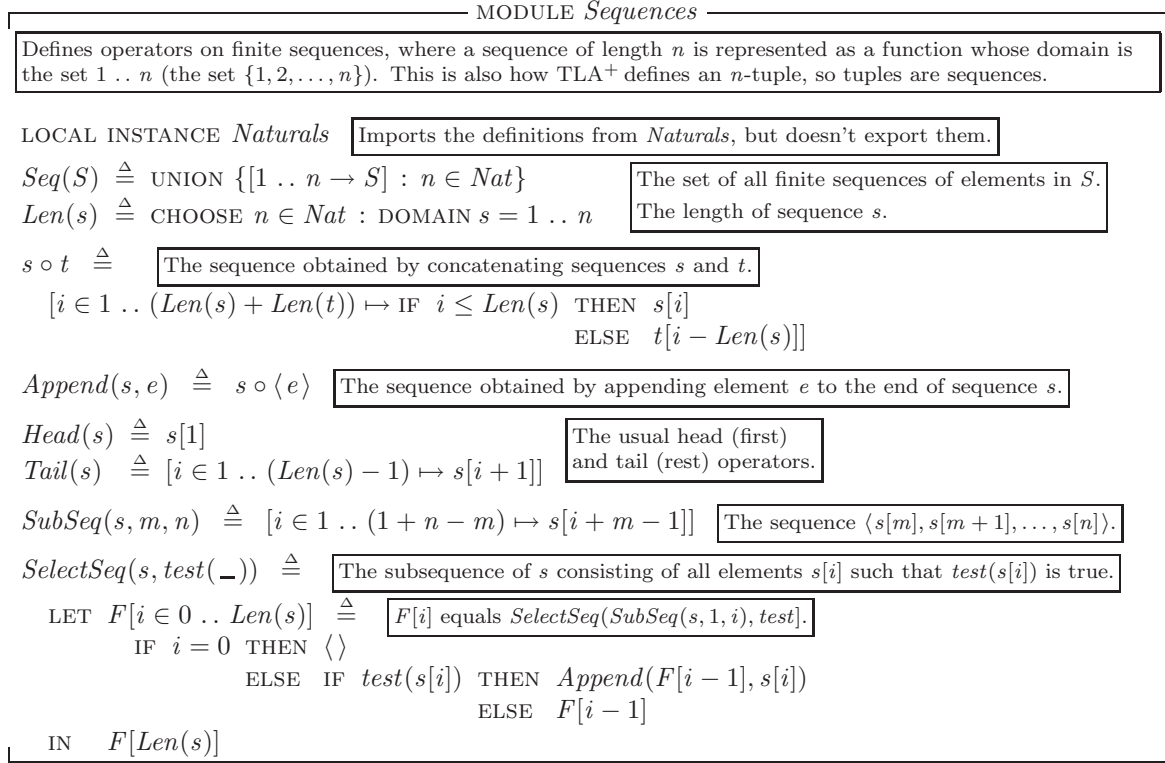
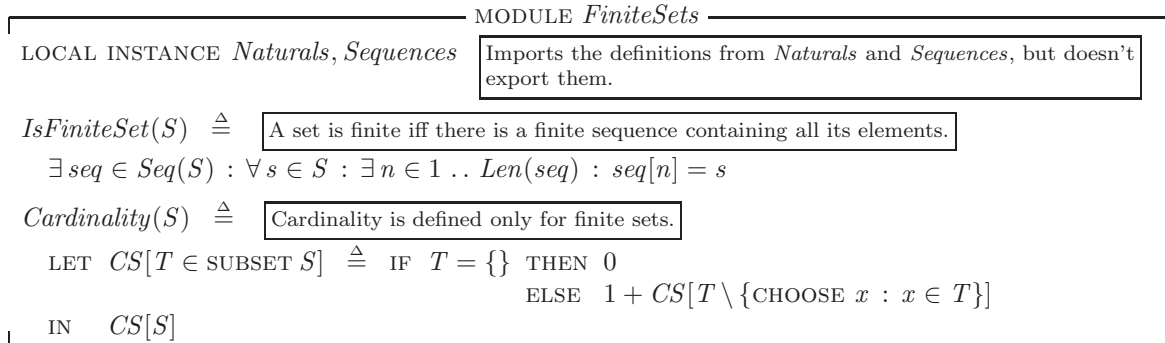
As described in Section 6.1 on page 66, the *FiniteSets* module defines the two operators *IsFiniteSet* and *Cardinality*. The definition of *Cardinality* is discussed on pages 69–70. The module itself is in Figure 17.2 on the next page.

## 17.3 Module *Bags*

A *bag*, also called a multiset, is a set that can contain multiple copies of the same element. A bag can have infinitely many elements, but only finitely many copies of any single element. Bags are sometimes useful for representing data structures. For example, the state of a network in which messages can be delivered in any order may be represented as a bag of messages in transit. Multiple copies of an element in the bag represent multiple copies of the same message in transit.

The *Bags* module defines a bag to be a function whose range is a subset of the positive integers. An element  $e$  belongs to bag  $B$  iff  $e$  is in the domain of  $B$ , in which case bag  $B$  contains  $B[e]$  copies of  $e$ . The module defines the following operators. In our customary style, we leave unspecified the value obtained by applying an operator on bags to something other than a bag.

|               |                                                                     |
|---------------|---------------------------------------------------------------------|
| $IsABag(B)$   | True iff $B$ is a bag.                                              |
| $BagToSet(B)$ | The set of elements at least one copy of which are in the bag $B$ . |

Figure 17.1: The standard *Sequences* module.Figure 17.2: The standard *FiniteSets* module.

|                     |                                                                                                                                                                                                                                                                           |
|---------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $SetToBag(S)$       | The bag that contains one copy of every element in the set $S$ .                                                                                                                                                                                                          |
| $BagIn(e, B)$       | True iff bag $B$ contains at least one copy of $e$ . $BagsIn$ is the $\in$ operator for bags.                                                                                                                                                                             |
| $EmptyBag$          | The bag containing no elements.                                                                                                                                                                                                                                           |
| $CopiesIn(e, B)$    | The number of copies of $e$ in bag $B$ . If $BagIn(e, B)$ is false, then $CopiesIn(e, B) = 0$ .                                                                                                                                                                           |
| $B1 \oplus B2$      | The union of bags $B1$ and $B2$ . The operator $\oplus$ satisfies<br>$CopiesIn(e, B1 \oplus B2) = CopiesIn(e, B1) + CopiesIn(e, B2)$ for any $e$ and any bags $B1$ and $B2$ .                                                                                             |
| $B1 \ominus B2$     | The bag $B1$ with the elements of $B2$ removed—that is, with one copy of an element removed from $B1$ for each copy of the same element in $B2$ . If $B2$ has at least as many copies of $e$ as $B1$ , then $B1 \ominus B2$ has no copies of $e$ .                        |
| $BagUnion(S)$       | The bag union of all elements of the set $S$ of bags. For example, $BagUnion(\{B1, B2, B3\})$ equals $B1 \oplus B2 \oplus B3$ .                                                                                                                                           |
| $B1 \sqsubseteq B2$ | True iff, for all $e$ , bag $B2$ has at least as many copies of $e$ as bag $B1$ does. Thus, $\sqsubseteq$ is the analog for bags of $\subseteq$ .                                                                                                                         |
| $SubBag(B)$         | The set of all subbags of bag $B$ . $SubBag$ is the bag analog of the SUBSET operator.                                                                                                                                                                                    |
| $BagOfAll(F, B)$    | The bag analog of the construct $\{F(x) : x \in B\}$ . It is the bag that contains, for each element $e$ of bag $B$ , one copy of $F(e)$ for every copy of $e$ in $B$ . This defines a bag iff, for any value $v$ , the set of $e$ in $B$ such that $F(e) = v$ is finite. |
| $BagCardinality(B)$ | If $B$ is a finite bag (one such that $BagToSet(B)$ is a finite set), then this is its cardinality—the total number of copies of elements in $B$ . Its value is unspecified if $B$ is not a finite bag.                                                                   |

The module appears in Figure 17.3 on the next page. Note the local definition of  $Sum$ , which makes  $Sum$  defined within the *Bags* module but not in any module that extends or instantiates it.

| MODULE <i>Bags</i>                                                                                                                                                                                                                                                                                                     |                                                                  |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------|
| LOCAL INSTANCE <i>Naturals</i>                                                                                                                                                                                                                                                                                         | Import definitions from <i>Naturals</i> , but don't export them. |
| $IsABag(B) \triangleq B \in [\text{DOMAIN } B \rightarrow \{n \in \text{Nat} : n > 0\}]$                                                                                                                                                                                                                               | True iff $B$ is a bag.                                           |
| $BagToSet(B) \triangleq \text{DOMAIN } B$                                                                                                                                                                                                                                                                              | The set of elements at least one copy of which is in $B$ .       |
| $SetToBag(S) \triangleq [e \in S \mapsto 1]$                                                                                                                                                                                                                                                                           | The bag that contains one copy of every element of the set $S$ . |
| $BagIn(e, B) \triangleq e \in BagToSet(B)$                                                                                                                                                                                                                                                                             | The $\in$ operator for bags.                                     |
| $EmptyBag \triangleq SetToBag(\{\})$                                                                                                                                                                                                                                                                                   |                                                                  |
| $B1 \oplus B2 \triangleq$                                                                                                                                                                                                                                                                                              | The union of bags $B1$ and $B2$ .                                |
| $[e \in (\text{DOMAIN } B1) \cup (\text{DOMAIN } B2) \mapsto$<br>$(\text{IF } e \in \text{DOMAIN } B1 \text{ THEN } B1[e] \text{ ELSE } 0) + (\text{IF } e \in \text{DOMAIN } B2 \text{ THEN } B2[e] \text{ ELSE } 0)]$                                                                                                |                                                                  |
| $B1 \ominus B2 \triangleq$                                                                                                                                                                                                                                                                                             | The bag $B1$ with the elements of $B2$ removed.                  |
| $\text{LET } B \triangleq [e \in \text{DOMAIN } B1 \mapsto \text{IF } e \in \text{DOMAIN } B2 \text{ THEN } B1[e] - B2[e] \text{ ELSE } B1[e]]$<br>$\text{IN } [e \in \{d \in \text{DOMAIN } B : B[d] > 0\} \mapsto B[e]]$                                                                                             |                                                                  |
| LOCAL $Sum(f) \triangleq$                                                                                                                                                                                                                                                                                              | The sum of $f[x]$ for all $x$ in $\text{DOMAIN } f$ .            |
| $\text{LET } DSum[S \in \text{SUBSET DOMAIN } f] \triangleq \text{LET } elt \triangleq \text{CHOOSE } e \in S : \text{TRUE}$<br>$\text{IN IF } S = \{\} \text{ THEN } 0$<br>$\text{ELSE } f[elt] + DSum[S \setminus \{elt\}]$                                                                                          |                                                                  |
| IN $DSum[\text{DOMAIN } f]$                                                                                                                                                                                                                                                                                            |                                                                  |
| $BagUnion(S) \triangleq$                                                                                                                                                                                                                                                                                               | The bag union of all elements of the set $S$ of bags.            |
| $[e \in \text{UNION } \{BagToSet(B) : B \in S\} \mapsto Sum([B \in S \mapsto \text{IF } BagIn(e, B) \text{ THEN } B[e] \text{ ELSE } 0])]$                                                                                                                                                                             |                                                                  |
| $B1 \sqsubseteq B2 \triangleq$                                                                                                                                                                                                                                                                                         | The subset operator for bags.                                    |
| $\wedge (\text{DOMAIN } B1) \subseteq (\text{DOMAIN } B2)$<br>$\wedge \forall e \in \text{DOMAIN } B1 : B1[e] \leq B2[e]$                                                                                                                                                                                              |                                                                  |
| $SubBag(B) \triangleq$                                                                                                                                                                                                                                                                                                 | The set of all subbags of bag $B$ .                              |
| $\text{LET } AllBagsOfSubset \triangleq$ The set of bags $SB$ such that $BagToSet(SB) \subseteq BagToSet(B)$ .<br>$\text{UNION } \{[SB \rightarrow \{n \in \text{Nat} : n > 0\}] : SB \in \text{SUBSET } BagToSet(B)\}$<br>$\text{IN } \{SB \in AllBagsOfSubset : \forall e \in \text{DOMAIN } SB : SB[e] \leq B[e]\}$ |                                                                  |
| $BagOfAll(F(-), B) \triangleq$                                                                                                                                                                                                                                                                                         | The bag analog of the set $\{F(x) : x \in B\}$ for a set $B$ .   |
| $[e \in \{F(d) : d \in BagToSet(B)\} \mapsto$<br>$Sum([d \in BagToSet(B) \mapsto \text{IF } F(d) = e \text{ THEN } B[d] \text{ ELSE } 0])]$                                                                                                                                                                            |                                                                  |
| $BagCardinality(B) \triangleq Sum(B)$                                                                                                                                                                                                                                                                                  | the total number of copies of elements in bag $B$ .              |
| $CopiesIn(e, B) \triangleq \text{IF } BagIn(e, B) \text{ THEN } B[e] \text{ ELSE } 0$                                                                                                                                                                                                                                  | The number of copies of $e$ in $B$ .                             |

Figure 17.3: The standard *Bags* module.

## 17.4 The Numbers Modules

The usual sets of numbers and operators on them are defined in the three modules *Naturals*, *Integers*, and *Reals*. These modules are tricky because their definitions must be consistent. A module  $M$  might extend both the *Naturals* module and another module that extends the *Reals* module. The module  $M$  thereby obtains two definitions of an operator such as  $+$ , one from *Naturals* and one from *Reals*. These two definitions of  $+$  must be the same. To make them the same, we have them both come from the definition of  $+$  in a module *ProtoReals*, which is locally instantiated by both *Naturals* and *Reals*.

The *Naturals* module defines the following operators:

|              |                |     |        |       |        |                  |
|--------------|----------------|-----|--------|-------|--------|------------------|
| $+$          | $*$            | $<$ | $\leq$ | $Nat$ | $\div$ | integer division |
| $-$          | $\wedge$       | $>$ | $\geq$ | $..$  | $\%$   |                  |
| binary minus | exponentiation |     |        |       |        |                  |

Except for  $\div$ , these operators are all either standard or explained in Chapter 2. We define integer division  $\div$  and modulus  $\%$  so that for any integer  $a$  and positive integer  $b$ :

$$a \% b \in 0 .. (b - 1) \quad a = b * (a \div b) + (a \% b)$$

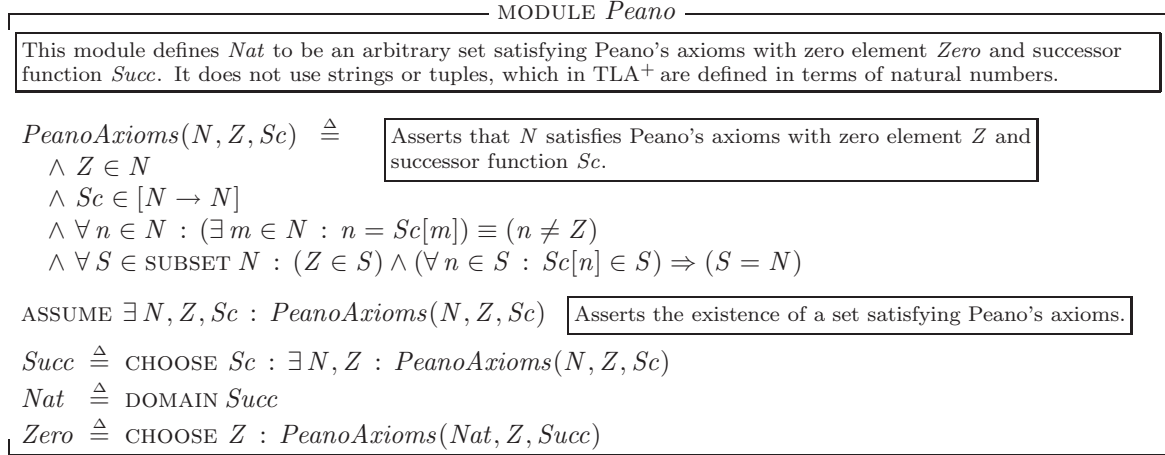
The *Integers* module extends the *Naturals* module and also defines the set *Int* of integers and unary minus ( $-$ ). The *Reals* module extends *Integers* and introduces the set *Real* of real numbers and ordinary division ( $/$ ). In mathematics, (unlike programming languages), integers are real numbers. Hence, *Nat* is a subset of *Int*, which is a subset of *Real*.

The *Reals* module also defines the special value *Infinity*. *Infinity*, which represents a mathematical  $\infty$ , satisfies the following two properties:

$$\forall r \in Real : -Infinity < r < Infinity \quad -(-Infinity) = Infinity$$

The precise details of the number modules are of no practical importance. When writing specifications, you can just assume that the operators they define have their usual meanings. If you want to prove something about a specification, you can reason about numbers however you want. Tools like model checkers and theorem provers that care about these operators will have their own ways of handling them. The modules are given here mainly for completeness. They can also serve as models if you want to define other basic mathematical structures. However, such definitions are rarely necessary for engineering specifications.

The set *Nat* of natural numbers, with its zero element and successor function is defined in the *Peano* module, which appears in Figure 17.4 on the next page. It simply defines the naturals to be a set satisfying Peano's axioms [?]. This definition is separated into its own module for the following reason. As explained in Section 15.1.9 (page 290) and Section 15.1.10 (page 291), the meanings of tuples and strings are defined in terms of the natural numbers. The *Peano* module,

Figure 17.4: The *Peano* module.

which defines the natural numbers, does not use tuples or strings. Hence, there is no circularity.

Most of the definitions in the number modules come from module *ProtoReals* in Figures 17.5 and 17.6 on the following two pages. To define the real numbers, it uses the well-known mathematical result that the reals are uniquely defined, up to isomorphism, as an ordered field in which every subset bounded from above has a least upper bound. The details will be of interest only to mathematically sophisticated readers who are curious about the formalization of ordinary mathematics. I hope that those readers will be as impressed as I am by how easy this formalization is—once you understand the mathematics.

Given the *ProtoReals* module, the rest is easy. The *Naturals*, *Integers*, and *Reals* modules appear in Figures 17.7–17.9 on page 332. Perhaps the most striking thing about them is the ugliness of an operator like  $R\!+$ , which is the version of  $+$  obtained by instantiating *ProtoReals* under the name *R*. It demonstrates that you should not use infix operators when writing a module that may be used with a named instantiation.

---

 MODULE *ProtoReals*


---

This module provides the basic definitions for the *Naturals*, *Integers*, and *Reals* module. It does this by defining the real numbers to be a complete ordered field containing the naturals.

EXTENDS *Peano*

*IsModelOfReals*(*R*, *Plus*, *Times*, *Leq*)  $\triangleq$

Asserts that *R* satisfies the properties of the reals with  $a + b = \text{Plus}[a, b]$ ,  $a * b = \text{Times}[a, b]$ , and  $(a \leq b) = (\langle a, b \rangle \in \text{Leq})$ . (We will have to quantify over the arguments, so they must be values, not operators.)

LET *IsAbelianGroup*(*G*, *Id*,  $- + -$ )  $\triangleq$

Asserts that *G* is an Abelian group with identity *Id* and group operation  $+$ .

$\wedge \text{Id} \in G$

$\wedge \forall a, b \in G : a + b \in G$

$\wedge \forall a \in G : \text{Id} + a = \text{Id}$

$\wedge \forall a, b, c \in G : (a + b) + c = a + (b + c)$

$\wedge \forall a \in G : \exists \text{minus} a \in G : a + \text{minus} a = \text{Id}$

$\wedge \forall a, b \in G : a + b = b + a$

$a + b \triangleq \text{Plus}[a, b]$

$a * b \triangleq \text{Times}[a, b]$

$a \leq b \triangleq \langle a, b \rangle \in \text{Leq}$

*RPos*  $\triangleq \{r \in R \setminus \{\text{Zero}\} : \text{Zero} \leq r\}$

*Dist*(*a*, *b*)  $\triangleq$  CHOOSE  $r \in RPos \cup \{\text{Zero}\} : \wedge a = b + r$   
 $\wedge b = a + r$  *D*(*a*, *b*) equals  $|a - b|$ .

IN  $\wedge \text{Nat} \subseteq R$

$\wedge \forall n \in \text{Nat} : \text{Succ}[n] = n + \text{Succ}[\text{Zero}]$

$\wedge \text{IsAbelianGroup}(R, \text{Zero}, +)$

$\wedge \text{IsAbelianGroup}(R \setminus \{\text{Zero}\}, \text{Succ}[\text{Zero}], *)$

$\wedge \forall a, b, c \in R : a * (b + c) = (a * b) + (a * c)$

$\wedge \forall a, b \in R : \wedge (a \leq b) \vee (b \leq a)$   
 $\wedge (a \leq b) \wedge (b \leq a) \equiv (a = b)$

$\wedge \forall a, b, c \in R : \wedge (a \leq b) \wedge (b \leq c) \Rightarrow (a \leq c)$   
 $\wedge (a \leq b) \Rightarrow \wedge (a + c) \leq (b + c)$   
 $\wedge (\text{Zero} \leq c) \Rightarrow (a * c) \leq (b * c)$

$\wedge \forall S \in \text{SUBSET } R :$

LET *SBound*(*a*)  $\triangleq \forall s \in S : s \leq a$

IN  $(\exists a \in R : \text{SBound}(a)) \Rightarrow$

$(\exists \text{sup} \in R : \wedge \text{SBound}(\text{sup})$

$\wedge \forall a \in R : \text{SBound}(a) \Rightarrow (\text{sup} \leq a))$

The first two conjuncts assert that *Nat* is embedded in *R*.

The next three conjuncts assert that *R* is a field.

The next two conjuncts assert that *R* is an ordered field.

The last conjunct asserts that every subset *S* of *R* bounded from above has a least upper bound *sup*.

THEOREM  $\exists R, \text{Plus}, \text{Times}, \text{Leq} : \text{IsModelOfReals}(R, \text{Plus}, \text{Times}, \text{Leq})$

*RM*  $\triangleq$  CHOOSE *RM* : *IsModelOfReals*(*RM.R*, *RM.Plus*, *RM.Times*, *RM.Leq*)

*Real*  $\triangleq \text{RM.R}$

---

Figure 17.5: The *ProtoReals* module (beginning).

We will define *Infinity*,  $<$ , and  $-$  so  $-Infinity < r < Infinity$ , for any  $r \in Real$ , and  $-(-Infinity) = Infinity$ .

$Infinity \triangleq \text{CHOOSE } x : x \notin Real$

$MinusInfinity \triangleq \text{CHOOSE } x : x \notin Real \cup \{Infinity\}$

*Infinity* and *MinusInfinity* (which will equal  $-Infinity$ ) are chosen to be arbitrary values not in *Real*.

$a + b \triangleq RM.Plus[a, b]$

$a * b \triangleq RM.Times[a, b]$

$a \leq b \triangleq \text{CASE } (a \in Real) \wedge (b \in Real) \rightarrow \langle a, b \rangle \in RM.Leq$   
 $\square (a = Infinity) \wedge (b \in Real \cup \{MinusInfinity\}) \rightarrow \text{FALSE}$   
 $\square (a \in Real \cup \{MinusInfinity\}) \wedge (b = Infinity) \rightarrow \text{TRUE}$   
 $\square a = b \rightarrow \text{TRUE}$

$a - b \triangleq \text{CASE } (a \in Real) \wedge (b \in Real) \rightarrow \text{CHOOSE } c \in Real : c + b = a$   
 $\square (a \in Real) \wedge (b = Infinity) \rightarrow MinusInfinity$   
 $\square (a \in Real) \wedge (b = MinusInfinity) \rightarrow Infinity$

$a/b \triangleq \text{CHOOSE } c \in Real : a = b * c$

$Int \triangleq Nat \cup \{Zero - n : n \in Nat\}$

We define  $a^b$  (exponentiation) for  $a > 0$ ,  $a \neq 0$  and  $b \in Int$ , or  $a \leq 0$  and  $b > 0$  by the four axioms:  
 $a^1 = a \quad a^{m+n} = a^m * a^n \text{ if } a \neq 0 \text{ and } m, n \in Int \quad 0^b = 0 \text{ if } b > 0 \quad a^{b*c} = (a^b)^c \text{ if } a > 0$   
 plus the continuity condition that  $0 < a$  and  $0 < b \leq c$  imply  $a^b \leq a^c$ .

$a^b \triangleq \text{LET } RPos \triangleq \{r \in Real \setminus \{Zero\} : Zero \leq r\}$

$exp \triangleq \text{CHOOSE } f \in [(RPos \times Real) \cup ((Real \setminus \{0\}) \times Int) \cup (\{0\} \times RPos) \rightarrow Real] :$   
 $\wedge \forall r \in Real : \wedge f[r, Succ[Zero]] = r$   
 $\wedge \forall m, n \in Int : f[r, m + n] = f[r, m] * f[r, n]$   
 $\wedge \forall r \in RPos : \wedge f[Zero, r] = 0$   
 $\wedge \forall s, t \in Real : f[r, s * t] = f[f[r, s], t]$   
 $\wedge \forall s, t \in RPos : (s \leq t) \Rightarrow (f[r, s] \leq f[r, t])$

IN  $exp[a, b]$

Figure 17.6: The *ProtoReals* module (end).

```

MODULE Naturals
LOCAL  $R \triangleq$  INSTANCE ProtoReals
 $Nat \triangleq R!Nat$ 
 $a + b \triangleq a R!+ b$ 
 $a - b \triangleq a R!- b$ 
 $a * b \triangleq a R!* b$ 
 $a^b \triangleq a R!^ b$ 
 $a \leq b \triangleq a R!\leq b$ 
 $a \geq b \triangleq b \leq a$ 
 $a < b \triangleq (a \leq b) \wedge (a \neq b)$ 
 $a > b \triangleq b < a$ 
 $a .. b \triangleq \{i \in Nat : (a \leq i) \wedge (i \leq b)\}$ 
 $a \div b \triangleq$  CHOOSE  $n \in R!Int : \exists r \in 0 .. (b - 1) : a = b * n + r$ 
 $a \% b \triangleq a - b * (a \div b)$ 

```

$R!+$  is the operator  $+$  defined in module *ProtoReals*.

$a^b$  is written in ASCII as `a^b`.

We define  $\div$  and  $\%$  so that  
 $a = b * (a \div b) + (a \% b)$   
for all integers  $a$  and  $b$  with  $b > 0$ .

Figure 17.7: The standard *Naturals* module.

```

MODULE Integers
EXTENDS Naturals
LOCAL  $R \triangleq$  INSTANCE ProtoReals
 $Int \triangleq R!Int$ 
 $- . a \triangleq 0 - a$ 

```

The *Naturals* module already defines operators like  $+$  to work on all real numbers.

Unary  $-$  is written `-.` when being defined or used as an operator argument.

Figure 17.8: The standard *Integers* module.

```

MODULE Reals
EXTENDS Integers
LOCAL  $R \triangleq$  INSTANCE ProtoReals
 $Real \triangleq R!Real$ 
 $a/b \triangleq a R!/ b$ 
 $Infinity \triangleq R!Infinity$ 

```

The *Integers* module already defines operators like  $+$  to work on all real numbers.

$R! /$  is the operator  $/$  defined in module *ProtoReals*.

Figure 17.9: The standard *Reals* module.

Part V

Appendix



## Appendix A

# The ASCII Specifications

### A.1 The Asynchronous Interface

```
----- MODULE AsynchInterface -----
EXTENDS Naturals
CONSTANT Data
VARIABLES val, rdy, ack

TypeInvariant == /\ val \in Data
                  /\ rdy \in {0, 1}
                  /\ ack \in {0, 1}
-----

Init == /\ val \in Data
         /\ rdy \in {0, 1}
         /\ ack = rdy

Send == /\ rdy = ack
        /\ val' \in Data
        /\ rdy' = 1 - rdy
        /\ UNCHANGED ack

Rcv  == /\ rdy # ack
        /\ ack' = 1 - ack
        /\ UNCHANGED <<val, rdy>>

Next == Send \/ Rcv

Spec == Init /\ [] [Next]_<<val, rdy, ack>>
-----
```

THEOREM Spec => []TypeInvariant

```
=====

----- MODULE Channel -----
EXTENDS Naturals
CONSTANT Data
VARIABLE chan

TypeInvariant ==
  chan \in [val : Data, rdy : {0, 1}, ack : {0, 1}]
-----

Init == /\ TypeInvariant
        /\ chan.ack = chan.rdy

Send(d) == /\ chan.rdy = chan.ack
            /\ chan' = [chan EXCEPT !.val = d, !.rdy = 1 - @]

Rcv == /\ chan.rdy # chan.ack
        /\ chan' = [chan EXCEPT !.ack = 1 - @]

Next == (\E d \in Data : Send(d)) \/ Rcv

Spec == Init /\ [] [Next]_chan
-----

THEOREM Spec => []TypeInvariant
=====
```

## A.2 A FIFO

```
----- MODULE InnerFIFO -----
EXTENDS Naturals, Sequences
CONSTANT Message
VARIABLES in, out, q
InChan == INSTANCE Channel WITH Data <- Message, chan <- in
OutChan == INSTANCE Channel WITH Data <- Message, chan <- out
-----

Init == /\ InChan!Init
        /\ OutChan!Init
        /\ q = << >>

TypeInvariant == /\ InChan!TypeInvariant
```

```

/\ OutChan!TypeInvariant
/\ q \in Seq(Message)

SSend(msg) == /\ InChan!Send(msg)
              /\ UNCHANGED <<out, q>>

BufRcv == /\ InChan!Rcv
          /\ q' = Append(q, in.val)
          /\ UNCHANGED out

BufSend == /\ q # << >>
           /\ OutChan!Send(Head(q))
           /\ q' = Tail(q)
           /\ UNCHANGED in

RRcv == /\ OutChan!Rcv
        /\ UNCHANGED <<in, q>>

Next == \/ \E msg \in Message : SSend(msg)
        \/ BufRcv
        \/ BufSend
        \/ RRcv

Spec == Init /\ [] [Next]_<<in, out, q>>
-----
THEOREM Spec => []TypeInvariant
=====

----- MODULE FIFO -----
CONSTANT Message
VARIABLES in, out
Inner(q) == INSTANCE InnerFIFO
Spec == \EE q : Inner(q)!Spec
=====

```

## A.3 A Caching Memory

```

----- MODULE MemoryInterface -----
VARIABLE memInt
CONSTANTS Send(_, _, _, _),
           Reply(_, _, _, _),

```

```

        InitMemInt,
        Proc,
        Adr,
        Val
ASSUME \A p, d, miOld, miNew :
        /\ Send(p,d,miOld,miNew) \in BOOLEAN
        /\ Reply(p,d,miOld,miNew) \in BOOLEAN
-----
MReq == [op : {"Rd"}, adr: Adr]
        \cup [op : {"Wr"}, adr: Adr, val : Val]

NoVal == CHOOSE v : v \notin Val
=====

----- MODULE InternalMemory -----
EXTENDS MemoryInterface
VARIABLES mem, ctl, buf
-----
IIInit == /\ mem \in [Adr->Val]
          /\ ctl = [p \in Proc |-> "rdy"]
          /\ buf = [p \in Proc |-> NoVal]
          /\ memInt \in InitMemInt

TypeInvariant ==
  /\ mem \in [Adr->Val]
  /\ ctl \in [Proc -> {"rdy", "busy", "done"}]
  /\ buf \in [Proc -> MReq \cup Val \cup {NoVal}]

```

```

Req(p) == /\ ctl[p] = "rdy"
          /\ \E req \in MReq :
              /\ Send(p, req, memInt, memInt')
              /\ buf' = [buf EXCEPT ![p] = req]
              /\ ctl' = [ctl EXCEPT ![p] = "busy"]
          /\ UNCHANGED mem

Do(p) ==
  /\ ctl[p] = "busy"
  /\ mem' = IF buf[p].op = "Wr"
            THEN [mem EXCEPT ![buf[p].adr] = buf[p].val]
            ELSE mem
  /\ buf' = [buf EXCEPT ![p] = IF buf[p].op = "Wr"
            THEN NoVal
            ELSE mem[buf[p].adr]]
  /\ ctl' = [ctl EXCEPT ![p] = "done"]
  /\ UNCHANGED memInt

Rsp(p) == /\ ctl[p] = "done"
          /\ Reply(p, buf[p], memInt, memInt')
          /\ ctl' = [ctl EXCEPT ![p] = "rdy"]
          /\ UNCHANGED <<mem, buf>>

INext == \E p \in Proc: Req(p) \/ Do(p) \/ Rsp(p)

ISpec == IInit /\ [] [INext]_<<memInt, mem, ctl, buf>>
-----
THEOREM ISpec => []TypeInvariant
=====

----- MODULE Memory -----
EXTENDS MemoryInterface
Inner(mem, ctl, buf) == INSTANCE InternalMemory
Spec == \EE mem, ctl, buf : Inner(mem, ctl, buf)!ISpec
=====

```

```

----- MODULE WriteThroughCache -----
EXTENDS Naturals, Sequences, MemoryInterface
VARIABLES mem, ctl, buf, cache, memQ
CONSTANT QLen
ASSUME (QLen \in Nat) /\ (QLen > 0)
M == INSTANCE InternalMemory
-----

Init == /\ M!Init
        /\ cache = [p \in Proc |-> [a \in Adr |-> NoVal] ]
        /\ memQ = << >>

TypeInvariant ==
  /\ mem   \in [Adr -> Val]
  /\ ctl   \in [Proc -> {"rdy", "busy", "waiting", "done"}]
  /\ buf   \in [Proc -> MReq \cup Val \cup {NoVal}]
  /\ cache \in [Proc -> [Adr -> Val \cup {NoVal}]]
  /\ memQ  \in Seq(Proc \X MReq)

Coherence == \A p, q \in Proc, a \in Adr :
              (NoVal \notin {cache[p][a], cache[q][a]})
              => (cache[p][a]=cache[q][a])
-----

Req(p) == M!Req(p) /\ UNCHANGED <<cache, memQ>>
Rsp(p) == M!Rsp(p) /\ UNCHANGED <<cache, memQ>>

RdMiss(p) == /\ (ctl[p] = "busy") /\ (buf[p].op = "Rd")
              /\ cache[p][buf[p].adr] # NoVal
              /\ Len(memQ) < QLen
              /\ memQ' = Append(memQ, <<p, buf[p]>>)
              /\ ctl' = [ctl EXCEPT ![p] = "waiting"]
              /\ UNCHANGED <<memInt, mem, buf, cache>>

DoRd(p) ==
  /\ ctl[p] \in {"busy", "waiting"}
  /\ buf[p].op = "Rd"
  /\ cache[p][buf[p].adr] # NoVal
  /\ buf' = [buf EXCEPT ![p] = cache[p][buf[p].adr]]
  /\ ctl' = [ctl EXCEPT ![p] = "done"]
  /\ UNCHANGED <<memInt, mem, cache, memQ>>

DoWr(p) ==
  LET r == buf[p]
  IN /\ (ctl[p] = "busy") /\ (r.op = "Wr")

```

```

/\ Len(memQ) < QLen
/\ cache' = [q \in Proc |->
    IF (p=q) /\ (cache[q][r.adr]#NoVal)
    THEN [cache[q] EXCEPT ![r.adr] = r.val]
    ELSE cache[q] ]
/\ memQ' = Append(memQ, <<p, r>>)
/\ buf' = [buf EXCEPT ![p] = NoVal]
/\ ctl' = [ctl EXCEPT ![p] = "done"]
/\ UNCHANGED <<memInt, mem>>

vmem ==
  LET f[i \in 0 .. Len(memQ)] ==
    IF i=0 THEN mem
    ELSE IF memQ[i][2].op = "Rd"
    THEN f[i-1]
    ELSE [f[i-1] EXCEPT ![memQ[i][2].adr] =
        memQ[i][2].val]
  IN f[Len(memQ)]

MemQWr == LET r == Head(memQ)[2]
  IN /\ (memQ # << >>) /\ (r.op = "Wr")
  /\ mem' = [mem EXCEPT ![r.adr] = r.val]
  /\ memQ' = Tail(memQ)
  /\ UNCHANGED <<memInt, mem, buf, ctl, cache>>

MemQRd ==
  LET p == Head(memQ)[1]
  r == Head(memQ)[2]
  IN /\ (memQ # << >>) /\ (r.op = "Rd")
  /\ memQ' = Tail(memQ)
  /\ cache' = [cache EXCEPT ![p][r.adr] = vmem[r.adr]]
  /\ UNCHANGED <<memInt, mem, buf, ctl>>

Evict(p,a) == /\ (ctl[p] = "waiting") => (buf[p].adr # a)
  /\ cache' = [cache EXCEPT ![p][a] = NoVal]
  /\ UNCHANGED <<memInt, mem, buf, ctl, memQ>>

Next == /\ \E p\in Proc : /\ Req(p) /\ Rsp(p)
  /\ RdMiss(p) /\ DoRd(p) /\ DoWr(p)
  /\ \E a \in Adr : Evict(p, a)
  /\ MemQWr /\ MemQRd

Spec ==

```

```

Init /\ [] [Next]_<<memInt, mem, buf, ctl, cache, memQ>>
-----
THEOREM Spec => [] (TypeInvariant /\ Coherence)
-----
LM == INSTANCE Memory
THEOREM Spec => LM!Spec
=====

```

## A.4 The Alternating Bit Protocol

```

----- MODULE AlternatingBit -----
EXTENDS Naturals, Sequences
CONSTANTS Data
VARIABLES msgQ, ackQ, sBit, sAck, rBit, sent, rcvd

ABInit == /\ msgQ = << >>
          /\ ackQ = << >>
          /\ sBit \in {0, 1}
          /\ sAck = sBit
          /\ rBit = sBit
          /\ sent = << >>
          /\ rcvd = << >>

TypeInv == /\ msgQ \in Seq({0,1} \X Data)
          /\ ackQ \in Seq({0,1})
          /\ sBit \in {0, 1}
          /\ sAck \in {0, 1}
          /\ rBit \in {0, 1}
          /\ sent \in Seq(Data)
          /\ rcvd \in Seq(Data)

SndNewValue(d) == /\ sAck = sBit
                 /\ sent' = Append(sent, d)
                 /\ sBit' = 1 - sBit
                 /\ msgQ' = Append(msgQ, <<sBit', d>>)
                 /\ UNCHANGED <<ackQ, sAck, rBit, rcvd>>

ReSndMsg ==
  /\ sAck # sBit
  /\ msgQ' = Append(msgQ, <<sBit, sent[Len(sent)]>>)
  /\ UNCHANGED <<ackQ, sBit, sAck, rBit, sent, rcvd>>

```

```

RcvMsg ==
  /\ msgQ # <<>>
  /\ msgQ' = Tail(msgQ)
  /\ rBit' = Head(msgQ)[1]
  /\ rcvd' = IF rBit' # rBit THEN Append(rcvd, Head(msgQ)[2])
              ELSE rcvd
  /\ UNCHANGED <<ackQ, sBit, sAck, sent>>

SndAck == /\ ackQ' = Append(ackQ, rBit)
          /\ UNCHANGED <<msgQ, sBit, sAck, rBit, sent, rcvd>>

RcvAck == /\ ackQ # << >>
          /\ ackQ' = Tail(ackQ)
          /\ sAck' = Head(ackQ)
          /\ UNCHANGED <<msgQ, sBit, rBit, sent, rcvd>>

Lose(c) ==
  /\ c # << >>
  /\ \E i \in 1..Len(c) :
    c' = [j \in 1..(Len(c)-1) |-> IF j \leq i THEN c[j]
                                             ELSE c[j+1] ]
  /\ UNCHANGED <<sBit, sAck, rBit, sent, rcvd>>

LoseMsg == Lose(msgQ) /\ UNCHANGED ackQ

LoseAck == Lose(ackQ) /\ UNCHANGED msgQ

ABNext == \/ \E d \in Data : SndNewValue(d)
          \/ ReSndMsg \/ RcvMsg \/ SndAck \/ RcvAck
          \/ LoseMsg \/ LoseAck

vars == << msgQ, ackQ, sBit, sAck, rBit, sent, rcvd>>
-----
Spec == ABInit /\ [] [ABNext]_vars
-----
Inv == /\ Len(rcvd) \in {Len(sent)-1, Len(sent)}
      /\ \A i \in 1..Len(rcvd) : rcvd[i] = sent[i]

THEOREM Spec => [] Inv
=====

----- MODULE MCAalternatingBit -----

```

```
EXTENDS AlternatingBit
CONSTANTS msgQLen, ackQLen, sentLen
```

```
ASSUME /\ msgQLen \in Nat
        /\ ackQLen \in Nat
        /\ sentLen \in Nat
```

```
SeqConstraint == /\ Len(msgQ) \leq msgQLen
                  /\ Len(ackQ) \leq ackQLen
                  /\ Len(sent) \leq sentLen
```

```
=====
```

# Bibliography

- [1] Donald E. Knuth. *The T<sub>E</sub>Xbook*. Addison-Wesley, Reading, Massachusetts, 1994.
- [2] Leslie Lamport. The temporal logic of actions. *ACM Transactions on Programming Languages and Systems*, 16(3):872–923, May 1994.
- [3] A. C. Leisenring. *Mathematical Logic and Hilbert's  $\varepsilon$ -Symbol*. Gordon and Breach, New York, 1969.