

Embedding an Object Calculus
in the
Unifying Theories of Programming

Michael Anthony Smith
Kellogg College

A thesis submitted in partial fulfilment
of the requirements for the degree of
Doctor of Philosophy in Software Engineering
at the University of Oxford,
Hilary Term 2009.

Abstract

Hoare and He's Unifying Theories of Programming (UTP) provides a rich model of programs as relational predicates. This theory is intended to provide a single framework in which any programming paradigms, languages, and features, can be modelled, compared and contrasted. The UTP already has models for several programming formalisms, such as imperative programming, higher-order programming (e.g. programming with procedures), several styles of concurrent programming (or reactive systems), class-based object-orientation, and transaction processing. We believe that the UTP ought to be able to represent all significant computer programming language formalisms, in order for it to be considered a unifying theory.

One gap in the UTP work is that of object-based object-orientation, such as that presented in Abadi and Cardelli's untyped object calculi (ζ -calculi). These ζ -calculi provide a prominent formalism of object-based object-oriented (OO) programs, which models programs as objects. We address this gap within this dissertation by presenting an embedding of an Abadi–Cardelli-style object calculus in the UTP. More formally, the thesis that this dissertation argues is that it is possible to provide an object-based object orientation to the UTP, with value- and reference-based objects, and a fully abstract model of references.

We have made three contributions to our area of study: first, to extend the UTP with a notion of object-based object orientation, in contrast with the existing class-based models; second, to provide an alternative model of pointers (references) for the UTP that supports both value-based compound values (e.g. objects) and references (pointers), in contrast to existing UTP models with pointers that have reference-based compound values; and third, to model an Abadi–Cardelli notion of an object in the UTP, and thus demonstrate that it can unify this style of object formalism.

Acknowledgements

I have had considerable support and encouragement whilst endeavouring to complete this part-time DPhil, since I started in September 1999. In particular I would like to thank:

- Professor Jim Woodcock (my original supervisor) and Professor Colin O'Halloran (my boss) for their support and encouragement to take on the challenge of a part-time DPhil in Software Engineering.
- Dr. Jeremy Gibbons (my current supervisor) for his patience, direction, and encouragement throughout the eight years I have been under his guidance.
- Nick Moffat my colleague – and fellow part-time DPhil student – for his willingness to provide detailed reviews of my paper submissions and discuss issues concerning my studies.
- Tom McCutcheon for his kind support and advice at a particularly low point in my DPhil. He kept my thoughts in perspective and enabled me to make the decision to go genuinely part-time at work.
- My work colleagues for their consideration and support. In particular, for allowing me to adapt my part-time working hours so that I could regularly study in blocks of one week, rather than a couple of days at a time.
- My church and bible study groups, for keeping me grounded.
- My parents (Pat and Dave Smith) for their faith in my ability and constant encouragement.
- The Oxford University Department for Continuing Education and Kellogg College staff, for their endless patience and willingness to help me through the University regulations and paper work.
- My examiners (Dr. Andrew Simpson and Dr. Andrew Butterfield) who provided many useful comments that have significantly improved the presentation of this thesis.

My DPhil fees were kindly paid for by the UK Ministry of Defence's "Defence Evaluation and Research Agency" and then by QinetiQ (following privatisation in July 2002). I would also like to thank Kellogg College for their bursary that enabled me to fund attendance and participation in three European conferences.

Contents

1	Introduction	1
1.1	Context and Motivation	1
1.2	What is Object Orientation?	2
1.3	Formal Models of Objects and Object Orientation	3
1.4	A Brief Introduction to the UTP and the Object Calculus	4
1.5	Thesis	4
1.5.1	Approach	4
1.5.2	Decisions taken	5
1.5.3	Contributions	5
1.5.4	Publications	6
1.5.5	Scope	6
2	Background	7
2.1	Common Notation and Conventions	7
2.1.1	The use of colour	7
2.1.2	Repeated terms	7
2.1.3	Identifying subterms	8
2.1.4	Meta-variable identifiers	8
2.2	Object Calculus	9
2.2.1	Abadi–Cardelli untyped object calculus	10
2.2.2	Basic extensions	11
2.2.3	Heap extensions	12
2.2.4	Remaining extensions	12
2.2.5	Substitution definition and object calculus examples	13
2.3	Unifying Theories of Programming	15
2.3.1	Relational model	15
2.3.2	Design model	17
2.3.3	UTP design commands	18
2.3.4	Hiding and filtering variables from programs	21
2.3.5	Higher order programming	23
2.4	Heap Storage Locations	24
3	Unifying Theories of Objects	25
3.1	Extending the UTP design	25
3.1.1	Objects and labels	26
3.1.2	Methods and functions	26
3.1.3	Abstract locations and the heap	27
3.1.4	Modelling notes	28
3.2	Operand Stack Model	30
3.2.1	Literal value programs	30
3.2.2	Function values	31

3.2.3	Object values and method definitions	31
3.2.4	Command expressions	31
3.2.5	Method updates and field assignments	32
3.2.6	Method and function invocation	33
3.2.7	Other programming mechanisms	35
3.2.8	Modelling the heap operations	36
3.3	Result Value Model	37
3.3.1	Operand values	38
3.3.2	Command expressions	38
3.3.3	Function and method invocation	39
3.3.4	Modelling the heap operations	39
3.3.5	Other mechanisms	41
3.4	Constant Map Model	41
3.4.1	Procedure definition	41
3.4.2	Method and function invocation	42
3.4.3	Alternative method and function definition compilation scheme	43
3.5	Summary	44
4	Object Model Consistency	46
4.1	Dynamic Operations	46
4.1.1	Variable substitution compilation scheme	47
4.1.2	The heap context compilation scheme	47
4.2	Structural Induction	48
4.2.1	Commuting Diagrams	48
4.2.2	Command expressions	49
4.2.3	Procedure invocation	50
4.2.4	Heap operations	52
4.2.5	Other programming mechanisms	53
4.2.6	Proof notation	54
4.3	Operand Stack Model	55
4.3.1	Command expressions	55
4.3.2	Procedure invocation	56
4.3.3	Heap operations	57
4.3.4	Other programming operations	58
4.4	Result Variable Model	60
4.4.1	Command expressions	60
4.4.2	Heap operations	62
4.4.3	Other programming operations	63
4.5	Constant Map Model	64
4.5.1	Procedure invocation	64
4.5.2	Other operations	66
4.6	Factorial Example	66
4.6.1	Notation	67
4.6.2	Object calculus iteration	67
4.6.3	UTP operand-stack design iteration	68
4.7	Summary	69
5	Unifying Theories of Locations	70
5.1	Locations	70
5.1.1	Terminology	71
5.1.2	Scope	71
5.1.3	Family tree	71

5.2	Concrete Representation	72
5.2.1	Concrete value notation	72
5.2.2	Concrete location heap	73
5.2.3	Concrete location model	74
5.2.4	Paths and their operations	76
5.3	Abstract Model	76
5.3.1	Graph abstraction	77
5.3.2	Traces	77
5.3.3	Trace-based node	78
5.3.4	Trace-based graph	80
5.4	UTP Model	83
5.4.1	Healthiness conditions	83
5.4.2	Operations	84
5.4.3	Full abstraction	84
5.5	Summary	85
6	Combining Objects and Locations	86
6.1	UTP Designs with Locations	86
6.1.1	Healthiness conditions	87
6.1.2	Operations	89
6.2	Location Graph Model of the Object Calculus	92
6.2.1	Value representation	92
6.2.2	Model definition	94
6.2.3	Cell – Method Invocation Example	98
6.2.4	Feasibility	104
6.3	Summary	105
7	Discussion	106
7.1	Summary	106
7.2	Predicative Objects	108
7.3	Abstract Location Model	110
7.4	Future Directions	112
7.5	Reflection	113
7.6	Conclusions	114
	References	115
A	Concepts and Terminology	125
A.1	Semantics	126
A.1.1	Semantic approaches	126
A.1.2	State values	127
A.1.3	Observable behaviour	128
A.1.4	Concurrency	129
A.2	Object-Oriented Infrastructure	130
A.2.1	Objects	130
A.2.2	Polymorphism	131
A.2.3	Program structure (composability)	132
A.2.4	Code reuse	133
A.3	Aliasing	134

B	Classes in the Object Calculus	136
B.1	Static Approaches to Sharing Methods	136
B.1.1	Prototype approach	136
B.1.2	Naive class-based approach	137
B.1.3	Pre-processed class-based approach	137
B.1.4	Sharing method instances	138
B.1.5	Caching shared methods	138
B.2	Dynamic Approaches to Sharing Methods	139
B.2.1	Simulating dynamic method sharing	139
B.2.2	Directly supporting dynamic method sharing	140
C	UTP Design Laws	142
C.1	Closed Design Laws	142
C.1.1	Conditional binding	142
C.1.2	Sequential composition	143
C.1.3	Variable introduction and elimination	145
C.1.4	Hiding	146
C.1.5	Block scoped variable declaration	147
C.2	General Design Laws	149
C.2.1	Skip the unit of sequential composition	149
C.2.2	Scoping variables	149
C.2.3	Assignment composition	150
C.2.4	Conjunctive implication	153
C.3	Object Model Design Laws and Example Proofs	153
C.3.1	Extract identity	153
C.3.2	Factorial example UTP correspondence proofs	153
D	Combined model extended example	157
	List of Abbreviations	161

This page is intentionally blank

Chapter 1

Introduction

1.1 Context and Motivation

QinetiQ’s Systems Assurance (SA) team uses a variety of formal method tools for reasoning about existing systems, such as the ProofPower theorem prover [Lem00] and the FDR (Failures-Divergence and Refinement) model checker [Fse05, Sca98]. The successful use of these tools on industrial scale projects has highlighted a variety of issues, such as the difficulty in translating results from one formal framework into another. One reason for this is that the notion of program correctness in one formal framework can be subtly different from the notion of correctness in another framework, as illustrated by Bolton’s work on the refinement of state-based and event-based models [Bol02, BL05].

Within the SA team it was proposed that Hoare and He’s Unifying Theories of Programming (UTP) [HH98] be adopted as a common underpinning framework. This solution is in principle fine for the majority of the team’s current tools, as they are designed to model either imperative or functional programs, possibly within a concurrent context (e.g. multiple threads). However, the support for object-oriented (OO) programs within the UTP is currently limited to class-based object orientation [HLL02a, HLL06, CSW03, CSW05], with the exception of our own work [SG07] (and this dissertation). These class-based models of object orientation support languages with single inheritance, where subtyping is subclassing.

In the slightly wider field of predicative programming, Kassios has introduced the notion of an object-specification [Kas04], which is a generalisation of a class, rather than an object. Here both objects and classes are defined as object-specifications that meet certain requirements. This work builds upon Hehner’s work on programming as “boolean expressions” [Heh93].

In contrast to the above work, our model of object orientation directly supports an object-based OO-paradigm, rather than a class-based OO-paradigm. Having said this, it is straightforward to use an object to define a class, as is discussed in [Appendix B](#). Further, object references are handled in our model via a *heap*, which is initially represented as a map from shareable locations to values.

Aside 1.1: It would be possible to base the modelling of objects on that of processes, where an object is seen as a special sort of process. Such objects would inherit the underlying process’ model of concurrency. Therefore, it is likely that a process-based model of objects would focus on one notion of concurrency, at the expense of other notions of concurrency. An alternative view is to separate the modelling of objects and concurrency. In principle, this would enable the appropriate combination of object and concurrency models to be chosen for a given problem. This is the view adopted; thus as concurrency is already well catered for within the UTP, the modelling of objects is considered.

□

The UTP framework is abstract, hence any model of objects for the UTP ought to be similarly abstract. This raises the issue of how to validate the UTP’s model of objects. What is required

is a formal notion of objects that has a simple semantic model that is straightforward to understand and can be executed for the purposes of validation.

1.2 What is Object Orientation?

There are currently no generally accepted notions of some of the key aspects of object orientation, such as precisely what is meant by the terms ‘object’ and ‘inheritance’. In general, the OO-paradigm can be thought of as describing a broad family of OO-frameworks, that provide facilities for creating, manipulating, and relating objects.

From our point of view, an *object* is an entity that has both *state* and *operations* for accessing and manipulating that state. Typically, the state is stored in *fields* (or *attributes*) and the operations are provided by *methods*. Conceptually, access to the fields is via the published methods (or interface). However, in practice, it is frequently possible to enable the state of an object to be directly visible (accessible) and updateable. The important point, from our perspective, is that an object conceptually owns both the fields and operations that can be applied to it. Further, the only way of updating its state is via invocation of one of its methods (operations); here unprotected fields are considered to have implicit ‘get’ and ‘set’ methods for reading and updating their data.

Inheritance is a more tricky issue. It is used for a variety of purposes, such as *code reuse* and *subtyping*. Meyer’s taxonomy of object orientation describes ten different forms of inheritance [Mey96, Mey97]. From our point of view inheritance is not an issue that we will consider in any depth, as the inheritance mechanisms can be defined (modelled) in terms of basic objects.

The most appropriate model of objects depends on what are considered to be *true* OO-features, as illustrated by the following example.

Example 1.2: Languages such as CLOS (Common Lisp Object System) [DG87] are considered by some, but not all, to be examples of the OO-paradigm. Essentially, this is because it is an example of a multi-method language, thus has “no receiver of a message” [Bru02]. Here the methods are treated more like procedures, which may be syntactically associated with a class of objects, but are conceptually collected into a central pool. This pool is referred to by the dynamic method dispatch algorithm, which is used to determine the method that *best fits* the call, where a call essentially consists of the method’s name and actual parameter types. The important point is that the method selection is not based on an individual (receiver) object, but rather on a collection of objects. Hence, standard OO-notions such as *self* (i.e. the current object) are meaningless. Further, the natural OO-concepts of encapsulation and design by contract (visible interface) are also difficult to model. But other things are easier – e.g. binary methods.

□

From our point of view, method invocations in an OO-language ought to specify the object that the method is going to be invoked on (i.e. the invocation message ‘receiver’). Therefore, we do not support multi-method flavours of object orientation; Castagna [Cas97] provides a detailed presentation of the underlying theory of multi-methods.

Having discarded multi-method languages, there are essentially two remaining families of OO-language, namely class-based and object-based languages. A pure class-based language organises its objects into classes of similar objects, where a class essentially provides a template for the construction of an object. In object-based languages, new objects have to be constructed from existing objects, and modified as necessary.

Appendix A provides a more detailed examination of typical OO-concepts and briefly discusses different styles of semantic model.

1.3 Formal Models of Objects and Object Orientation

Reynold’s article on user defined types and procedural data structures [Rey75] is considered by some to be the first ‘important’ contribution to the formalisation of OO-programming [GM94, Bru02]. Having said this, there are earlier works that contribute to this area including:

- the formalisation of type systems, such as the simply typed lambda calculus [Chu40], and polymorphic lambda calculus [Rey74];
- the defacto operational formalisation of compilers (and interpreters) for existing object oriented languages, such as Simula [BDMN79] and Smalltalk [Lew95].

Other important works in the formalisation of object-orientation include Cardelli and Wegner’s on types and data abstraction [CW85], Cardelli’s semantics of inheritance [Car88], and Cook’s denotational semantics of inheritance [Coo89]. The last of these works is reprinted in [GM94], which collects what the editors consider to be the fifteen most important works on the formalisation of OO-programming prior to 1993; note that [Rey75] is also reprinted in this collection.

Currently there is an extensive collection of approaches to the formalisation of the OO-paradigm, from relatively intuitive operation descriptions based on abstract machines, through to more abstract mathematics such as higher-order logics and category theory. Here, the intuitive models can be used to validate the more abstract mathematical models, which are better suited for reasoning and proof.

A special sort of algebra, known as an initial algebra, is a traditional way of modelling an abstract data type [GTWW77, MT92]. Essentially each algebra defines an interface, which explicitly states what constants and operations are provided to generate values of that algebra’s type (i.e. carrier set). Algebraic formalisms of objects and OO-features include: Rees’ algebraic theory of software interfaces [Ree00]; and Wei’s framework for studying the semantics of object-oriented programs [Wei01].

The problem with the algebraic approach is that it focuses on how to build (or construct) an entity, rather than how the entity behaves (responds to requests for information). Co-algebras address this issue, by modelling an entity’s behaviour, rather than how it can be constructed. Co-algebraic formalisms include: Reichel’s work on object semantics [Rei95]; Jacobs and colleagues’ work on reasoning about classes and objects [Jac96, Jac98, HHJT98]; and Poll’s work on the semantics of subtyping [Pol01].

Various calculi have been proposed for the modelling of the OO-paradigm. Some are adaptations of the untyped lambda calculus (λ -calculus), such as: Castagna’s work on the unified foundation of OO-programming [Cas97]; Pierce’s work on the typed programming languages [Pie02]; and Bruce, Pierce and Cardelli’s work on comparing object encodings [BCP99]. Others are built on process calculi, notably the pi-calculus, such as: Sangiorgi and Walker’s work on objects as mobile processes [SW01, Part 7]; Jones’ work on the Pobl ($\pi o \beta \lambda$) language [Jon93]; and the PhD theses of Turner [Tur95], Hordhagen [Nor98], Lumpe [Lum99], and Merro [Mer00]. Then there are the various object calculi based on Abadi and Cardelli’s theory of objects [AC96], such as: Gordon, Hankin and Lassen’s work on imperative objects [GHL99]; Bugliesi and Crafa’s work on dynamic messages [BC99]; and Clarke’s work on object ownership and containment [Cla02].

Logics for modelling OO-programs and OO-specification languages include: Abadi and Leino’s logic of object-oriented programs [AL98]; Lano, Bicarregui and Evans’ structured axiomatic semantics for UML (Unified Modelling Language) models [LBE00]; and Smith’s logic for object-Z [Smi95].

All of the above approaches have merit, although none are considered to be the definitive definition of an object or object orientation. The Abadi and Cardelli style of object calculus was chosen as the starting point of the *base-line* semantics for validating the proposed UTP model of object orientation, as discussed in Section 1.5.2.

1.4 A Brief Introduction to the UTP and the Object Calculus

Hoare and He’s Unifying Theories of Programming (UTP) [HH98] can be used to formally define how results produced in one formal model can be translated as assumptions to another formal model. Essentially, programs are considered to be predicates that relate the values of their observable input and output variables (their alphabet). For example, the increment program $x := x + 1$ is typically defined by the relational predicate $x' = x + 1$, where: the predicate variables x and x' denote the input and output values of the program variable x ; and the set $\{x, x'\}$ denotes its alphabet. One key observation is that predicates with the same alphabets can be partially ordered by logical implication (e.g. $x' = 3 \Rightarrow x' > 0$). Here, the predicates $x' = 3$ and $x' > 0$ can be considered to represent the implementation and specification respectively. In general, this forms an implication-based refinement ordering, where **false** is the strongest element in the order and **true** the weakest.

This basic relational model has been specialised to reflect the semantics of various programming paradigms and languages, such as: imperative programs without subroutines; reactive systems for simple message-based concurrency; and class-based object orientation (as discussed in Section 1.1). Each of these specialisations introduce some *observational variables* to model various aspects of the program’s context, such as: whether the program has started or terminated; the history of past communication events; and the inheritance relationship between class definitions. Alongside these observational variables, the specialisations also include some *healthiness conditions*, which are used to constrain a healthy program’s behaviour. For example, a healthy program should not be able to alter the history of past communication events; a formalisation of this constraint is one of the healthiness conditions that applies to all the models of concurrency presented in [HH98]. The UTP is presented in more detail in Section 2.3.

Abadi and Cardelli’s untyped object calculus (ζ -calculus) [AC96] provides a mathematical foundation for objects that is defined in terms of a collection of reduction rules. These rules are typically used to form the basis of an operational semantics. One characteristic of the ζ -calculus is the explicit representation of the *self* parameter; this differs from some class-based languages, such as Java and C++, where the *self* is implicitly accessible via the *this* keyword. The self parameter mechanism allows several notions of self to be visible within the same scope (e.g. that of an inner method definition); it provides the calculus with some of its expressive power.

The ζ -calculus has been extended, by various authors, to include several other notions, such as literal values and their operations, typing, object references, functions, and object ownership. We provide our own object calculus – the $\mathcal{O}$ -calculus – which directly supports value-based objects, eagerly evaluated functions, references, literal values and their operations, conditional evaluation, local declaration blocks, and sequential composition – but not types, classes, or other class-based notions such as inheritance. The $\mathcal{O}$ -calculus is presented in Section 2.2.

1.5 Thesis

The thesis that this dissertation argues is as follows: it is possible to provide an object-based object orientation to the UTP with value- and reference-based objects, and a fully abstract model of references.

1.5.1 Approach

The high level approach we take is to: identify an existing object-based framework that can be embedded into the UTP; construct a corresponding denotational model of this framework within the UTP; demonstrate that the models are consistent; and then demonstrate that the model of references is fully abstract. An Abadi–Cardelli-style of object calculus [AC96] was chosen as the existing object framework (Section 2.2). Three UTP models of this framework

are presented in [Chapter 3](#), and shown to be consistent in [Chapter 4](#). A model of locations is introduced in [Chapter 5](#), which is used to define a fully abstract model of references (using a heap). This abstract heap is then used to construct a fourth model of the chosen object framework ([Section 6.2](#)).

1.5.2 Decisions taken

We now present seven key decisions taken during our studies. First, we limit our consideration of object-based object orientation to those languages in which one invokes a method on a specified ‘receiver’ object, rather than passing an object as the first parameter of a more general method invocation mechanism. Hence, we do not cover multi-method flavours of object orientation – such as presented by [\[Cas97\]](#).

Second, the untyped Abadi–Cardelli-style object calculus [\[AC96\]](#) was chosen as the object-based framework to embed. Essentially, this is because it provides a small framework that is computationally complete, reasonably clear, and straightforward to extend. The small framework was thought to be the best way of ensuring the feasibility of providing both a UTP encoding and a demonstration that it is consistent. The computational completeness is demonstrated by a straightforward encoding of the λ -calculus [\[AC96\]](#). The clarity is important for illustrating the language-specific definitions of OO-features. This is assisted by being able to add some straightforward extensions, which can be used to remove some of the more obscure encoding such as the Abadi–Cardelli object encoding of both conditional evaluation and the natural numbers.

Third, our variant of the Abadi–Cardelli calculus directly represents both value and reference objects. Such distinctions are important for modelling the atomicity of assignment, and for being able to control which elements of a compound structure can be directly referenced. For example, it mirrors the distinction between boxed and unboxed values with the C# language.

Fourth, the operational semantics of the object calculus is presented in a deterministic small-step reduction-rule style [\[Pie02\]](#). The consistency proofs between this small-step semantics and the UTP semantics is done as a structural induction over these steps.

Fifth, we present multiple UTP encodings of the object calculus. Here the idea is to start with simple encodings that mirror the nature of the object calculus, and gradually adapt them to remove unwanted features. This approach provides a mechanism for informally explaining why the UTP models are consistent with the object calculus.

Sixth, the last of the UTP models completely changes the way in which a program is represented, from a combination of local variables and a map from heap locations to values, to a trace-based graph of program state. In order to facilitate this change, a concrete model of the trace-based graph has been introduced to aid understanding and validate the modelling approach.

Seventh (and last) we consider the model of references to be fully abstract iff every modelled location is reachable from the root (program) node and every location’s value is uniquely defined by its position within the graph, and not by some arbitrary label.

1.5.3 Contributions

The work in this dissertation makes three significant contributions to our area of study (i.e. object-based object orientation within the UTP). First, we provide the UTP with an object-based model of object-oriented programs, in contrast to the existing class-based models.

Second, we provide an alternative model of pointers (references) for UTP that handles both value- and reference-based compound values (e.g. objects), in contrast to existing models, which handle one or the other. This is intended to contribute to the current discussions in the community as to what a general UTP pointer model should include and how one could generalise it.

Third, we provide an encoding of a prominent formalism of object orientation in the UTP (namely the Abadi–Cardelli-style of object calculus [AC96]) in contrast to developing our own notion of objects. Note that the encoding is performed in a systematic manner, with the hope that this will enable our work to be straightforwardly adapted for use with other Abadi–Cardelli-style object calculi.

1.5.4 Publications

This dissertation builds on two of our previously published works, [SG07] and [SG08]. The [SG07] paper introduces a cut down version of our chosen object calculus, a corresponding UTP model, and an outline of the consistency proof. Both the model and the consistency proofs are extended in this dissertation, in Chapter 3 and Chapter 4 respectively.

The [SG08] paper presents a model of locations that can be used to provide an abstract layout of memory, where each location is identified by precisely one path-set (i.e. the one that defines it). This paper also presents a brief outline of how this model can be linked to the basic relational model of the UTP, which is extended to the UTP model of designs in Chapter 6 of this dissertation.

Between them these publications provide evidence for the originality of the thesis being presented. They form the basis of work being presented in this dissertation. Here [SG07] corresponds to Section 2.2, Chapter 3 and Chapter 4, and [SG08] corresponds to Chapter 5.

1.5.5 Scope

At this point it is worth stating some areas that are beyond the scope of our consideration. First, we do not consider typing. This is consistent with the untyped nature of the UTP, at least as presented by Hoare and He in their seminal book [HH98], which introduces the subject. A significant body of research has gone into typed versions of the object calculus, e.g. [AC96, GH98, Cla02]. It ought to be possible to take advantage of this work, such as by assuming that the corresponding UTP model is provided with type-correct programs. In this circumstance it would be reasonable to use properties that follow from this type correctness in our UTP model proofs – for example, that named object members always exist when selected.

Second, we do not consider concurrency. It ought to be possible to add the UTP notion of reactive processes to the UTP model of objects. Having done this, it would be a matter of choosing a reactive process model that could straightforwardly support the operational semantics of an Abadi–Cardelli-style calculi.

Third, we do not consider refinement. The models were not intended to be used for the basis of a refinement calculus. It is not clear that it would be straightforward to use these models for the purposes of a refinement calculus without significant adaptation. Current work on providing the UTP with an OO-refinement calculus includes [HLL06, Kas06].

Fourth, we do not consider object ownership or reference encapsulation techniques. For example, object *representation exposure* is not considered. There is a significant amount of work in this area [NVP98, NCP99, Cla02, Par05, PNCB06, PB08] (and indeed, an entire workshop series on ‘Aliasing, Confinement and Ownership’). These works use a variety of formalisms including the object calculus (e.g. [Cla02]) and separation logic (e.g. [Par05]).

Chapter 2

Background

This chapter introduces the mathematical frameworks and notations that are used throughout this dissertation. In particular, it presents:

- an overview of Abadi and Cardelli’s untyped object calculus (ς -calculus) [AC96], along with a definition of our $\mathcal{O}$ -calculus [SG07]; and
- an overview of Hoare and He’s Unifying Theories of Programming (UTP) [HH98].

Before we introduce the Abadi–Cardelli calculus and the UTP, it is worth presenting some of the common notations that are used throughout this dissertation.

2.1 Common Notation and Conventions

Within this dissertation various notations and conventions are used to clarify the presentation. They will typically be introduced as they are needed. The remainder of this section contains some basic notation and conventions that are used throughout the dissertation.

2.1.1 The use of colour

So far we have only used colour to highlight cross-references and citations, both of which are coloured in [blue](#) (and hyper-linked in the electronic version). **Brown** is used to highlight meta-data within our $\mathcal{O}$ -calculus reduction rules; this meta-data status is also highlighted using a Quine’s corner notation ($\ulcorner \dots \urcorner$) [Qui40] when it is not clear from context. Lastly, **purple** is used to highlight literal label values (e.g. an object’s member names); a literal label is also typeset in `teletype` font.

2.1.2 Repeated terms

Comma-separated lists: It is frequently useful to denote an arbitrary comma-separated list of repeated terms. For example, $t_{i=1}^k$ is used to denote the list of terms $t_1, t_2, \dots, t_k$ whenever $k \geq 1$, and the empty list of terms whenever $k < 1$. Such an empty list of terms is explicitly denoted by $\perp$.

It may be necessary to scope the range of a repeated term, and this is denoted in the usual manner, i.e. $(t)_{i=1}^k$; in such cases, the outer brackets are not included in its expansion, which is $t_1, t_2, \dots, t_k$. When outer brackets are to be kept, the following alternative notation may be used:

$$\begin{array}{ccc} (t_{i=1}^k) & \widehat{=} & (t_{i=1}^k) \\ \{t_{i=1}^k\} & \widehat{=} & \{t_{i=1}^k\} \\ [t_{i=1}^k] & \widehat{=} & [t_{i=1}^k] \end{array}$$

Occasionally, it is useful to start counting at a different number, index over sets, or count backwards. These cases are denoted by $t_{i=j}^k$, $t_{i \in I}$, and $t_j^{i=k}$ respectively, where j is an integer and I is a set of indexes (such as program variable names). Note that the indexed terms $t_2^{i=3}$ and $t_{i=3}^2$ produce the results ‘ t_3, t_2 ’ and $\perp$ respectively; the latter empty case follows from the observation that there is no index i that is simultaneously greater than or equal to 3 and less than 2; i.e. $\neg(\exists i \bullet i \geq 3 \wedge i < 2)$.

Indexed operators: It is frequently useful to be able to repeatedly apply an associative binary operator to a list of arguments. For example, the indexed sequential composition operator $\circ_{i=1}^k t_i$ denotes $t_1 \circ t_2 \circ \dots \circ t_k$ whenever $k \geq 1$, and **skip** (the unit of composition) whenever $k < 1$. Note that such operators are only defined for empty lists of arguments if they have a *unit* (e.g. 0 for addition, 1 for multiplication, **true** for conjunction, and **false** for disjunction).

2.1.3 Identifying subterms

It is sometimes useful to write functions that operate over the abstract syntax tree representation of a program. For example, it is sometimes possible to identify the *free* variables of a program in this manner. Here the only interesting cases are the variable identifier terms, and the variable binding terms. All other cases simply merge the results of applying the *free* variable function to their subterms. Therefore, it would be convenient to be able to define these other cases by a single generic case. This is achieved via the use of a notation for formally denoting the subterms of an abstract syntax tree term.

Generic subterms: Let a generic abstract syntax tree term t containing k subterms $t_{i=1}^k$ be denoted by $t\{t_{i=1}^k\}$, where a leaf term ($k = 0$) has the empty set of subterms.

In principle the *free* variable function (FV) can now be defined by the following three case equations, where: the first case whose left hand side (LHS) pattern matches the provided argument is taken; and the meta-term $bind(x, t)$ represents any term that binds the variable identifier x to the term t .

$$\begin{aligned} FV(bind(x, t)) &\hat{=} FV(t) \setminus \{x\} \\ FV(x) &\hat{=} \{x\} \\ FV(t\{t_{i=1}^k\}) &\hat{=} \bigcup_{i=1}^k FV(t_i) \end{aligned}$$

This style of function definition places certain restrictions on the structure of the language’s abstract syntax tree. Specifically, in the case of the free-variable function that has just been defined, it must not be applied to the syntactic form of the standard UTP terms for representing variable introduction (and elimination); here, a variable is introduced (and terminated) as a subterm of sequential composition, so the scope of the variable is not contained as a subterm of the variable introduction operation, thus cannot be easily modelled using this approach. Fortunately, this issue can be avoided by constraining a program to use higher level *variable block* commands, which have the appropriate form, and whose semantics is defined in terms of the standard lower-level variable introduction (and elimination) commands.

2.1.4 Meta-variable identifiers

Within this dissertation, various meta-variables will be used to represent different types of entity within a syntactic description or formula. For example, the meta-variables l and o are used to represent a label and an object entity respectively. When more than one label or object is required, new label and object meta-variables can be constructed by the addition of subscripts and prime marks (e.g. l'_3 , a primed label meta-variable with subscript 3). The following table summaries the typical mapping of meta-variables to entity types.

Meta-Variable		Meta-Variable	
<i>Ids</i>	Type	<i>Ids</i>	Type
<i>b</i>	Boolean	<i>p, q</i>	Conditional (predicate)
<i>cv</i>	Compound Value	<i>P, Q</i>	(Relational) predicate
<i>d</i>	Definition	<i>lp</i>	(Labelled) path
<i>D</i>	Design	<i>Lp</i>	(Labelled) path-set
<i>e</i>	Expression	<i>pt</i>	Program text
<i>f</i>	Function	<i>r</i>	Reference Value
<i>i, j</i>	Integer (Index)	<i>t</i>	Term
<i>k</i>	Integer (Size)	<i>tl</i>	Trace label
<i>l</i>	Label	<i>tr</i>	Trace
<i>lv</i>	Literal Value	<i>Tr</i>	Trace-set
<i>m</i>	Method	<i>u</i>	List of variables
<i>n</i>	Node	<i>v</i>	Value
<i>N</i>	Set of nodes	<i>w</i>	List of variables
<i>o</i>	Object	<i>x, y</i>	(Program) variables
<i>ov</i>	Operand Value	<i>z</i>	Fix-point identifier

Table 2.1: Meta-variable identifiers

2.2 Object Calculus

The object calculus used in this dissertation is an extension of the ς -calculus presented in Chapter 6 of Abadi and Cardelli’s book on objects [AC96]. Our object calculus is known as the $\mathcal{O}$ -calculus; it is essentially the same as the heap extended core object calculus we presented in [SG07]. The main differences are in the order of presentation and the inclusion of two further extensions, which directly support conditional evaluation and untyped lambda calculus (λ -calculus) functions.

The operational semantics of the object calculus operations is defined in terms of a collection of small-step evaluation rules. These rules are similar to those in [Pie02], except that they include a notion of a general context (Γ), which is essentially used to denote those *specific* contexts that are irrelevant to a given rule. A specific context, such as the heap in Section 2.2.3, can be selected and set as follows:

$$\begin{array}{l|l} \{\text{HEAP} \mapsto H\} & \text{Let } H \text{ denote the HEAP context.} \\ \{\text{HEAP} \leftarrow H\} & \text{Set the HEAP context to the value of } H. \end{array}$$

The objective of the evaluation rules is to provide the circumstances under which a term t in a context Γ can evaluate in one step to a term t' in context Γ' ; such one-step evaluations are denoted by $\Gamma \bullet t \longrightarrow \Gamma' \bullet t'$, where $_ \bullet _$ denotes the context-term pair binder and $_ \longrightarrow _$ denotes an individual evaluation step. It is now possible to define the rule representation as follows:

$$\frac{\langle \text{condition}_1 \rangle \quad \dots \quad \langle \text{condition}_n \rangle}{\langle \text{concluded term evaluation step} \rangle} \text{RULENAME} \quad \boxed{\begin{array}{l} \langle \text{side condition}_1 \rangle \\ \dots \\ \langle \text{side condition}_m \rangle \end{array}}$$

Here, a condition_i is either a logical constraint or an evaluation step, where as a side condition_j is either a meta-level logical constraint or a meta-variable definition. Note that if there are no *conditions* (i.e. $n = 0$), then the concluded term evaluation step can be used whenever the term on its left hand side pattern matches with that of the actual term.

We now define our $\mathcal{O}$ -calculus in five stages. The first stage presents an overview of the ς -calculus (Section 2.2.1). The second stage extends this with a few basic enhancements, for

handling literal values and field assignment (Section 2.2.2). The third stage adds an abstract heap (Section 2.2.3). Finally, the last two stages provide direct support for handling conditional evaluation and eagerly evaluated λ -calculus functions respectively (Section 2.2.4).

2.2.1 Abadi–Cardelli untyped object calculus

The Abadi–Cardelli ς -calculus introduces the notion of an object as a collection of labelled methods that can be updated and selected as follows.

$\left[\begin{smallmatrix} k \\ i=1 \end{smallmatrix} l_i = m_i \right]$	denotes an object value – a partial map from labels to methods, where method m_i is identified by label l_i .
$\varsigma(x) e$	denotes a method whose body is defined by the expression e , which may contain one or more instances of the self variable (identifier) x .
$o.l \Leftarrow m$	denotes a method update operation, which generates a new object by taking a copy of the object o and replacing the method identified by label l with the method m .
$o.l$	denotes a method selection operation, which evaluates the body of the method with label l in object o , after each instance of the method's <i>self</i> variable has been replaced by a copy of the invoking object o .

Here, the meta-variables

o, l	denote an object value and a label respectively;
m	denotes a method;
e, x	denote an expression and the self identifier (variable/expression).

Note that a ς -calculus expression is either an object value, variable identifier, or an application of the method selection or update operators. In particular, neither a label nor a method is considered to be a value-expression.

Method update: The base case for the method update operations can now be defined by the following small-step evaluation rule.

$$\frac{l \in o}{\Gamma \bullet o.l \Leftarrow m \longrightarrow \Gamma \bullet \textcolor{brown}{o} \oplus \{l \mapsto m\}^\top} \text{UPDM}$$

Here, the meanings of the components are as follows

$l \in o$	label l is in the domain of object o .
$o_1 \oplus o_2$	object map o_2 overrides object map o_1 .
$\textcolor{brown}{e}^\top$	the meta-expression e .

There is one other small-step evaluation rule, which ensures that the evaluable argument (i.e. expression argument) of the method update operation is evaluated prior to the operation being applied.

$$\frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet e.l \Leftarrow m \longrightarrow \Gamma' \bullet e'.l \Leftarrow m} \text{UPDM-1}$$

Method invocation (or selection): The base case for the method invocation operations of the ς -calculus can be defined by the following rule.

$$\frac{l \in o}{\Gamma \bullet o.l \longrightarrow \Gamma \bullet \textcolor{brown}{e}\{x \leftarrow o\}^\top} \text{INVM} \quad \boxed{\begin{array}{l} m \hat{=} o(l) \\ \varsigma(x) e \doteq m \end{array}}$$

Here:

$o(l)$	is the method of object o with label l .
$m \hat{=} e$	defines variable m to be the evaluation of meta-expression e .
$\varsigma(x) e \doteq m$	binds x and e to the self-variable and body of method m .
$e\{x \leftarrow o\}$	the substitution of object o for free variable x in expression e (it is defined in Section 2.2.5).

The rule for evaluating an evaluable argument before the base rule can be applied is defined in a similar manner to that of the method update operation.

$$\frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet e.l \longrightarrow \Gamma' \bullet e'.l} \text{ INV M-1}$$

2.2.2 Basic extensions

The basic extensions introduce four sorts of – non-object – literal value: the booleans, the integers, the strings, and the unset unit sort. Here the integers are denoted by their usual Arabic numbers, the booleans by the constants `true` and `false`, the strings by double-quoted text (e.g. "Some text"), and the unset sort by the constant `!`.

Aside 2.1: The unset value serves the same role as the Perl language's `undef` value; e.g. it provides a mechanism for testing whether a variable contains, or a function returns, a defined value.

□

The basic extension also provides some logical, arithmetic, and text operations, for working with the booleans, integers, and strings respectively. The precise collection of operations is not important, so they are abstracted by a generalised literal value operation ($\Box_k$), which takes k literal arguments and returns a literal value; here k is omitted when it is clear from context. Lastly, the basic extension provides a field assignment operation, which is used to assign the evaluation of an expression to a given object member.

$\Box(e_{i=1}^k)$	denotes a k -ary operation on literal values (e.g. binary '+' and '≤', with $k = 2$).
$o.l := e$	denotes the field assignment operation, which evaluates the expression e to a value v , then applies the method update operation $o.l \leftarrow \varsigma(-) v$.

The following rules specify the base cases for a generic literal operation (LITOP) and the field assignment operation (FLDA).

$$\frac{\Gamma \bullet \Box(v_{i=1}^k) \longrightarrow \Gamma \bullet \Box(v_{i=1}^k)}{\Gamma \bullet \Box(v_{i=1}^k) \longrightarrow \Gamma \bullet \Box(v_{i=1}^k)} \text{ LITOP} \quad \boxed{v_{i=1}^k \in \text{dom}(\Box)}$$

$$\frac{l \in o}{\Gamma \bullet o.l := v \longrightarrow \Gamma \bullet o \oplus \{l \mapsto \varsigma(-) v\}} \text{ FLDA}$$

The other cases for these operations ensure that their evaluable arguments are processed in a left to right order.

$$\frac{\Gamma \bullet e_k \longrightarrow \Gamma' \bullet e'_k}{\Gamma \bullet \Box(v_{i=1}^{k-1}, e_{i=k}^k) \longrightarrow \Gamma' \bullet \Box(v_{i=1}^{k-1}, e'_k, e_{i=k+1}^k)} \text{ LITOP-K}$$

$$\frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1.l := e_2 \longrightarrow \Gamma' \bullet e'_1.l := e_2} \text{ FLDA-1} \quad \frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet o.l := e \longrightarrow \Gamma' \bullet o.l := e'} \text{ FLDA-3}$$

Note that the digit distinguishing the field assignment rules FLDA-1 and FLDA-3 identifies the argument to evaluate.

2.2.3 Heap extensions

We now introduce a copy-based heap storage model [Pie02], where the heap is a partial map from abstract locations to values. Here the contents of an abstract location can be read (dereferenced) or updated (assigned) via atomic operations that take copies of the source values. The new constants and operators introduced by this model now follow.

- ℓ_i denotes an abstract location on the heap and an allocated reference value.
- null** denotes the null (i.e. unallocated) reference value.
- fresh** denotes the operation that results in the location of a newly allocated heap entry, whose contents are unset.
- $*r$ denotes the operation that takes a copy of the contents in heap location r .
- $r * = v$ denotes the assignment, by copy, of value v to location r .

Here, the meta-variables r , v , and i denote the reference, general, and integer values respectively.

The following rules specify the base cases for the **fresh**, **dereference**, and **assignment** (reference update) operators; these rules select and manipulate the **heap** component of the rule-context Γ . The other cases for these operators are defined to follow the usual left to right evaluation order.

$$\begin{array}{c}
 \frac{}{\Gamma\{\mathbf{heap} \mapsto H\} \bullet \mathbf{fresh} \longrightarrow \Gamma\{\mathbf{heap} \mapsto H'\} \bullet r} \text{FRESH} \quad \boxed{\begin{array}{l} r \triangleq \mathbf{fresh}_{\text{LOC}}(\text{dom } H) \\ H' \triangleq H \oplus \{r \mapsto \mathbf{i}\} \end{array}} \\
 \frac{r \in H}{\Gamma\{\mathbf{heap} \mapsto H\} \bullet *r \longrightarrow \Gamma \bullet \lceil H(r) \rceil} \text{DEREF} \\
 \frac{r \in H}{\Gamma\{\mathbf{heap} \mapsto H\} \bullet r * = v \longrightarrow \Gamma\{\mathbf{heap} \mapsto H \oplus \{r \mapsto v\}\} \bullet r} \text{UPDL}
 \end{array}$$

Here, $\mathbf{fresh}_{\text{LOC}}$ is a meta-function that takes a set of location values and returns a location that is not within this set.

The other cases for the heap operations ensure that their arguments are processed in a left to right order.

$$\begin{array}{c}
 \frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet *e \longrightarrow \Gamma' \bullet *e'} \text{DEREF-1} \\
 \frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1 * = e_2 \longrightarrow \Gamma' \bullet e'_1 * = e_2} \text{UPDL-1} \quad \frac{\Gamma \bullet e_2 \longrightarrow \Gamma' \bullet e'_2}{\Gamma \bullet r * = e_2 \longrightarrow \Gamma' \bullet r * = e'_2} \text{UPDL-2}
 \end{array}$$

2.2.4 Remaining extensions

The remaining extensions introduce some support for directly handling conditional evaluation, sequential composition, eagerly evaluated λ -calculus functions and local blocks.

- if b then e_1 else e_2** denotes the operation that evaluates expression e_1 when the boolean b is **true**, and expression e_2 otherwise.
- $o?l$** denotes the membership test operation, which evaluates to **true** whenever object o contains a method with label l , and **false** otherwise.
- $e_1 \circ e_2$** denotes the sequential composition of the expressions e_1 and e_2 . Here, the result of evaluating e_1 is ignored; hence, e_1 's contribution is via its side-effects (e.g. updating the heap).
- $\lambda(x) e$** denotes a function with formal parameter x whose body is defined by the expression e .

$f \ v$ denotes the function application operation, which evaluates the body of the function f , once its parameter has been replaced by a copy of value v , wherever it occurs (free within the body).

$\text{let } d_{i=1}^k \text{ in } e$ denotes the local block that introduces a comma separated list of local declarations $d_{i=1}^k$ that can be used in expression e .

Here, the meta-variables f and d denote a function and a local declaration respectively, and the local declarations have one of two forms:

$x \equiv t$ denotes the binding of an unevaluated term t to the identifier x .

$x \hat{=} e$ denotes the binding of the evaluation of an expression e to the identifier x .

The following conditional operation rules have a slightly different form to those previously defined. Specifically, they have two base cases; and these cases may contain some potentially unevaluated evaluable arguments. This is important to model the usual notion of conditional evaluation, where only the branch to be selected is evaluated. An alternative would be to evaluate both branches, before selecting the relevant branch. Note that this alternative scheme provides the usual semantics for conditional evaluation so long as a branch terminates and does not contain any globally visible side effects.

$$\frac{}{\Gamma \bullet \text{if true then } e_1 \text{ else } e_2 \longrightarrow \Gamma \bullet e_1} \text{IFT} \quad \frac{}{\Gamma \bullet \text{if false then } e_1 \text{ else } e_2 \longrightarrow \Gamma \bullet e_2} \text{IFF}$$

$$\frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet \text{if } e \text{ then } e_1 \text{ else } e_2 \longrightarrow \Gamma' \bullet \text{if } e' \text{ then } e_1 \text{ else } e_2} \text{IF}$$

The operations for membership test, function invocation and sequential composition can be specified in the usual manner. Their base cases are now defined by the following rules.

$$\frac{}{\Gamma \bullet o?l \longrightarrow \Gamma \bullet \ulcorner l \in o \urcorner} \text{HASM} \quad \frac{}{\Gamma \bullet v \ ; \ e \longrightarrow \Gamma \bullet e} \text{SEQC}$$

$$\frac{}{\Gamma \bullet f \ v \longrightarrow \Gamma \bullet \ulcorner e \{x \leftarrow v\} \urcorner} \text{INV F} \quad \boxed{\lambda(x) \ e \doteq f}$$

The other cases for these operations ensure that their evaluable arguments are processed in a left to right order.

$$\frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet e?l \longrightarrow \Gamma' \bullet e'?l} \text{HASM-1} \quad \frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1 \ ; \ e_2 \longrightarrow \Gamma' \bullet e'_1 \ ; \ e_2} \text{SEQC-1}$$

$$\frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1 \ e_2 \longrightarrow \Gamma' \bullet e'_1 \ e_2} \text{INV F-1} \quad \frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet f \ e \longrightarrow \Gamma' \bullet f \ e'} \text{INV F-2}$$

Finally, the local block is introduced as a form of function definition and invocation, as follows:

$$\frac{}{\Gamma \bullet \text{let } x \equiv t \text{ in } e \longrightarrow \Gamma \bullet \ulcorner e \{x \leftarrow t\} \urcorner} \text{LETM} \quad \frac{}{\Gamma \bullet \text{let } x \hat{=} e_1 \text{ in } e_2 \longrightarrow \Gamma \bullet (\lambda(x) \ e_2) \ e_1} \text{LETD}$$

$$\frac{}{\Gamma \bullet \text{let } d_{i=1}^k \text{ in } e \longrightarrow \Gamma \bullet \text{let } d \text{ in let } d_{i=2}^k \text{ in } e} \text{LETL} \quad \boxed{k > 1}$$

Here, $d_{i=1}^k$ denotes a list of local declarations with at least two elements.

2.2.5 Substitution definition and object calculus examples

Substitution: Three of the reduction rules for our $\mathcal{O}$ -calculus have used the substitution operation to replace free occurrences of an identifier within a term with another term. This

can be formally defined as follows:

$$\begin{array}{lll}
x \llbracket x \leftarrow t \rrbracket & \hat{=} & t \\
(\varsigma(x) e) \llbracket x \leftarrow t \rrbracket & \hat{=} & \varsigma(x) e \\
(\lambda(x) e) \llbracket x \leftarrow t \rrbracket & \hat{=} & \lambda(x) e \\
(\text{let } x \equiv t_1 \text{ in } e) \llbracket x \leftarrow t \rrbracket & \hat{=} & \text{let } x \equiv t_1 \text{ in } e \\
(\text{let } x \hat{=} e_1 \text{ in } e) \llbracket x \leftarrow t \rrbracket & \hat{=} & \text{let } x \hat{=} e_1 \text{ in } e \\
(t_0 \{t_{i=1}^k\}) \llbracket x \leftarrow t \rrbracket & \hat{=} & t_0 \{t_{i=1}^k \llbracket x \leftarrow t \rrbracket\}
\end{array}$$

Here, the first of the pattern matching definition lines matches only when the term being substituted is an identifier which is syntactically identical to that of the identifier being substituted. The following four definition lines show the four ways of introducing a local identifier (variable); these prevent the substitution of bound variables with the same name. And the last of the definition lines shows the general case, where the substitution of any other term amounts to substituting its subterms (if there are any). For example:

$$\begin{aligned}
(\lambda(y) y + (\lambda(x) \dots)(x + 1)) \llbracket x \leftarrow 3 \rrbracket &= \lambda(y \llbracket x \leftarrow 3 \rrbracket) (y + (\lambda(x) \dots)(x + 1)) \llbracket x \leftarrow 3 \rrbracket \\
&= \lambda(y) (y + (\lambda(x) \dots)(x + 1)) \llbracket x \leftarrow 3 \rrbracket \\
&= \lambda(y) y \llbracket x \leftarrow 3 \rrbracket + ((\lambda(x) \dots)(x + 1)) \llbracket x \leftarrow 3 \rrbracket \\
&= \lambda(y) y + ((\lambda(x) \dots)(x + 1)) \llbracket x \leftarrow 3 \rrbracket \\
&= \lambda(y) y + (\lambda(x) \dots) \llbracket x \leftarrow 3 \rrbracket (x + 1) \llbracket x \leftarrow 3 \rrbracket \\
&= \lambda(y) y + (\lambda(x) \dots)(x + 1) \llbracket x \leftarrow 3 \rrbracket \\
&= \lambda(y) y + (\lambda(x) \dots)(x \llbracket x \leftarrow 3 \rrbracket + 1 \llbracket x \leftarrow 3 \rrbracket) \\
&= \lambda(y) y + (\lambda(x) \dots)(3 + 1 \llbracket x \leftarrow 3 \rrbracket) \\
&= \lambda(y) y + (\lambda(x) \dots)(3 + 1)
\end{aligned}$$

Value example: The simplest $\mathcal{O}$ -calculus programs (expressions) are the values, as they are already evaluated, and thus represent their own results. The integer 42 is an example of such a program. This example is used to illustrate how to interpret the results of compiling a $\mathcal{O}$ -calculus program into one of our UTP object models.

Restorable cell example: This example is essentially the third variant of the Abadi–Cardelli’s untyped storage cell [AC96, Section 6.5.5]. Here we have an uninitialised (unset) **data** field, whose value can be accessed via **get** and **set** methods. The **set** method also stores the previous value of the **data** field in a dynamically updated **undo** method.

$$\begin{aligned}
UnsetRestCell \quad = \quad & [\text{data} = \mathbf{i}, \\
& \text{get} = \varsigma(x) x.\text{data}, \\
& \text{set} = \varsigma(x) \lambda(y) (\\
& \quad x.\text{undo} \Leftarrow \varsigma(z) z.\text{data} := x.\text{data} \\
& \quad).\text{data} := y, \\
& \text{undo} = \varsigma(x) x.\text{data} := \mathbf{i} \\
&]
\end{aligned}$$

Note that the **set** method makes use of two distinct *self* variables x and z , which represent the current and future values respectively. Further, both of these variables are in scope in **set**’s inner method (i.e. $x.\text{undo} \Leftarrow \varsigma(z) z.\text{data} := x.\text{data}$). This is why the Abadi–Cardelli calculus uses a named self parameter, rather than the *self* (or *this*) keyword that is used in many object-oriented (OO) languages.

Factorial example: A simple recursive algorithm to calculate the factorial of a number can be written in the $\mathcal{O}$ -calculus as follows:

$$\text{Factorial} \quad \hat{=} \quad [\text{fact} = \varsigma(x) \lambda(i) \text{ if } i < 2 \text{ then } 1 \text{ else } i \times x.\text{fact}(i - 1)]$$

Here, the *Factorial* object provides a single method that contains an anonymous function, which takes a natural number and outputs its factorial value. It is later used as a concrete example of how one of the UTP encodings of the $\mathcal{O}$ -calculus is compiled and reasoned with.

Sharing methods example: [Appendix B](#) provides a collection of $\mathcal{O}$ -calculus programs that enable the methods of an object to be shared. These techniques can be used when modelling a class or an inheritance scheme, in terms of $\mathcal{O}$ -calculus objects.

2.3 Unifying Theories of Programming

Various notations and conventions are adopted in the UTP book [HH98]. Only some of these are required for the purposes of this dissertation; in general, they will be introduced when needed. However, the following notation and conventions, concerning the alphabets of a predicate, are worth highlighting at this point.

First, it is useful to be able to partition the predicate P 's alphabet of *free*, i.e. not quantified, variables into inputs and outputs; convention dictates that only the output variables are dashed, thus the output value of a variable x is denoted by x' , whereas the input value is denoted by x itself. The alphabet associated with a predicate P , its inputs, and its outputs are denoted by αP , $in\alpha P$, and $out\alpha P$ respectively.

Second, as variables in real programming languages have both input and output values, even if they are the same, it is frequently assumed that the set of output variables can be derived from the set of input variables by dashing the input variables. In such cases, the output alphabet of a program represented by predicate P is $out\alpha P = \{x' \mid x \in in\alpha P\}$, and is denoted by $in\alpha' P$.

Last, the letters w and w' represent the list of all input and output variables of a predicate respectively, where $w(i)$ represents the i^{th} variable in w that has $\#w$ elements. This is done so that it is possible to do mass renaming of these variables, for an arbitrary predicate P , using the notation $P[w_1/w_2]$, to stand for the predicate that is formed by replacing all free occurrences of variable $w_2(i)$ with variable $w_1(i)$ in P , where $i \in 1 \dots \#w_1$ and $\#w_1 = \#w_2$.

2.3.1 Relational model

The semantics of the UTP is based on the notion of a relation between variables, specifically input and output variables, at observable points in a program. Arguably the simplest program is the one that does nothing. This is modelled in the UTP by the *skip* program, which is defined by the predicate $w' = w$, where: w represents both the list of program variables and their initial (input) values; and w' the final (output) values of the same list of program variables. The usual UTP convention of denoting the input and output values of a program variable x , by logical variables x and x' is followed.

Formally, a UTP program is a relational predicate, denoted by the pair $(\alpha P, P)$, where P is a predicate that relates inputs and output variables within its alphabet αP , which is at least the set of *free*, i.e. not quantified, variables in P . Normally, the alphabet of a predicate is taken to be those variables that are free within it (or become required in order to compose it with another UTP program). Hence, the predicate part of the relation is often used to stand for the relation as a whole. This practise is adopted in this dissertation.

Aside 2.2: Most predicate operators assume that the alphabets of the predicates they combine are the same. Hence, there is a special operation, P_{+A} , which expands a predicate P for a given a set of input variable names A , by creating both an associated set of output variables, i.e. $\{a : A \bullet a'\}$, and insisting that they are bound to their associated inputs. Formally this operation is defined by $P_{+A} \hat{=} P \wedge (\bigwedge_{a \in A} \bullet a' = a)$, provided the sets of input names and free variables are disjoint, i.e. $A \cap FV(P) = \emptyset$.

□

Having said that a program can be represented as a predicate relation, there must be a way of representing the semantics of a program's constructs as predicate operators (i.e. operators that take and return predicates). Chapters 1 and 2 of the UTP book [HH98] demonstrate how the semantics of basic imperative constructs, such as assignment, conditional execution, and loops, can be written as predicates. This is straightforward for all the constructs except loops, which require an understanding of the *fix-point* theory associated with lattices [DP02, Chapters 2 and 4] and [BCG00].

Example 2.3: The sequential composition predicate operator $(-; -)$ constructs a new program fragment predicate from two existing program fragment predicates, as follows:

$$\begin{aligned} P ; Q &\hat{=} \exists w_0 \bullet P[w_0/w'] \wedge Q[w_0/w] && \text{provided } out\alpha P = in\alpha' Q = \{w'\} \\ in\alpha(P ; Q) &\hat{=} in\alpha P \\ out\alpha(P ; Q) &\hat{=} out\alpha Q \end{aligned}$$

□

Correctness: Having introduced the notion of representing relations between input and output variables as predicates, what is the method of stating that a program correctly implements its specification? It is natural to choose the logical implication ordering to determine program correctness, where a program P is considered to correctly implement its specification S if it implies its specification, i.e. $P \Rightarrow S$; one caveat to this statement is that both the program and specification must have the same alphabet, i.e. $\alpha P = \alpha S$. The formal notion of program correctness is the universal quantification of $P \Rightarrow S$, and this is denoted $[P \Rightarrow S]$ in accordance with Dijkstra and Scholten's convention [DS90].

$$[P \Rightarrow S] = \forall x_{i=1}^k \bullet P \Rightarrow S \quad \text{provided } \alpha P = \alpha S = \{x_{i=1}^k\}$$

Note that the implication ordering also defines what it means for one program to refine another in the UTP relational model.

Refinement: A program P is said to refine a specification S whenever it implies the specification.

$$P \sqsubseteq S \hat{=} P \Rightarrow S \quad \text{provided } \alpha P = \alpha S$$

The transitive nature of the refinement (implication) ordering allows for a staged development of concepts. Critically, the various steps in the development form a complete lattice of programs with alphabet A , where **false** _{A} is the strongest element of the predicate lattice ordered by implication and **true** _{A} is the weakest element. Here the notion of truth and falsehood is in relation to a specific alphabet of free variables. Essentially, **true** _{A} is modelled by the relation that does not restrict the values of its outputs, and is sometimes referred to in programming languages as *abort* or *chaos*. Unfortunately, **false** _{A} does not have such an intuitive description, but it is technically defined to be the result of the disjunction of an empty set of relations over alphabet A , i.e. $\bigvee_{P \in \emptyset}$, where:

$$\begin{aligned} \bigvee_{P \in \emptyset} &\hat{=} \text{false} \\ \bigvee_{P \in \{P_{i=1}^k\}} &\hat{=} P_1 \vee \bigvee_{P \in \{P_{i=2}^k\}} \end{aligned}$$

The **false** relation is sometimes called the *miracle*, as it implements an arbitrary specification S ; i.e. $\forall S \bullet [\text{false} \Rightarrow S]$.

2.3.2 Design model

A UTP *design* is a special form of relation that introduces the notions of successful initialisation and termination, via a couple of special variables, namely Π_{ok} and Π_{ok}' respectively. These variables are *free* in every design, and represent a program's execution context; as such they cannot be referred to by any of the program's variables or constructs. Instead, they are used to change some of the constraints that the basic relational model imposes on the implementation context of a program, such as backwards execution; that is only evaluating those inputs that are actually needed for the output. Though this style of execution can be supported, it is not supported by typical imperative languages, such as C, Pascal and Ada. Hence, the UTP notion of a design provides for the optimisation of forward execution typical of imperative languages; that is the evaluation of inputs as they appear, rather than as they are needed.

Aside 2.4 – Context variable naming conventions: The UTP models of the $\mathcal{O}$ -calculus use special, non-program, variables to model various aspects of the OO-context. Such contextual variables are denoted by identifiers that start with the capital Greek letter pi (Π), so that they can be easily distinguished from normal program variables. In particular, the usual name of the special imperative programming (design) variable, used to model program initialisation and termination of a program, is changed from *ok* to Π_{ok} so that it can conform to this convention. Another example of such a contextual variable is Π_{res} , which is used to denote the *result* of a program (or subprogram). $\square$

All of the UTP models of the $\mathcal{O}$ -calculus presented in this thesis extend the notion of a UTP design, which captures the semantics of first-order imperative programs without subroutines (e.g. procedures and functions). However, instead of using the notion of the design defined in the seminal book on UTP [HH98], we use He's updated variant of a design [HLL02a]. It updates the original definition in [HH98] by ensuring that “the assumption is a precondition, containing only undashed [input] variables [... which ...] corresponds exactly to the third healthiness condition” of [HH98, page 84].

Design definition: A design predicate $p \vdash P$ states that the program represented by the relational predicate P will successfully terminate whenever it has been started in an input state that satisfies input assumption (precondition) predicate p .

$$\begin{aligned} p \vdash P &\hat{=} (\Pi_{ok} \wedge p) \Rightarrow (\Pi_{ok}' \wedge P) \\ \alpha(p \vdash P) &\hat{=} \alpha P \cup \{\Pi_{ok}, \Pi_{ok}'\} \\ &\text{provided } \alpha p \subseteq in \alpha P \text{ and } \Pi_{ok}, \Pi_{ok}' \notin \alpha P \end{aligned}$$

Two consequences of this design model are: it is impossible to say what the initial values of a variable are prior to the starting of a program; and it is impossible to say what the final values of a variable are prior to the termination of a program. Here $\neg \Pi_{ok}'$ represents a situation where a program has yet to terminate. Having said this, if the program is a design where both Π_{ok} and p hold, then it is guaranteed to terminate. The issue is in demonstrating (proving) that the precondition p is met. This is non-trivial and, in general, at least as hard as the halting problem.

Correctness: The notion of correctness for a design is the same as that for a relation, namely, that of an implementation-design implying the specification-design. Equivalently, the notion of program correctness can be defined as the conjunction of: the specification's pre-condition implying the implementation's pre-condition; and the implementation's post-condition implying the specification's post-condition, within the context of the specification's pre-condition.

In other words, supposing that implementation and specification designs are represented by $D_I = p_I \vdash P_I$ and $D_S = p_S \vdash P_S$ respectively, then the notion of program correctness is either $[D_I \Rightarrow D_S]$ or equivalently $[p_S \Rightarrow p_I] \wedge [(p_S \wedge P_I) \Rightarrow P_S]$. The proof of the equivalence of these two notions of design correctness is given in theorem 3.1.2 of the UTP book [HH98].

Aside 2.5: The notion of implication is used for both the definition of a design and the correctness of a design. In the first case, the pre-condition implies the post-condition; and in the second case, an implementation design implies the specification design. It is all too easy to confuse or merge these distinct usages of the implication ordering. Having said this, it is precisely the interaction between the usage of implication in the definition and correctness of a design which leads to the alternative (second) definition of design correctness in the preceding paragraph.

□

2.3.3 UTP design commands

In this section we present a slightly extended version of the standard commands for a UTP design; i.e. those operations that can be used to construct a UTP program.

We start by introducing the notion of a design *frame*, which is used to clarify the definition of the *basic* UTP commands, by focusing attention on only those elements of the alphabet that the commands actually change. Here, a *basic* command is considered to be any command that does not take a UTP program as an argument.

Each of the *non-basic* commands is defined in terms of how it affects the UTP programs that it consumes, which may be designs. We also show that if all the program arguments are designs, then the result is a design. Hence, the commands that we present are closed, when applied to designs. This closure is important for the use of Tarski's fix-point theory, which is used to define the semantics of loops (and recursion – see Chapter 3).

Design frame: Let V be a set of program variables. A design with frame V has the form $V : (p \vdash P)$, denoting the predicate $p \vdash P \wedge w' = w$ with alphabet $\{V, V', w, w'\}$, where the vector w contains the variables within the input alphabet of $p \vdash P$ except those in the set V (i.e. $\{x \mid x \in w\} = \text{in}\alpha(p \vdash P) \setminus V$).

Skip command: The skip_A^D command does nothing and terminates successfully.

$$\begin{aligned} \text{skip}_A^D & \hat{=} \emptyset : (\text{true} \vdash \text{true}) \\ \alpha(\text{skip}_A^D) & \hat{=} A \end{aligned}$$

Miracle command: The *miracle* places an impossible collection of restrictions on the outputs of the variables within its alphabet (i.e. *false*).

$$\begin{aligned} \text{miracle}_A & \hat{=} \emptyset : (\text{true} \vdash \text{false}) \\ \alpha \text{miracle}_A & \hat{=} A \end{aligned}$$

Note that the *miracle* element simplifies to the predicate $\neg \Pi_{\text{OK}}$, and is the top of the design refinement ordering ($\top_D$). For a program to be implementable its result must not be reducible to a miracle. However, miracles are useful for modelling contractual guarantees, as failure of the guarantee to hold is not the fault of the design (or program).

Chaos command: The chaos_A command, also known as *abort*, behaves unpredictably.

$$\begin{aligned} \text{chaos}_A & \hat{=} \emptyset : (\text{false} \vdash \text{true}) \\ \alpha(\text{chaos}_A) & \hat{=} A \end{aligned}$$

Note that the **chaos** element simplifies to the predicate **true**, and is the bottom of the design refinement ordering ($\perp_D$). It could also have been defined by (**false** $\vdash$ **false**), which is the last of the four programs involving only truth and falsehood.

Conditional command: The conditional command ($P \triangleleft b \triangleright Q$) behaves like the design P if the boolean expression b is true, otherwise it behaves like Q .

$$\begin{aligned} P \triangleleft b \triangleright Q & \hat{=} b \wedge P \vee \neg b \wedge Q, \quad \text{provided } \alpha P = \alpha Q \\ \alpha(P \triangleleft b \triangleright Q) & \hat{=} \alpha P \end{aligned}$$

When both of a conditional command's program arguments are designs, and its boolean argument is defined, then the result of the conditional is also a design. This is proved in [Appendix C.1.1](#) – specifically it is shown that:

$$(p \vdash P) \triangleleft b \triangleright (q \vdash Q) = (p \triangleleft b \triangleright q) \vdash (P \triangleleft b \triangleright Q)$$

In English, the resultant design states that whenever the boolean variable is defined and the selected precondition is met, then the selected postcondition will also hold.

Non-deterministic choice command: The non-deterministic choice ($P \sqcap Q$) behaves either like design P or design Q .

$$\begin{aligned} P \sqcap Q & \hat{=} P \vee Q, \quad \text{provided } \alpha P = \alpha Q \\ \alpha(P \sqcap Q) & \hat{=} \alpha P \end{aligned}$$

When both of a non-deterministic command's arguments are designs, then the result of the non-deterministic command is also a design. It is straightforward to prove this using the propositional calculus and the definitions of the non-determinism command and the UTP design. Here the only interesting point is that the disjunction of designs ($p \vdash P$) and ($q \vdash Q$) forms a conjunction of their pre-conditions.

$$(p \vdash P) \sqcap (q \vdash Q) = p \wedge q \vdash P \sqcap Q$$

Sequential composition command: The sequential composition ($P \circledast Q$) behaves like design P and then design Q , where the final state of P is also the initial state of Q .

$$\begin{aligned} P \circledast Q & \hat{=} \exists w_0 \bullet P[w_0/w'] \wedge Q[w_0/w], \quad \text{provided } out\alpha P = in\alpha' Q = \{w'\} \\ in\alpha(P \circledast Q) & \hat{=} in\alpha P \\ out\alpha(P \circledast Q) & \hat{=} out\alpha Q \end{aligned}$$

The proof that the composition of two designs is also a design is in [Appendix C.1.2](#). This differs from the proof given in the UTP book, as we are using He's more recent version of a design, which is slightly more restrictive. In particular, it is possible to simplify the representation of the precondition for the composed design, by side-stepping the need for sequential composition with the bottom element in the standard relational calculus (i.e. **true**).

$$(p \vdash P) \circledast (q \vdash Q) = p \wedge \neg(P \circledast \neg q) \vdash P \circledast Q$$

In English, this says that the composition of two designs is the design that performs $P \circledast Q$ whenever: the first design's precondition p is met; and the result of the first design produces a context which ensures that the second design's precondition q is met.

Assignment command: The assignment command $(x := e)$ assigns the evaluation of expression e to the variable x so long as the expression e is defined and the variable x is in the alphabet of the assignment ($\alpha(_ := _) = \{x\}$). We adopt the usual UTP practice of assuming that such expressions are well defined and thus do not put in a definedness pre-condition check in the definition of assignment. The assignment command can also be generalised to the simultaneous assignment command, which takes vectors of distinct variables $(x_{i=1}^k)$ and general expressions $(e_{i=1}^k)$.

$$\begin{aligned} x :=_A e & \hat{=} \{x\} : (\text{true} \vdash x' = e), \quad \text{provided } x' \in A \\ \alpha(x :=_A e) & \hat{=} A \\ x_{i=1}^k :=_A e_{i=1}^k & \hat{=} \{x_{i=1}^k\} : (\text{true} \vdash x_{i=1}'^k = e_{i=1}^k), \quad \text{provided } x_{i=1}'^k \in A \\ \alpha(x_{i=1}^k :=_A e) & \hat{=} A \end{aligned}$$

This assignment operation differs from that of the standard higher-order programming found in Chapter 9 of [HH98], which defines assignment in terms of refinement. Here, the assignment is defined by the relational predicate $x := e \hat{=} (x' \sqsubseteq e \wedge w' = w)$, where a first order value (v) is only refined by itself (i.e. $x' \sqsubseteq v \Leftrightarrow x' = v$). In particular, it allows an implementation of a procedure P to refine its specification; i.e. $(P \sqsubseteq S) \Rightarrow (x := P \sqsubseteq x := S)$.

Iteration command: The iteration command $b * P$ repeats the program P as long as the condition b is true.

$$b * P \hat{=} \mu X \bullet P \circ X \triangleleft b \triangleright \text{skip}^D,$$

Here, $\mu X \bullet F(X)$ stands for the *Tarski weakest fix-point* [Tar55] of the equation $X = F(X)$, and X is a predicate from the lattice of UTP relational predicates [HH98].

The program $b * (p \vdash P)$ results in a design because either: it infinitely loops, in which case it is equivalent to the *chaos* design; or it sequentially composes a finite number of designs, which is a design.

Variable introduction command: The local variables $x_{i=1}^k$ are introduced by the command $(\text{var } x_{i=1}^k)$, so long as they are not already in the alphabet of the program.

$$\begin{aligned} \text{var } x_{i=1}^k & \hat{=} \exists x_{i=1}^k \bullet \text{skip}_A^D \\ & = \text{true} \vdash \exists x_{i=1}^k \bullet \text{skip}_{A1} \end{aligned}$$

Here, the alphabet of the command as a whole, i.e. $\alpha(\text{var } x_{i=1}^k)$, is $A \setminus \{x_{i=1}^k\}$ and the alphabet $A1 = A \setminus \{\Pi_{\text{OK}}, \Pi_{\text{OK}}'\}$. If $k < 1$ then an empty variable list is introduced; this is considered to have the same effect as skip_A^D (i.e. the no operation command). The proof that this command produces a design is presented in [Appendix C.1.3](#).

When the introduction of variable x is composed with program Q , x must be in the alphabet of Q (i.e. $x \in \alpha Q$), and the scope of the existential quantification is extended to include program Q ; that is $(\text{var } x) \circ Q = (\exists x \bullet Q)$. This result is proved in [Appendix C.2.2](#); it relies on the (existential) renaming in the definition of sequential composition.

Initialised variable introduction: It is sometimes convenient to initialise the variables as they are introduced.

$$\text{var } x_{i=1}^k := e_{i=1}^k \hat{=} (\text{var } x_{i=1}^k) \circ (x_{i=1}^k := e_{i=1}^k)$$

Variable completion command: The local variables $x_{i=1}^k$ are removed from the current scope by the command $(\text{end } x_{i=1}^k)$.

$$\begin{aligned} \text{end } x_{i=1}^k &\hat{=} \exists x'_{i=1} \bullet \text{skip}_A^D \\ &= \text{true} \vdash \exists x'_{i=1} \bullet \text{skip}_{A1} \end{aligned}$$

Here, the alphabet of the command as a whole, i.e. $\alpha(\text{end } x_{i=1}^k)$, is $A \setminus \{x_{i=1}^k\}$ and the alphabet $A1 = A \setminus \{\Pi_{\text{OK}}, \Pi_{\text{OK}}'\}$. If $k < 1$ then an empty variable list is completed in a similar manner to that of variable introduction; i.e. it is considered to have the same effect as skip_A^D . The proof that this command produces a design is presented in [Appendix C.1.3](#).

When a program Q is composed with the command representing the completion of variable x 's scope, then the scope of the existential quantification is extended to include program Q ; i.e. $Q \circ (\text{end } x) = (\exists x' \bullet Q)$. This result is proved in [Appendix C.2.2](#), in a similar manner to that of variable introduction.

Variable declaration command: In this dissertation local variables are used in a *traditional* block-scoping role; i.e. variables are introduced for a known limited scope. The following two block-scoping variable declaration commands are used to introduce an uninitialised collection and an initialised map of local variables respectively.

$$\begin{aligned} \text{decl } x_{i=1}^k \text{ in } P &\hat{=} (\text{var } x_{i=1}^k) \circ P \circ (\text{end } x_{i=1}^k) \\ \text{decl } \{x_{i=1}^k \mapsto v_i\} \text{ in } P &\hat{=} (\text{var } x_{i=1}^k := v_{i=1}^k) \circ P \circ (\text{end } x_{i=1}^k) \end{aligned}$$

Note that both of the syntactic forms of block scoping are suitable for a top-down structural free variable identification algorithm as previously discussed in [Section 2.1.3](#).

Given that the block variable declaration commands are applied to designs, then the resulting program is also a design, as follows, where u and V stand for $x_{i=1}^k$ and $v_{i=1}^k$ respectively.

$$\begin{aligned} \text{decl } u \text{ in } p \vdash P &= (\forall u \bullet p) \vdash (\exists u, u' \bullet P) \\ \text{decl } \{x_{i=1}^k \mapsto v_i\} \text{ in } p \vdash P &= p[V/u] \vdash (\exists u' \bullet P[V/u]) \end{aligned}$$

These results are proved in [Appendix C.1.5](#).

2.3.4 Hiding and filtering variables from programs

It is sometimes useful to hide some variables from a subprogram. For example, if a subprogram happens to declare a local variable with the same name as a global variable that it does not use, then, according to the standard UTP rules on variable introduction, either the local or global variable must be renamed. Within this context such hiding would provide a convenience mechanism for not having to rename variables. However, there are contexts where it is important to reduce the alphabet of a subprogram so that it can conform to external constraints. For example, in the context of recursion it is important that the alphabets either side of an unwinding of the fix-point are the same. The ability to hide unwanted local variables from the fix-point iteration then becomes a rather more important issue.

Hide variables command: The hide variables command $\text{hide } x_{i=1}^k \text{ from } P$ removes the variables $x_{i=1}^k$ from the alphabet of P , but following P 's completion restores these variables with their previous values.

$$\begin{aligned} \text{hide } _ \text{ from } P &\hat{=} P \\ \text{hide } x_{i=1}^k \text{ from } P &\hat{=} \exists y_{i=1}^k \mid y_{i=1}^k = x_{i=1}^k \bullet (\text{end } x_{i=1}^k) \circ P \circ (\text{var } x_{i=1}^k := y_{i=1}^k) \\ \alpha(\text{hide } \{x_{i=1}^k\} \text{ from } P) &= \{x_{i=1}^k, x'_i\} \cup \alpha P \end{aligned}$$

where

$$\begin{aligned} \forall_{i=1}^k x_i \notin \alpha P \wedge y_i \notin \alpha P \\ out \alpha P = (in \alpha P)' \end{aligned}$$

Note that the hide variables command is similar to that of the standard UTP alphabet extension command (P_{+A}), in the sense that both can be used to enable the sequential composition of a program with a smaller alphabet. The difference is that the hide variables command is specialised to work with designs, rather than general relations; applying the hide operator to a design $p \vdash P$ returns a design of the form $p \vdash P \wedge u' = u$, rather than $(p \vdash P) \wedge u' = u$, where u represents the vector of variables being hidden.

Lemma 2.6: Hide design

$$\text{hide } u \text{ from } p \vdash P \quad = \quad p \vdash P \wedge u' = u$$

□

This lemma is proved in [Appendix C.1.4](#).

Sometimes it is useful to limit a subprogram to a specific set of variables. This can be achieved by hiding all variables in the alphabet except for those of interest. For example, it is possible to limit the scope of subprogram P to the alphabet $\{u, u'\}$, within a program that has an alphabet $\{w, w'\}$, using the command

$$\text{hide } w|u \text{ from } P$$

where $w|u$ is the vector w without the variables contained in vector u . One issue with this approach is that it is not possible to infer what the alphabet of the surrounding program should be. Therefore, any *compilation scheme* that used this form of hiding would become context dependant, in the sense that it would have to maintain the program's current alphabet. Our original UTP constant-map-design ($\mathcal{U}_M$ -design) of the $\mathcal{O}$ -calculus did precisely this, in order to enable the definition of the method invocation command to appropriately reduce the scope of the program representing the target method (the current $\mathcal{U}_M$ -design is presented in [Section 3.4](#)). This significantly complicated the consistency proofs for this model, and thus support for a context independent compilation scheme was sought.

The adopted solution to the context independent compilation issue was inspired by the observation that scoped (hidden) subprograms were being sequentially composed with subprograms that could be compiled in a context independent manner – specifically that the alphabets of these independent subprograms could be used to infer the alphabet of a scoped subprogram. Here, the new UTP command would take three arguments: the independent subprogram; the list of variables to be kept; and the scoped subprogram.

Filter-in composition command: The left- and right-hand filter-in composition commands perform a sequential composition of subprograms P and Q , where only the variables contained in the vector u are allowed in the subprogram P .

$$\begin{aligned} P \upharpoonright u \circ Q &\widehat{=} (\text{hide } w|u \text{ from } P) \circ Q \\ Q \circ u \upharpoonright P &\widehat{=} Q \circ (\text{hide } w|u \text{ from } P) \end{aligned}$$

Here, vector w contains all the variables in the input alphabet of Q , and $w|u$ is the vector w without the variables contained in vector u .

It is also possible to define left- and right-hand filter-out composition commands in a similar manner, though in this case it is not required for simplifying the compilation process.

Filter-out composition command: The left- and right-hand filter-out composition commands perform a sequential composition of subprograms P and Q , where the variables contained in the vector u are removed from (hidden in) the subprogram P .

$$\begin{aligned} P \downarrow u \circ Q &\hat{=} (\text{hide } u \text{ from } P) \circ Q \\ Q \circ u \downarrow P &\hat{=} Q \circ (\text{hide } u \text{ from } P) \end{aligned}$$

2.3.5 Higher order programming

Church’s λ -calculus provides a simple, yet very powerful, model for the whole of computing, based on an abstract idea of what a function is [Chu36, Chu85, Pie02]. Briefly, the main idea is that of the lambda abstraction, i.e. “ $f = \lambda(x) e$ ”, where f is a single parameter function, whose formal parameter x may be used in the definition of its body (expression) e . Having provided a syntax for defining a function, the λ -calculus also provides a syntax for applying it, namely “ $f x$ ”, where the function f is applied to the argument x , which itself can be a function. It is this concept of applying a function to a function (i.e. treating a function as data or a value) that is given the term *higher order* in this context.

In general, the concept of higher order programming can be used to enable programs – such as macros, procedures and subroutines – to be treated as data. The important point is that suitable program variables can have inner programs assigned to them. Chapter 9 of the UTP book [HH98] provides a succinct introduction to these concepts, from both imperative and functional language points of view, which can be adapted to model OO-methods. Having said this, some would argue that [HH98] simulates, rather than provides, higher ordered programming to the UTP. Here, functions (and procedures) are modelled in terms of predicate relations, rather than being a first class concept. Specifically they are modelled as UTP predicate texts that are extracted when needed.

Before continuing with the discussion on higher order programming we illustrate how higher order programming can be simulated in a set-theoretic context, such as provided by the Z specification language [ISO02, Spi92, WD96].

Example 2.7: Suppose that a programming language has two basic data types, the natural numbers ($\mathbb{N}$) and the booleans ($\mathbb{B}$). Further, let there be two operations for constructing new data types from old, namely, the usual power-set ($\mathbb{P}_-$) and Cartesian product ($_- \times \dots \times _-$) type constructors. This enables the type of an anonymous record (a tuple) that contains a natural number followed by a boolean to be represented by $\mathbb{N} \times \mathbb{B}$. Now a function *even* can be modelled by a set of such tuples, one for each natural number n , where: the first element of the tuple is set to n , and the second is set to $true_{\mathbb{B}}$ when n is even and the value $false_{\mathbb{B}}$ otherwise. Formally this can be defined by $even \hat{=} \{(n, n \bmod 2 = 0) \mid n \in \mathbb{N}\}$, where the *even* function can be represented as an element of the type $\mathbb{P}(\mathbb{N} \times \mathbb{B})$.

□

The previous example demonstrates how a function can be modelled by a set of pairs. In general this modelling technique can be extended to any relation, and illustrates the key insight behind the simulation of higher order functionality in the UTP.

Object-oriented modelling: Two teams have made significant progress with adding class-based OO-support to the UTP. They have both used higher order programming concepts, in the sense that object values essentially contain the methods (functions) that can be invoked on them. The first team’s work [HLL02a, HLL02b, HLXS04, HLL06, CHH⁺07] has resulted in the development of the Refinement Calculus of Object Systems (rCOS) language; it supports object references. The second team’s work [CSW05, SCS06, CHW06, HCW08] is being used to extend the Circus language; here the work on objects and pointers has been done independently, but with a view to merging the results.

Our work is similar to these in that we use higher order concepts to represent methods as designs, whose semantics is determined by an appropriate weakest fix-point construction. The key difference, is that in our approach methods are not associated with a class, but with an object.

2.4 Heap Storage Locations

Sharing (*aliasing*) objects is supported by most mainstream OO-languages, such as C++ and Java. Here objects are dynamically created via the `new ClassName(parameters)` constructor, which allocates space on the heap for the object of type *ClassName*, initialises its fields using the supplied *parameters*, and then returns a pointer (object reference) to the location (address) on the heap where the object is stored. Such object references can then be assigned (copied) to other variables, and thus the objects that they are pointing at can be aliased (shared). [Appendix A.3](#) contains a slightly more detailed discussion on aliasing.

An object's location can be considered to represent its unique *identity*, as once created an object retains its heap location (identity) for the duration of its existence; in other words the object's identity is immutable. Some consider that “object identity is the foundation of object-oriented programming” [NVP98]. Our $\mathcal{O}$ -calculus does not support this notion of unique object identity. Instead, it supports an explicit notion of *self* that is independent of where the object is located. From our point of view, heap storage provides a means for sharing objects, such as required by typical implementations of the doubly-linked-list.

Example 2.8: A record-like object for representing both an unset node and an example middle node in a doubly-linked-list could be defined as follows in our $\mathcal{O}$ -calculus.

$$\begin{array}{ll} \text{UnsetNode} \equiv [\text{data} = \varsigma(-) \text{!} & \text{MidNode} \equiv [\text{data} = \varsigma(-) 42 \\ \text{next} = \varsigma(-) \text{nil} & \text{next} = \varsigma(-) \ell_{57} \\ \text{prev} = \varsigma(-) \text{nil} & \text{prev} = \varsigma(-) \ell_{15} \\] &] \end{array}$$

Here, the *MidNode* node in the middle of the list is shared by its two adjacent nodes (at locations ℓ_{57} and ℓ_{15}), where the preceding node's (ℓ_{15} 's) **next** field and the following node's (ℓ_{57} 's) **prev** field both contain an object reference to the *MidNode* node.

□

Semantics: Two broad categories of semantic modelling approaches are value semantics and location (or reference) semantics. Here a value semantics is typically written in terms of a map from variables to values, whereas a location semantics is typically written in terms of two maps, from variables to locations and locations to values. Having said this, some variants of a location semantics simply add a map from locations to values and treat a location as a value.

Alternatively, the layout of storage can be viewed as a directed graph. One possibility is to view variables and values as nodes within the graph. Here the node representing a compound value, such as a record, would have an outbound edge for each of its fields, which is labelled by the distinct field name, to the value for that field. This corresponds to the concrete graph of locations presented in [Section 5.2](#). Another possibility is to view the nodes as locations and the edges as the values. This corresponds to the abstract graph of locations presented in [Section 5.3](#). Further, these abstract graphs underpin the fully abstract model of memory, that is used by our last UTP model $\mathcal{O}$ -calculus.

Chapter 3

Unifying Theories of Objects

This chapter presents a family of three Unifying Theories of Programming (UTP) models of our $\mathcal{O}$ -calculus:

1. the UTP operand–stack–design ($\mathcal{U}_s$ -design), which stores the result of evaluating an operand, i.e. a subprogram, on an operand stack;
2. the UTP result–value–design ($\mathcal{U}_r$ -design), which replaces this stack with a combination of local variables and a result value; and
3. the UTP constant–map–design ($\mathcal{U}_m$ -design), which extends the use of variables so that they represent the direct and inherited parameters of function calls and method invocations.

These models extend the notion of a UTP design with values and commands for modelling objects, methods, and pointers (references). Here the approach is to convert each $\mathcal{O}$ -calculus expression into a UTP command that denotes its semantics.

Aside 3.1: The reason that an $\mathcal{O}$ -calculus expression is transformed into a UTP command, rather than a UTP expression, is because programs are represented by expressions in the $\mathcal{O}$ -calculus and commands in the UTP.

□

The notion of a *command-expression* is introduced to systematically transform the majority of the $\mathcal{O}$ -calculus's expressions into equivalent UTP programs, as follows:

1. transform each argument in the expression to a UTP program that calculates its result;
2. evaluate the arguments in the same order as that of the $\mathcal{O}$ -calculus's operational semantics;
3. store the evaluated arguments (e.g. in an operand stack or some local variables); and
4. apply the underlying $\mathcal{O}$ -calculus operation on the stored values, and store the result.

Item 4 is handled by a model specific *transformation-command*, which applies the $\mathcal{O}$ -calculus operation to the evaluated arguments and appropriately stores the result as a UTP program. Specifically, a UTP program that produces the given result.

Some operations such as method invocation do not follow this model, as they require more sophisticated handling, in this case the use of fix-points.

3.1 Extending the UTP design

The UTP notion of a design is extended so that it can model $\mathcal{O}$ -calculus entities such as objects, labels, methods, functions, and abstract locations (along with their associated heap).

3.1.1 Objects and labels

An object is modelled by a partial map from labels to method definitions, where a label is denoted by an alphanumeric identifier and a method as a tuple of program-texts (Section 3.1.2). This mirrors the approach of the $\mathcal{O}$ -calculus, thus has similar benefits and limitations. The main benefit is that all members of an object have the same form: they are methods. The main limitation is that it is difficult to distinguish a method that represents a field from a method whose body consists of a single – possibly compound – value.

Aside 3.2: Originally we planned to model an object as a partial map from object member-labels to method-labels, which were used to index a global method table. The idea was that this table would be useful for the fix-point definition of a mutually recursive method’s semantics. A benefit of this approach was that an object only contained scalar values, i.e. member and method labels, rather than the compound method definitions and object values. However, the indirection through a global method table proved to be both unnecessary and a hindrance. First, it introduced the problem of naming methods, and managing these names. Second, the extra indirection complicated the definition of method invocation, by introducing two extra steps in the following invocation process – specifically steps 2 and 3.

1. check whether the object contained the member-label;
2. get the corresponding global method-label;
3. check whether this method-label is in the global method table;
4. get the corresponding method;
5. call the method with the object (in step 1) as its self-parameter.

Only stages 1, 4, and 5 are required in our current model of method invocation.

□

3.1.2 Methods and functions

Within the higher-order chapter of the UTP book [HH98] subroutines (e.g. functions) are essentially modelled by *program texts* that are unpacked (extracted) and then sequentially composed with their calling environment. An object’s methods can be viewed in a similar manner. One important observation is that the calling context sets the method’s initial alphabet. Therefore care must be taken to ensure that we know what the calling context’s alphabet is. A simple approach to address this problem is to constrain the program alphabet at the point of a method (or function) call so that it only contains the UTP model’s observation variables (and not any program variables). This approach is sufficient for our purposes and is implemented by each of our systematic compilation schemes. For convenience when we wish to talk about both a method and a function, we shall use the term *procedure*.

One significant consequence of having a constant alphabet at the point of a procedure call is that the initial alphabet of every procedure is identical. This uniformity is essential for the way in which mutually recursive procedures have been handled. Here, each model of the $\mathcal{O}$ -calculus provides a single command for invoking a procedure, which identifies both the argument value and the program-text from the UTP model’s state variables; that is, the values of the variables in the program’s alphabet. This enables the definition of mutually recursive procedures to be given in terms of a fix-point over program states.

Aside 3.3: An alternative way of modelling the $\mathcal{O}$ -calculus method would be to treat it in a similar manner to that of a single argument untyped lambda calculus (λ -calculus) function – specifically, that the method’s object is passed by value as the argument to the function, which is then treated as an immutable parameter (constant) within its definition. Inner methods could then be defined in a similar manner to the way in which inner functions are defined in

the λ -calculus. Here, the parameters containing the immutable object arguments of the outer scope methods are available for use within a inner method, as read-only global variables. These global variables must be included within the alphabets of the methods that use them, but not in the alphabets of the methods that declare (introduce) them. The important point is that this approach leads to methods which have different alphabets, and thus do not conform to the limitations that we have chosen to adopt.

□

Within each of our UTP object models, non-local variables (e.g. from an outer method self-variables or function parameters) must be either substituted for their values prior to their procedure's invocation or passed as part of their procedure's invocation. This ensures that any references to the parent's formal parameter are available for use within the inner procedure invocation.

Example 3.4: In [Section 2.2.5](#) we introduced a *Factorial* object that contained a single method, which made use of an anonymous inner function. This is now reproduced for our convenience.

$$\text{Factorial} \quad \widehat{=} \quad [\text{fact} = \varsigma(x) \lambda(i) \text{ if } i < 2 \text{ then } 1 \text{ else } i \times x.\text{fact}(i - 1)]$$

Note that the anonymous inner function body (i.e. $\text{if } i < 2 \text{ then } 1 \text{ else } i \times x.\text{fact}(i - 1)$) contains a reference to the outer method's self parameter (x). It is this outer variable (x) that needs to be substituted for its value or passed to the inner procedure as part of its invocation.

□

The simplest approach to the non-local variable issue is to ensure that the procedure invocation processes substitute (a syntactic variant of) their argument's value for their formal parameter in their *program-text*, prior to its extraction. Having said this, there are a couple of hurdles to overcome:

1. the ability to identify all *free* occurrences of the variable (parameter) within a program-text, so that only these get substituted;
2. the ability to transform any semantic value, which can be passed as an argument to a method or function, into a suitable program text.

The standard UTP technique for overcoming the first hurdle is to ensure that inner bound variables never have the same name as a free (alphabet) variable. This constraint causes problems with the $\mathcal{O}$ -calculus compilation process to the UTP. It can be easily overcome by restricting our program texts to a UTP syntax that has a convenient block scoping structure. Having restricted our program texts to this convenient form, it is straightforward to define a free-variable substitution processes, for each of our UTP models, in a manner similar to that of the $\mathcal{O}$ -calculus ([Section 2.2.5](#)).

The second hurdle is also straightforward to overcome due to the form of the values within our model. Here, most values are their own texts (e.g. the literal 3). Functions (and methods) are represented by a tuple of texts, which already conform to the variable scoping compilation restriction, thus their textual representation can be this tuple. Objects are represented as a map from labels to methods, which we textually denote by a set of pairs, where labels are their own texts and methods are tuples of texts.

3.1.3 Abstract locations and the heap

The $\mathcal{O}$ -calculus directly models the heap as a map from abstract locations to values. This map is stored as an environmental context variable, which is updated as a *side effect* of reducing a term, using the small step reduction rules. Having said this, the main purpose of the rules for creating, updating, and dereferencing heap locations is to manage such side effects.

Before continuing with our discussion on the modelling of the $\mathcal{O}$ -calculus, it is worth recalling that its **fresh** operation has the side effect of updating the **HEAP** context variable. This side effect has important implications on the ordering of evaluation as illustrated by the following example.

Example 3.5: Consider creating a zero-initialised fresh abstract location, whose value (not its contents) is used to initialise two fields of a *new* object.

$$\text{let } loc \hat{=} (\text{fresh } *= 0) \text{ in } [\mathbf{f1} = \varsigma(-) \text{ loc}, \mathbf{f2} = \varsigma(-) \text{ loc}]$$

Note that it is not possible to replace the definition of the location variable *loc* in the body of the local block, as this change in execution order would result in fields **f1** and **f2** pointing to different abstract locations, rather than sharing the same abstract location. Essentially, this is because the **fresh** command takes an environmental parameter, which is updated on each use. Therefore, the UTP predicate that formally models the operational semantics of the **fresh** command needs to be parametrised by this environmental information.

If two distinct fresh locations were required then the following, subtly different, $\mathcal{O}$ -calculus program could be used.

$$\text{let } loc \equiv (\text{fresh } *= 0) \text{ in } [\mathbf{f1} = \varsigma(-) \text{ loc}, \mathbf{f2} = \varsigma(-) \text{ loc}]$$

Here, the only difference is in the **let** variable-binder which changes from definition ($\hat{=}$) to macro-expansion ($\equiv$).

□

The simple heap approach presented in this chapter mirrors that of the $\mathcal{O}$ -calculus. Therefore, this chapter's UTP models of a heap all inherit the *garbage* detection and collection issues of the $\mathcal{O}$ -calculus; that is, the issues with being able to detect and delete locations on the (abstract) heap that cannot be referenced by the program's execution state. Note that the operational semantics of the $\mathcal{O}$ -calculus ignores such garbage collection issues, as this complication was not necessary for illustrating the usefulness of the calculus. Within the context of the UTP models this issue is perhaps more important, as we wish to be able to say when one program is equal to another, and this will involve a comparison of abstract heap locations.

A garbage collection mechanism can be added to the $\mathcal{O}$ -calculus, via a function that examines the current usage of the heap by the terms representing a program and its context, and removes the unreferenceable locations. A similar approach could be adopted for the first of our UTP models, but not the latter two, as they use existential quantification to temporarily hide references to the heap (whereas the first does not). [Chapter 5](#) provides an alternative – fully abstract – model of a heap, which is used to address these issues. It is based on the idea of representing store as a directed graph, which has anonymous nodes with distinctly labelled outbound edges. [Section 3.4.3](#) presents an alternative compilation scheme which ensures that all compound terms (e.g. objects and methods) have distinctly named components, and thus can be directly represented by such a directed graph.

3.1.4 Modelling notes

Badly behaved programs: A program is considered to be *badly behaved* if it does not successfully terminate. Infinitely recursing (looping) programs are badly behaved and are denoted by the **chaos** command in each of our UTP object models, which is the weakest fix-point of a non-terminating design [HH98, WC02]. Therefore, **chaos** appears to be a good choice for modelling badly behaved programs. However, an $\mathcal{O}$ -calculus program can also be badly behaved by becoming *stuck*.

Aside 3.6: Recall that an $\mathcal{O}$ -calculus program's operational semantics are written in terms of a collection of – non overlapping – small-step reduction rules. Therefore, given a program term, at most one rule can be applied to the whole term (rather than one of its subterms). If no rules can be applied, one of two conditions can hold: either the term representing the program is now a value, in which case it has successfully terminated, or else the program has become stuck (and thus is badly behaved).

□

The standard UTP approach is to not distinguish between different types of badness, and thus model stuck $\mathcal{O}$ -calculus programs by **chaos**. This is what we have done.

Note that the dynamic conditions under which a program becomes stuck are identifiable, as they correspond to when the arguments of an $\mathcal{O}$ -calculus term have the wrong form (e.g. type). In principle, the form of the arguments could be explicitly checked for within our UTP models, which in turn would enable them to generate a different sort of failure. This might be useful in the future, if the distinction between stuck and non-terminating is required. Currently, some of the checks are implicitly captured in the pre-conditions of the designs we use to model the $\mathcal{O}$ -calculus operations.

UTP program texts and compilation: Both UTP operations and arithmetic expressions are typically considered to directly denote semantic entities, rather than specific syntactic forms, which are then provided with a semantics.

Example 3.7: The arithmetic expressions $x \times 2$ and $x + x$ both denote the same mathematical entity, thus they can be substituted for each other in any (semantic) context. Here, it is assumed that the syntactic presentation of such a substitution automatically adds brackets as necessary (in order to prevent change of meaning due to the precedence of surrounding operators).

The UTP relational predicates $x := 4$ and $x := 2 \ ; \ x := x + x$ both denote the same mathematical entity, thus they can be substituted for each other in any (semantic) context, in a similar manner to that of the arithmetic example.

□

Within our models of the $\mathcal{O}$ -calculus methods are represented by UTP program texts whose syntactic form has the block scoping property; i.e. the scope of a bound variable can be limited to the children of the term that introduces the variable.

One consequence of introducing the concept of a UTP program text, is that we need to be able to convert between such texts and their semantic entities. Converting a UTP program text to its corresponding semantic entity, is straight forward; it is denoted by the usual semantic meaning brackets as follows

$$\llbracket t \rrbracket \hat{=} t$$

Here, t is a *valid* output of the program compilation process, which ensures that the syntactic form of its terms have the block scoping property. The reverse process of taking a semantic entity and producing a syntactic representation of is defined, in [Section 3.2.5](#), by the text_J operation (where J is the UTP model identifier).

Having provided some notation for converting between syntax and semantics, we introduce a notation for defining the systematic transformation (i.e. a compilation) of an $\mathcal{O}$ -calculus program into the corresponding UTP program. This compilation process is denoted by $\langle\langle t \rangle\rangle_J^M$, where t is the $\mathcal{O}$ -calculus term to be compiled, m is the optional – staged compilation – mode, and J is the UTP model identifier.

We also provide a shorthand for the presenting the UTP semantics of an $\mathcal{O}$ -calculus term as follows:

$$\llbracket t \rrbracket_J^M \hat{=} \llbracket \langle\langle t \rangle\rangle_J^M \rrbracket$$

3.2 Operand Stack Model

The UTP operand-stack-design ($\mathcal{U}_s$ -design) extends the notion of a UTP design with an operand stack for storing intermediate results, and a heap map for storing dynamically allocated values in abstract heap locations. Formally this stack and heap are denoted by the UTP context variables Π_{STK} , Π_{HEAP} , Π_{STK}' , and Π_{HEAP}' , which represent the input and output states (values) of the operand stack and heap storage respectively.

The contents of the stack are the values (e.g. the integers, objects, functions, and heap locations) and other semantic entities (i.e. labels and methods), which between them represent the operands of the $\mathcal{O}$ -calculus operations. The idea is that following the execution of a subprogram the top value on the stack represents its result.

The heap storage (map) context is essentially taken from the $\mathcal{O}$ -calculus, the main differences being in the changes to its name and the precise representation of its contents (values). In particular, the restriction that the heap can only contain values is kept; thus unlike the stack a heap cannot contain labels or methods.

3.2.1 Literal value programs

The simplest object calculus program is represented by a literal value, which is also the final result value of the program. Therefore, the compilation of such an $\mathcal{O}$ -calculus program must result in the $\mathcal{U}_s$ -design that pushes a single value onto the operand stack. This is achieved by the following $\mathcal{U}_s$ -design evaluation ($\mathcal{E}_s$), which pushes an operand value (ov) onto the operand stack.

$$\mathcal{E}_s \text{ } ov \quad \widehat{=} \quad \{ \Pi_{\text{STK}} \} : (\text{true} \vdash \Pi_{\text{STK}}' = \langle ov \rangle \wedge \Pi_{\text{STK}})$$

Note that the evaluation command has been subscripted by the operand stack model identifier (i.e. s). The evaluation command in the other models will be subscripted by their model identifiers, so that these commands can be distinguished. In general, this form of subscripting is applied to all the model specific commands (and utility operations).

The compilation of a literal value lv to the $\mathcal{U}_s$ -design is now defined in stages by the following two compilation rules. Here, the first rule compiles the value to a UTP program, whereas the second rule compiles the value to a UTP expression.

$$\llbracket lv \rrbracket_s \quad \widehat{=} \quad \mathcal{E}_s \llbracket lv \rrbracket_s^E \qquad \llbracket lv \rrbracket_s^E \quad \widehat{=} \quad lv$$

Note that these compilation rules produce $\mathcal{U}_s$ -design texts, which can then be converted into a $\mathcal{U}_s$ -design program by applying the semantic meaning brackets as follows.

$$\llbracket t \rrbracket \quad \widehat{=} \quad t$$

Here, t is a *valid* output of the program compilation process, which ensures that the syntactic forms of its terms have the *block scoping* property; i.e. the scope of a bound variable can be limited to the children of the term that introduces the variable. Therefore, the texts for standard UTP variable introduction and elimination commands (i.e. `var` and `end`) cannot be output as part of the compilation process. This is important for defining the free-variable substitution function (Section 3.2.6) using a structural induction over UTP program texts.

Before continuing with the model description it is worth illustrating the UTP semantics of a simple literal value program.

Example 3.8: The $\mathcal{U}_s$ -design for the $\mathcal{O}$ -calculus program in the ‘value example’ on page 14 is:

$$\begin{aligned} \llbracket 42 \rrbracket_s &= \mathcal{E}_s \llbracket 42 \rrbracket_s^E = \mathcal{E}_s 42 \\ &= \Pi_{\text{STK}} := \langle 42 \rangle \wedge \Pi_{\text{STK}} = \{ \Pi_{\text{STK}} \} : (\text{true} \vdash \Pi_{\text{STK}}' = \langle 42 \rangle \wedge \Pi_{\text{STK}}) \end{aligned}$$

Recall that the resultant value is that contained on the top of the stack.

□

3.2.2 Function values

In the $\mathcal{U}_s$ -design a function value is defined as a pair of compiled program texts that represent the function's parameter and body. It is denoted by $\langle\!\langle x, P \rangle\!\rangle$, where the scope of the parameter x is the program-text P .

$$\langle\!\langle \lambda(x) e \rangle\!\rangle_s \hat{=} \mathcal{E}_s \langle\!\langle \lambda(x) e \rangle\!\rangle_s^E \quad \langle\!\langle \lambda(x) e \rangle\!\rangle_s^E \hat{=} \langle\!\langle \langle\!\langle x \rangle\!\rangle_s^E, \langle\!\langle e \rangle\!\rangle_s \rangle\!\rangle$$

Here, a variable is represented by itself in the declaration (evaluation) context and the evaluation of itself in the general program context.

$$\langle\!\langle x \rangle\!\rangle_s \hat{=} \mathcal{E}_s \langle\!\langle x \rangle\!\rangle_s^E \quad \langle\!\langle x \rangle\!\rangle_s^E \hat{=} x$$

Note that the evaluation ($\mathcal{E}_s$) of a variable is not defined by the $\mathcal{U}_s$ -design, but it may be contained in the function's body (i.e. a compiled program text). Such variables are substituted by their values, prior to the program text being extracted to a $\mathcal{U}_s$ -design subprogram. Details of the program text variable-substitution and extraction processes are presented in the discussion of function and method invocation ([Section 3.2.6](#)).

3.2.3 Object values and method definitions

In the $\mathcal{U}_s$ -design an object value is defined as a map from labels to methods. This is denoted by $\{_{i=1}^k l_i \mapsto m_i\}$, where k represents the number of object methods m_i with distinct labels l_i . The compilation of an object value is similar to that of literal values.

$$\begin{aligned} \langle\!\langle [_{i=1}^k l_i = m_i] \rangle\!\rangle_s &\hat{=} \mathcal{E}_s \langle\!\langle [_{i=1}^k l_i = m_i] \rangle\!\rangle_s^E \\ \langle\!\langle [_{i=1}^k l_i = m_i] \rangle\!\rangle_s^E &\hat{=} \{_{i=1}^k \langle\!\langle l_i \rangle\!\rangle_s^E \mapsto \langle\!\langle m_i \rangle\!\rangle_s^E \} \end{aligned}$$

A method is defined as a pair of compiled program texts that represent the method's *self* variable and body. Like the function, it is denoted by $\langle\!\langle x, P \rangle\!\rangle$, where the scope of the *self* variable (parameter) x is the program text P . However, methods differ from functions in their compilation; specifically, a method cannot occur as a top level program as it is not considered to be a value, thus it only has an evaluation moded compilation scheme.

$$\langle\!\langle \varsigma(x) e \rangle\!\rangle_s^E \hat{=} \langle\!\langle \langle\!\langle x \rangle\!\rangle_s^E, \langle\!\langle e \rangle\!\rangle_s \rangle\!\rangle$$

3.2.4 Command expressions

In the $\mathcal{O}$ -calculus almost all the programming operations are expressions. Such expressions are converted into UTP commands that evaluate each of their arguments, store the results on the operand stack, and then apply an appropriate stack transformation command. This is modelled by the `cmdExps` command, which applies a transformation program (command) P to the result of executing a list of programs $P_{i=1}^k$ – representing the eager evaluation of its arguments – in a left to right order.

$$\text{cmdExp}_s(P, (P_{i=1}^k)) \hat{=} (\wp_{i=1}^k P_i) \circ P$$

Note that some of the $\mathcal{O}$ -calculus operations that we wish to transform have *unevaluable* arguments, such as the method update operation $e.l \Leftarrow m$, where the expression e should evaluate to an object containing an unevaluable label l , which indexes the method that is going to be replaced by the unevaluable method m . In general, the contents of unevaluable arguments are considered to be already evaluated, and are thus loaded on to the operand stack.

An example of a $\mathcal{U}_s$ -design stack transformation program is the `transs( $f, k$ )` command, which takes a k -parameter function f – that defines the operation being modelled – and constructs

a UTP program that applies this function to the top k elements of the operand stack. Care must be taken to ensure that the parameters are in the order that they are going to appear on the operand stack, as this may not be the same as the left-to-right declaration order.

Given that the meta-variables $x_{i=1}^k$ represent the arguments for function f , then the updated stack can be modelled by $\langle f(x_{i=1}^k) \rangle \cap (\text{tail}^k \Pi_{\text{STK}})$, assuming that: it has started ($\Pi_{\text{OK}} = \text{true}$); there are sufficient arguments ($k \leq \#\Pi_{\text{STK}}$); and these arguments are in the domain of the function being modelled ($x_{i=1}^k \in f$).

$$\begin{aligned} \text{trans}_s(f, k) &\hat{=} \begin{array}{c} k \leq \#\Pi_{\text{STK}} \wedge (\Pi_{\text{STK}} \text{ proj } k) \in f \\ \vdash \\ \Pi_{\text{STK}}' = \langle f((\Pi_{\text{STK}} \text{ proj } k)) \rangle \cap (\text{tail}^k \Pi_{\text{STK}}) \end{array} \\ (s \text{ proj } k) &\hat{=} (\overset{k}{i=1} \text{ head}(\text{tail}^{k-i} s)) \end{aligned}$$

Example 3.9: The $\mathcal{O}$ -calculus addition operation can be modelled in terms of the cmdExp_s operation as follows:

$$\begin{aligned} \langle \langle e_1 + e_2 \rangle \rangle_s &\hat{=} \text{cmdExp}_s(\\ &\quad \text{trans}_s((- + -), 2), \\ &\quad (\langle \langle e_1 \rangle \rangle_s, \langle \langle e_2 \rangle \rangle_s) \\ &\quad) \\ &= \langle \langle e_1 \rangle \rangle_s \circ \langle \langle e_2 \rangle \rangle_s \circ \text{trans}_s((- + -), 2) \end{aligned}$$

□

Note, the definition of the cmdExp_s generalises the one presented in [SG07] by removing the limitation that only the trans_s command could be used to specify the appropriate transformation. It was generalised in order to minimised the difference between the UTP object models in this chapter and the one in Section 6.2.

3.2.5 Method updates and field assignments

Method update in the $\mathcal{O}$ -calculus is compiled in two parts: first, the terms representing the arguments are compiled; and second, they are combined by an appropriate method update transformation function.

$$\begin{aligned} \langle \langle e.l \Leftarrow m \rangle \rangle_s &\hat{=} \text{cmdExp}_s(\\ &\quad \text{trans}_s(\text{methUpd}, 3), \\ &\quad (\langle \langle e \rangle \rangle_s, \mathcal{E}_E \langle \langle l \rangle \rangle_s^E, \mathcal{E}_E \langle \langle m \rangle \rangle_s^E) \\ &\quad) \\ \text{methUpd} &\hat{=} \{ (o, l, m) \mapsto o \oplus \{l \mapsto m\} \mid \\ &\quad (o, l, m) \in \text{Object} \times \text{Label} \times \text{Method} \wedge l \in o \\ &\quad \} \end{aligned}$$

Note that the methUpd function can never be called outside its domain, as it is always used in the context of the trans_s UTP program, which checks to see if the arguments are in the domain of the transformation function to be used – in this case methUpd ; if the arguments are in the domain, then the transformation function is applied and the result stored on the top of the operand stack, otherwise the program is defined to become chaotic (i.e. become the **chaos** program).

A field assignment is compiled in a similar manner, except for the use of a utility function (text_s) for converting a variable, a value, a method, and some programs to their corresponding program texts.

$$\begin{aligned} \llbracket e_1.l := e_2 \rrbracket_s &\hat{=} \text{cmdExp}_s(\\ &\quad \text{trans}_s(\text{fldUpd}, 3), \\ &\quad (\llbracket e_1 \rrbracket_s, \mathcal{E}_s \llbracket l \rrbracket_s^E, \llbracket e_2 \rrbracket_s) \\ &\quad) \\ \text{fldUpd} &= \{ (o, l, v) \mapsto o \oplus \{l \mapsto \langle _, \text{text}_s(v) \rangle\} \mid \\ &\quad (o, l, v) \in \text{Object} \times \text{Label} \times \text{Value} \wedge l \in o \\ &\quad \} \end{aligned}$$

Recall that in [Section 3.1.4](#) we introduced the text_s helper operation, and said that it takes an UTP semantic entity and produces a textual representation of it. What we did not say was what such a textual representation would be.

In [Section 3.1.2](#) we provided an outline of how operand values could be converted back into their textual forms. Here, variables and atomic operand values (which include object labels) are their own texts. Procedures are represented by a tuple of texts, which already conform to the variable scoping compilation restriction, thus their textual representation can be this tuple. Objects are represented as a map from labels to methods, which we textually denote by a set of label-method pairs.

In addition to these operand-values we provide textual representations for those programs that are supplied as a fix-point argument (as discussed in [Section 3.2.6](#)). For the moment it is sufficient to say that these can be constructed solely from the call_s , apply_s and chaos programs, each of which can be represented by its own name, within a program text. Further, as we do not need to convert other semantic entities into program texts, no definition for these semantic entities is given.

It is now possible to indirectly define the text_s operation by a set of equations it must obey, as follows:

$$\begin{aligned} \llbracket \text{text}_s(x) \rrbracket &= x \\ \llbracket \text{text}_s(av) \rrbracket &= av \\ \llbracket \text{text}_s(\langle x, t \rangle) \rrbracket &= \langle x, t \rangle \\ \llbracket \text{text}_s(\{_{i=1}^k l_i \mapsto m_i\}) \rrbracket &= \{_{i=1}^k \text{text}_s(\llbracket l_i \rrbracket) \mapsto \text{text}_s(\llbracket m_i \rrbracket)\} \\ \llbracket \text{text}_s(\text{call}_s) \rrbracket &= \text{call}_s \\ \llbracket \text{text}_s(\text{apply}_s(z)) \rrbracket &= \text{apply}_s(\text{text}_s(\llbracket z \rrbracket)) \\ \llbracket \text{text}_s(\text{chaos}) \rrbracket &= \text{chaos} \end{aligned}$$

Note that the reason that this operation is not defined constructively is because more than one program text can represent the same semantic entity. For example, the identity function $\langle x, \mathcal{E}_s x \rangle$, can also be written as $\langle y, \mathcal{E}_s y \rangle$. Here, the proof that these two functions are semantically equal would involve demonstrating that for every context Γ and for all values v the equation $\langle x, \mathcal{E}_s x \rangle(v) = \langle y, \mathcal{E}_s y \rangle(v)$ holds. In practise this is not a problem for our proofs, as our textual denotation of the semantic entities mirrors that of the above equations, and we choose the obvious identity like conversion to a program text; that is for $\text{text}_s(\llbracket pt \rrbracket)$, we choose the program text pt . We could have chosen any other program text pt' , so long as $\llbracket pt' \rrbracket = \llbracket pt \rrbracket$, but then we would have to prove that this equation holds.

3.2.6 Method and function invocation

Method invocation in the $\mathcal{O}$ -calculus is compiled in two parts. First an object-member pair is constructed from the invocation arguments: an expression (e) representing the target object (o) and a label (l). This is achieved by retrieving the method with label l from object o . The second

part performs the actual method invocation, using a generic method call command (call_s). It executes the body of the method where the method's self variable has been instantiated with the calling object's value.

$$\begin{aligned} \langle\langle e.l \rangle\rangle_s &\hat{=} \text{cmdExp}_s(\\ &\quad \text{trans}_s(\text{omPair}, 2), \\ &\quad (\langle\langle e \rangle\rangle_s, \langle\langle l \rangle\rangle_s) \\ &\quad) \circ \text{call}_s \\ \text{omPair} &\hat{=} \{(o, l) \mapsto (o, o(l)) \mid (o, l) \in \text{Object} \times \text{Label} \wedge l \in o\} \end{aligned}$$

Function invocation is compiled in a similar manner. The key difference is that the function's argument is a general expression, rather than a label.

$$\begin{aligned} \langle\langle e_1 e_2 \rangle\rangle_s &\hat{=} \text{cmdExp}_s(\\ &\quad \text{trans}_s(\text{vtPair}, 2), \\ &\quad (\langle\langle e_1 \rangle\rangle_s, \langle\langle e_2 \rangle\rangle_s) \\ &\quad) \circ \text{call}_s \\ \text{vtPair} &\hat{=} \{(v, pt) \mapsto (pt, v) \mid (pt, v) \in \text{ProgramText} \times \text{Value}\} \end{aligned}$$

Before formally defining the generic call_s command, it is worth presenting three helper functions for operating with program texts – specifically, one for substituting value-texts for variable-texts, one for substituting a call-text with a fix-point value-text, and one for extracting a program text (i.e. converting a program text into a $\mathcal{U}_s$ -design). These functions are defined by cases, where the first case that matches is taken.

The program text substitution function ($t\{x \leftarrow ov\}$) performs the same role as that of its $\mathcal{O}$ -calculus counterpart, in that it replaces all free occurrences of the program variable x with the operand value ov in the program text t .

$$\begin{aligned} x\{x \leftarrow ov\}_s &\hat{=} ov \\ \langle\langle x, t \rangle\rangle_s\{x \leftarrow ov\}_s &\hat{=} \langle\langle x, t \rangle\rangle_s \\ t\{_{i=1}^k t_i\}\{x \leftarrow ov\}_s &\hat{=} t\{_{i=1}^k t_i\{x \leftarrow ov\}_s\} \end{aligned}$$

The fix_s function constructs a program text by substituting the terms representing a procedure invocation (call_s) with the program text representing the instantiation of the fix-point variable z . Note that this does not directly substitute the call_s texts within inner procedures; inner procedures texts are appropriately updated by the procedure application command apply_s , which is defined at the end of this section.

$$\begin{aligned} \text{fix}_s(\text{call}_s, z) &\hat{=} z \\ \text{fix}_s(\langle\langle x, pt \rangle\rangle_s) &\hat{=} \langle\langle x, pt \rangle\rangle_s \\ \text{fix}_s(t\{_{i=1}^k t_i\}, z) &\hat{=} t\{_{i=1}^k \text{fix}_s(t_i, z)\} \end{aligned}$$

The following ext_s function extracts a method's program text (t) given that there is a program text to represent the value of fix-point variable (z).

$$\text{ext}_s(t, z) \hat{=} \llbracket \text{fix}_s(t, \text{text}_s(z)) \rrbracket$$

It turns out that the value of the fix-point variable can always be represented in the form $\text{apply}_s^i(\text{chaos})$, where i is a natural number. Here, non-terminating programs are represented by $i = 0$ (i.e. chaos) and terminating programs by $n \geq 1$, where n is the number of method and function invocations that are actually invoked. In either case, programs of this form are handled by the existing program text generation function (text_s).

We are now in a position to define the call_s command. It is defined as the least fix-point of the apply_s function, which substitutes the argument v for the parameter x in program text pt .

$$\begin{aligned}
\text{call}_s &\hat{=} \mu z \bullet \text{apply}_s(z) \\
\text{apply}_s(z) &\hat{=} (\exists v, x, pt \mid (v, (\llbracket x, pt \rrbracket_s)) = \text{head } \Pi_{\text{STK}} \bullet \\
&\quad \text{pop}_s \circ \text{ext}_s(pt\{x \leftarrow \text{text}_s(v)\})_s, z) \\
&\quad) \\
&\quad \triangleleft \# \Pi_{\text{STK}} > 0 \wedge (\text{head } \Pi_{\text{STK}}) \in \text{Value} \times \text{ProgramText} \triangleright \\
&\quad \text{chaos} \\
\text{pop}_s &\hat{=} \{ \Pi_{\text{STK}} \} : (\# \Pi_{\text{STK}} > 0 \vdash \Pi_{\text{STK}}' = \text{tail } \Pi_{\text{STK}})
\end{aligned}$$

3.2.7 Other programming mechanisms

Conditional Execution: The UTP provides a notion of conditional execution, where a boolean expression is used to select the appropriate subprogram. It differs from the $\mathcal{O}$ -calculus's conditional execution operator in two significant ways:

1. The $\mathcal{O}$ -calculus allows the execution of an arbitrary program to determine the value of its boolean control argument, whereas the UTP conditional command only allows a predicate that is completely defined in terms of the conditional's input alphabet (i.e. a UTP conditional).
2. The evaluation of the $\mathcal{O}$ -calculus conditional expression need not produce a boolean value (or for that matter any result), whereas the UTP conditional expression is required to be both a predicate and defined.

The first of these differences is catered for by extracting the program that determines the value of the conditional expression and composing it with the UTP switch_s program.

$$\llbracket \text{if } e_1 \text{ then } e_2 \text{ else } e_3 \rrbracket_s \hat{=} \llbracket e_1 \rrbracket_s \circ \text{switch}_s(\llbracket e_2 \rrbracket_s, \llbracket e_3 \rrbracket_s)$$

The switch_s program essentially selects one of a pair of subprograms depending on the value on the top of the operand-stack. However, it also has to cater for the possibility that the evaluation of the conditional program did not produce a boolean value. In such cases, the corresponding $\mathcal{O}$ -calculus operational semantics would become *stuck*, and thus fail to produce a useful result; recall that we have chosen to model all non-useful programs by the UTP *chaos* command.

$$\begin{aligned}
\text{switch}_s(P_1, P_2) &\hat{=} (\text{pop}_s \circ P_1 \triangleleft \text{head } \Pi_{\text{STK}} \triangleright \text{pop}_s \circ P_2) \\
&\quad \triangleleft \# \Pi_{\text{STK}} > 0 \wedge (\text{head } \Pi_{\text{STK}}) \in \text{Boolean} \triangleright \\
&\quad \text{chaos}
\end{aligned}$$

Membership test: An example of the test that can precede a conditional execution is the object membership test $o?l$. Such a test is compiled using a command expression as follows.

$$\begin{aligned}
\llbracket o?l \rrbracket_s &\hat{=} \text{cmdExp}_s(\\
&\quad \text{trans}_s((_ \in _), 2), \\
&\quad (\llbracket l \rrbracket_s, \llbracket e \rrbracket_s) \\
&\quad)
\end{aligned}$$

Note that the order of the parameter evaluation has been swapped. This is fine because the program representing the evaluation of a label ($\mathcal{E}_s l$) does not depend on or alter the environment; i.e. it is independent of the program representing the expression e .

Sequential Composition: The sequential composition of two expressions can be straightforwardly represented by the sequential composition of their compiled programs, where the result of the first program is ignored (i.e. popped from the operand stack).

$$\langle\langle e_1; e_2 \rangle\rangle_s \hat{=} \langle\langle e_1 \rangle\rangle_s ; \text{pop}_s ; \langle\langle e_2 \rangle\rangle_s$$

Local declaration: The compilation of a local declaration block uses a similar rewriting strategy as that of our $\mathcal{O}$ -calculus (Section 2.2.4). Here, the *macro* declaration block is expanded (at compilation time), the *definition* declaration block is rewritten in terms of a function and its application, and the *list* declaration block is expanded.

$$\begin{aligned} \langle\langle \text{let } x \equiv t \text{ in } e \rangle\rangle_s &\hat{=} \langle\langle e_0 \rangle\rangle_s \\ \langle\langle \text{let } x \hat{=} e_1 \text{ in } e \rangle\rangle_s &\hat{=} \langle\langle (\lambda(x) e) e_1 \rangle\rangle_s \\ \langle\langle \text{let } d_{i=1}^k \text{ in } e \rangle\rangle_s &\hat{=} \langle\langle \text{let } d_1 \text{ in} \\ &\quad \text{let } d_{i=2}^k \text{ in } e \rangle\rangle_s \end{aligned}$$

where $e_0 \hat{=} e\{x \leftarrow t\}$ and $k \geq 2$

Note that we precompute the macro substitution before its body is compiled to ensure that this does not get confused with the compilation of the substitution scheme that is presented in Chapter 4.

3.2.8 Modelling the heap operations

The $\mathcal{U}_s$ -design of a heap mirrors that of the $\mathcal{O}$ -calculus, where the location, unset and null values are shared semantic entities between the models.

Fresh Operator: The $\mathcal{O}$ -calculus fresh operation is compiled to its corresponding $\mathcal{U}_s$ -design command.

$$\langle\langle \text{fresh} \rangle\rangle_s \hat{=} \text{fresh}_s$$

The fresh command creates a new location on the heap, which is initialised to the explicit unset value; it then pushes the value of this new location onto the operand stack.

$$\begin{aligned} \text{fresh}_s &\hat{=} \exists r \mid r = \text{fresh}_{\text{LOC}}(\text{dom } \Pi_{\text{HEAP}}) \bullet \\ &\quad \Pi_{\text{STK}}, \Pi_{\text{HEAP}} := \langle r \rangle \wedge \Pi_{\text{STK}}, \Pi_{\text{HEAP}} \oplus \{r \mapsto \mathbf{i}\} \end{aligned}$$

Note that this operation deliberately uses the same fresh location generation function as in the $\mathcal{O}$ -calculus, as it simplifies the consistency proof between the models. Without this we would have to have a notion of heap equivalence.

Dereference Operator: The dereference $\mathcal{O}$ -calculus operation is compiled by evaluating the expression representing the heap location, then applying the $\mathcal{U}_s$ -design's command for dereferencing the current result.

$$\langle\langle *e \rangle\rangle_s \hat{=} \langle\langle e \rangle\rangle_s ; \text{deref}_s$$

The heap dereference command (deref_s) takes the heap location on the top of the stack and replaces it with a copy of the associated heap value.

$$\begin{aligned} \text{deref}_s &\hat{=} \{ \Pi_{\text{STK}} \} : (\\ &\quad \# \Pi_{\text{STK}} > 0 \wedge (\text{head } \Pi_{\text{STK}}) \in \Pi_{\text{HEAP}} \\ &\quad \vdash \\ &\quad \Pi_{\text{STK}}' = \langle \Pi_{\text{HEAP}}(\text{head } \Pi_{\text{STK}}) \rangle \wedge (\text{tail } \Pi_{\text{STK}}) \\ &\quad) \end{aligned}$$

Heap Update Operator: The heap update $\mathcal{O}$ -calculus operation is compiled by evaluating its arguments in a left to right order, storing their results into a single location-value pair, and then applying the model of the heap update operation.

$$\langle\langle e_1 * e_2 \rangle\rangle_s \hat{=} \text{cmdExp}_s(\text{trans}_s(lvPair, 2), (\langle\langle e_1 \rangle\rangle_s, \langle\langle e_2 \rangle\rangle_s) \circ \text{update}_s)$$

Here, $lvPair$ is a variant of the identity function whose domain elements are defined to be the location-value pairs.

$$lvPair \hat{=} \{(r, v) \mapsto (r, v) \mid (r, v) \in Location \times Value\}$$

The reason for combining the location and value into a pair is so that it has the same form as the heap extended result-value and constant-map UTP models of the $\mathcal{O}$ -calculus. This helps to highlight the semantic, rather than syntactic, differences between the models.

The heap update command consumes the location-value pair on the top of the stack, assigns the new value to the existing heap location, and then pushes the heap location onto the stack.

$$\begin{aligned} \text{update}_s \hat{=} & \left(\exists r, v \mid (r, v) = \text{head } \Pi_{\text{STK}} \bullet \right. \\ & \Pi_{\text{STK}}, \Pi_{\text{HEAP}} := \langle r \rangle \frown (\text{tail } \Pi_{\text{STK}}), \Pi_{\text{HEAP}} \oplus \{r \mapsto v\} \\ & \triangleleft r \in \Pi_{\text{HEAP}} \triangleright \\ & \text{chaos} \\ & \left. \right) \\ & \triangleleft \# \Pi_{\text{STK}} > 0 \wedge (\text{head } \Pi_{\text{STK}}) \in Location \times Value \triangleright \\ & \text{chaos} \end{aligned}$$

3.3 Result Value Model

The $\mathcal{U}_R$ -design removes the need for an operand stack by storing intermediate results in local variables, which are introduced in the standard UTP way. The cost of this modelling approach is that the alphabets of subprograms are no longer uniform, and thus have to be managed to ensure that they can be composed. In particular, care has to be taken to ensure that the alphabet at the point of a recursive method invocation is the same as the alphabet of all existing method invocations. This restriction ensures that the sequential composition of the current program state with the unwinding of a method invocation is well defined.

In the $\mathcal{U}_R$ -design an object value is represented as a map from object labels to methods, where a method definition is a compiled – but unevaluated – program text. Such program texts are processed as part of the method invocation command.

In this model we introduce the Π_{OK} and Π_{RES} special modelling variables, where:

- Π_{OK} represents the starting and termination of a program;
- Π_{RES} represents the result of executing a program (or method).

The remainder of this section presents only the differences between this model and the previous $\mathcal{U}_S$ -design. Specifically, when the only action required to use a previous definition is to replace each old model identifier subscript (s) with this model's identifier subscript (R), then the definition is not reproduced.

3.3.1 Operand values

The values, labels, and methods used within a $\mathcal{U}_R$ -design are the same as those used within a $\mathcal{U}_S$ -design. One consequence of this is that both the definitions and compilation scheme for these operand-values are the same, up to renaming of model identifiers, for both models, and thus are not repeated here. However, the definition of the evaluation command ($\mathcal{E}_R$), which is used to model the simplest object calculus program (i.e. a literal value), needs to be updated to use the result-value (Π_{RES}) instead of the operand-stack (Π_{STK}). This is achieved by the following $\mathcal{U}_R$ -design evaluation command, which assigns an operand value (ov) to the result-value (Π_{RES}).

$$\mathcal{E}_S \text{ } ov \quad \widehat{=} \quad \Pi_{\text{RES}} := ov$$

Example 3.10: The $\mathcal{U}_R$ -design for the $\mathcal{O}$ -calculus program in the ‘value example’ on [page 14](#) is:

$$\begin{aligned} \llbracket 42 \rrbracket_R &= \mathcal{E}_R \llbracket 42 \rrbracket_R^E = \mathcal{E}_R 42 \\ &= \Pi_{\text{RES}} := 42 = \{ \Pi_{\text{RES}} \} : (\text{true} \vdash \Pi_{\text{RES}}' = 42) \end{aligned}$$

□

3.3.2 Command expressions

In the $\mathcal{O}$ -calculus almost all the programming operations are expressions. As previously discussed, such operations are modelled by a command-expression $\text{cmdExp}_R(P, P_{i=1}^k)$ that applies the transformation command P , which represents the operation, to the results of executing the subprograms $P_{i=1}^k$. Here, each of the subprogram results is stored in an *intermediate* variable, before being consumed by the *transformation* command that applies the operation. Specifically, the $\mathcal{U}_R$ -design’s command-expression is defined in three stages, as follows:

1. introduce a local variable for each operand, which is scoped to the next two commands;
2. calculate each argument and store its value in the appropriate local variable;
3. takes the contents of these local variables and apply the transformation command to them.

$$\begin{aligned} \text{cmdExp}_R(P, (P_{i=1}^k)) &\widehat{=} \text{decl } x_{i=1}^k \text{ in} & (1) \\ &(\circ_{j=1}^k P_j \downarrow x_{i=1}^k \circ x_j := \Pi_{\text{RES}}) \circ & (2) \\ &P & (3) \end{aligned}$$

Note that Item 2 uses one of the scoped sequential composition commands of [Section 2.3.4](#) – specifically, the form $P \downarrow x_{i=1}^k \circ Q$ that has the effect of hiding the variables $x_{i=1}^k$ from the scope of a subprogram P before sequentially composing it with subprogram Q .

Item 3’s result transformation command can be systematically constructed by supplying a function f to the $\text{trans}_R(f, (x_{i=1}^k))$ command, where $x_{i=1}^k$ are the names of the local variables storing the argument values. The result of this command is stored in the special result variable (Π_{RES}), so long as the values of the parameters are in the domain of the function f ; if this conditions fails to hold then the resulting program behaves chaotically.

$$\text{trans}_R(f, 2) \quad \widehat{=} \quad \{ \Pi_{\text{RES}} \} : ((x_{i=1}^k) \in f \vdash \Pi_{\text{RES}}' = f(x_{i=1}^k))$$

Example 3.11: The $\mathcal{O}$ -calculus addition operation can be modelled in terms of the trans_R function by sequentially processing the programs represented by expressions e_1 and e_2 , storing

their results in some temporary local variables, and then applying the binary addition function to this result.

$$\begin{aligned}
\langle\!\langle e_1 + e_2 \rangle\!\rangle_{\mathbf{R}} &\widehat{=} \text{cmdExp}_{\mathbf{R}}(\\
&\quad \text{trans}_{\mathbf{R}}((- + -), 2), \\
&\quad (\langle\!\langle e_1 \rangle\!\rangle_{\mathbf{R}}, \langle\!\langle e_2 \rangle\!\rangle_{\mathbf{R}}) \\
&\quad) \\
&= \text{decl } x_1, x_2 \text{ in} \\
&\quad (\langle\!\langle e_1 \rangle\!\rangle_{\mathbf{R}} \downarrow x_1, x_2 \circ x_1 := \Pi_{\text{RES}}) \circledast \\
&\quad (\langle\!\langle e_2 \rangle\!\rangle_{\mathbf{R}} \downarrow x_1, x_2 \circ x_2 := \Pi_{\text{RES}}) \circledast \\
&\quad \text{trans}_{\mathbf{R}}((- + -), 2) \\
&= \text{decl } x_1, x_2 \text{ in} \\
&\quad ((\text{hide } x_1, x_2 \text{ from } \langle\!\langle e_1 \rangle\!\rangle_{\mathbf{R}}) \circledast x_1 := \Pi_{\text{RES}}) \circledast \\
&\quad ((\text{hide } x_1, x_2 \text{ from } \langle\!\langle e_2 \rangle\!\rangle_{\mathbf{R}}) \circledast x_2 := \Pi_{\text{RES}}) \circledast \\
&\quad \text{trans}_{\mathbf{R}}((- + -), 2)
\end{aligned}$$

□

3.3.3 Function and method invocation

The procedure invocation $\mathcal{O}$ -calculus operations are compiled in a similar manner to that of their $\mathcal{U}_s$ -design counterparts.

Aside 3.12: Recall that as the $\mathcal{U}_r$ -design's command expression ($\text{cmdExp}_{\mathbf{R}}$) hides the variables it introduces to store intermediate results from the evaluation of its argument subprograms, these variables will not be in scope of any procedure invocation. Further, as this is the only way of introducing program variables, the alphabet of the method invocation will only contain the six model variables (i.e. Π_{OK} , Π_{RES} , Π_{HEAP} , and their dashed variants).

□

The only aspect of function and method invocation that has changed from the previous model is the definition of the $\text{call}_{\mathbf{R}}$ command, or more specifically, its fix-point function. Here, the $\text{call}_{\mathbf{R}}$ command is defined as the least fix-point of the $\text{apply}_{\mathbf{R}}$ function, which substitutes the value v in the program-text t for its parameter x .

$$\begin{aligned}
\text{call}_{\mathbf{R}} &\widehat{=} \mu z \bullet \text{apply}_{\mathbf{R}}(z) \\
\text{apply}_{\mathbf{R}}(z) &\widehat{=} (\exists v, x, t \mid (v, (\downarrow x, t)) = \Pi_{\text{RES}} \bullet \\
&\quad \text{ext}_{\mathbf{R}}(t \{x \leftarrow \text{text}_{\mathbf{R}}(v)\}_{\mathbf{R}}, z) \\
&\quad) \\
&\triangleleft \Pi_{\text{RES}} \in \text{Value} \times \text{ProgramText} \triangleright \\
&\text{chaos}
\end{aligned}$$

3.3.4 Modelling the heap operations

The $\mathcal{U}_r$ -design is a straightforward update of the $\mathcal{U}_s$ -design. The only change is to replace the definitions of the fresh_s , deref_s and update_r , with versions of the commands that use the result value instead of the operand stack.

Fresh Operator: The fresh command creates a new location on the heap, which is initialised to the explicit unset value; it sets the result value to the new heap location.

$$\begin{aligned}
\text{fresh}_{\mathbf{R}} &\widehat{=} \exists r \mid r = \text{fresh}_{\text{LOC}}(\text{dom } \Pi_{\text{HEAP}}) \bullet \\
&\quad \Pi_{\text{RES}}, \Pi_{\text{HEAP}} := r, \Pi_{\text{HEAP}} \oplus \{r \mapsto \mathbf{i}\}
\end{aligned}$$

Dereference Operator: The heap dereference command (deref_R) takes the location stored in the result variable, and replaces it with a copy of the associated heap value.

$$\text{deref}_R \hat{=} \{ \Pi_{\text{RES}} \}_R : (\begin{array}{l} \Pi_{\text{RES}} \in \Pi_{\text{HEAP}} \\ \vdash \\ \Pi_{\text{RES}}' = \Pi_{\text{HEAP}} \Pi_{\text{RES}} \end{array})$$

Heap Update Operator: The heap update command assumes that the input result value (Π_{RES}) contains a location-value pair, where the location already exists on the heap. If this is the case, then the value of that location is updated appropriately.

$$\text{update}_R \hat{=} (\begin{array}{l} \exists r, v \mid (r, v) = \Pi_{\text{RES}} \bullet \\ \Pi_{\text{RES}}, \Pi_{\text{HEAP}} := r, \Pi_{\text{HEAP}} \oplus \{ r \mapsto v \} \\ \triangleleft r \in \Pi_{\text{HEAP}} \triangleright \\ \text{chaos} \end{array}) \\ \triangleleft \Pi_{\text{RES}} \in \text{Location} \times \text{Value} \triangleright \\ \text{chaos}$$

Garbage collection revisited: When discussing the modelling of the heap in [Section 3.1.3](#), we stated that the hiding of local variables, via existential quantification, could lead to situations where it is not possible to determine if a heap location is *garbage* by just examining the state variables within a program. In such situations, knowing when a heap location is genuinely unreachable is very difficult. This is illustrated by the following example.

Example 3.13: Consider the following $\mathcal{O}$ -calculus program, which creates a fresh location on the heap, assigns the value 42 to it, and then dereferences that location.

$*(\text{fresh} \text{ } *= 42)$

It can be compiled to the following UTP program, using the $\mathcal{U}_R$ -design compilation scheme.

$(\text{cmdExp}_R(\begin{array}{l} \text{trans}_R(lvPair, 2), \\ (\text{fresh}_R, \mathcal{E}_R 42) \end{array}) \circ \text{update}_R \circ \text{deref}_R$

Having done this, we expand the definition of the cmdExp_R command to generate the following program, which illustrates the hidden variable issue.

$(\text{decl } x_1, x_2 \text{ in } \begin{array}{l} (\text{hide } x_1, x_2 \text{ from } \text{fresh}_R) \circ x_1 := \Pi_{\text{RES}} \circ \\ (\text{hide } x_1, x_2 \text{ from } \mathcal{E}_R 42) \circ x_2 := \Pi_{\text{RES}} \circ \\ \text{trans}_R(lvPair, 2) \end{array}) \circ \text{update}_R \circ \text{deref}_R$

We focus on the evaluation of the subprogram $\mathcal{E}_R 42$, in the context where the subprogram fresh_R has evaluated to the abstract location ℓ_{17} . First, we observe that the alphabet of this subprogram does not include either of the hidden local variables x_1 or x_2 . Second, we observe that just prior to the evaluation of our nominated subprogram, only the Π_{RES} result variable and the hidden local variable x_1 contain references to location ℓ_{17} . Third, we observe that

the evaluation overwrites the Π_{RES} result variable with the value 42, thus leaving the hidden variable x_1 with the only reference to location ℓ_{17} . Therefore, from our nominated subprograms point of view the location ℓ_{17} is not referenced by any variable within its alphabet, and thus is garbage and can be removed.

□

3.3.5 Other mechanisms

Control flow mechanisms: The conditional evaluation command is compiled in the same manner as before. It executes subprogram P_1 , P_2 , or **chaos** depending on whether the result variable (Π_{RES}) contains **true**, **false**, or a non-boolean value respectively.

$$\text{switch}_{\text{R}}(P_1, P_2) \hat{=} ((P_1 \triangleleft \Pi_{\text{RES}} \triangleright P_2) \triangleleft \Pi_{\text{RES}} \in \text{Boolean} \triangleright \text{chaos})$$

Sequential Composition: The sequential composition of two expressions can be straightforwardly represented by the sequential composition of their compiled programs. It differs from the $\mathcal{U}_{\text{S}}$ -design in that it implicitly ignores the result of the first program, rather than explicitly popping the unwanted result from the operand stack.

$$\langle\langle e_1; e_2 \rangle\rangle_{\text{S}} \hat{=} \langle\langle e_1 \rangle\rangle_{\text{S}} ; \langle\langle e_2 \rangle\rangle_{\text{S}}$$

3.4 Constant Map Model

The $\mathcal{U}_{\text{M}}$ -design is a variant of the $\mathcal{U}_{\text{R}}$ -design, which represents procedures as triples rather than pairs: in addition to the formal parameter and procedure body text there is a constant variable map. This map is used to represent the environmental bindings of the procedure body's free variables (which may include the formal parameter). The idea is that:

1. each of the variables within such a map is defined before the completion of the procedure invocation command, which itself establishes the binding of its formal parameter to both the procedure being invoked and any nested procedures;
2. the procedure invocation command introduces the variable-value pairs in the constant map as read-only variables for the scope of this procedure's body, but excluding both procedure invocations and inner-procedure definitions.

The two exclusions in the scope of a procedure's context variables follow from the constraints of modelling recursive procedure invocations, specifically, that the alphabet before (and after) each procedure invocation call must be the same. This can be achieved by ensuring that the current procedure's context variables are hidden before each procedure invocation. Following the completion of the procedure invocation the hidden context variables are reinstated. The inner procedure exclusion follows from the observation that the invocation of any procedure removes the context of the calling (enclosing) procedure context.

3.4.1 Procedure definition

A procedure is defined as a triple of compiled program texts that represent the procedure's *parameter*, *constant-map* and *body*, where the constant map contains an entry for each free-variable within the procedure's body. These variables are set – once – prior to or during the

loading of the procedure's body. A procedure definition is denoted by $\langle x, M, pt \rangle$, where the scope of the parameter x and the constant-map M is the program text pt .

In contrast to the previous UTP models of the $\mathcal{O}$ -calculus the tuple denoting the procedure is not considered to be its meaning. Instead, we define its meaning by a pair $\langle x, body \rangle$, where the second and third arguments of the tuple are combined to form its body – specifically the variable bindings in the constant map are used to provide a local declaration block for the program text.

$$\langle x, M, pt \rangle \hat{=} \langle x, \text{decl } M \text{ in } pt \rangle$$

There are two reasons for denoting a procedure by a triple and defining it as a pair. First, the triple represents the conceptual aspects of the procedure definition, which are later used to structure the layout of a graph based model of state (Section 6.2.2). Second, the pair provides a representation of a method that simplifies the consistency proofs, which follow in the next chapter.

A method only has an evaluation context compilation mode, as it is not considered to be a value. It is defined as follows:

$$\langle \langle \varsigma(x) e \rangle_M^E \rangle \hat{=} \left(\begin{array}{l} x, \\ \{y \mapsto \text{free} \mid y \in Y\}, \\ \langle e \rangle_M \end{array} \right) \begin{array}{l} \text{Self variable} \\ \text{Constant map} \\ \text{Body} \end{array}$$

Here, $Y = \text{FV}(e) \cup \{x\}$ and **free** is a special marker value, for unbound variables. Note, the reason we use **free** to denote an unbound variable rather than the existing unset value ($\mathfrak{!}$) is because $\mathfrak{!}$ is a legitimate value for a variable to contain.

A function is compiled in a similar fashion to that of the method; the main difference is that functions are considered to be values, thus have a normal compilation mode (as well).

$$\begin{aligned} \langle \langle \lambda(x) e \rangle_M \rangle &\hat{=} \mathcal{E}_M \langle \langle \lambda(x) e \rangle_M^E \rangle \\ \langle \langle \lambda(x) e \rangle_M^E \rangle &\hat{=} \langle \langle \varsigma(x) e \rangle_M^E \rangle \end{aligned}$$

The method update and field assignment are essentially the same as those defined for the $\mathcal{U}_R$ -design, which in turn are inherited from the $\mathcal{U}_S$ -design (see Section 3.2.5 for details).

3.4.2 Method and function invocation

The method invocation $\mathcal{O}$ -calculus operation is compiled in a similar manner to that of the $\mathcal{U}_R$ -design, which is inherited from the $\mathcal{U}_S$ -design. The key difference is in the use of the ternary variable filtering command, which ensures that only the special program variables Π_{OK} and Π_{RES} are in the alphabet of the call_M command.

$$\langle \langle e.l \rangle_M \rangle \hat{=} \text{cmdExp}_M(\begin{array}{l} \text{trans}_M(\text{omPair}, 2), \\ (\langle e \rangle_M, \mathcal{E}_M \langle \langle l \rangle_M^E \rangle) \end{array}) \circ \Pi_{\text{OK}}, \Pi_{\text{RES}} \upharpoonright \text{call}_M$$

Here, the *omPair* function is as defined in the operand stack model (Section 3.2.6).

Function invocation $\mathcal{O}$ -calculus command is compiled in a similar manner to that of the method invocation.

$$\langle \langle e_1 e_2 \rangle_M \rangle \hat{=} \text{cmdExp}_M(\begin{array}{l} \text{trans}_M(\text{vtPair}, 2), \\ (\langle e_1 \rangle_M, \langle e_2 \rangle_M) \end{array}) \circ \Pi_{\text{OK}}, \Pi_{\text{RES}} \upharpoonright \text{call}_M$$

Here, the *vtPair* function is as defined in the operand stack model (Section 3.2.6).

Before formally defining the generic call_M command it is worth presenting an updated helper function for parameter substitution, which handles the new definition of a procedure. The notion of variable substitution is similar to that of the previous models, in the sense that it replaces a *free* variable for its definition. The difference is that a variable is only considered to be *free* if it is contained in the domain of a method's constant map and has the unbound value (*free*). This leads to the following algorithm for substituting the invoking object for the self variable (with name x):

1. If the method's constant map contains an unset version of the named constant x , then set it to this constant.
2. For each *direct* inner method, check whether the self variable shares the name x . If it does then this inner method's processing is complete; otherwise go to step 1.

Unfortunately this algorithm is awkward to formally define using our structural induction technique, due to the special (contextual) handling of the first method invocation. Therefore, we take the inner-method's point of view (step 2 in the algorithm), and defer the special case processing of the invoking method to the call_M command.

$$\begin{aligned} (x \mapsto \text{free}) \llbracket x \leftarrow v \rrbracket_M &\hat{=} (x \mapsto v) \\ (\llbracket x, M, t \rrbracket) \llbracket x \leftarrow v \rrbracket_M &\hat{=} (\llbracket x, M, t \rrbracket) \\ t\{t_{i=1}^k t_i\} \llbracket x \leftarrow v \rrbracket_M &\hat{=} t\{t_{i=1}^k t_i \llbracket x \leftarrow v \rrbracket_M \} \end{aligned}$$

We are now in a position to define the call_M command. It is defined as the least fix-point of the apply_M function, which substitutes the argument value v in the program-text t for its formal parameter x . However, if either the argument value or the program-text does not exist, then the resultant program behaves chaotically.

$$\begin{aligned} \text{call}_M &\hat{=} \mu z \bullet \text{apply}_M(z) \\ \text{apply}_M(z) &\hat{=} (\exists v, x, t \mid (v, (x, t)) = \Pi_{\text{RES}} \bullet \\ &\quad \text{ext}_M(t \llbracket x \leftarrow \text{text}_M(v) \rrbracket, z) \\ &\quad) \\ &\quad \triangleleft \Pi_{\text{RES}} \in \text{Value} \times \text{ProgramText} \triangleright \\ &\quad \text{chaos} \end{aligned}$$

3.4.3 Alternative method and function definition compilation scheme

Section 6.2 presents the trace-based model of the $\mathcal{O}$ -calculus, which has a requirement to uniquely name both the inner procedures and the objects that a procedure directly contains. It is straightforward to update the $\mathcal{U}_M$ -design being presented here to meet this requirement, once we observe that such inner procedures and object values could be bound to free variables outside the scope of this method's definition. However, we do not want to insist that methods are written in this fashion; instead we simulate the external declaration of these values as part of the compilation process.

The first stage of this simulation is to identify those elements (subterms) of a procedure that need to be extracted; namely, the directly defined inner procedure definitions and object values. The following objsAndProcs function extracts the set of terms that representing these inner procedures and objects, where o , m , f , and t , have their usual meanings; namely an object value term, a method definition term, a function definition term, and a general term respectively.

$$\begin{aligned} \text{objsAndProcs}(o) &\hat{=} \{o\} \\ \text{objsAndProcs}(m) &\hat{=} \{m\} \\ \text{objsAndProcs}(f) &\hat{=} \{f\} \\ \text{objsAndProcs}(t\{t_{i=1}^k t_i\}) &\hat{=} \bigcup_{i=1}^k \text{objsAndProcs}(t_i) \end{aligned}$$

The second stage is to generate a *fresh* name to identify each of these extracted terms.

$$\begin{aligned} \text{freshMap}(Y, \emptyset) &\stackrel{\cong}{=} \emptyset \\ \text{freshMap}(Y, \{t_{i=1}^k\}) &\stackrel{\cong}{=} \{y \mapsto t_1\} \cup \text{freshMap}(\{y\} \cup Y, \{t_{i=2}^k\}) \end{aligned}$$

Here, $y = \text{fresh}_{\text{VAR}}(Y)$ generates a variable that is not in Y .

The third (and last) stage is to use these functions in the compilation of a procedure. Here, the identified inner procedures and objects are replaced by their *freshly* generated names within the *body* of the procedure, and the constant map is updated to include the entries for these inner procedures and objects. The compilation of a $\mathcal{O}$ -calculus method now follows.

$$\begin{array}{l} \langle\langle \varsigma(x) \, e \rangle\rangle_M^E \stackrel{\cong}{=} \\ \left(\begin{array}{l} x, \\ M_1 \cup M_2, \\ \langle\langle e \, o \leftarrow z \mid z \mapsto o \in M_2 \rangle\rangle_M \end{array} \right) \end{array} \quad \left| \begin{array}{l} \text{Self variable} \\ \text{Constant map} \\ \text{Body} \end{array} \right.$$

where

$$\begin{array}{l} Y = \text{FV}(e) \cup \{x\} \\ M = \text{freshMap}(Y, \text{objsAndProcs}(e)) \\ M_1 = \{y \mapsto \text{free} \mid y \in Y\} \\ M_2 = \{y \mapsto \langle\langle e \rangle\rangle_M \mid y \mapsto e \in M\} \end{array} \quad \left| \begin{array}{l} \text{The method body's free variables} \\ \text{The named inner procedure and object map} \\ \text{The map of, un-set (free), external variables} \\ \text{The map of simulated external variables} \end{array} \right.$$

Note that as the compilation of an $\mathcal{O}$ -calculus function is defined in terms of the compilation of an $\mathcal{O}$ -calculus method, its definition does not need to be updated.

3.5 Summary

We have presented three UTP models of the $\mathcal{O}$ -calculus, along with their corresponding systematic compilation schemes. These models differ in the way that they handle both the final and intermediate values within a calculation (i.e. the execution of a program), as follows:

1. the $\mathcal{U}_s$ -design uses an operand stack to store the results of both the final and intermediate calculations;
2. the $\mathcal{U}_r$ -design replaces the operand stack with a combination of a result variable and some temporary local variable declaration blocks, which are used to store the final and intermediate calculations respectively; and
3. the $\mathcal{U}_m$ -design extends the use of local variable declaration blocks, so that they also model procedure parameters.

All three of these models share the same straightforward model of a heap, which is a partial map from abstract locations to values. Such a model of a heap is not fully-abstract as discussed in [Section 1.5.2](#). Here, we have focused on the core modelling of object-based object orientation, which in turn is focused on the modelling of potentially mutually recursive method (and function) invocation.

Potentially mutually recursive procedure calls are defined in terms of the UTP weakest fix-point operator, where each procedure in the recursive loop is essentially sequentially composed with the subprogram that invoked it and the subprogram that follows it. One consequence of this is that the procedure invocation mechanism must guarantee that the alphabets of the procedure and the calling contexts will always agree. Our solution to this issue, was to ensure that the alphabet at the point of procedure invocation was guaranteed to contain just the model's observational variables (i.e. the non program variables).

It is now possible to argue that the $\mathcal{U}_s$ -design is the simplest of the UTP models, as it does not use the UTP's variable management facilities (e.g. introduction and elimination), and thus the alphabet of all subprograms within this UTP model are guaranteed to be the same. The $\mathcal{U}_r$ -design is only slightly more complicated as the local variables it uses to store intermediate results of calculations (i.e. subprograms) are guaranteed to be outside the scope of the subprograms themselves, and thus the alphabet at point of method invocation will always be the same. Lastly, the $\mathcal{U}_m$ -design guarantees that the alphabet of a procedure invocation is always the same, by temporarily hiding all local program variables from scope just prior to invoking a procedure.

These models form the basis for two of our three contributions (as presented in [Section 1.5.3](#) and discussed in [Chapter 7](#)). Specifically, we have encoded – via a compilation scheme – an Abadi–Cardelli-style object calculus in the UTP. What we have not shown (so far) is that this compilation scheme produces UTP programs that have the same semantics as that of their source (i.e. an $\mathcal{O}$ -calculus program). This is the topic of [Chapter 4](#).

Chapter 4

Object Model Consistency

The three Unifying Theories of Programming (UTP) object models in [Chapter 3](#) each provide a semantics for our $\mathcal{O}$ -calculus, which was presented in [Section 2.2](#). We demonstrate that the denotational semantics of each UTP model is consistent with the operational semantics of the $\mathcal{O}$ -calculus, via a structural induction over the $\mathcal{O}$ -calculus's terms; i.e. the UTP denotational semantics of an $\mathcal{O}$ -calculus operation is the same as that of its result.

One complication is that the $\mathcal{O}$ -calculus's operational semantics introduces some dynamic terms, such as the *fresh location values*. These values index (point at) entries on the abstract heap, which is considered to be empty at the beginning of a program's execution. Another dynamic term is that of the substitution meta-operation, which replaces an $\mathcal{O}$ -calculus parameter with its value. The fresh location terms have already been included in our UTP models of the $\mathcal{O}$ -calculus, via the modelling of the three heap operations (**fresh**, *****, and ***=**). However, the substitution operation has not been directly modelled, as it was not required for defining the UTP semantics of the $\mathcal{O}$ -calculus. Having said this, the substitution operation does appear in the resultant part of several $\mathcal{O}$ -calculus small-step reduction rules, and thus could be considered as part of the structure that the consistency proof is inducting over. This is the view that we take.

This chapter presents the dynamic operations for the three UTP models ([Section 4.1](#)), an outline of the structural induction argument ([Section 4.2](#)), the main proofs for each of the models ([Section 4.3](#) to [Section 4.5](#)), and an example ([Section 4.6](#)). Here, the example provides an illustration of the correspondence between the operational $\mathcal{O}$ -calculus and denotational UTP semantics via one iteration of a factorial program. Additional laws and proofs are contained in [Appendix C](#) and are referenced where appropriate.

4.1 Dynamic Operations

The previous chapter defined three UTP models of the $\mathcal{O}$ -calculus ([Section 2.2](#)), where each model had at least one compilation scheme for taking an $\mathcal{O}$ -calculus program to its corresponding UTP model. However, these compilation schemes did not include translations for the *run-time* information that the $\mathcal{O}$ -calculus's operational rules introduce. In particular, we do not provide compilation schemes for

- the variable substitution meta-operation ($t\{x \leftarrow ov\}$); and
- the execution context (Γ) (i.e. the heap).

We now update our compilation schemes to handle this dynamic run-time information, so that it is straightforward to define the UTP semantics of an $\mathcal{O}$ -calculus term at any point (step) during its operational evaluation (execution).

4.1.1 Variable substitution compilation scheme

The meta-function for substituting variables in an $\mathcal{O}$ -calculus program is recursively defined over the syntactic structure of an $\mathcal{O}$ -calculus term (program). In particular, it cannot be applied to the meanings of its arguments, just their syntactic form. For the UTP operand–stack–design ($\mathcal{U}_s$ -design) this amounts to compiling the term representing the object calculus program, and then applying the $\mathcal{U}_s$ -design’s syntactic variable substitution process.

$$\langle\!\langle t \{x \leftarrow ov\} \rangle\!\rangle_s \hat{=} \langle\!\langle t \rangle\!\rangle_s \{x \leftarrow \langle\!\langle ov \rangle\!\rangle_s^E\}_s$$

The same compilation scheme is used for the UTP result–value–design ($\mathcal{U}_r$ -design), but not for the UTP constant–map–design ($\mathcal{U}_m$ -design). The $\mathcal{U}_m$ -design explicitly models parameters via the introduction of UTP variables. Therefore, the variable substitution compilation scheme is adapted to use this approach. Specifically, it introduces a variable to hold the *constant* value of the variable substitution, for all terms within the current program-text, except for (inner) procedures. Such inner procedures contain a *constant map*, which records the values of its *inherited* parameters as they are set, via the use of the $\mathcal{U}_m$ -design’s variable substitution function.

$$\langle\!\langle t \{x \leftarrow ov\} \rangle\!\rangle_M \hat{=} \text{decl } \{x \mapsto \langle\!\langle ov \rangle\!\rangle_M^E\} \text{ in } \langle\!\langle t \rangle\!\rangle_M \{x \leftarrow \langle\!\langle ov \rangle\!\rangle_M^E\}_M$$

4.1.2 The heap context compilation scheme

The heap in the environmental context is a, potentially empty, finite map from $\mathcal{O}$ -calculus dynamic locations to values. Each of the three UTP models in Chapter 3 mirrors the $\mathcal{O}$ -calculus model, in the sense that it introduces a Π_{HEAP} observation variable to store the map from dynamic locations to values; in these cases, the locations and values are of the relevant UTP model. Therefore, the compilation scheme for the heap context straightforwardly transforms each of the entries in the $\mathcal{O}$ -calculus map to the corresponding entries in the UTP model’s map, as follows:

$$\langle\!\langle \{_{i=1}^k r_i \mapsto v_i\} \rangle\!\rangle_J \hat{=} \{_{i=1}^k \langle\!\langle r_i \rangle\!\rangle_J \mapsto \langle\!\langle v_i \rangle\!\rangle_J\}$$

Here, J stands for one of the three UTP design model identifiers S , R , and M .

The dynamic aspects of the $\mathcal{O}$ -calculus are introduced by the application of the small-step reduction rules to an $\mathcal{O}$ -calculus expression that represents a program. These rules are applied to (and return) dynamic terms of the form $C \bullet e$, where C denotes a context and e an expression. The UTP semantics for such dynamic terms is the sequential composition of the UTP designs that denote the context C and the expression e respectively, as follows:

$$\langle\!\langle C \bullet e \rangle\!\rangle_J \hat{=} \langle\!\langle C \rangle\!\rangle_J \circ \langle\!\langle e \rangle\!\rangle_J$$

Here, the context C is in one of the following three styles:

$$\begin{array}{lll} \langle\!\langle \Gamma \{\text{heap} \mapsto H\} \rangle\!\rangle_J & \hat{=} & \Pi_{\text{HEAP}} :=_J \langle\!\langle H \rangle\!\rangle_J & \text{Read context heap value} \\ \langle\!\langle \Gamma \{\text{heap} \leftarrow H\} \rangle\!\rangle_J & \hat{=} & \Pi_{\text{HEAP}} :=_J \langle\!\langle H \rangle\!\rangle_J & \text{Write context heap value} \\ \langle\!\langle \Gamma \rangle\!\rangle_J & \hat{=} & \text{skip}^D & \text{Context not accessed} \end{array}$$

Recall that the last form of the context occurs in $\mathcal{O}$ -calculus small-step reduction rules that do not use (or alter) the contextual information, and thus can be modelled by the unit of composition.

4.2 Structural Induction

Each valid $\mathcal{O}$ -calculus expression (program) is either a value or reducible to a value using the deterministic small-step reduction rules. From our structural induction perspective, the interesting cases are those terms (structures) that correspond to the left hand side (LHS) of a reduction rule, as the UTP semantics of the before and after states of these rules can be compared and shown to be consistent (i.e. equal). Those rules that enable a subterm (i.e. an argument) to be reduced via the application of another rule to the subterm are already catered for by the induction hypothesis; i.e. that a sub-term has a consistent UTP semantics. Therefore, we focus our attention on the remaining rules, which typically apply to fully evaluated arguments; one notable exception is that of the IF_T and IF_F conditional evaluation rules, both of which only require one of their arguments to be evaluated (Section 2.2.4).

$\mathcal{O}$ -calculus expressions can be invalid for two main reasons. First, a non-value expression may not match any of the small-step reduction rules, in which case the program's execution becomes *stuck*. For example, there is no rule that enables the logical negation of an integer (e.g. $\neg 3$ cannot be reduced). Second, the expression may be *infinitely reducible*, as the program that it represents does not terminate (e.g. $[l = \varsigma(x) x.l].l$). Our UTP model does not distinguish between these two forms of invalidity; they are both modelled by *chaos* (Section 3.1.4). The stuck behaviours are modelled by precondition checks within each of the UTP models; if such preconditions fail the resulting UTP designs are equivalent to *chaos*, as required. The infinitely looping terms result from recursive (or mutually recursive) method and function calls that never reduce to a value. The UTP weakest fix-point of such terms is equivalent to *chaos* due to Tarski's theorem as discussed in [HH98, WC04], as required. We still have to show that the reduction of an individual method and function call is consistent, so as to preserve the non-terminating sequence of calls. This consistency check is already being performed as part of the valid term consistency check, so no further work is required.

Another constraint that we add to the structural induction is that it can only be used to show the consistency of closed $\mathcal{O}$ -calculus programs; i.e. $\mathcal{O}$ -calculus programs that have no free variables. This syntactically checkable constraint is required for the assumption that all variables within the program are bound to a value before they are used. The other requirement for this assumption to hold is that all the variables are initialised before they are used. This is guaranteed by the syntactic structure of the program and small-step reduction rules, which ensure: that variables are initialised as part of function or method invocation; and that let expressions are reduced to function invocations.

The remainder of this section introduces the use of commuting diagrams to present one step of the consistency structural induction argument, and then applies such diagrams to the reducible terms. Here, we merely state the lemmas; the proofs for the three UTP models follow in Sections 4.3 to 4.5.

4.2.1 Commuting Diagrams

The commuting diagram in Figure 4.1 illustrates the structure of the proof that the semantic models for an object calculus operation op with k subterms are consistent, where:

- st_i is the i^{th} subterm of the original operation.
- ov_i is the i^{th} subterm of the operation after its arguments (i.e. operand values) have been evaluated in the correct order.
- $\llbracket t \rrbracket_J^M$ is the combination of the semantic meaning and compilation functions (i.e. $\llbracket \llbracket t \rrbracket_J^M \rrbracket$), where M is the compilation mode and J is the model identifier.
- Γ_i is the i^{th} object calculus run-time context variable.

- A0 is the assumption that the subterms can be evaluated in the correct order. Note this assumption also guarantees the consistency of the subterm mappings (i.e. $\forall_{i=1}^k \llbracket ov_i \rrbracket_J = \mathcal{E}_J ov_i$).
- A1 is the assumption that the arguments are in the domain of the operation being modelled (i.e. $(ov_{i=1}^k) \in op$).
- A2 is the assumption that the result of executing the operation with arguments $ov_{i=1}^k$ is the expression e .
- $\rightarrow, \twoheadrightarrow$ are the one-step and multi-step $\mathcal{O}$ -calculus operational semantics.
- $\downarrow$ is a compilation and/or semantic evaluation function.
- $=, \parallel$ are two different representations of the equality relation, for horizontal and vertical display contexts respectively.

$$\begin{array}{ccccc}
\Gamma_0 \bullet op\{st_{i=1}^k\} & \xrightarrow{A0} & \Gamma_1 \bullet op\{ov_{i=1}^k\} & \xrightarrow{A1, A2} & \Gamma_2 \bullet e \\
\downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma_0 \rrbracket_J \mathbin{\dot{;}} \llbracket op\{st_{i=1}^k\} \rrbracket_J & \xrightarrow{A0} & \llbracket \Gamma_1 \rrbracket_J \mathbin{\dot{;}} \llbracket op\{ov_{i=1}^k\} \rrbracket_J & \xrightarrow{A1, A2} & \llbracket \Gamma_2 \rrbracket_J \mathbin{\dot{;}} \llbracket e \rrbracket_J
\end{array}$$

Figure 4.1: Commuting diagram principle

The left hand square of the commuting diagram in Figure 4.1 is essentially the same for every operator being checked, as it mirrors the use of the induction hypothesis, that an operation's arguments (subterms) can be evaluated successfully. Therefore, in practice this aspect of the diagram is omitted, as illustrated in Figure 4.2.

$$\begin{array}{ccc}
\Gamma_1 \bullet op\{ov_{i=1}^k\} & \xrightarrow{A1, A2} & \Gamma_2 \bullet e \\
\downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma_1 \rrbracket_J \mathbin{\dot{;}} \llbracket op\{ov_{i=1}^k\} \rrbracket_J & \xrightarrow{A1, A2} & \llbracket \Gamma_2 \rrbracket_J \mathbin{\dot{;}} \llbracket e \rrbracket_J
\end{array}$$

Figure 4.2: Commuting diagram practice

4.2.2 Command expressions

The $\mathcal{O}$ -calculus provides a variety of operations that are non-recursively defined, such as the simple arithmetic and logical operations. As these operations are systematically compiled, it is possible to present a generic proof that these operations are consistently modelled. The commuting diagram in Figure 4.3 outlines the structure of the proof that a function f with k scalar parameters has a consistent denotational semantics. Here we assume that:

- A3 The scalar values $ov_{i=1}^k$ are in the domain of the k -ary function; i.e. $(ov_{i=1}^k) \in f$.
- A4 The scalar value ov is the result of evaluating the function f ; i.e. $ov = f(ov_{i=1}^k)$.

Lemma 4.1 is the key step in this consistency proof, which demonstrates that a command expression has the expected semantics. This is denoted by **L4.1** in Figure 4.3's commuting diagram, which in turn forms a template for all the $\mathcal{O}$ -calculus operations that are defined as command expressions within the three UTP models. In particular, field assignment and method update are also covered by this proof template; the commuting diagram that specifically corresponds to the method update operation is supplied at the end of this section (in Figure 4.4).

$$\begin{array}{ccc}
\Gamma \bullet f(ov_{i=1}^k) & \xrightarrow{A3, A4} & \Gamma \bullet ov \\
\downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma \rrbracket_J \mathbin{;} \text{cmdExp}_J(\text{trans}_J(f, k), (\mathop{\mathcal{E}}_J ov_i)_{i=1}^k)) & & \llbracket \Gamma \rrbracket_J \mathbin{;} \mathcal{E}_J ov \\
\downarrow A3, L4.1 & & \downarrow \\
\llbracket \Gamma \rrbracket_J \mathbin{;} \mathcal{E}_J f(ov_{i=1}^k) & \xlongequal{A4} & \llbracket \Gamma \rrbracket_J \mathbin{;} \mathcal{E}_J ov
\end{array}$$

Figure 4.3: Generic command expression commuting diagram

Lemma 4.1 – Command Expression Lemma: The effect of applying a cmdExp_J command to a function f , with *pre-evaluated* arguments $ov_{i=1}^k$, is the same as the effect of applying the evaluation command to the result of the function f on its arguments, assuming that the arguments are in the domain of the function (i.e. $(ov_{i=1}^k) \in f$).

$$(ov_{i=1}^k) \in f \quad \Rightarrow \quad \text{cmdExp}_J(\text{trans}_J(f, k), (\mathop{\mathcal{E}}_J ov_i)_{i=1}^k) = \mathcal{E}_J(f(ov_{i=1}^k))$$

Recall that a *pre-evaluated* argument is an operand-value (i.e. a label, method definition or a value).

□

Member update: Figure 4.4 specialises the command expression commuting diagram template of Figure 4.3 to show the consistency argument for the method update operation.

$$\begin{array}{ccc}
\Gamma \bullet o.l \Leftarrow m & \xrightarrow{A3, A4} & \Gamma \bullet o_0 \\
\downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma \rrbracket_J \mathbin{;} \text{cmdExp}_J(\text{trans}_J(\text{methUpd}, 3), (\llbracket o \rrbracket_J, \mathcal{E}_J l, \mathcal{E}_J \llbracket m \rrbracket_J^E)) & & \llbracket \Gamma \rrbracket_J \mathbin{;} \mathcal{E}_J o_0 \\
\parallel A3, \text{C.E. lemma} & & \downarrow \\
\llbracket \Gamma \rrbracket_J \mathbin{;} \mathcal{E}_J \text{methUpd}(\llbracket o \rrbracket_J^E, l, \llbracket m \rrbracket_J^E) & \xlongequal{A4} & \llbracket \Gamma \rrbracket_J \mathbin{;} \mathcal{E}_J \llbracket o_0 \rrbracket_J^E
\end{array}$$

Figure 4.4: Method update commuting diagram

4.2.3 Procedure invocation

There are two similar forms of procedure invocation: method invocation and function invocation. The commuting diagram in Figure 4.5 demonstrates that the $\mathcal{O}$ -calculus method invocation is consistent with one unwinding of the fix-point function (call_J) that defines it, where assumption A5 is that the label l of object o has the method $\varsigma(x) e$. Similarly, the commuting diagram in Figure 4.6 demonstrates that the $\mathcal{O}$ -calculus function invocation is consistent with one unwinding of the fix-point function (call_J) that defines it, where assumption A6 is that the function f is denoted by $\lambda(x) e$.

The consistency diagrams in Figure 4.5 and Figure 4.6 rely on three lemmas: the command expression lemma (Lemma 4.1); the unwinding lemma (Lemma 4.2), which uses the fix-point

$$\begin{array}{ccc}
\Gamma \bullet o.l & \xrightarrow{A5} & \Gamma \bullet e\{x \leftarrow o\} \\
\downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma \rrbracket_J \mathbin{\text{\textcircled{;}}} \text{cmdExp}_J(\text{trans}_J(\text{omPair}, 2), (\mathcal{E}_J \llbracket o \rrbracket_J, \mathcal{E}_J l)) \mathbin{\text{\textcircled{;}}} \text{call}_J & & \\
\parallel \text{L4.1, L4.2, A5} & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma \rrbracket_J \mathbin{\text{\textcircled{;}}} \mathcal{E}_J(\llbracket o \rrbracket_J^E, \llbracket \varsigma(x) e \rrbracket_J^E) \mathbin{\text{\textcircled{;}}} \text{apply}_J(\text{call}_J) & \xrightarrow{\text{L4.3}} & \llbracket \Gamma \rrbracket_J^E \mathbin{\text{\textcircled{;}}} \llbracket e\{x \leftarrow o\} \rrbracket_M
\end{array}$$

Figure 4.5: Method invocation commuting diagram

$$\begin{array}{ccc}
\Gamma \bullet f v & \xrightarrow{A6} & \Gamma \bullet e\{x \leftarrow v\} \\
\downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma \rrbracket_J \mathbin{\text{\textcircled{;}}} \text{cmdExp}_J(\text{trans}_J(\text{vtPair}, 2), (\mathcal{E}_J \llbracket f \rrbracket_J, \mathcal{E}_J \llbracket v \rrbracket_J)) \mathbin{\text{\textcircled{;}}} \text{call}_J & & \\
\parallel \text{L4.1, L4.2, A6} & & \downarrow \llbracket - \rrbracket_J \\
\llbracket \Gamma \rrbracket_J \mathbin{\text{\textcircled{;}}} \mathcal{E}_J(\llbracket v \rrbracket_J^E, \llbracket \lambda(x) e \rrbracket_J^E) \mathbin{\text{\textcircled{;}}} \text{apply}_J(\text{call}_J) & \xrightarrow{\text{L4.3}} & \llbracket \Gamma \rrbracket_J^E \mathbin{\text{\textcircled{;}}} \llbracket e\{x \leftarrow v\} \rrbracket_M
\end{array}$$

Figure 4.6: Function invocation commuting diagram

definition to provide a single unwinding; and the procedure call lemma (Lemma 4.3), which demonstrates that this unwinding is correct.

Recall that the procedure call operation is defined as the fix-point of an apply function ($\mu z \bullet \text{apply}_J(z)$). The structure of this definition leads to the following unwinding lemma.

Lemma 4.2 – Unwinding Lemma:

$$\text{call}_J = \text{apply}_J(\text{call}_J)$$

□

Proof

$$\begin{aligned}
& \text{call}_J \\
= & \dots \dots \dots \text{Defn. of call}_J \\
& (\mu z \bullet \text{apply}_J(z)) \\
= & \dots \dots \dots \text{Defn of fix-point } \mu \\
& \text{apply}_J(\mu z \bullet \text{apply}_J(z)) \\
= & \dots \dots \dots \text{Defn. of call}_J \\
& \text{apply}_J(\text{call}_J)
\end{aligned}$$

□

The following procedure call lemma does not directly meet the requirements of the commuting diagram, in that it does not start with the pre-compiled form of a function or method. Instead, it starts with the compiled form of the function or method, which depends on the UTP model being used. A similar comment can be applied to the dynamic substitution operation. Both of these definitional transformations are straightforward, and thus the procedure call lemma indirectly meets the requirements of the commuting diagram.

Lemma 4.3 – Procedure Call Lemma: The effect of applying the call_J command is equivalent to the effect of applying one iteration of this command to itself.

$$\begin{aligned} \mathcal{E}_J(\llbracket v \rrbracket_J^E, (\downarrow x, \llbracket e \rrbracket_J^E) \circ \text{apply}_J(\text{call}_J)) &= \llbracket \llbracket e \rrbracket_J \{x \leftarrow \llbracket v \rrbracket_J^E\} \rrbracket_J \\ \mathcal{E}_M(\llbracket v \rrbracket_M^E, (\downarrow x, M, pt) \circ \text{apply}_M(\text{call}_M)) &= \llbracket \text{decl}\{x \mapsto v\} \text{ in} \\ &\quad (\downarrow pt) \{x \leftarrow \llbracket v \rrbracket_M^E\} \rrbracket_M \end{aligned}$$

where $J \in \{S, R\}$ and

$$\begin{aligned} M &= \{x \mapsto \text{free}, (y_i \mapsto ov_i)_{i=1}^k\} \\ pt &= \llbracket e \rrbracket_M^E \{y_i \mapsto ov_i\}_M^k \end{aligned}$$

Note that operand values ov_1^k represent the values of *all* the inherited parameters y_1^k respectively; this follows from limiting our consistency analysis to closed $\mathcal{O}$ -calculus programs (as if one of the parent parameters was not set, then it would be free in the $\mathcal{O}$ -calculus program, and thus contradict our limitation). $\square$

4.2.4 Heap operations

Fresh heap locations: The commuting diagram in Figure 4.7 outlines the structure of the proof that the **fresh** operator in the $\mathcal{O}$ -calculus has a consistent denotational semantics. Here we assume that:

- A7 The expression $\text{fresh}_{\text{LOC}}(\text{dom } H_0)$ evaluates to ℓ_j .
- A8 The context-heap value H_1 is $H_0 \oplus \{\ell_j \mapsto \mathbf{i}\}$.

$$\begin{array}{ccc} \Gamma\{\text{heap} \mapsto H_0\} \bullet \text{fresh} & \xrightarrow{A7, A8} & \Gamma\{\text{heap} \leftarrow H_1\} \bullet \ell_j \\ \llbracket - \rrbracket_J \downarrow & & \llbracket - \rrbracket_J \downarrow \\ (\Pi_{\text{HEAP}} := \llbracket H_0 \rrbracket_J^E) \circ \text{fresh}_J & \xrightarrow{L4.4, A7, A8} & (\Pi_{\text{HEAP}} := \llbracket H_1 \rrbracket_J^E) \circ (\mathcal{E}_J \ell_j) \end{array}$$

Figure 4.7: Fresh location commuting diagram

Lemma 4.4 – Fresh location lemma:

$$(\Pi_{\text{HEAP}} := H) \circ \text{fresh}_J = (\Pi_{\text{HEAP}} := H \oplus \{r \mapsto \mathbf{i}\}) \circ (\mathcal{E}_J r)$$

where $r = \text{fresh}_{\text{LOC}}(\text{dom } \Pi_{\text{HEAP}})$

$\square$

Dereferencing a heap location: The commuting diagram in Figure 4.8 outlines the structure of the proof that the dereference operator $(*__)$ in the $\mathcal{O}$ -calculus has a consistent denotational semantics. Here, assumption A9 is that $r \in \Pi_{\text{HEAP}}$ and $v = \Pi_{\text{HEAP}}(r)$.

$$\begin{array}{ccc} \Gamma\{\text{heap} \mapsto H\} \bullet *_r & \xrightarrow{A9} & \Gamma\{\text{heap} \leftarrow H\} \bullet v \\ \llbracket - \rrbracket_J \downarrow & & \llbracket - \rrbracket_J \downarrow \\ (\Pi_{\text{HEAP}} := \llbracket H \rrbracket_J^E) \circ \llbracket r \rrbracket_J \circ \text{deref}_J & \xrightarrow{L4.5, A9} & (\Pi_{\text{HEAP}} := \llbracket H \rrbracket_J^E) \circ \llbracket v \rrbracket_J \end{array}$$

Figure 4.8: Dereferencing location commuting diagram

Lemma 4.5 – Dereferencing location lemma:

$$(\Pi_{\text{HEAP}} := H) \circledast (\mathcal{E}_J r) \circledast \text{deref}_J = (\Pi_{\text{HEAP}} := H) \circledast (\mathcal{E}_J v)$$

where $r \in \Pi_{\text{HEAP}}$ and $v = \Pi_{\text{HEAP}}(r)$

□

Updating a heap location: The commuting diagram in Figure 4.9 outlines the structure of the proof that the reference update operator ($_ \ast _$) in the $\mathcal{O}$ -calculus has a consistent denotational semantics. Here, assumption A10 is that the context-heap value $H_1 = H_0 \oplus \{r \mapsto v\}$.

$$\begin{array}{ccc}
 \Gamma\{\text{heap} \mapsto H_0\} \bullet r \ast v & \xrightarrow{A10} & \Gamma\{\text{heap} \mapsto H_1\} \bullet r \\
 \downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
 (\Pi_{\text{HEAP}} := \llbracket H_0 \rrbracket_J^E) \circledast & & \\
 \text{cmdExp}_J(\text{trans}_J(rvPair, 2), (\llbracket r \rrbracket_J, \llbracket v \rrbracket_J)) \circledast & & \\
 \text{update}_J & & \\
 \downarrow L4.1 & & \downarrow L4.6, A10 \\
 (\Pi_{\text{HEAP}} := \llbracket H_0 \rrbracket_J^E) \circledast \mathcal{E}_J(\llbracket r \rrbracket_J^E, \llbracket v \rrbracket_J^E) \circledast \text{update}_J & \xlongequal{\quad} & (\Pi_{\text{HEAP}} := \llbracket H_1 \rrbracket_J^E) \circledast \llbracket r \rrbracket_J
 \end{array}$$

Figure 4.9: Updating location commuting diagram

Lemma 4.6 – Updating location lemma:

$$(\Pi_{\text{HEAP}} := H) \circledast \mathcal{E}_J(r, v) \circledast \text{update}_J = (\Pi_{\text{HEAP}} := H \oplus \{r \mapsto v\}) \circledast (\mathcal{E}_J r)$$

where $r \in \Pi_{\text{HEAP}}$

□

4.2.5 Other programming mechanisms

Sequential composition: The composition of an $\mathcal{O}$ -calculus value and an expression results in the evaluation of the expression. The two commuting diagrams in Figure 4.10 illustrate the consistency argument for the $\mathcal{U}_S$ -design and the other UTP models respectively. In the latter

$$\begin{array}{ccc}
 \Gamma \bullet v \circledast e & \xrightarrow{\quad} & \Gamma \bullet e \\
 \downarrow \llbracket - \rrbracket_S & & \downarrow \llbracket - \rrbracket_S \\
 \llbracket \Gamma \rrbracket_S \circledast \mathcal{E}_S v \circledast \text{pop}_S \circledast \llbracket e \rrbracket_S & \xlongequal{L4.7} & \llbracket \Gamma \rrbracket_S \circledast \llbracket e \rrbracket_S
 \end{array}
 \qquad
 \begin{array}{ccc}
 \Gamma \bullet v \circledast e & \xrightarrow{\quad} & \Gamma \bullet e \\
 \downarrow \llbracket - \rrbracket_J & & \downarrow \llbracket - \rrbracket_J \\
 \llbracket \Gamma \rrbracket_J \circledast \mathcal{E}_J v \circledast \llbracket e \rrbracket_J & \xlongequal{L4.7, A11} & \llbracket \Gamma \rrbracket_J \circledast \llbracket e \rrbracket_J
 \end{array}$$

Figure 4.10: Sequential composition commuting diagram

– right hand side (RHS) – case, the assumption A11 is that the semantics of expression e is independent of the value v . This assumption holds whenever the LHS of the sequential composition does not update the evaluation context (Γ), thus as the evaluation of a value can never update the context, assumption A11 holds in this case.

Lemma 4.7 – Sequential Composition:

$$\begin{aligned} \mathcal{E}_S \text{ ov } \mathbin{\text{\textcircled{;}}} \text{pop}_S \mathbin{\text{\textcircled{;}}} P &= P \\ \mathcal{E}_J \text{ ov } \mathbin{\text{\textcircled{;}}} P &= P \end{aligned}$$

Here, for J in $\{R, M\}$ we also require that the side condition $\forall \text{ov} \bullet P = P[\text{ov}/\Pi_{\text{RES}}]$, which is assumption A11.

□

Conditional choice: The conditional term is split into two cases, depending on the evaluation of its first argument (the conditional), which is either **true** or **false**. The commuting diagrams in Figure 4.11 illustrate the consistency argument for both of these cases, where:

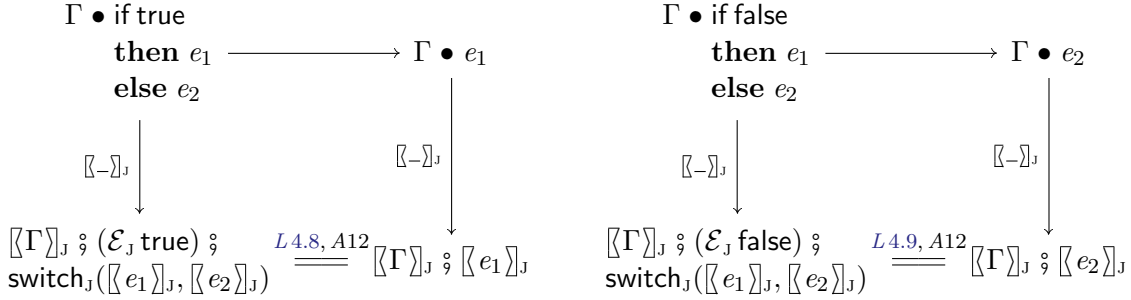


Figure 4.11: Conditional commuting diagrams

assumption A12 is that neither expression e_1 or e_2 are dependent on the evaluation of the conditional argument; and L4.8 and L4.9 denote the following true and false conditional lemmas respectively.

Lemma 4.8 – Conditional True:

$$\mathcal{E}_J \text{ true } \mathbin{\text{\textcircled{;}}} \text{switch}_J(P, Q) = P$$

Here, for J in $\{R, M\}$, we also require that the side conditions $\forall b \bullet P = P[b/\Pi_{\text{RES}}]$ and $\forall b \bullet Q = Q[b/\Pi_{\text{RES}}]$ hold, which is assumption A12.

□

Lemma 4.9 – Conditional False:

$$\mathcal{E}_J \text{ false } \mathbin{\text{\textcircled{;}}} \text{switch}_J(P, Q) = Q$$

Here, for J in $\{R, M\}$, we also require that the side conditions $\forall b \bullet P = P[b/\Pi_{\text{RES}}]$ and $\forall b \bullet Q = Q[b/\Pi_{\text{RES}}]$ hold, which is assumption A12.

□

The proofs for these lemmas follow almost immediately from the definition of the switch_J command, which mirrors that of the $\mathcal{O}$ -calculus conditional expression.

4.2.6 Proof notation

We introduce an alternative notation for the design framing mechanism presented in Section 2.3.3, to simplify the presentation of the proofs that follow in the remainder of this chapter.

$$p \vdash^V P \hat{=} V : (p \vdash P)$$

4.3 Operand Stack Model

We now prove the lemmas that require model-specific steps for the $\mathcal{U}_s$ -design of the $\mathcal{O}$ -calculus.

4.3.1 Command expressions

The command expression lemma (Lemma 4.1) is now specialised for the $\mathcal{U}_s$ -design model as follows

$$(ov_{i=1}^k) \in f \Rightarrow \text{cmdExp}_s(\text{trans}_s(f, k), (\mathcal{E}_s ov_i)) = \mathcal{E}_s(f(ov_{i=1}^k))$$

Within the following proof, the LHS of the initial implication is added as an assumption to the proof context.

Proof

$$\begin{aligned}
& \text{cmdExp}_s(\text{trans}_s(f, k), (\mathcal{E}_s ov_i)) \\
= & \dots \text{Defn. of cmdExp}_s \\
& (\mathcal{E}_s ov_i) \circ \text{trans}_s(f, k) \\
= & \dots \text{Defn. of } \mathcal{E}_s \\
& (\mathcal{E}_s (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \langle ov_i \rangle \cap \Pi_{\text{STK}})) \circ \text{trans}_s(f, k) \\
= & \dots \text{Defn. of } _ \circ _ \text{ and predicate logic} \\
& (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \langle \overset{i=k}{1} ov_i \rangle \cap \Pi_{\text{STK}}) \circ \text{trans}_s(f, k) \\
= & \dots \text{Defn. of trans}_s \\
& (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \langle \overset{i=k}{1} ov_i \rangle \cap \Pi_{\text{STK}}) \circ \\
& (\# \Pi_{\text{STK}} \geq k \wedge (\Pi_{\text{STK}} \text{proj } k) \in f \\
& \quad \vdash \{\Pi_{\text{STK}}\} \\
& \quad \Pi_{\text{STK}}' = \langle f(\Pi_{\text{STK}} \text{proj } k) \rangle \cap (\text{tail}^k \Pi_{\text{STK}}) \\
&) \\
= & \dots \text{Law C.13 assign comp.} \\
& (\# \Pi_{\text{STK}} \geq k \wedge (\Pi_{\text{STK}} \text{proj } k) \in f)[\langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}} / \Pi_{\text{STK}}] \\
& \vdash \{\Pi_{\text{STK}}\} \\
& (\Pi_{\text{STK}}' = \langle f(\Pi_{\text{STK}} \text{proj } k) \rangle \cap (\text{tail}^k \Pi_{\text{STK}}))[\langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}} / \Pi_{\text{STK}}] \\
= & \dots \text{Semantic substitution} \\
& \#(\langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}}) \geq k \wedge ((\langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}}) \text{proj } k) \in f \\
& \vdash \{\Pi_{\text{STK}}\} \\
& \Pi_{\text{STK}}' = \langle f((\langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}}) \text{proj } k) \rangle \cap (\text{tail}^k \langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}}) \\
= & \dots \text{Simplification} \\
& ((\langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}}) \text{proj } k) \in f \quad \# \langle \overset{i=k}{1} ov_1 \rangle = k \\
& \vdash \{\Pi_{\text{STK}}\} \quad \text{tail}^k(\langle \overset{i=k}{1} ov_1 \rangle = k) = \langle \rangle \\
& \Pi_{\text{STK}}' = \langle f(\langle \overset{i=k}{1} ov_1 \rangle \cap \Pi_{\text{STK}}) \text{proj } k \rangle \cap \Pi_{\text{STK}} \\
= & \dots \langle \overset{i=k}{1} t \rangle \cap t \text{proj } k = t_{i=1}^k. \\
& (ov_{i=1}^k) \in f \\
& \vdash \{\Pi_{\text{STK}}\} \\
& \Pi_{\text{STK}}' = \langle f(ov_{i=1}^k) \rangle \cap \Pi_{\text{STK}} \\
= & \dots \text{A3, i.e. } (ov_{i=1}^k) \in f \\
& \text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \langle f(ov_{i=1}^k) \rangle \cap \Pi_{\text{STK}} \\
= & \dots \text{Defn. of } \mathcal{E}_s \\
& \mathcal{E}_s(f(ov_{i=1}^k))
\end{aligned}$$

□

4.3.2 Procedure invocation

The only lemma that has not already been proven for showing consistency of either the method or function invocation operations is the procedure call lemma (Lemma 4.3), which is now specialised for the $\mathcal{U}_s$ -design as follows

$$\mathcal{E}_s(\llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle) \mathbin{\text{\textcircled{;}}} \text{apply}_s(\text{call}_s) = \text{ext}_s(\langle \llbracket e \rrbracket_s \{x \leftarrow \llbracket v \rrbracket_s^E\}_s, \text{call}_s)$$

Proof

$$\begin{aligned}
& \mathcal{E}_s(\llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle) \mathbin{\text{\textcircled{;}}} \text{apply}_s(\text{call}_s) \\
= & \dots \text{Defn. of } \mathcal{E}_s \\
& (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \langle \llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle \rangle \wedge \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
& \text{apply}_s(\text{call}_s) \\
= & \dots \text{Defn. of } \text{apply}_s \\
& (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \langle \llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle \rangle \wedge \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
& (\ (\ \exists v_1, x_1, t_1 \mid (v_1, \langle x_1, t_1 \rangle) = \text{head } \Pi_{\text{STK}} \bullet \\
& \quad \text{pop}_s \mathbin{\text{\textcircled{;}}} \text{ext}_s(\langle t_1 \{x_1 \leftarrow \text{text}_s(v_1)\}_s, \text{call}_s) \\
& \quad) \\
& \quad \triangleleft \# \Pi_{\text{STK}} > 1 \ \wedge \ (\text{head } \Pi_{\text{STK}}) \in \text{Value} \times \text{ProgramText} \triangleright \\
& \quad \text{chaos} \\
& \) \\
= & \dots \text{Defn. of } \text{pop}_s \\
& (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \langle \llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle \rangle \wedge \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
& (\ (\ \exists v_1, x_1, t_1 \mid (v_1, \langle x_1, t_1 \rangle) = \text{head } \Pi_{\text{STK}} \bullet \\
& \quad (\# \Pi_{\text{STK}} > 1 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \text{tail } \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
& \quad \text{ext}_s(\langle t_1 \{x_1 \leftarrow \text{text}_s(v_1)\}_s, \text{call}_s) \\
& \quad) \\
& \quad \triangleleft \# \Pi_{\text{STK}} > 1 \ \wedge \ (\text{head } \Pi_{\text{STK}}) \in \text{Value} \times \text{ProgramText} \triangleright \\
& \quad \text{chaos} \\
& \) \\
= & \dots \text{Defn. of } _ \mathbin{\text{\textcircled{;}}} _ \\
& (\ \exists v_1, x_1, t_1 \mid (v_1, \langle x_1, t_1 \rangle) = \langle \llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle \rangle \bullet \\
& \quad (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
& \quad \text{ext}_s(\langle t_1 \{x_1 \leftarrow \text{text}_s(v_1)\}_s, \text{call}_s) \\
& \) \\
& \triangleleft \text{true} \ \wedge \ (\llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle) \in \text{Value} \times \text{ProgramText} \triangleright \\
& \text{chaos} \\
= & \dots \text{Defn. of } _ \triangleleft _ \triangleright _ \\
& \exists v_1, x_1, t_1 \mid (v_1, \langle x_1, t_1 \rangle) = \langle \llbracket v \rrbracket_s^E, \langle x, \llbracket e \rrbracket_s \rangle \rangle \bullet \\
& \quad (\text{true} \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
& \quad \text{ext}_s(\langle t_1 \{x_1 \leftarrow \text{text}_s(v_1)\}_s, \text{call}_s) \\
= & \dots \text{One point rule} \\
& \text{skip}^D \mathbin{\text{\textcircled{;}}} \\
& \text{ext}_s(\langle \llbracket e \rrbracket_s \{x \leftarrow \text{text}_s(\llbracket v \rrbracket_s^E)\}_s, \text{call}_s) \\
= & \dots \text{skip}^D \text{ unit of } _ \mathbin{\text{\textcircled{;}}} _ \\
& \text{ext}_s(\langle \llbracket e \rrbracket_s \{x \leftarrow \llbracket v \rrbracket_s^E\}_s, \text{call}_s) \\
& \text{and defn. of } \text{text}_s
\end{aligned}$$

□

4.3.3 Heap operations

Fresh heap locations: The commuting diagram in Figure 4.7 outlines the structure of the proof that the fresh operator in the $\mathcal{O}$ -calculus has a consistent denotational semantics. Here we assume that:

- A7 The expression $\text{fresh}_{\text{LOC}}(\text{dom } H_0)$ evaluates to ℓ_j .
- A8 The context-heap value H_1 is $H_0 \oplus \{\ell_j \mapsto \mathbf{i}\}$.

The fresh location lemma (Lemma 4.4) is specialised for the $\mathcal{U}_s$ -design as follows

$$(\Pi_{\text{HEAP}} := H) \mathbin{\text{\textcircled{;}}} \text{fresh}_s = (\Pi_{\text{HEAP}} := H \oplus \{r \mapsto \mathbf{i}\}) \mathbin{\text{\textcircled{;}}} (\mathcal{E}_s r)$$

where $r = \text{fresh}_{\text{LOC}}(\text{dom } \Pi_{\text{HEAP}})$.

Proof

$$\begin{aligned}
& (\Pi_{\text{HEAP}} := H) \mathbin{\text{\textcircled{;}}} \text{fresh}_s \\
= & \dots \dots \dots \text{Defn. of } \text{fresh}_s \\
& (\Pi_{\text{HEAP}} := H) \mathbin{\text{\textcircled{;}}} \\
& (\exists r \mid r = \text{fresh}_{\text{LOC}}(\text{dom } \Pi_{\text{HEAP}}) \bullet \\
& \quad \Pi_{\text{STK}}, \Pi_{\text{HEAP}} := \langle r \rangle \frown \Pi_{\text{STK}}, \Pi_{\text{HEAP}} \oplus \{r \mapsto \mathbf{i}\} \\
&) \\
= & \dots \dots \dots \text{Law C.13 assign comp.} \\
& \exists r \mid r = \text{fresh}_{\text{LOC}}(\text{dom } H) \bullet \\
& \quad \Pi_{\text{STK}}, \Pi_{\text{HEAP}} := \langle r \rangle \frown \Pi_{\text{STK}}, H \oplus \{r \mapsto \mathbf{i}\} \\
= & \dots \dots \dots \text{A7} \\
& \exists r \mid r = \ell_j \bullet \\
& \quad \Pi_{\text{STK}}, \Pi_{\text{HEAP}} := \langle r \rangle \frown \Pi_{\text{STK}}, H \oplus \{r \mapsto \mathbf{i}\} \\
= & \dots \dots \dots \text{1-point rule} \\
& \Pi_{\text{STK}}, \Pi_{\text{HEAP}} := \langle \ell_j \rangle \frown \Pi_{\text{STK}}, H \oplus \{\ell_j \mapsto \mathbf{i}\} \\
= & \dots \dots \dots \text{Defns. of } _ \mathbin{\text{\textcircled{;}}} _ \text{ and } \mathcal{E}_s \\
& (\Pi_{\text{HEAP}} := H \oplus \{\ell_j \mapsto \mathbf{i}\}) \mathbin{\text{\textcircled{;}}} (\mathcal{E}_s \ell_j) \\
\Box &
\end{aligned}$$

Dereferencing heap locations: The commuting diagram in Figure 4.8 outlines the structure of the proof that the dereference operator in the $\mathcal{O}$ -calculus has a consistent denotational semantics. Here we assume that $r \in H$ and $v = H(r)$.

The dereference location lemma (Lemma 4.5) is specialised for the $\mathcal{U}_s$ -design as follows

$$(\Pi_{\text{HEAP}} := H) \mathbin{\text{\textcircled{;}}} \text{fresh}_s = (\Pi_{\text{HEAP}} := H \oplus \{r \mapsto \mathbf{i}\}) \mathbin{\text{\textcircled{;}}} (\mathcal{E}_s r)$$

where $r = \text{fresh}_{\text{LOC}}(\text{dom } \Pi_{\text{HEAP}})$.

Proof

$$\begin{aligned}
& \Pi_{\text{HEAP}} := H \mathbin{\text{\textcircled{;}}} \mathcal{E}_s r \mathbin{\text{\textcircled{;}}} \text{deref}_s \\
= & \dots \dots \dots \text{Defn. of } \mathcal{E}_s \\
& \Pi_{\text{HEAP}} := H \mathbin{\text{\textcircled{;}}} \Pi_{\text{STK}} := \langle r \rangle \frown \Pi_{\text{STK}} \mathbin{\text{\textcircled{;}}} \text{deref}_s \\
= & \dots \dots \dots \text{Independent assignments} \\
& \Pi_{\text{HEAP}}, \Pi_{\text{STK}} := H, \langle r \rangle \frown \Pi_{\text{STK}} \mathbin{\text{\textcircled{;}}} \text{deref}_s
\end{aligned}$$

$$\begin{aligned}
 &= \dots \dots \dots \text{Defn. of } \text{deref}_s \\
 &\quad \Pi_{\text{HEAP}}, \Pi_{\text{STK}} := H, \langle r \rangle \frown \Pi_{\text{STK}} \text{;} \\
 &\quad (\# \Pi_{\text{STK}} > 0 \wedge (\text{head } \Pi_{\text{STK}}) \in \Pi_{\text{HEAP}} \\
 &\quad \quad \vdash \{ \Pi_{\text{STK}} \} \\
 &\quad \quad \Pi_{\text{STK}}' = \langle \Pi_{\text{HEAP}}(\text{head } \Pi_{\text{STK}}) \rangle \frown (\text{tail } \Pi_{\text{STK}}) \\
 &\quad) \\
 &= \dots \dots \dots \text{Law C.13 assign comp.} \\
 &\quad (\#(\langle r \rangle \frown \Pi_{\text{STK}}) > 0 \wedge \text{head}(\langle r \rangle \frown \Pi_{\text{STK}}) \in H \\
 &\quad \quad \vdash \\
 &\quad \quad \Pi_{\text{HEAP}}' = H \wedge \\
 &\quad \quad \Pi_{\text{STK}}' = \langle H(\text{head}(\langle r \rangle \frown \Pi_{\text{STK}})) \rangle \frown \text{tail}(\langle r \rangle \frown \Pi_{\text{STK}}) \\
 &\quad) \\
 &= \dots \dots \dots \text{Simplification} \\
 &\quad (1 + \# \Pi_{\text{STK}} > 0 \wedge r \in H \\
 &\quad \quad \vdash \\
 &\quad \quad \Pi_{\text{HEAP}}' = H \wedge \Pi_{\text{STK}}' = \langle H(r) \rangle \frown \Pi_{\text{STK}} \\
 &\quad) \\
 &= \dots \dots \dots \text{Assumption A9} \\
 &\quad (\text{true} \\
 &\quad \quad \vdash \\
 &\quad \quad \Pi_{\text{HEAP}}' = H \wedge \Pi_{\text{STK}}' = \langle v \rangle \frown \Pi_{\text{STK}} \\
 &\quad) \\
 &= \dots \dots \dots \text{Defn. of } _ := _ \\
 &\quad \Pi_{\text{HEAP}}, \Pi_{\text{STK}} := H, \langle v \rangle \frown \Pi_{\text{STK}} \\
 &= \dots \dots \dots \text{Independent assignments} \\
 &\quad \Pi_{\text{HEAP}} := H \text{;} \Pi_{\text{STK}} := \langle v \rangle \frown \Pi_{\text{STK}} \\
 &= \dots \dots \dots \text{Defn. of } \mathcal{E}_s \\
 &\quad \Pi_{\text{HEAP}} := H \text{;} \mathcal{E}_s v \\
 &\square
 \end{aligned}$$

4.3.4 Other programming operations

Sequential composition operation: The sequential composition operation structural induction (of Section 4.2.5) had one outstanding lemma (Lemma 4.7), which is specialised for the $\mathcal{U}_s$ -design as follows

$$\mathcal{E}_s \text{ ov } \text{;} \text{pop}_s \text{;} P = P$$

Proof

$$\begin{aligned}
 &\mathcal{E}_s \text{ ov } \text{;} \text{pop}_s \text{;} P \\
 &= \dots \dots \dots \text{Defn. of } \mathcal{E}_s \\
 &\quad (\Pi_{\text{STK}} := \langle \text{ov} \rangle \frown \Pi_{\text{STK}}) \text{;} \text{pop}_s \text{;} P \\
 &= \dots \dots \dots \text{Defn. of } \text{pop}_s \\
 &\quad (\Pi_{\text{STK}} := \langle \text{ov} \rangle \frown \Pi_{\text{STK}}) \text{;} \\
 &\quad (\# \Pi_{\text{STK}} > 0 \vdash \{ \Pi_{\text{STK}} \} \Pi_{\text{STK}}' = \text{tail } \Pi_{\text{STK}}) \text{;} \\
 &\quad P
 \end{aligned}$$

$$\begin{aligned}
&= \dots\dots\dots \text{Law C.13 assign comp.} \\
&\quad (\#(\langle ov \rangle \cap \Pi_{\text{STK}}) > 0 \\
&\quad \quad \vdash \{\Pi_{\text{STK}}\} \\
&\quad \quad \Pi_{\text{STK}}' = \text{tail}(\langle ov \rangle \cap \Pi_{\text{STK}}) \\
&\quad) \mathbin{\text{\textcircled{;}}} P \\
&= \dots\dots\dots \text{Simplification} \\
&\quad (1 + \#\Pi_{\text{STK}} > 0 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} P \\
&= \dots\dots\dots \text{Simplification} \\
&\quad (\text{true} \vdash^\emptyset \text{true}) \mathbin{\text{\textcircled{;}}} P \\
&= \dots\dots\dots \text{Defn. skip}^D \\
&\quad \text{skip}^D \mathbin{\text{\textcircled{;}}} P \\
&= \dots\dots\dots \text{skip}^D \text{ unit of } _ \mathbin{\text{\textcircled{;}}} _ \\
&\quad P \\
&\square
\end{aligned}$$

Conditional operation: The conditional operation structural induction (of [Section 4.2.5](#)) left two undischarged lemmas, which corresponded to the **true** and **false** branches of the conditional statement. We now prove the **false** conditional lemma ([Lemma 4.9](#)), which is specialised for the $\mathcal{U}_S$ -design as follows

$$\mathcal{E}_S \text{ false} \mathbin{\text{\textcircled{;}}} \text{switch}_S(P, Q) = Q$$

Proof

$$\begin{aligned}
&\mathcal{E}_S \text{ false} \mathbin{\text{\textcircled{;}}} \text{switch}_S(P, Q) \\
&= \dots\dots\dots \text{Defn. of } \mathcal{E}_S \\
&\quad \Pi_{\text{STK}} := \langle \text{false} \rangle \cap \Pi_{\text{STK}} \mathbin{\text{\textcircled{;}}} \text{switch}_S(P, Q) \\
&= \dots\dots\dots \text{Defn. of } \text{switch}_S \\
&\quad \Pi_{\text{STK}} := \langle \text{false} \rangle \cap \Pi_{\text{STK}} \mathbin{\text{\textcircled{;}}} \\
&\quad ((\text{pop}_S \mathbin{\text{\textcircled{;}}} P \triangleleft \text{head } \Pi_{\text{STK}} \triangleright \text{pop}_S \mathbin{\text{\textcircled{;}}} Q) \\
&\quad \quad \triangleleft \#\Pi_{\text{STK}} > 0 \wedge (\text{head } \Pi_{\text{STK}}) \in \text{Boolean} \triangleright \\
&\quad \quad \text{chaos} \\
&\quad) \\
&= \dots\dots\dots \text{Defn. of } \text{pop}_S \\
&\quad \Pi_{\text{STK}} := \langle \text{false} \rangle \cap \Pi_{\text{STK}} \mathbin{\text{\textcircled{;}}} \\
&\quad (((\#\Pi_{\text{STK}} > 0 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \text{tail } \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} P \\
&\quad \quad \triangleleft \text{head } \Pi_{\text{STK}} \triangleright \\
&\quad \quad (\#\Pi_{\text{STK}} > 0 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \text{tail } \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} Q \\
&\quad) \\
&\quad \triangleleft \#\Pi_{\text{STK}} > 0 \wedge (\text{head } \Pi_{\text{STK}}) \in \text{Boolean} \triangleright \\
&\quad \text{chaos} \\
&\quad) \\
&= \dots\dots\dots \text{Law C.13 assign comp.} \\
&\quad ((\#(\langle \text{false} \rangle \cap \Pi_{\text{STK}}) > 0 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \text{tail}(\langle \text{false} \rangle \cap \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
&\quad \quad P \\
&\quad \quad \triangleleft \text{head}(\langle \text{false} \rangle \cap \Pi_{\text{STK}}) \triangleright \\
&\quad \quad (\#(\langle \text{false} \rangle \cap \Pi_{\text{STK}}) > 0 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \text{tail}(\langle \text{false} \rangle \cap \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} \\
&\quad \quad \quad Q \\
&\quad) \\
&\quad \triangleleft \#(\langle \text{false} \rangle \cap \Pi_{\text{STK}}) > 0 \wedge \text{head}(\langle \text{false} \rangle \cap \Pi_{\text{STK}}) \in \text{Boolean} \triangleright \\
&\quad \text{chaos}
\end{aligned}$$

$$\begin{aligned}
 &= \dots\dots\dots \text{Simplification} \\
 &\quad ((1 + \#\Pi_{\text{STK}} > 0 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} P \\
 &\quad \quad \triangleleft \text{false} \triangleright \\
 &\quad \quad (1 + \#\Pi_{\text{STK}} > 0 \vdash \{\Pi_{\text{STK}}\} \Pi_{\text{STK}}' = \Pi_{\text{STK}}) \mathbin{\text{\textcircled{;}}} Q \\
 &\quad) \\
 &\quad \triangleleft 1 + \#\Pi_{\text{STK}} > 0 \wedge \text{false} \in \text{Boolean} \triangleright \\
 &\quad \text{chaos} \\
 &= \dots\dots\dots \text{Simplification} \\
 &\quad ((\text{true} \vdash^\emptyset \text{true}) \mathbin{\text{\textcircled{;}}} P \\
 &\quad \quad \triangleleft \text{false} \triangleright \\
 &\quad \quad (\text{true} \vdash^\emptyset \text{true}) \mathbin{\text{\textcircled{;}}} Q \\
 &\quad) \\
 &\quad \triangleleft \text{true} \triangleright \\
 &\quad \text{chaos} \\
 &= \dots\dots\dots \text{Defn. of } _ \triangleleft _ \triangleright _ \\
 &\quad (\text{true} \vdash^\emptyset \text{true}) \mathbin{\text{\textcircled{;}}} Q \\
 &= \dots\dots\dots \text{Defn. of skip}^D \\
 &\quad \text{skip}^D \mathbin{\text{\textcircled{;}}} Q \\
 &= \dots\dots\dots \text{skip}^D \text{ unit of } _ \mathbin{\text{\textcircled{;}}} _ \\
 &\quad Q \\
 &\square
 \end{aligned}$$

The **true** conditional lemma (Lemma 4.8) is specialised and proved in a similar manner; the $\mathcal{U}_s$ -design specialised variant is

$$\mathcal{E}_s \text{ true } \mathbin{\text{\textcircled{;}}} \text{switch}_s(P, Q) = P$$

4.4 Result Variable Model

We now prove the lemmas that require model specific steps for the $\mathcal{U}_R$ -design of the $\mathcal{O}$ -calculus.

4.4.1 Command expressions

The command expression lemma (Lemma 4.1) is now specialised for the $\mathcal{U}_R$ -design model, as follows

$$(ov_{i=1}^k) \in f \Rightarrow \text{cmdExp}_R(\text{trans}_R(f, k), (\mathop{\mathcal{E}}_R \mathbin{\text{\textcircled{;}}} ov_i)) = \mathcal{E}_R(f(ov_{i=1}^k))$$

Within the following proof, the LHS of the initial implication is added as an assumption to the proof context.

Proof

$$\begin{aligned}
 &\text{cmdExp}_R(\text{trans}_R(f, k), (\mathop{\mathcal{E}}_R \mathbin{\text{\textcircled{;}}} ov_i)) \\
 &= \dots\dots\dots \text{Defn. of cmdExp}_R \\
 &\quad \text{decl } x_{i=1}^k \text{ in} \\
 &\quad \quad (\mathop{\mathcal{E}}_R \mathbin{\text{\textcircled{;}}} ov_j \mid x_{i=1}^k \mathbin{\text{\textcircled{;}}} x_j := \Pi_{\text{RES}}) \mathbin{\text{\textcircled{;}}} \\
 &\quad \quad \text{trans}_R(f, k) \\
 &= \dots\dots\dots \text{Defn. of } _ \mid _ \mathbin{\text{\textcircled{;}}} _ \\
 &\quad \text{decl } x_{i=1}^k \text{ in} \\
 &\quad \quad (\mathop{\mathcal{E}}_R \mathbin{\text{\textcircled{;}}} (\text{hide } x_{i=1}^k \text{ from } \mathop{\mathcal{E}}_R \mathbin{\text{\textcircled{;}}} ov_j) \mathbin{\text{\textcircled{;}}} x_j := \Pi_{\text{RES}}) \mathbin{\text{\textcircled{;}}} \\
 &\quad \quad \text{trans}_R(f, k)
 \end{aligned}$$

$$\begin{aligned}
&= \dots \dots \dots \text{Defn. of } \mathcal{E}_R \text{ and } _ := _ \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad (\circ_{j=1}^k \text{ (hide } x_{i=1}^k \text{ from} \\
&\quad \quad \quad (\text{true} \vdash \Pi_{\text{RES}}' = \text{ov}_j \wedge \Pi_{\text{HEAP}}' = \Pi_{\text{HEAP}}) \\
&\quad \quad \quad) \circ \\
&\quad \quad \quad x_j := \Pi_{\text{RES}} \\
&\quad \quad \quad) \circ \\
&\quad \quad \text{trans}_R(f, k) \\
&= \dots \dots \dots \text{Law C.7 hide design} \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad (\circ_{j=1}^k \text{ (true} \\
&\quad \quad \quad \vdash \\
&\quad \quad \quad \Pi_{\text{RES}}' = \text{ov}_j \wedge \Pi_{\text{HEAP}}' = \Pi_{\text{HEAP}} \wedge x_{i=1}'^k = x_{i=1}^k \\
&\quad \quad \quad) \circ \\
&\quad \quad \quad x_j := \Pi_{\text{RES}} \\
&\quad \quad \quad) \circ \\
&\quad \quad \text{trans}_R(f, k) \\
&= \dots \dots \dots \text{defn. } _ := _ \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad (\circ_{j=1}^k \Pi_{\text{RES}} := \text{ov}_j \circ x_j := \Pi_{\text{RES}}) \circ \\
&\quad \quad \text{trans}_R(f, k) \\
&= \dots \dots \dots \text{Transitivity of assignment} \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad (\circ_{j=1}^k \Pi_{\text{RES}} := \text{ov}_j \circ x_j := \text{ov}_j) \circ \\
&\quad \quad \text{trans}_R(f, k) \\
&= \dots \dots \dots \text{Independent assignments} \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad (\circ_{j=1}^k x_j := \text{ov}_j \circ \Pi_{\text{RES}} := \text{ov}_j) \circ \\
&\quad \quad \text{trans}_R(f, k) \\
&= \dots \dots \dots \text{Simplification} \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad x_{j=1}^k, \Pi_{\text{RES}} := \text{ov}_{j=1}^k, \text{ov}_k \circ \\
&\quad \quad \text{trans}_R(f, k) \\
&= \dots \dots \dots \text{Defn. of } \text{trans}_R \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad x_{j=1}^k, \Pi_{\text{RES}} := \text{ov}_{j=1}^k, \text{ov}_k \circ \\
&\quad \quad (x_{i=1}^k \in f \vdash \{\Pi_{\text{RES}}\} \Pi_{\text{RES}}' = f(x_{i=1}^k)) \\
&= \dots \dots \dots \text{Law C.13 assign comp.} \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad (\text{ov}_{i=1}^k \in f \vdash \{\Pi_{\text{RES}}\} \Pi_{\text{RES}} = f(\text{ov}_{i=1}^k)) \\
&= \dots \dots \dots \text{A3, i.e. } (\text{ov}_{i=1}^k) \in f \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad \text{true} \vdash \{\Pi_{\text{RES}}\} \Pi_{\text{RES}}' = f(\text{ov}_{i=1}^k) \\
&= \dots \dots \dots \text{Defn. of } \mathcal{E}_R \\
&\quad \text{decl } x_{i=1}^k \text{ in} \\
&\quad \quad \Pi_{\text{RES}} := f(\text{ov}_{i=1}^k) \\
&= \dots \dots \dots \text{Defn. of decl} \\
&\quad \text{var } x_{i=1}^k \circ \Pi_{\text{RES}} := f(\text{ov}_{i=1}^k) \circ \text{end } x_{i=1}^k
\end{aligned}$$

$$\begin{aligned}
 &= \dots\dots\dots x_{i=1}^k \text{ not in assignment} \\
 &\quad \text{var } x_{i=1}^k \text{ ; end } x_{i=1}^k \text{ ; } \Pi_{\text{RES}} := f(ov_{i=1}^k) \\
 &= \dots\dots\dots \text{Unused local vars.} \\
 &\quad \text{skip}^D \text{ ; } \Pi_{\text{RES}} := f(ov_{i=1}^k) \\
 &= \dots\dots\dots \text{Defn. of } \mathcal{E}_R, \text{ and skip}^D \text{ unit} \\
 &\quad \mathcal{E}_R f(ov_{i=1}^k) \text{ of } - \text{ ; } - \\
 \square
 \end{aligned}$$

4.4.2 Heap operations

Dereferencing a heap location: The dereference location lemma (Lemma 4.5) is specialised for the $\mathcal{U}_S$ -design as follows

$$(\Pi_{\text{HEAP}} := H) \text{ ; fresh}_R = (\Pi_{\text{HEAP}} := H \oplus \{r \mapsto \mathfrak{!}\}) \text{ ; } (\mathcal{E}_R r)$$

where $r = \text{fresh}_{\text{LOC}}(\text{dom } \Pi_{\text{HEAP}})$.

Proof

$$\begin{aligned}
 &\Pi_{\text{HEAP}} := H \text{ ; } \mathcal{E}_R r \text{ ; deref}_R \\
 &= \dots\dots\dots \text{Defn. of } \mathcal{E}_R \\
 &\quad \Pi_{\text{HEAP}} := H \text{ ; } \Pi_{\text{RES}} := r \text{ ; deref}_R \\
 &= \dots\dots\dots \text{Independent assignments} \\
 &\quad \Pi_{\text{HEAP}}, \Pi_{\text{RES}} := H, r \text{ ; deref}_R \\
 &= \dots\dots\dots \text{Defn. of deref}_R \\
 &\quad \Pi_{\text{HEAP}}, \Pi_{\text{RES}} := H, r \text{ ;} \\
 &\quad (\Pi_{\text{RES}} \in \Pi_{\text{HEAP}} \vdash \{\Pi_{\text{RES}}\} \Pi_{\text{RES}}' = \Pi_{\text{HEAP}}(\Pi_{\text{RES}})) \\
 &= \dots\dots\dots \text{Law C.13 assign comp.} \\
 &\quad r \in H \vdash \Pi_{\text{HEAP}}' = H \wedge \Pi_{\text{RES}}' = H(r) \\
 &= \dots\dots\dots \text{Assumption A9} \\
 &\quad \text{true} \vdash \Pi_{\text{HEAP}}' = H \wedge \Pi_{\text{RES}}' = v \\
 &= \dots\dots\dots \text{Defn. of } - := - \\
 &\quad \Pi_{\text{HEAP}}, \Pi_{\text{RES}} := H, v \\
 &= \dots\dots\dots \text{Independent assignments} \\
 &\quad \Pi_{\text{HEAP}} := H \text{ ; } \Pi_{\text{RES}} := v \\
 &= \dots\dots\dots \text{Defn. of } \mathcal{E}_R \\
 &\quad \Pi_{\text{HEAP}} := H \text{ ; } \mathcal{E}_R v \\
 \square
 \end{aligned}$$

Updating a heap location: The updating location lemma (Lemma 4.6) is specialised for the $\mathcal{U}_R$ -design as follows

$$(\Pi_{\text{HEAP}} := H) \text{ ; } \mathcal{E}_R(r, v) \text{ ; update}_R = (\Pi_{\text{HEAP}} := H \oplus \{r \mapsto v\}) \text{ ; } (\mathcal{E}_R r)$$

where $r \in \Pi_{\text{HEAP}}$.

Proof (for the $\mathcal{U}_R$ -design).

$$\begin{aligned}
 &\Pi_{\text{HEAP}} := H \text{ ; } \mathcal{E}_R(r, v) \text{ ; update}_R \\
 &= \dots\dots\dots \text{Defn. of } \mathcal{E}_R \\
 &\quad \Pi_{\text{HEAP}} := H \text{ ; } \Pi_{\text{RES}} := (r, v) \text{ ; update}_R \\
 &= \dots\dots\dots \text{Independent assignments} \\
 &\quad \Pi_{\text{HEAP}}, \Pi_{\text{RES}} := H, (r, v) \text{ ; update}_R
 \end{aligned}$$

$$\begin{aligned}
&= \dots \text{Defn. of } \text{update}_R \\
&\quad \Pi_{\text{HEAP}}, \Pi_{\text{RES}} := H, (r, v) \circ \\
&\quad (\exists r_1, v_1 \mid (r_1, v_1) = \Pi_{\text{RES}} \bullet \\
&\quad \quad \Pi_{\text{RES}}, \Pi_{\text{HEAP}} := r_1, \Pi_{\text{HEAP}} \oplus \{r_1 \mapsto v_1\} \\
&\quad \quad \triangleleft r_1 \in \Pi_{\text{HEAP}} \triangleright \\
&\quad \quad \text{chaos} \\
&\quad) \\
&\quad \triangleleft \Pi_{\text{RES}} \in \text{Location} \times \text{Value} \triangleright \\
&\quad \text{chaos} \\
&= \dots \text{Law C.13 assign comp.} \\
&\quad (\exists r_1, v_1 \mid (r_1, v_1) = (r, v) \bullet \\
&\quad \quad \Pi_{\text{RES}}, \Pi_{\text{HEAP}} := r_1, H \oplus \{r_1 \mapsto v_1\} \\
&\quad \quad \triangleleft r_1 \in H \triangleright \\
&\quad \quad \text{chaos} \\
&\quad) \\
&\quad \triangleleft (r, v) \in \text{Location} \times \text{Value} \triangleright \\
&\quad \text{chaos} \\
&= \dots (r, v) \in \text{Location} \times \text{Value} \\
&\quad \exists r_1, v_1 \mid (r_1, v_1) = (r, v) \bullet \\
&\quad \quad \Pi_{\text{RES}}, \Pi_{\text{HEAP}} := r_1, H \oplus \{r_1 \mapsto v_1\} \\
&\quad \quad \triangleleft r_1 \in H \triangleright \\
&\quad \quad \text{chaos} \\
&= \dots \text{One point rule} \\
&\quad \Pi_{\text{RES}}, \Pi_{\text{HEAP}} := r, H \oplus \{r \mapsto v\} \\
&\quad \triangleleft r \in H \triangleright \\
&\quad \text{chaos} \\
&= \dots \text{Assumption A10, } r \in H \\
&\quad \Pi_{\text{RES}}, \Pi_{\text{HEAP}} := r, H \oplus \{r \mapsto v\} \\
&= \dots \text{Independent assignments} \\
&\quad \Pi_{\text{HEAP}} := H \oplus \{r \mapsto v\} \circ \Pi_{\text{RES}} := r \\
&= \dots \text{Defn. of } \mathcal{E}_R \\
&\quad \Pi_{\text{HEAP}} := H \oplus \{r \mapsto v\} \circ \mathcal{E}_R r \\
&\square
\end{aligned}$$

4.4.3 Other programming operations

Sequential composition operation: The sequential composition operation structural induction (of Section 4.2.5) had one outstanding lemma (Lemma 4.7), which is specialised for the $\mathcal{U}_R$ -design as follows

$$\mathcal{E}_R \text{ ov } \circ P = P$$

where $\forall \text{ ov } \bullet P = P[\text{ov}/\Pi_{\text{RES}}]$, which is assumption A9.

Proof

$$\begin{aligned}
&\mathcal{E}_R \text{ ov } \circ P \\
&= \dots \text{Defn. of } \mathcal{E}_R \\
&\quad \Pi_{\text{RES}} := \text{ov } \circ P \\
&= \dots \text{Law C.13 assign comp.} \\
&\quad P[\text{ov}/\Pi_{\text{RES}}] \\
&= \dots \text{A11,} \\
&\quad P \qquad \qquad \qquad \forall \text{ ov } \bullet P = P[\text{ov}/\Pi_{\text{RES}}] \\
&\square
\end{aligned}$$

Conditional operation: The conditional operation structural induction (of [Section 4.2.5](#)) left two undischarged lemmas, which corresponded to the **true** and **false** branches of the conditional statement. We now prove the **false** conditional lemma ([Lemma 4.9](#)), which is specialised for the $\mathcal{U}_R$ -design as follows

$$\mathcal{E}_R \text{ false} \circ \text{switch}_R(P, Q) = Q$$

where $\forall b \bullet P = P[b/\Pi_{\text{RES}}]$ and $\forall b \bullet Q = Q[b/\Pi_{\text{RES}}]$, which is assumption A10.

Proof

$$\begin{aligned} & \mathcal{E}_R \text{ false} \circ \text{switch}_R(P, Q) \\ = & \dots \dots \dots \text{Defn. of } \mathcal{E}_R \\ & \Pi_{\text{RES}} := \text{false} \circ \text{switch}_R(P, Q) \\ = & \dots \dots \dots \text{Defn. of } \text{switch}_R \\ & \Pi_{\text{RES}} := \text{false} \circ \\ & ((P \triangleleft \text{head } \Pi_{\text{RES}} \triangleright Q) \\ & \quad \triangleleft \Pi_{\text{RES}} \in \text{Boolean} \triangleright \\ & \quad \text{chaos} \\ &) \\ = & \dots \dots \dots \text{Law C.13 assign comp.} \\ & ((P[\text{false}/\Pi_{\text{RES}}] \triangleleft \text{false} \triangleright Q[\text{false}/\Pi_{\text{RES}}]) \\ & \quad \triangleleft \text{false} \in \text{Boolean} \triangleright \\ & \quad \text{chaos}[\text{false}/\Pi_{\text{RES}}] \\ &) \\ = & \dots \dots \dots \text{Defn. of } _ \triangleleft _ \triangleright _ \\ & Q[\text{false}/\Pi_{\text{RES}}] \\ = & \dots \dots \dots \text{A12, } Q = Q[\text{false}/\Pi_{\text{RES}}] \\ & Q \end{aligned}$$

□

The **true** conditional lemma ([Lemma 4.8](#)) is specialised and proved in a similar manner; the $\mathcal{U}_R$ -design specialised variant is

$$\mathcal{E}_R \text{ true} \circ \text{switch}_R(P, Q) = P$$

where $\forall b \bullet P = P[b/\Pi_{\text{RES}}]$ and $\forall b \bullet Q = Q[b/\Pi_{\text{RES}}]$, which is assumption A10.

4.5 Constant Map Model

We now prove the lemmas that require model specific steps for the $\mathcal{U}_M$ -design of the $\mathcal{O}$ -calculus.

4.5.1 Procedure invocation

The proof for the $\mathcal{U}_M$ -design's procedure call lemma ([Lemma 4.3](#)) makes use of the following call iteration lemma, which reduces the proof to that of a single iteration (step) of a potentially recursive (or mutually recursive) procedure call.

Lemma 4.10 – Call iteration:

$$\mathcal{E}_M(v, \langle x, M, pt \rangle) \circ \text{apply}_M(\text{call}_M) = \text{decl } \llbracket M \{x \leftarrow \text{text}_M(v)\} \rrbracket_M \text{ in } \llbracket pt \{x \leftarrow \text{text}_M(v)\} \rrbracket_M$$

□

The proof of this lemma follows at the end of this section.

We now specialise the procedure call lemma for the $\mathcal{U}_M$ -design as follows

$$\mathcal{E}_M(\llbracket v \rrbracket_M^E, (\downarrow x, M, pt \downarrow)) \circ \text{apply}_M(\text{call}_M) = \text{decl } \{x \mapsto \llbracket v \rrbracket_M^E\} \text{ in } \llbracket \llbracket e \rrbracket_M^k y_i \leftarrow v_i \rrbracket_M \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket$$

where

$$pt \hat{=} e \{ \{ \rrbracket_{i=1}^k y_i \leftarrow v_i \} \}$$

$$M \hat{=} \{x \mapsto \text{free}, (y_i \mapsto \llbracket v_i \rrbracket_M^E)_{i=1}^k\}$$

Proof

$$\begin{aligned} & \mathcal{E}_M(\llbracket v \rrbracket_M^E, (\downarrow x, M, pt \downarrow)) \circ \text{apply}_M(\text{call}_M) \\ = & \dots \text{Lemma 4.10 call iter.} \\ & \text{decl } \llbracket M \{x \leftarrow \text{text}_M(\llbracket v \rrbracket_M^E)\}_M \rrbracket \text{ in } \\ & \quad \llbracket pt \{x \leftarrow \text{text}_M(\llbracket v \rrbracket_M^E)\}_M \rrbracket \\ = & \dots \text{Defn. text}_M \\ & \text{decl } \llbracket M \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \text{ in } \\ & \quad \llbracket pt \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \\ = & \dots \text{Defn. } M \\ & \text{decl } \llbracket \{x \mapsto \text{free}, (y_i \mapsto \llbracket v_i \rrbracket_M^E)_{i=1}^k\} \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \text{ in } \\ & \quad \llbracket pt \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \\ = & \dots \text{Defn. of } \llbracket - \leftarrow - \rrbracket_M \\ & \text{decl } \{x \mapsto \llbracket v \rrbracket_M^E, (y_i \mapsto \llbracket \llbracket v_i \rrbracket_M^E \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket)_{i=1}^k\} \text{ in } \\ & \quad \llbracket pt \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \\ = & \dots \text{Decl. by cases} \\ & \text{decl } \{x \mapsto \llbracket v \rrbracket_M^E\} \text{ in } \\ & \quad \text{decl } \{ \{ \rrbracket_{i=1}^k y_i \mapsto \llbracket \llbracket v_i \rrbracket_M^E \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \} \text{ in } \\ & \quad \quad \llbracket pt \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \\ = & \dots \text{Defn. of } pt \\ & \text{decl } \{x \mapsto \llbracket v \rrbracket_M^E\} \text{ in } \\ & \quad \text{decl } \{ \{ \rrbracket_{i=1}^k y_i \mapsto \llbracket \llbracket v_i \rrbracket_M^E \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \} \text{ in } \\ & \quad \quad \llbracket \llbracket e \rrbracket_M^k y_i \leftarrow v_i \rrbracket_M \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \\ = & \dots \text{Defn. of } \llbracket - \leftarrow - \rrbracket_M \\ & \text{decl } \{x \mapsto \llbracket v \rrbracket_M^E\} \text{ in } \\ & \quad \llbracket (\text{decl } \{ \{ \rrbracket_{i=1}^k y_i \mapsto \llbracket v_i \rrbracket_M^E \} \text{ in } \\ & \quad \quad \llbracket e \rrbracket_M^k y_i \leftarrow v_i \rrbracket_M \\ & \quad) \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \\ = & \dots \text{Defn. of } \llbracket - \leftarrow - \rrbracket_M \\ & \text{decl } \{x \mapsto \llbracket v \rrbracket_M^E\} \text{ in } \\ & \quad \llbracket \llbracket e \rrbracket_M^k y_i \leftarrow v_i \rrbracket_M \{x \leftarrow \llbracket v \rrbracket_M^E\}_M \rrbracket \\ \square \end{aligned}$$

Proof of the call iteration lemma (Lemma 4.10).

$$\begin{aligned} & \mathcal{E}_M(v, (\downarrow x, M, pt \downarrow)) \circ \text{apply}_M(\text{call}_M) \\ = & \dots \text{Defn. of } \mathcal{E}_M \\ & \Pi_{\text{RES}} := (v, (\downarrow x, M, pt \downarrow)) \circ \\ & \quad \text{apply}_M(\text{call}_M) \\ = & \dots \text{Defn. of } (\downarrow -, -, -) \\ & \Pi_{\text{RES}} := (v, (\downarrow x, \text{decl } M \text{ in } pt \downarrow)) \circ \\ & \quad \text{apply}_M(\text{call}_M) \end{aligned}$$

$$\begin{aligned}
&= \dots\dots\dots \text{Law C.13 assign comp.} \\
&\quad \text{apply}_M(\text{call}_M)[(v, \langle x, \text{decl } M \text{ in } pt \rangle / \Pi_{\text{RES}})] \\
&= \dots\dots\dots \text{Defn. of apply}_J \\
&\quad (\langle \exists v_1, x_1, t_1 \mid (v_1, (x_1, t_1)) = \Pi_{\text{RES}} \bullet \\
&\quad \quad \text{ext}_M(t_1 \{x_1 \leftarrow \text{text}_M(v_1)\}_M, \text{call}_M) \\
&\quad) \\
&\quad \langle \Pi_{\text{RES}} \in \text{Value} \times \text{ProgramText} \rangle \\
&\quad \text{chaos} \\
&\quad)[(v, \langle x, \text{decl } M \text{ in } pt \rangle / \Pi_{\text{RES}})] \\
&= \dots\dots\dots \text{Law C.18 extract id.} \\
&\quad (\langle \exists v_1, x_1, t_1 \mid (v_1, (x_1, t_1)) = \Pi_{\text{RES}} \bullet \\
&\quad \quad \llbracket t_1 \{x_1 \leftarrow \text{text}_M(v_1)\}_M \rrbracket \\
&\quad) \\
&\quad \langle \Pi_{\text{RES}} \in \text{Value} \times \text{ProgramText} \rangle \\
&\quad \text{chaos} \\
&\quad)[(v, \langle x, \text{decl } M \text{ in } pt \rangle / \Pi_{\text{RES}})] \\
&= \dots\dots\dots \text{--[-/-] distributes} \\
&\quad (\exists v_1, x_1, t_1 \mid (v_1, (x_1, t_1)) = (v, \langle x, \text{decl } M \text{ in } pt \rangle) \bullet \quad \text{through } (- \triangleleft - \triangleright -) \\
&\quad \quad \llbracket t_1 \{x_1 \leftarrow \text{text}_M(v_1)\}_M \rrbracket \\
&\quad) \\
&\quad \triangleleft (v, \langle x, \text{decl } M \text{ in } pt \rangle \in \text{Value} \times \text{ProgramText} \triangleright \\
&\quad \text{chaos}[(v, \langle x, \text{decl } M \text{ in } pt \rangle / \Pi_{\text{RES}})] \\
&= \dots\dots\dots \text{Select true branch} \\
&\quad \exists v_1, x_1, t_1 \mid (v_1, (x_1, t_1)) = (v, \langle x, \text{decl } M \text{ in } pt \rangle) \bullet \\
&\quad \quad \llbracket t_1 \{x_1 \leftarrow \text{text}_M(v_1)\}_M \rrbracket \\
&= \dots\dots\dots \text{One point rule} \\
&\quad \llbracket (\text{decl } M \text{ in } pt) \{x \leftarrow \text{text}_M(v)\}_M \rrbracket \\
&= \dots\dots\dots \text{Defn. of } -\{_\}_M \\
&\quad \text{decl } \llbracket M \{x \leftarrow \text{text}_M(v)\}_M \rrbracket \text{ in } \llbracket pt \{x \leftarrow \text{text}_M(v)\}_M \rrbracket \\
&\square
\end{aligned}$$

4.5.2 Other operations

The $\mathcal{U}_M$ -design inherits most of its functionality from the $\mathcal{U}_R$ -design; therefore, its proofs of the other operations are almost identical to those of the $\mathcal{U}_R$ -design. The main issue in updating these proofs would be to correctly pass through (but not change) the alphabet variables representing a procedure's parameters. The only changes to such variables occur during the procedure invocation process, which we have already presented (in [Section 4.5.1](#)).

4.6 Factorial Example

We now present a worked example of one iteration of a factorial calculation in both the $\mathcal{O}$ -calculus and its associated $\mathcal{U}_S$ -design. First, we provide our $\mathcal{O}$ -calculus definition of the mathematical factorial operation, as the only member of the following *sFact* object.

$$sFact \equiv [\text{fact} = \varsigma(x) \lambda(i) \text{ if } i < 2 \text{ then } 1 \text{ else } i \times x.\text{fact}(i - 1)]$$

Note that the **fact** method returns a function that takes any integer parameter i and calculates the factorial of i , if i is a natural number; if i is not a natural number this function returns the value 1.

In the remainder of this section we introduce some extra notation for clarifying the presentation of the examples, we illustrate one iteration of the factorial function using the $\mathcal{O}$ -calculus's

deterministic reduction rules (in four steps), and we demonstrate that the $\mathcal{U}_s$ -design semantics is consistent for each step of the reduction.

4.6.1 Notation

In order to simplify the presentation of this example we introduce some extra notation. First, we introduce a named variant of the operational calculus's single-step evaluation arrow, which is denoted by $_ \xrightarrow{\text{R-NAM}} _$, where R-NAM is the name of the rule being used to perform the evaluation step. Note that in general, this might be a sequence of nested rules (which correspond to the rules used to evaluate an argument). When this situation arises we use a comma-separated list of rule names, starting at the outer-most rule.

Second, we introduce a new syntax for the `cmdExps` command for those cases when we want to apply either an infix relation (e.g. $_ < _$) or an infix operation (e.g. $_ \times _$) to literal values.

$$P \text{ rel}_s^{\text{CE}} Q \equiv \text{cmdExp}_s(\text{trans}_s(_ \text{ rel } _), (P, Q))$$

$$P \text{ op}_s^{\text{CE}} Q \equiv \text{cmdExp}_s(\text{trans}_s(_ \text{ op } _), (P, Q))$$

This notation is then generalised for any binary function fn whose arguments are either a literal value or program-text.

$$fn_s^{\text{CE}}(P, Q) \equiv \text{cmdExp}_s(\text{trans}_s(fn), (P, Q))$$

Third (and last), we introduce a notation for providing a staged compilation of an $\mathcal{O}$ -calculus program into its $\mathcal{U}_s$ -design counterpart, via providing parametrised macro definitions of the form $\text{MACRO_NAME}(x_{i=1}^k) \equiv P[x_{i=1}^k]$, where the arguments associated with the parameters $x_{i=1}^k$ can replace the occurrences of these parameters within the compiled predicate P .

4.6.2 Object calculus iteration

The factorial of a number n can be calculated by evaluating the $sFact.\text{fact}(n)$ $\mathcal{O}$ -calculus program. Further, in order to perform at least one iteration (which in this case is a recursive call) of the `fact` method, n must be greater than 1. Therefore, let us assume that n is a literal integer value that is greater than 1, in the following $\mathcal{O}$ -calculus evaluation of one iteration of the factorial calculation.

$$\begin{aligned} \Gamma \bullet sFact.\text{fact}(n) &\xrightarrow{\text{INV-M}} \Gamma \bullet (\lambda(i) \text{ if } i < 2 \text{ then } 1 \text{ else } i \times sFact.\text{fact}(i - 1))(n) \\ &\xrightarrow{\text{INV-F}} \Gamma \bullet \text{ if } n < 2 \text{ then } 1 \text{ else } n \times sFact.\text{fact}(n - 1) \\ &\xrightarrow{\text{IF,LITOP}} \Gamma \bullet \text{ if false then } 1 \text{ else } n \times sFact.\text{fact}(n - 1) \\ &\xrightarrow{\text{IFF}} \Gamma \bullet n \times sFact.\text{fact}(n - 1) \end{aligned}$$

It is useful to name each of these steps, as this enables a clearer presentation of the correspondence between this operation evaluation, and the UTP denotational simplification.

$$sFactFn \equiv \lambda(i) \text{ if } i < 2 \text{ then } 1 \text{ else } i \times sFact.\text{fact}(i - 1)$$

$$sFactCond1 \equiv \text{ if } n < 2 \text{ then } 1 \text{ else } n \times sFact.\text{fact}(n - 1)$$

$$sFactCond2 \equiv \text{ if false then } 1 \text{ else } n \times sFact.\text{fact}(n - 1)$$

$$sFactIter \equiv n \times sFact.\text{fact}(n - 1)$$

Note that only the function value (rather than) its call is named in the first step; therefore the $\mathcal{O}$ -calculus program at the end of the first step is denoted by $sFactFn(n)$.

4.6.3 UTP operand-stack design iteration

We now demonstrate that each operational reduction step is consistent with the $\mathcal{U}_s$ -design's semantics. First, we introduce our staged compilation scheme, then we present a correspondence lemma for each of the four steps; these lemmas are proved in [Appendix C.3.2](#).

Compilation: We now present a staged compilation of the five named $\mathcal{O}$ -calculus programs and three additional programs which invoke either an object's method, a function, or both.

$$\begin{aligned}
\llbracket sFact \rrbracket_s &= O & \llbracket sFact.fact \rrbracket_s &= OF(x) \\
\llbracket sFactFn \rrbracket_s &= F(O) & \llbracket sFactFn(n) \rrbracket_s &= SF(F(O), n) \\
\llbracket sFactCond1 \rrbracket_s &= E1(O, n) & \llbracket sFact.fact(n) \rrbracket_s &= SF(OF(O), n) \\
\llbracket sFactCond2 \rrbracket_s &= E2(O, n, \text{false}) \\
\llbracket sFactIter \rrbracket_s &= ITER(O, n)
\end{aligned}$$

where

$$\begin{aligned}
O &\equiv \{\text{fact} \mapsto \langle x, F(x) \rangle\} & SW(x, i) &\equiv \text{switch}_s(\mathcal{E}_s 1, ITER(x, i)) \\
F(x) &\equiv \mathcal{E}_s \langle i, E1(x, i) \rangle & ITER(x, i) &\equiv \mathcal{E}_s i \times_s^{\text{CE}} SF(OF(x), i) \\
E1(x, i) &\equiv (\mathcal{E}_s i <_s^{\text{CE}} \mathcal{E}_s 2) \ ; \ & SF(f, i) &\equiv \text{vfPair}_s^{\text{CE}}(f, \mathcal{E}_s i) \ ; \\
&\quad SW(x, i) & &\quad \text{call}_s \\
E2(x, i, b) &\equiv \mathcal{E}_s b \ ; \ & OF(x) &\equiv \text{omPair}_s^{\text{CE}}(\mathcal{E}_s x, \mathcal{E}_s \text{fact}) \ ; \\
&\quad SW(x, i) & &\quad \text{call}_s
\end{aligned}$$

Step 1 – Method invocation to function invocation:

Lemma 4.11 – Factorial Example Step 1:

$$\llbracket sFact.\text{fact}(n) \rrbracket_s = \llbracket sFactFn(n) \rrbracket_s$$

□

Step 2 – Function invocation to conditional expression:

Lemma 4.12 – Factorial Example Step 2:

$$\llbracket sFactFn(n) \rrbracket_s = \llbracket sFactCond1 \rrbracket_s$$

□

Step 3 – Conditional expression decision evaluation:

Lemma 4.13 – Factorial Example Step 3:

$$\llbracket sFactCond1 \rrbracket_s = \llbracket sFactCond2 \rrbracket_s$$

□

Step 4 – Selection of else conditional expression branch:

Lemma 4.14 – Factorial Example Step 4:

$$\llbracket sFactCond2 \rrbracket_s = \llbracket sFactIter \rrbracket_s$$

□

4.7 Summary

We have proved that the UTP semantics of the $\mathcal{O}$ -calculus is consistent with that of its operational semantics. Here, the UTP semantics of an $\mathcal{O}$ -calculus program is equal to the UTP semantics of the same $\mathcal{O}$ -calculus program after any number of operational execution steps. The only additional requirement of a faithful encoding is that the semantics of the resultant values are the same, as it might be possible to construct a consistent mapping that changes the semantics. Here, our compilation schemes straightforwardly map both literal and object values on to their corresponding UTP values; and the function values and methods are compiled into program texts that we have already proved to be consistent. Therefore, the UTP models of our $\mathcal{O}$ -calculus, in [Chapter 3](#), are faithful.

The main contribution of this chapter to our overall thesis is to provide further evidence that the contributions claimed in the previous chapter are well founded.

Chapter 5

Unifying Theories of Locations

This chapter presents an abstract model of locations that can be used to provide a fully abstract heap for a Unifying Theories of Programming (UTP) representation of our $\mathcal{O}$ -calculus. It starts by discussing the notion of a location and some of the constraints that we place upon them [Section 5.1](#). Having done this, we present a concrete model of locations ([Section 5.2](#)), which is abstracted ([Section 5.3](#)) and then integrated into the basic relational model of the UTP ([Section 5.4](#)).

5.1 Locations

The UTP model in [\[CHW06\]](#) supports the notion of a compound data structure via the introduction of a record datatype, which essentially maps distinct labels to values. These labels are also used when unambiguously specifying the location of a value and determining whether it is shared.

An object can be modelled in a similar manner to that of a record. For example, in C^{++} and $C^\#$ the object and record types are defined by the `class` and `struct` datatype constructors respectively. Here, a variable of an object type contains a pointer to an object, whereas a variable of a record type contains the record itself. It is this distinction between variables of object and record types that we believe is important to explicitly model in a general theory of UTP pointers. Specifically, the contents of a record are duplicated, whereas the contents of an object are aliased (shared).

The UML class diagram in [Figure 5.1](#) provides a high level overview of our model; it illustrates the relationships between abstract locations and the values contained within them. Here an abstract location is either shareable or containable, but not both. All such locations store precisely one value, which can be a reference, a literal or a compound value. Reference values can contain either the null value or the address of a shareable location. Compound values can contain any number of distinctly named fields, where each field is stored in a separate containable location.

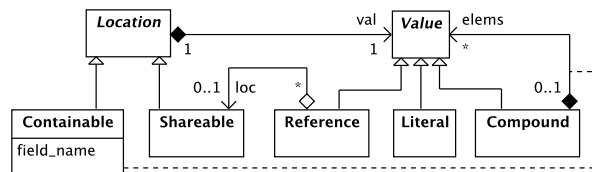


Figure 5.1: Location Model – Class Diagram

5.1.1 Terminology

The notions of a location, pointer, reference and object-identity are used to provide a means for storing or locating a value. Precisely what is meant by each of these terms varies; some typical definitions are:

Location a location in memory. It may be tagged with attributes, such as the type of data it can store, the owner of the data, and whether the data is read-only. For the purposes of the model being presented a location is essentially tagged with whether its contents can be directly addressed by a pointer.

Pointer a value that points to a directly addressable memory location. It may be tagged with attributes in a similar manner to that of a location. Note that a pointer's attributes need not match those of the location, they just ought to be consistent with them.

Reference a pointer whose value cannot be used in, or the result of, pointer arithmetic.

Object identity the identity of an object, which may be its address in memory, a name in a Java Remote Method Invocation (RMI) registry, a name in a CORBA (Common Object Request Broker Architecture) naming service, etc.

5.1.2 Scope

The model of locations we present is not intended to support the concepts of object ownership or reference containment, such as discussed in ownership models [Cla02, NCP99] and separation logics [Rey02]. Nor is this model of locations intended to support low-level pointer operations, such as those operations that create a new pointer by adding an arbitrary offset to an existing pointer (e.g. $ptr = ptr + 2$) or get the address of a record's element (e.g. $ptr = \&(rec.fld)$). Having said this, it is straightforward to write C++ and C# programs that do not directly use such low-level pointer operations and this ought to be syntactically checkable. For example, in C# this could be achieved by banning the use of the `unsafe` keyword.

5.1.3 Family tree

Within this chapter we use instantiations of data-types associated with a family tree, as illustrated by the class diagram in Figure 5.2, to provide examples of abstract data-structures. Essentially a family tree contains a relationship between a set of people, where each person has a name, a gender and a set of relations; a person may also have a record of dates associated with them to record either or both of their birth and death dates. Further, each entry in a person's relations set contains the type of relationship (as indicated by the 'rel' entry), the set of people that this is with, and a link back to itself (i.e. the containing person entry). We give a

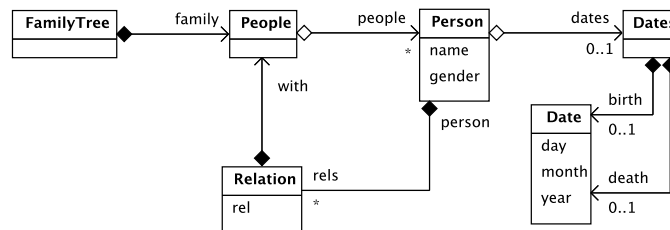


Figure 5.2: Family Tree Example – Class Diagram

specific interpretation to the aggregations within Figure 5.2 to link them to the concepts within our abstract location model (as previously illustrated in Figure 5.1). Here, shareable (hollow

diamond) and composite (solid diamond) aggregations are used to distinguish between shareable and containable locations, respectively; aggregations that have any number of instances are represented by lists.

5.2 Concrete Representation

5.2.1 Concrete value notation

The two types of *literal value* used within this chapter are the integers (e.g. -32) and the strings (e.g. "Some text"). There is also a special *unset* literal constant, denoted by $\mathbf{i}$; this is used to represent the contents of a freshly created location, and the value of a missing element.

The two remaining types of concrete value are the *compound* and *reference values*. A compound value is represented by a partial map from field names (which we identify with containable locations) to concrete values. It is denoted by $\{_{i=1}^k nm_i = v_i\}$, where the name nm_i indexes the concrete value v_i . A name is denoted by an alpha-numeric word starting with a letter or the dollar symbol. The name represented purely by a single dollar symbol, which we refer to as the ‘dollar name’, is reserved for denoting a shareable location, and thus cannot be used as a compound value’s field name.

A reference value is either null or an index to a shareable location. Such values are denoted by $\ominus$ and ℓ_i respectively, where two non-null reference values ℓ_i and ℓ_j index distinct shareable locations whenever $i \neq j$.

Figure 5.3 provides both an instance of the family tree’s **Dates** class and a concrete value representation of this instance (object). Here, the **Dates** object explicitly sets only one of its two optional **Date** fields, **birth**, to 12 Aug 1980. The other optional field, **death**, is left unset. This data structure can be drawn as a graph, as illustrated in Figure 5.4, where: a literal

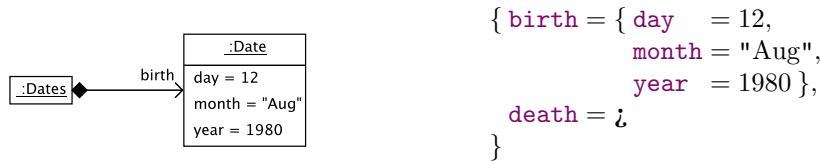


Figure 5.3: Dates – object diagram and concrete representation

value is denoted by a boxed node containing the literal; and a compound value is denoted by a circular node, whose outgoing edges are labelled with its distinct field names. Reference values

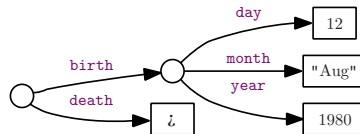


Figure 5.4: Dates – concrete graph representation

are denoted by a diamond node that contains the reference (Figure 5.6 and Figure 5.8).

In addition to defining the concrete representations of the values we use meta-variables for representing each of the different types of value, in accordance with Section 2.1.4. Here i, j and k represent the integers, and s, lv, cv, rv, v , and t represent strings, literal values, compound values, reference values, values and terms respectively. Note that the meta-variable t includes the concrete values, field names, and the yet-to-be terms such as location graphs. These meta-variables are typically used to define functions, as illustrated by the following example, which extracts the shareable locations contained within a term. This example also uses the generalised term notation that was introduced in Section 2.1.3, where $t\{_{i=1}^k t_i\}$ denotes

a term t with k subterms, $t_{i=1}^k$, where a leaf term ($k = 0$) has the empty set of subterms.

$$\begin{aligned} sLocs \ell_i & \hat{=} \{ \ell_i \} \\ sLocs t \{_{i=1}^k t_i \} & \hat{=} \bigcup \{_{i=1}^k sLocs t_i \} \end{aligned}$$

We read such definitions by pattern-matching from top to bottom, accepting the first equation that matches an actual argument. Thus, the order in which the lines of a function are presented may affect its meaning. In this case, swapping the order would produce a function that returns the empty set.

The $sLocs$ function is applied to terms that are yet to be defined, such as the heap value term in Section 5.2.2. Note that this does not require an update to the $sLocs$ function as these terms are already handled by the second definitional line, which can be applied to any term (i.e. a general term).

5.2.2 Concrete location heap

A *location heap* is a partial map from shareable locations to values; it is denoted by $\{_{i:N} \ell_i \mapsto v_i\}$, where N is some finite subset of the natural numbers. For example, the object diagram in Figure 5.5 illustrates that Jane Doe is married to John Doe, where only the instances of the Person class are considered to be shareable. It can be represented by the following concrete

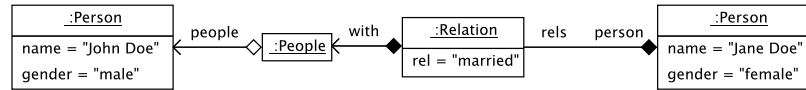


Figure 5.5: Marriage example – object diagram

location heap, where the contents of shared locations ℓ_3 and ℓ_5 contain the John Doe and Jane Doe Person objects respectively.

$$\begin{aligned} \{ \ell_3 \mapsto \{ \text{name} = \text{"John Doe"}, \text{gender} = \text{"male"}, \text{dates} = \emptyset, \text{rels} = \mathcal{I} \}, \\ \ell_5 \mapsto \{ \text{name} = \text{"Jane Doe"}, \text{gender} = \text{"female"}, \text{dates} = \emptyset, \\ \text{rels} = \{ \$1 = \{ \text{rel} = \text{"married"}, \text{person} = \ell_5, \text{with} = \{ \text{people} = \{ \$1 = \ell_3 \} \} \} \} \} \} \end{aligned}$$

Note that the **rels** fields within the Jane Doe Person object contains a singleton list of relations, whose first (and only) entry is labelled with **\$1**. Further, this entry's **with** field contains a list of People objects with a single entry, again labelled with **\$1**, that contains a pointer to the John Doe Person object. In general, the n^{th} element of a list would be labelled with **\$n**.

Figure 5.6 provides the alternative graph representation of the example, where the dashed edges are used to link a reference value to its contents. Note that these edges are labelled with the dollar name (\$) as discussed in Section 5.2.1.

Before moving on to present the concrete location model, we observe that a concrete heap can reference a shareable location that it does not define; i.e. a concrete heap can contain reference values that are not in the domain of the heap's partial map. In order to classify location heaps that do not have this undesirable property, we introduce a healthiness condition, which considers a heap to be healthy whenever all references to shared locations within the graph's values are defined by the graph itself.

$$\text{HC}_H H \hat{=} sLocs H \subseteq \text{dom } H$$

where

H is the meta variable representing a concrete location heap
 $\text{dom } _$ returns the domain of a relation or function

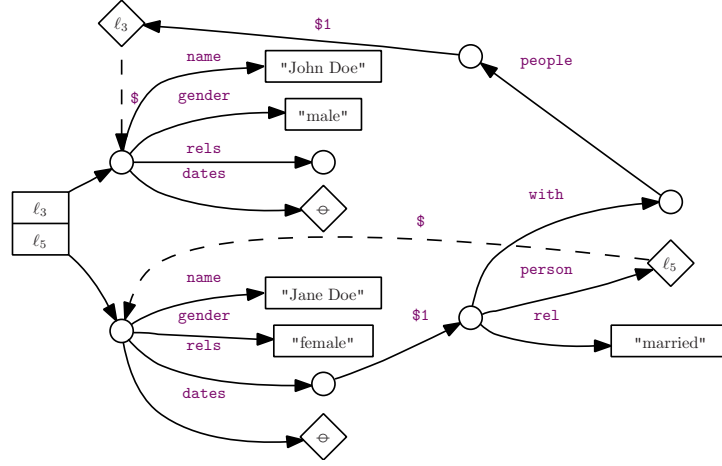


Figure 5.6: Marriage example – concrete heap graph

Further, any heap can be made healthy by adding an entry for each missing shared location and setting that location's value to the unset value ($\mathfrak{!}$), as follows:

$$\text{MH}_H H \hat{=} H \cup \{r \mapsto \mathfrak{!} \mid r \in (sLocs H) \setminus (\text{dom } H)\}$$

Note that a healthy heap is unaffected by the application of MH_H (and vice-versa), because all the shareable locations in the heap are contained within its domain; i.e. $(\text{MH}_H H = H) \Leftrightarrow \text{HC}_H H$.

The notion of equivalence ($\equiv$) between location heaps is more complex than that for concrete values, which is mathematical equality (where the ordering of elements within a set or map is not significant). Here, the equivalence relationship between heaps allows for the renaming of shareable locations. Specifically, two heaps are considered to be equivalent if there exists a bijective map (f) that can be applied to one heap to produce the other.

$$H_1 \equiv H_2 \Leftrightarrow \exists f \bullet H_1 = \text{rename}(H_2, f)$$

where

$$\begin{aligned} \text{rename}(\ell_i, f) &\hat{=} \ell_{f(i)} \\ \text{rename}(t\{_{i=1}^k t_i\}, f) &\hat{=} t\{_{i=1}^k \text{rename}(t_i, f)\} \end{aligned}$$

5.2.3 Concrete location model

Location models extend this notion of the heap by adding a starting point, which is represented by a concrete value. Therefore, a location model is denoted by a value-heap pair (v, H) . Here, the idea is that the value v represents the root of a computational unit, such as a program, whose elements can share data via the shareable locations in the heap H .

Like the location heap that preceded it, location models have a healthiness condition which ensures that the heap is valid; that is, all the shareable locations referenced within a model are defined by the heap.

$$\text{HC1}_L(v, H) \hat{=} sLocs(v, H) = \text{dom } H$$

A location model can be made HC1_L -valid in a similar manner to a heap.

$$\text{MH1}_L(v, H) \hat{=} (v, H \cup \{r \mapsto \mathfrak{!} \mid r \in sLocs(v, H) \setminus (\text{dom } H)\})$$

Locations in the model are considered to be *reachable* if they are either contained within the starting value v or indirectly contained within the contents of v 's reference values. For HC1_L -healthy models, this can be formalised by the following functions, where R and R' represent the shareable locations that have already been taken into account and are contained within a value respectively.

$$\begin{aligned} \text{reachable}(v, H) &\hat{=} \text{reachValue}(v, H, \emptyset) \\ \text{reachValue}(v, H, R) &\hat{=} sLocs\ v \cup \text{reachDeref}(sLocs\ v, H, R) \\ \text{reachDeref}(R', H, R) &\hat{=} \bigcup \{ \text{reachValue}(H\ r, H, R' \cup R) \mid r \in R' \setminus R \} \end{aligned}$$

The following normal-form healthiness condition ensures that there is no unreachable information within the model; i.e. every shareable location that is defined by a model's heap is reachable.

$$\text{HC2}_L(v, H) \hat{=} \text{dom } H = \text{reachable}(\text{MH1}_L(v, H))$$

Note that we apply the MH1_L healthiness constructor prior to performing the reachability calculation, in order to ensure that HC2_L calculation is defined. If a location model is not in normal form (i.e. HC2_L -healthy), it can be made so by ensuring that it is HC1_L -healthy and then removing all the unreachable locations.

$$\text{MH2}_L(v, H) \hat{=} \text{let } (v_1, H_1) \hat{=} \text{MH1}_L(v, H) \text{ in } (v_1, \{rv \mid rv \in H_1 \wedge (\text{first } rv) \in \text{reachable}(v_1, H_1)\})$$

$$\text{where } \text{first}(x, y) \hat{=} x$$

We are now in a position to define an equivalence relation over location models. It is similar to that of heaps, except that we first ensure that models are made healthy before performing the check, as we only want to consider reachable elements in a model's heap. In other words, two location models are equivalent iff there exists some bijective shareable-location-renaming function f that enables two normalised heaps to be made equal.

$$\begin{aligned} (v_1, H_1) &\equiv (v_2, H_2) \\ \Leftrightarrow \\ \exists f \bullet \text{MH2}_L(v_1, H_1) &= \text{rename}(\text{MH2}_L(v_2, H_2), f) \end{aligned}$$

It is this notion of equivalence up to which our UTP model of locations is fully abstract, as described in [Section 5.4.3](#).

The family tree example can now be extended to illustrate a concrete location model, by adding an object to represent the family tree, as illustrated in [Figure 5.7](#). The concrete graph

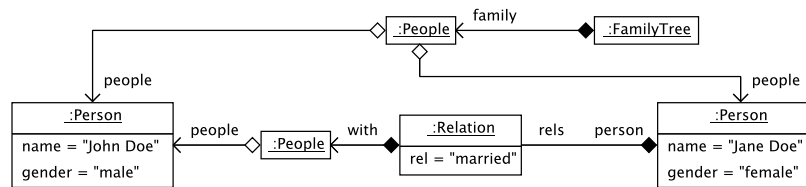


Figure 5.7: Family tree example – object diagram

representation of this example is provided by [Figure 5.8](#), where the explicit visualisation of the heap has been removed, as it is no longer required for representing the shareable locations. Such locations are now represented by the dashed edges within the graph, which are now guaranteed to exist due to the reachability healthiness condition.

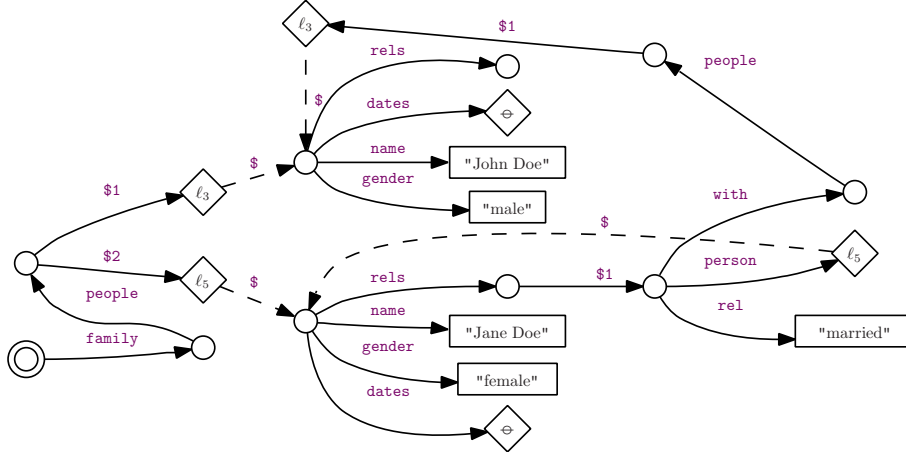


Figure 5.8: Family tree example – concrete model graph

5.2.4 Paths and their operations

A *compound value path* describes a route from a compound value to one of its elements, via a non-empty dot-separated sequence of field names (labels). Such labelled paths are essentially used to describe routes to contained locations, which we can access and update by using the following functions

$$\begin{aligned}
 *cv.nm & \hat{=} cv \ nm \\
 cv.nm := v & \hat{=} cv \oplus \{nm \mapsto v\} \\
 cv.nm.lp := v & \hat{=} cv \oplus \{nm \mapsto (*cv.nm).lp := v\}
 \end{aligned}$$

where lp is the meta-variable for labelled paths, $(- _)$ is the function or map application operation, and $(- \oplus -)$ is the function override operation. This notion of a path is extended to define location model update and access functions as follows:

$$\begin{aligned}
 *(v, L) & \hat{=} v \\
 *(v, L).lp & \hat{=} *v.lp \\
 *(v, L).\ell_i & \hat{=} L \ \ell_i \\
 *(v, L).\ell_i.lp & \hat{=} (L \ \ell_i).lp \\
 (v, L) := v' & \hat{=} (v', L) \\
 (v, L).lp := v' & \hat{=} (v, lp := v', L) \\
 (v, L).\ell_i := v' & \hat{=} (v, L \oplus \{\ell_i \mapsto v'\}) \\
 (v, L).\ell_i.lp := v' & \hat{=} (v, L \oplus \{\ell_i \mapsto (L \ \ell_i).lp := v'\})
 \end{aligned}$$

Further, it is possible to extend this notion to copy a value from one location to another, as follows

$$(v, L).lop := (v, L).llp' \hat{=} (v, L).lop := *(v, L).llp'$$

where: lop denotes either a labelled path (lp), a shareable location index (ℓ_i), or a shareable location index followed by a path ($\ell_i.lp$); llp denotes a lop or the empty path (which is itself denoted by $\emptyset$); and $(v, L).\emptyset$ denotes (v, L) .

It is also straightforward to define other operations, such as for deleting elements from compound values and the heap.

5.3 Abstract Model

In [Section 5.2.3](#), graphs represented healthy concrete location models, where:

- the solid and dashed edges denote distinct compound and potentially shared shareable locations, respectively;
- the rectangular, circular, and diamond nodes denote literal, compound, and reference values, respectively.

This section presents: a brief overview of the trace-based graph abstraction; some utility operations for manipulating traces; a model of nodes as a set of traces; and an overview of the trace-based location graph model.

5.3.1 Graph abstraction

We can determine the value of dereferencing a reference node of a concrete location graph by following that node's outbound edge (as shown in Figure 5.8). That is, the shareable location index contained within a reference node is not required. Thus, this unused data can, and will, be ignored in our abstraction.

We observe that the outbound edges of each node within a healthy concrete location graph have distinct labels. Therefore, we can use a finite non-empty sequence of names to unambiguously define a path from a graph's root node to any other node. Such a path is from now on referred to as an *absolute path*.

The location of a node within a concrete location graph can be modelled by the set of all absolute paths to that node, which we from now on refer to as an *absolute path-set*. Hence, one way of providing a UTP model of locations would be as a partial map from such an absolute path-set to an appropriate abstraction of the data directly associated with its corresponding node. For example, the data associated with:

- a literal or null-reference node could be modelled by its concrete value;
- a compound node could be modelled by its set of outbound edge labels;
- a non-null reference node could be modelled by its outbound edge label.

Such a model of locations is similar to that presented in [CHW06], which uses the idea of an entity group to model shared locations. Here, each group contains the set (equivalence class) of fully qualified variables that share the same location.

Another approach is to change the notion of an absolute path-set, from representing the location of a node to representing both the location and contents of a node. To avoid confusion, we refer to such paths as *traces*. Here, the idea is that the last value in a trace represents its content, and the front of the trace its location. In other words, a trace is a path lp followed by a trace label l , where l represents either a name, a literal value, or the null-reference value; it is denoted by $lp.l$. This is the basis of our UTP model of locations, which we present in Section 5.4. Such a model of locations is similar to that presented in [HH99], which uses trace-sets to model both locations and values. Here, the main difference is in our introduction of a containable location and its effect on assignment (which [HH99] refers to as 'pointer swing'). Specifically, within our model the contents of contained locations need to be duplicated, whereas the contents of shared locations are referenced.

5.3.2 Traces

As previously stated, a trace is denoted by $lp.l$, where lp is a path and l is a trace label (i.e. a name, a literal, a null-reference). One consequence of this is that it is only possible to concatenate two traces (denoted by $tr_1.tr_2$) when the last label in the first trace tr_1 is a name, as only names are allowed within a path.

The remainder of this section defines some utility operations on traces and trace-sets, that are used in the construction of our abstract model of locations. First we introduce two

operations **front** and **last** for extracting the location and content components of a non-empty trace.

$$\text{front } lp.l \hat{=} lp \quad \text{last } lp.l \hat{=} l$$

The front operation can be used to generate the set of locations visited by a trace, as characterised by their paths, where each path within this set is considered to be a *prefix* of the original trace. Such a set of paths is referred to as the *proper prefixes* of the given trace. The function prefixes_T defines the non-proper version of the prefix set.

$$\begin{aligned} \text{prefixes}_T \emptyset &\hat{=} \{\emptyset\} \\ \text{prefixes}_T tr &\hat{=} \{tr\} \cup \text{prefixes}_T(\text{front } tr) \end{aligned}$$

The prefixes also provide a natural ordering over traces.

$$tr_1 <_T tr_2 \hat{=} \text{prefixes}_T(tr_1) \subset \text{prefixes}_T(tr_2)$$

Having defined an ordering over traces, it is now possible to use that ordering to define a subtraction operation. This is eventually used to define the relative paths between nodes in a set.

$$(- -_T -) \hat{=} \lambda tr_1, tr_2 \mid tr_2 \leq_T tr_1 \bullet \text{pick}\{tr \mid tr_1 = tr_2.tr\}$$

Here, the **pick** function picks the singleton element from a set (i.e. $\text{pick}\{x\} \hat{=} x$).

Before leaving the trace utilities, we lift the definitions of the **front**, **last**, and prefixes_T operations to trace-sets. The first two are lifted by applying their definitions to each non-empty trace with the set. The latter one is lifted to a trace-set (denoted by TR) by applying the prefixes operation to each trace within the set and merging the results.

$$\begin{aligned} \text{frontTraces } TR &\hat{=} \{\text{front } tr \mid tr \in (TR \setminus \{\emptyset\})\} \\ \text{lastLabels } TR &\hat{=} \{\text{last } tr \mid tr \in (TR \setminus \{\emptyset\})\} \\ \text{prefixes}_N TR &\hat{=} \bigcup \{\text{prefixes}_T tr \mid tr \in TR\} \end{aligned}$$

5.3.3 Trace-based node

A trace-based graph node is modelled by a set of traces that satisfies two healthiness conditions. Both of these conditions follow from the observation that the only way a concrete location graph node may have more than one incoming edge, is if all these edges are labelled with the dollar name. Consequently, every trace to a node is guaranteed to end with the same label, except for the root node which has no label. This is modelled by the first healthiness condition, which states that all incoming edges to a node (trace-set n) have the same label.

$$\text{HC1}_N(n) \hat{=} \# \text{lastLabels}(n) \leq 1$$

Another consequence of the observation is that a node may only have multiple parents if it is stored in a shareable location. This is modelled by the second healthiness condition, which states that the trace to any node that has more than one parent must end with the special shareable location label.

$$\text{HC2}_N(n) \hat{=} \# \text{lastLabels}(\text{frontTraces}(n)) > 1 \Rightarrow \text{lastLabels}(n) = \{\$ \}$$

Any healthy node can be denoted by $Lp.l$, where each path in the path-set Lp is extended by the trace-label l to form the trace-set $\{lp.l \mid lp \in Lp\}$. The remainder of this section now presents some useful utility relations and operations on nodes.

Node relations: The *child-of* and *descendant-of* relations test whether one node is an immediate child of or a descendant of another node. These tests assume that the nodes come from a healthy graph, where all the routes to the parent are contained within the child.

$$\begin{aligned} n_1 \text{ childOf } n_2 & \hat{=} n_2 \subseteq \text{frontTraces}(n_1) \\ n_1 \text{ descendantOf } n_2 & \hat{=} n_2 \subseteq \text{prefixes}_N(\text{frontTraces}(n_1)) \end{aligned}$$

In addition to knowing whether two nodes are related, it is sometimes useful to identify the relative traces from a parent to child node.

$$\text{traces}_N(n_1, n_2) \hat{=} \{tr_2 -_T tr_1 \mid tr_1 \in n_1 \wedge tr_2 \in n_2 \wedge tr_1 \leq_T tr_2\}$$

Such trace-sets are used to determine whether two nodes are related via shareable or via containable locations. Here, two nodes are related by a shareable location if one of the traces within the trace-set includes the dollar label. Similarly they are related by a containable location if one of the traces within the trace-set does not contain the dollar label.

$$\begin{aligned} n_1 \text{ shareDescOf } n_2 & \hat{=} n_1 \text{ descendantOf } n_2 \wedge \$ \in_T \text{traces}_N(n_2, n_1) \\ n_1 \text{ containDescOf } n_2 & \hat{=} n_1 \text{ descendantOf } n_2 \wedge \$ \notin_T \text{traces}_N(n_2, n_1) \end{aligned}$$

where $l \in_T tr \hat{=} \exists lp \bullet lp.l \in \text{prefixes}_T tr$

Note that the only way a node can be both a shareable and a containable descendant of another node, is if the nodes are both contained in the same cycle. In this case, all the containable descendants are also shareable descendants.

Node unlinking (deletion): Part of the assignment process involves the removal of previously held data. This is the purpose of the following unlinking operations, which remove all traces of either a node (n_1) or its children from the specified target node (n_2).

$$\begin{aligned} \text{unlink}_N n_1 \text{ from } n_2 & \hat{=} n_2 \setminus \{tr_2 \mid tr_1 \in n_1 \wedge tr_2 \in n_2 \wedge tr_1 \leq_T tr_2\} \\ \text{unlinkChildren}_N n_1 \text{ from } n_2 & \hat{=} n_2 \setminus \{tr_2 \mid tr_1 \in n_1 \wedge tr_2 \in n_2 \wedge tr_1 <_T tr_2\} \end{aligned}$$

These operations can be lifted to the graph context by unlinking a given node from a node-set.

$$\begin{aligned} \text{unlink}_G(N, n) & \hat{=} \{(\text{unlink}_N n \text{ from } n') \mid n' \in N\} \\ \text{unlinkChildren}_G(N, n) & \hat{=} \{(\text{unlinkChildren}_N n \text{ from } n') \mid n' \in N\} \end{aligned}$$

Node duplication (replacement): It is sometimes useful to construct a new node from a pair of existing nodes, a source node (n_2) and one of its descendants (n_3). Here the idea is to extract the traces between the source and descendant nodes, and then append them to a new source node (n_1), which is the target of the duplication.

$$\text{replace}_N n_1 \text{ for } n_2 \text{ in } n_3 \hat{=} \{tr_1.tr \mid tr_1 \in n_1 \wedge tr \in \text{traces}_N(n_2, n_3)\}$$

For example, consider the two graphs presented in Figure 5.9. Here, the traces between the source and destination node in the initial graph is the singleton set $\{\text{month."Aug"}\}$. The traces within this set are appended to each trace in the new source node (n_3) to produce the trace-set $\{\text{death.month."Aug"}\}$. This trace-set could be considered to define a new node within the graph, as illustrated in the updated part of Figure 5.9. Having said this, care needs to be taken to ensure that any formal graph duplication operation also appropriately copies the intermediate nodes (see later discussion on assignment in Section 5.3.4).

Instead of replacing one parent for another, we may want to add a parent; for example, when copying a reference to a shareable location. This is essentially achieved by performing the replacement operation and merging in the original data.

$$\text{add}_N n_1 \text{ to } n_2 \text{ in } n_3 \hat{=} (\text{replace}_N n_1 \text{ for } n_2 \text{ in } n_3) \cup n_3$$

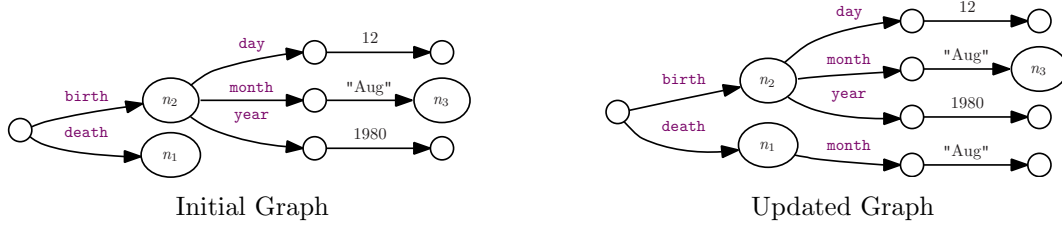


Figure 5.9: Simple node duplication

These operations can then be lifted so that they operate on node-sets, by reparenting each node in the set.

$$\begin{aligned} \text{replace}_G n_1 \text{ for } n_2 \text{ in } N &\hat{=} \{(\text{replace}_N n_1 \text{ for } n_2 \text{ in } n) \mid n \in N\} \\ \text{add}_G n_1 \text{ to } n_2 \text{ in } N &\hat{=} \{(\text{add}_N n_1 \text{ to } n_2 \text{ in } n) \mid n \in N\} \end{aligned}$$

We use these operations to prepare a subgraph for being moved or copied to a new location.

5.3.4 Trace-based graph

A *trace-based graph* is a set of trace-based nodes that satisfies four healthiness conditions. The first healthiness condition states that each of the graph's nodes is healthy.

$$\text{HC1}_G(G) \hat{=} \forall n \mid n \in G \bullet \text{HC1}_N(n) \wedge \text{HC2}_N(n)$$

The second healthiness condition states that the nodes of a graph are disjoint. This ensures that an absolute trace can be used to identify a single node.

$$\text{HC2}_G(G) \hat{=} \forall n_1, n_2 \mid \{n_1, n_2\} \subseteq G \wedge n_1 \neq n_2 \bullet n_1 \cap n_2 = \emptyset$$

For a graph (G) that satisfies condition HC2_G , it is possible to define an operation for extracting the node (n) that has an absolute trace (lp), so long as the trace is within the graph.

$$\text{node}_G(G, tr) \hat{=} \lambda G, tr \mid \text{HC2}_G(G) \wedge tr \in_G G \bullet \text{pick}(\{n \mid n \in G \wedge tr \in n\})$$

Here, the $(- \in_G -)$ relation determines whether a trace is in the graph:

$$lp \in_G G \hat{=} lp \in (\bigcup G)$$

The third healthiness condition states that each of a node's traces is consistently extended; i.e. if it is possible to take an edge with label l from node n_1 to node n_2 , then the trace-set formed by appending the label l to each of n_1 's traces is a subset of n_2 's trace-set.

$$\begin{aligned} \text{HC3}_G(G) \hat{=} \forall n_1, n_2, tr_1, tr_2, l \mid \\ \{n_1, n_2\} \subseteq G \wedge \{tr_1, tr_2\} \in n_1 \bullet \\ tr_1.l \in n_2 \Rightarrow tr_2.l \in n_2 \end{aligned}$$

The fourth healthiness condition states that the parents of a node are contained within the graph; in other words, the traces within a graph are prefix closed.

$$\text{HC4}_G(G) \hat{=} \forall tr, l \mid tr.l \in_G G \bullet tr \in_G G$$

The combination of the first three graph healthiness conditions defines what it means for the trace model to have a consistent, but not necessarily complete, set of nodes. Thus, these conditions should be satisfied by any healthy subgraph.

The remainder of this section provides operations for manipulating the contents of a location graph model, such as operations for: extracting a subgraph; extracting the value at a location; and assigning a value to a location.

Children and descendants subgraphs: Subgraphs can be formed by selecting only some of a graph's nodes. The `childOf` and `descendantOf` relations can be used to filter a graph to form children and descendants subgraphs respectively.

$$\begin{aligned} \text{children} & \hat{=} \lambda G, n \mid n \in G \bullet \{n' \mid n' \in G \wedge n' \text{ childOf } n\} \\ \text{descendants} & \hat{=} \lambda G, n \mid n \in G \bullet \{n' \mid n' \in G \wedge n' \text{ descendantOf } n\} \end{aligned}$$

Note that a node can be a descendant of itself if, and only if, there is a non-empty sequence of edges back to itself.

Dereferencing a location's value: A location is represented by a healthy node whose last label is either a field name or the dollar name. Such a node can be represented by a path-set, as each trace within this node may only contain names. The value of a location node is determined by recursively examining its children, or more specifically its child labels. There are three cases to consider.

1. There is a single null-reference or literal value (nlv) child label. In this case, the label value is returned as the location's value.
2. There is a single child label that contains the dollar name. In this case the path-set (reference value) that models the child node is returned.
3. There is a set of child-labels that contain field names. In this case a compound value is recursively constructed from its children.

$$\begin{aligned} *_G(G, lp) & \hat{=} *_G(G, \text{node}_G(G, lp)) \\ *_G(G, Lp) & \hat{=} *_G(G, \text{children}(G, Lp)) \\ *_G(G, \{Lp.nlv\}) & \hat{=} nlv \\ *_G(G, \{Lp.\$ \}) & \hat{=} Lp.\$ \\ *_G(G, \{_{i=1}^k Lp.nm_i\}) & \hat{=} \{_{i=1}^k nm_i = *_G(G, Lp.nm_i)\} \end{aligned}$$

Recall that we introduced $Lp.l$ as an alternative notation for denoting a healthy node, in [Section 5.3.3](#), where $Lp.l \hat{=} \{lp.l \mid lp \in Lp\}$.

Preparing a location for assignment: The preparation required for assigning a value to a location depends on a number of factors, such as whether the location already exists. We could limit assignments to existing locations, but then this would not mirror our concrete model, which defined assignment in terms of the map overriding operation ($- \oplus -$). Instead we categorise a potential location as either *existing* (E_{pm}), *freshly containable* (C_{pm}), *freshly shareable* (S_{pm}), or *invalid* (U_{pm}), as follows

$$\text{prepMode}(G, Lp.nm) \hat{=} \begin{cases} E_{pm}, & \text{if } Lp.nm \in G \\ S_{pm}, & \text{if } Lp.nm \notin G \wedge nm = \$ \\ C_{pm}, & \text{if } Lp.nm \notin G \wedge nm \neq \$ \wedge *_G(G, Lp) \in CV \\ U_{pm}, & \text{otherwise} \end{cases}$$

where CV denotes the set of compound values (i.e. the compound value type). Note that the above definition of freshly created locations ensures that a compound value may only contain containable locations (and vice versa). In general, a path-set Lp is considered to represent an assignable location within a graph G whenever it has a valid assignable location mode.

$$Lp \text{ assignableIn}_G G \hat{=} \text{prepMode}(G, Lp) \neq U_{pm}$$

It is now possible to define the preparation for an assignable location by ensuring that it exists and contains no contents. This can involve the clearing (unlinking) of an existing node's contents and the creation of a new location node.

$$\begin{array}{lll}
 \text{prep}_G(G, Lp) & \widehat{=} & \text{prep}(G, Lp, \text{prepMode}(G, Lp)) \\
 \text{prep}(G, Lp, E_{pm}) & \widehat{=} & \text{unlinkChildren}_G(G, Lp) \\
 \text{prep}(G, Lp.l, S_{pm}) & \widehat{=} & \text{unlinkChildren}_G(G, Lp) \cup \{Lp.\$\} \\
 \text{prep}(G, Lp.l, C_{pm}) & \widehat{=} & G \cup \{Lp.l\}
 \end{array}$$

Assigning a null-reference or literal value: A null-reference or literal value (nlv) can be assigned to a graph location by preparing the location and setting its contents to the given value.

$$\begin{array}{lll}
 (G, lp) :=_G nlv & \widehat{=} & (G, \text{node}_G(G, lp)) :=_G nlv \\
 (G, Lp) :=_G nlv & \widehat{=} & \text{prep}_G(G, Lp) \cup \{Lp.nlv\}
 \end{array}$$

Assigning an encapsulated compound value: An encapsulated compound value is a concrete compound value that contains no shareable locations (i.e. a compound value in the set $\{cv \mid sLocs\ cv = \emptyset\}$). Such values are represented by the meta-variable ecv . It can be assigned to a location by preparing the location and setting its contents to the subtree that represents the compound value.

$$\begin{array}{lll}
 (G, lp) :=_G ecv & \widehat{=} & (G, \text{node}_G(G, lp)) :=_G ecv \\
 (G, Lp) :=_G ecv & \widehat{=} & \text{prep}_G(G, Lp) \cup \bigcup \{Lp.tr \mid tr \in (cvTrs\ ecv)\}
 \end{array}$$

Here, the $cvTrs$ function converts an encapsulated compound value into a prefix closed set of traces, representing each trace through the compound value's structure.

Assigning the contents of an existing location: In the concrete model, we referred to this as the copying of a location's value. This is more tricky than the previous cases for a number of reasons. One significant reason is that the location we are copying may be contained within the target location that we are assigning to. In such a case, the location preparation process could remove (clear) the location we want to copy. This limitation can be overcome by a three-step process. First, copy the value to a fresh temporary location, which is not contained within the contents of the target location. Second, prepare the target location and copy the value of the temporary location to it. Last, remove the temporary location.

What would make a good temporary location is dependant on what the location graph is being used to model, so in general we cannot specify this. Having said that, what we can do is specify how to assign the contents of a location to a prepared location node.

Assigning to a cleared location node: When assigning the contents of a cleared location, care has to be taken to ensure that the contents of reference values are pointed to rather than duplicated. In order to facilitate this, two utility operations are defined: one for identifying the referenced nodes (copyRefSG); and the other to add the copied pointer (path) to these identified nodes (copyRefNodes).

$$\begin{array}{lll}
 \text{copyRefSG}(G, n) & \widehat{=} & \{n' \mid n' \in G \wedge n' \text{ shareDescOf } n\} \\
 \text{copyRefNodes } n_1 \text{ to } n_2 \text{ in } G & \widehat{=} & \text{add}_G n_2 \text{ for } n_1 \text{ in } \text{copyRefSG}(G, n_1)
 \end{array}$$

Care also has to be taken to ensure that a duplicate of the value nodes are added to the copy node. This is facilitated by two utility operations: one for identifying the nodes to be

duplicated (`copyValSG`); and the other to perform the duplication (`copyValNodes`) using the node replacement operation.

$$\begin{aligned} \text{copyValSG}(G, n) &\hat{=} \{n' \mid n' \in G \wedge n' \text{ containDescOf } n\} \\ \text{copyValNodes } n_1 \text{ to } n_2 \text{ in } G &\hat{=} \text{replace}_G n_2 \text{ for } n_1 \text{ in } \text{copyValSG}(G, n_1) \end{aligned}$$

It is now possible to define the graph transformation operation of copying the contents of a source node to the empty location as the union of: the appropriately updated reference nodes; the descendant nodes that were not updated; the non-descendant nodes; and the duplicated value nodes.

$$\begin{aligned} \text{copy}_G n_1 \text{ to } n_2 \text{ in } G &\hat{=} (\text{copyRefNodes } n_1 \text{ to } n_2 \text{ in } G) \\ &\cup \text{descendants}(G, n_1) \setminus \text{copyRefSG}(G, n_1) \\ &\cup G \setminus \text{descendants}(G, n_1) \\ &\cup (\text{copyValNodes } n_1 \text{ to } n_2 \text{ in } G) \end{aligned}$$

Now given that the location with path `$copy` is an assignable location that does not exist within the graph, the copy assignment can be defined as follows:

$$\begin{aligned} (G, lp) :=_G lp' &\hat{=} (G, lp) :=_G \text{node}_G(G, lp') \\ (G, lp) :=_G Lp' &\hat{=} (G, \text{node}_G(G, lp)) :=_G Lp' \\ (G, Lp) :=_G Lp' &\hat{=} \text{let } G_1 \hat{=} (\text{copy}_G Lp' \text{ to } \{\text{\textcolor{violet}{\$copy}}\} \text{ in } (G \cup \{\text{\textcolor{violet}{\$copy}}\})) \\ &\quad G_2 \hat{=} (\text{copy}_G \{\text{\textcolor{violet}{\$copy}}\} \text{ to } Lp \text{ in } \text{prep}_G(G_1, Lp)) \\ &\quad \text{in } \text{unlink}_G(G_2, \{\text{\textcolor{violet}{\$copy}}\}) \end{aligned}$$

5.4 UTP Model

Our UTP model of locations uses the Abstract Location Trace Graph (ALTG) of [Section 5.3.4](#) to provide a semantics of locations, where the special logical variables Π_{ALTG} and Π_{ALTG}' represent the before and after states of the graph. The contents of this ALTG are then linked to the normal UTP program variables, using a technique inspired by [\[CHW06\]](#). In our case, the values of normal program variables are mirrored by correspondingly named first-level nodes in the graph. For example, the logical input and output variables for a UTP program variable x are represented by the node $\{x\}$ in the Π_{ALTG} and Π_{ALTG}' graphs respectively. Note that whenever there could be confusion between whether a variable is being used to denote its name rather than its value, we prefix the variable with a dash to get its name. For example, the predicate $x = *_G(\Pi_{\text{ALTG}}, 'x)$ holds whenever the value of variable x equals the value of extracting its corresponding element from the graph Π_{ALTG} (i.e. the one with the path name $'x$).

The remainder of this section introduces the healthiness conditions on the UTP model of locations, provides the definitions for a few operations, such as assignment, and relates the abstract and concrete models. Here the meta variable Q denotes a relational predicate that defines a UTP location model program.

5.4.1 Healthiness conditions

Before we formalise the relationship between a program's variables and the ALTG, it is worth introducing a healthiness condition to ensure that both the Π_{ALTG} and Π_{ALTG}' graphs are healthy (as defined in [Section 5.3.4](#)).

$$\begin{aligned} \text{HC}_{1_U}(Q) &\hat{=} Q = (Q \wedge \text{HC}_U(\Pi_{\text{ALTG}}) \wedge \text{HC}_U(\Pi_{\text{ALTG}}')) \\ \text{HC}_U(G) &\hat{=} \text{HC}_{1_G}(G) \wedge \text{HC}_{2_G}(G) \wedge \text{HC}_{3_G}(G) \wedge \text{HC}_{4_G}(G) \end{aligned}$$

The first step in formalising the link between the graph and program variables is by insisting that the first-level nodes within the graph correspond precisely to the UTP program variables other than the observational variables (i.e. Π_{ALTG} and Π_{ALTG}').

$$\begin{aligned} \text{HC2a}_U(Q) &\hat{=} Q = (Q \wedge \{ 'x \mid x \in \text{inv}\alpha Q \} = \text{labels}_G(\Pi_{\text{ALTG}}, \emptyset)) \\ \text{HC2b}_U(Q) &\hat{=} Q = (Q \wedge \{ 'x \mid x' \in \text{outv}\alpha Q \} = \text{labels}_G(\Pi_{\text{ALTG}}', \emptyset)) \\ \text{HC2}_U(Q) &\hat{=} \text{HC2a}_U(Q) \wedge \text{HC2b}_U(Q) \end{aligned}$$

Here, $\text{inv}\alpha Q$ and $\text{outv}\alpha Q$ represent the input and output alphabets of program Q except for the observational variables Π_{ALTG} and Π_{ALTG}' respectively; and the child labels of graph path are defined by $\text{labels}_G(G, p) \hat{=} \{ l \mid P.l \in \text{children}(G, \text{node}_G(G, p)) \}$.

The second, and last, step in formalising the link between the graph and program variables is to ensure that the value of a variable is the same as the value stored within the ALTG.

$$\begin{aligned} \text{HC3a}_U(Q) &\hat{=} Q = (Q \wedge (\bigwedge_{x \in \text{inv}\alpha Q} \bullet x = *_G(\Pi_{\text{ALTG}}, 'x'))) \\ \text{HC3b}_U(Q) &\hat{=} Q = (Q \wedge (\bigwedge_{x' \in \text{outv}\alpha Q} \bullet x' = *_G(\Pi_{\text{ALTG}}', 'x'))) \\ \text{HC3}_U(Q) &\hat{=} \text{HC3a}_U(Q) \wedge \text{HC3b}_U(Q) \end{aligned}$$

5.4.2 Operations

We only present those operations that significantly differ from those of the standard UTP model, as presented in Chapter 2 of [HH98]: specifically, the assignment and program variable management operations. A more detailed account is provided in Section 6.1, where we adapt these ideas to work with UTP designs.

The assignment operation is broken down into three cases, depending on the type of the r-value (i.e. the value to be assigned). These mirror the three cases presented in the trace-based graph model, except that the location is always defined in terms of a possibly empty path from a UTP program variable. It is defined as follows:

$$\begin{aligned} x.lp := nlv &\hat{=} \text{HCS}_U(\Pi_{\text{ALTG}}' = ((\Pi_{\text{ALTG}}, 'x.lp) :=_G nlv)) \\ x.lp := ecv &\hat{=} \text{HCS}_U(\Pi_{\text{ALTG}}' = ((\Pi_{\text{ALTG}}, 'x.lp) :=_G ecv)) \\ x.lp := y.lp_1 &\hat{=} \text{HCS}_U(\Pi_{\text{ALTG}}' = ((\Pi_{\text{ALTG}}, 'x.lp) :=_G 'y.lp_1)) \end{aligned}$$

where $x.\emptyset = x$ and $\text{HCS}_U(Q) \hat{=} \text{HC3}_U(\text{HC2}_U(\text{HC1}_U(Q)))$

Note that the combined healthiness condition ensures that the consequences of updating shared values can be seen by all participating UTP program variables.

The variable introduction and elimination operations are also defined in terms of their effects on the ALTG. Here the variable introduction operation provides a default unset value to the introduced variable, and the variable elimination operation removes all references to the value from the graph.

$$\begin{aligned} \text{var } x &\hat{=} \exists x \bullet \text{HCS}_U(\Pi_{\text{ALTG}}' = ((\Pi_{\text{ALTG}}, 'x) :=_G \mathbf{i})) \\ \text{end } x &\hat{=} \exists x' \bullet \text{HCS}_U(\Pi_{\text{ALTG}}' = \text{unlink}_G(\Pi_{\text{ALTG}}, 'x')) \end{aligned}$$

5.4.3 Full abstraction

The ALTG-based UTP model of locations, outlined here, is fully abstract in the sense described earlier: two concrete location graphs are equivalent, as defined in Section 5.2.3, iff their corresponding ALTGs are equal. This is essentially because the underpinning ALTG model is fully abstract by design; it removes the need for explicitly indexed shareable locations. Here, each location has precisely one path-set that represents it.

5.5 Summary

We have introduced a UTP model of locations that has two types of location: shareable locations, which can be dereferenced by a pointer and are aliased when copied; and containable locations, which cannot be dereferenced by a pointer and are duplicated when copied. This differs from the other UTP models of pointers, where every location is considered to be shareable, as discussed in [Section 7.3](#).

An advantage of our location approach is that it directly supports languages that have compound entities that are duplicated – rather than aliased – on assignment, without losing the ability to share data via pointer aliasing. For example, both a C# struct and an $\mathcal{O}$ -calculus object are duplicated on assignment, but can be shared by storing the whole value in a shareable location. Further, the model of locations presented here is fully-abstract, as required by our thesis ([Section 1.5](#)); it is the basis of the fully-abstract model of objects presented in [Section 6.2](#).

Chapter 6

Combining Objects and Locations

This chapter illustrates how the Unifying Theories of Programming (UTP) model of locations presented in [Section 5.4](#) can in principle be extended to support both the UTP theory of designs and our UTP theory of objects ([Chapter 3](#)). This model of locations uses an Abstract Location Trace Graph (ALTG) to model the state of a program ([Section 5.3](#)), where:

1. each location is modelled by a *path-set* that contains the potentially infinite set of paths from the root node to the node representing that location; and
2. each UTP program variable x corresponds to a first-level node in the ALTG with the same name (i.e. the node denoted by the path-set $\{x\}$).

The second of these design choices is slightly updated within this chapter in order to simplify the presentation of the combined object and location model. Here, we associate the UTP program variables with the children of the first-level *program variables* node $\{\text{var}\}$. Therefore, each program variable x – within the alphabet of the UTP program – now corresponds to the node $\{\text{var}.x\}$, where x denotes the name of program variable x rather than its value (as in the previous chapter).

Aside 6.1: The first-level nodes within the ALTG are also used to represent some meta-variables that are outside the alphabet of a UTP program; for example, those variables that are hidden from the scope of a subprogram that calculates the value of a command-expression’s argument. Here, the $\{\text{hide}\}$ node is used to store a stack of such hidden local variable blocks (frames). Note that this stack is also used to store the call-frames of procedure invocation. Variable hiding is discussed in more detail towards the end of [Section 6.1.2](#).
□

6.1 UTP Designs with Locations

We now illustrate a way of adapting the UTP model of locations presented in [Section 5.4](#) so that it supports the notion of a design. Our outline of a UTP design-based model of locations uses an Abstract Location Trace Graph (ALTG) to provide a semantics of locations, where the logical observation variables Π_{ALTG} and Π_{ALTG}' represent the before and after states of the graph. The contents of this ALTG are linked to the UTP program variables (in [Section 6.1.1](#)) as discussed in the introduction to this chapter. Having done this, we then (in [Section 6.1.2](#)) illustrate how those UTP design operations – that directly update the state of a program – can be adapted to work with this hybrid model. Operations, such as sequential composition and conditional choice, that do not directly update the state of a program, require no adaptation, as they do not directly change the values held in any program or observational variable. Hence, they do not change the contents of the ALTG.

6.1.1 Healthiness conditions

The healthiness conditions presented in this section provide a conservative extension to those of the UTP design, in the sense that when they are applied to a design they return a design. Our strategy is to define the healthiness conditions via the composition of two designs: one that specifies what it is to be healthy; and the other the program to be made healthy. Therefore, we have to transform (adapt) our existing – conjunctive – healthiness conditions for the UTP model of locations (Section 5.4.1) into this form.

Aside 6.2: All of our location healthiness conditions have the property that healthy designs can be written as the sequential composition of two parts: one that is written solely in terms of the input alphabet; and the other that is written solely in terms of the output alphabet. This is not possible for arbitrary conjunctive healthiness conditions in general.

□

Healthiness prefixing: We now consider those conjunctive healthiness conditions, of the form $P = P \wedge q$, where: the UTP program P is made healthy by being conjunctively combined with the – healthiness – predicate q ; and the predicate q is written solely in terms of the input alphabet of P . Such a healthiness predicate can be systematically transformed into the first of the two sequentially composed designs by the following **preHc** utility function (UTP command).

$$\text{preHc}(q) \hat{=} \emptyset : (q \vdash \text{true})$$

Having systematically transformed the healthiness predicate into a healthiness design, the next step is to show that when it is sequentially combined with another design it returns a design of the expected form; i.e. a design where the healthiness predicate q is a precondition for the correctness of the design as a whole.

Lemma 6.3 – Left healthiness design lemma:

$$\text{preHc}(q) \circ (p \vdash P) = q \wedge p \vdash P$$

provided that q only contains input variables (i.e. undashed variables).

□

Proof

$$\begin{aligned}
 & \text{preHc}(q) \circ (p \vdash P) \\
 = & \dots\dots\dots \text{Defn. of preHc} \\
 & (q \vdash w' = w) \circ (p \vdash P) \\
 = & \dots\dots\dots \text{Law C.2 comp. design} \\
 & q \wedge \neg(w' = w \circ \neg p) \vdash w' = w \circ P \\
 = & \dots\dots\dots \text{skip rel. unit of } (- \circ -) \\
 & q \wedge \neg(\neg p) \vdash P \\
 = & \dots\dots\dots \text{Law of the excluded middle} \\
 & q \wedge p \vdash P
 \end{aligned}$$

□

Note that it is also straightforward to demonstrate that any design $p \vdash P$ is equal to $p \vdash p \wedge P$. Therefore, the healthiness predicate implicitly prefixes the resultant part (i.e. the post condition) of the design as well.

Healthiness postfixing: Healthiness postfixing is similar to prefixing, except that the conjunctive healthiness predicate is of the form $P = P \wedge Q$, where: the UTP program P is made healthy by being conjunctively combined with the healthiness predicate Q ; and the predicate Q is written solely in terms of the output alphabet of P . Such a healthiness predicate can be systematically transformed into the second of the two sequentially composed designs by the following `postHc` utility function (UTP command).

$$\text{postHc}(Q) \hat{=} \emptyset : (\text{true} \vdash Q)$$

Having systematically transformed the healthiness predicate into a healthiness design, the next step is to show that when it is sequentially combined with another design it returns a design of the expected form; i.e. a design where the healthiness predicate Q is a postcondition for the correctness of the design as a whole.

Lemma 6.4 – Right healthiness design lemma:

$$(p \vdash P) \mathbin{\circ} \text{postHc}(Q) = p \vdash P \wedge Q$$

provided that Q only contains output variables (i.e. dashed variables).

□

Proof

$$\begin{aligned} & (p \vdash P) \mathbin{\circ} \text{postHc}(Q) \\ = & \dots\dots\dots \text{Defn. of postHc} \\ & (p \vdash P) \mathbin{\circ} (\text{true} \vdash Q \wedge w' = w) \\ = & \dots\dots\dots \text{Corol. C.3 comp. pre-true} \\ & p \vdash P \mathbin{\circ} (Q \wedge w' = w) \\ = & \dots\dots\dots \text{Defn. of } (- \mathbin{\circ} -) \\ & p \vdash \exists w_0 \bullet P[w_0/w'] \wedge Q[w_0/w] \wedge w' = w_0 \\ = & \dots\dots\dots \text{No input vars in } Q \\ & p \vdash \exists w_0 \bullet P[w_0/w'] \wedge (Q \wedge w' = w_0) \\ = & \dots\dots\dots \text{One point rule } (w_0 = w') \\ & p \vdash P \wedge Q \end{aligned}$$

□

Location design healthiness conditions: We are now in a position to systematically transform the three pairs of location healthiness conditions of the UTP location model presented in [Section 5.4.1](#). The first pair of healthiness conditions are that the graphs representing both the input and output states of a program are healthy. That is the graph satisfies the four ALTG healthiness conditions presented in [Section 5.3.4](#).

$$\begin{aligned} \text{HC1a}_D(P) & \hat{=} \text{preHc}(\text{HC}_G(\Pi_{\text{ALTG}})) \mathbin{\circ} P \\ \text{HC1b}_D(P) & \hat{=} P \mathbin{\circ} \text{postHc}(\text{HC}_G(\Pi_{\text{ALTG}}')) \\ \text{HC1}_D(P) & \hat{=} \text{HC1a}_D(\text{skip}^D) \mathbin{\circ} \text{HC1b}_D(P) \end{aligned}$$

$$\text{where } \text{HC}_G(G) = \text{HC1}_G(G) \wedge \text{HC2}_G(G) \wedge \text{HC3}_G(G) \wedge \text{HC4}_G(G)$$

Note that the healthiness conditions are deliberately split into pre and post stages. This enables the post-stage healthiness conditions to be applied to the output of a UTP program – that specifies a change to its observational Π_{ALTG} variable – and produce a new UTP operation that healthily maps the UTP program (and other observational) variables to their Π_{ALTG} counterparts.

Having specified the healthiness conditions for ALTG itself, the next step is to formalise the link between the graph and the program variables. This is achieved in two steps, in a similar

manner to that of the UTP location model; they differ in the details of which graph nodes are used to represent the program variables (as discussed in the introduction to this chapter).

The first step in formalising the link between the graph and program variables is by insisting that the children of the $\{\mathbf{var}\}$ node within the graph correspond precisely to the UTP program variables other than the observational variables (i.e. Π_{OK} , $\Pi_{\text{OK}'}$, Π_{ALTG} and $\Pi_{\text{ALTG}'}$).

$$\begin{aligned} \text{HC2a}_D(P) &\hat{=} \text{preHc}(\{x \mid x \in \text{inv}\alpha P\} = \text{labels}_G(\Pi_{\text{ALTG}}, \mathbf{var})) \ ; P \\ \text{HC2b}_D(P) &\hat{=} P \ ; \text{postHc}(\{x \mid x' \in \text{outv}\alpha P\} = \text{labels}_G(\Pi_{\text{ALTG}'}, \mathbf{var})) \\ \text{HC2}_D(P) &\hat{=} \text{HC2a}_D(\text{skip}^D) \ ; \text{HC2b}_D(P) \end{aligned}$$

Here, $\text{inv}\alpha P$ and $\text{outv}\alpha P$ represent the input and output alphabets of program P except for the observational variables; and $\text{labels}_G(G, p)$ returns the child labels of graph path (as in Section 5.4.1).

The second step in formalising the link between the graph and program variables is to ensure that the value of a variable is the same as the value stored within the ALTG.

$$\begin{aligned} \text{HC3a}_D(P) &\hat{=} \text{preHc}(\bigwedge_{x \in \text{inv}\alpha P} x = *_G(\Pi_{\text{ALTG}}, \mathbf{var}.x)) \ ; P \\ \text{HC3b}_D(P) &\hat{=} P \ ; \text{postHc}(\bigwedge_{x' \in \text{outv}\alpha P} x' = *_G(\Pi_{\text{ALTG}'}, \mathbf{var}.x)) \\ \text{HC3}_D(P) &\hat{=} \text{HC3a}_D(\text{skip}^D) \ ; \text{HC3b}_D(P) \end{aligned}$$

6.1.2 Operations

We now present a sample of four operations (UTP commands) that directly change the state of a program. The first two commands ($\mathbf{var_}$) and ($\mathbf{end_}$) introduce and remove a local variable from the scope of the program respectively. The third command ($\mathbf{_} :=_D \mathbf{_}$) defines the assignment of a value to a variable (multiple assignment is left for future work). The fourth command ($\mathbf{hide_from_}$) temporarily hides a collection of variables from the scope of a subprogram.

Each of the sample commands is written in terms of operations on the ALTG, which transform the input ALTG (Π_{ALTG}) into the resultant output ALTG ($\Pi_{\text{ALTG}'}$). They do not however specify the required updates to the linking of UTP variables (and their values) to the nodes in the ALTG. This is achieved by applying the latter two of the post healthiness condition constructors (i.e. HC2b_D and HC3b_D). For convenience these healthiness condition constructs are combined into the following output healthiness condition constructor.

$$\text{HC23b}_D(P) \hat{=} \text{HC2b}_D(P) \ ; \text{HC3b}_D(\text{skip})$$

Note that the first of the post healthiness conditions is not applied, as this is assumed to be guaranteed by the application of the ALTG graph operations to a healthy graph. Further, the pre healthiness conditions are not explicitly mentioned within the definition, in order to keep the definition as straightforward as possible. Given that these preconditions are met (and the operations are correctly specified), then it ought to be possible to show that the output design is guaranteed to be healthy. Here, we are only giving an illustration of how our previous work could be combined; it is for future work to iron out the details, and perform the necessary consistency proofs.

Local variable introduction and completion: The local variable x can be introduced by the command ($\mathbf{var\ } x$), so long as it is not already in this command's initial alphabet (i.e. $x \notin A$). When introduced its value is set to the explicit *unset* value $\mathbf{i}$.

$$\begin{aligned} \mathbf{var}_A x &\hat{=} \exists x \bullet \text{HC23b}_D(\\ &\quad \text{true} \\ &\quad \vdash \\ &\quad \Pi_{\text{ALTG}'} = ((\Pi_{\text{ALTG}}, \mathbf{var}.x) :=_G \mathbf{i}) \\ &\quad) \end{aligned}$$

Here, the alphabet of the command as a whole $\alpha(\text{var } x)$ is $A \setminus \{x\}$.

The local variable x can be removed from the current scope by the command $(\text{end } x)$.

$$\begin{aligned} \text{end}_A x \quad \widehat{=} \quad & \exists x' \bullet \text{HC23b}_D(\\ & \text{true} \\ & \vdash \\ & \Pi_{\text{ALTG}}' = \text{unlink}_G(\Pi_{\text{ALTG}}, \text{var. } 'x) \\ &) \end{aligned}$$

Here, the alphabet of the command as a whole $\alpha(\text{end } x)$ is $A \setminus \{x'\}$.

Assignments: The assignment command is split into three cases, depending on the type of the rvalue (i.e. the value to be assigned). These mirror the three cases presented in the ALTG, except that the location is always defined in terms of a possibly empty path from a UTP program variable. Each of the assignments has the form $x.lp :=_{DA} vp$, where: x represents a program variable in the alphabet A ; lp is the path from x to assignable location in Π_{ALTG} ; and vp is either a value that contains no references to shareable locations or a path to a location in Π_{ALTG} . For example, $\Pi_{\text{RES}} :=_D 42$ assigns the literal 42 to the special result variable. Note that in this case the empty path ($\emptyset$) has been omitted as it is the unit of path concatenation ($-\cdot-$).

Assignment of a scalar: The scalar assignment command $(x.lp :=_{DA} nlv)$ assigns a null or literal value (nlv) to the location $x.lp$, where A is the alphabet of the command, x is a program variable (i.e. $x \in A$) and lp is a possibly empty path from x to an assignable location in the ALTG.

$$\begin{aligned} x.lp :=_{DA} nlv \quad \widehat{=} \quad & \text{HC23b}_D(\\ & \text{true} \\ & \vdash \\ & \Pi_{\text{ALTG}}' = (\Pi_{\text{ALTG}}, \text{var. } 'x.lp) :=_G nlv \\ &) \\ \alpha(x.lp :=_{DA} e) \quad \widehat{=} \quad & A \end{aligned}$$

Assignment of an encapsulated compound value: The compound assignment command $(x.lp :=_{DA} ecv)$ assigns an encapsulated compound value to the variable x so long as the variable x is in the alphabet of the assignment. (Recall that on [page 82](#) an ‘encapsulated compound value’ was defined as a compound value that does not contain shareable locations.)

$$\begin{aligned} x.lp :=_{DA} ecv \quad \widehat{=} \quad & \text{HC23b}_D(\\ & \text{true} \\ & \vdash \\ & \Pi_{\text{ALTG}}' = (\Pi_{\text{ALTG}}, \text{var. } 'x.lp) :=_G ecv \\ &) \\ \alpha(x.lp :=_{DA} e) \quad \widehat{=} \quad & A \end{aligned}$$

Assignment of a program variable: The variable assignment command $(x.lp_1 :=_{DA} y.lp_2)$ assigns the contents of location $y.lp_2$ to the location $x.lp_1$ so long as both x and y are program variables in the alphabet of the assignment and the locations $y.lp_2$ and $x.lp_1$ exist and are

assignable respectively.

$$\begin{aligned}
 x.lp :=_{DA} y & \quad \widehat{=} \quad \text{HC23b}_D(\\
 & \quad \text{true} \\
 & \quad \vdash \\
 & \quad \Pi_{\text{ALTG}}' = (\Pi_{\text{ALTG}}, \text{var}. 'x.lp) :=_G \text{var}. 'y \\
 & \quad) \\
 \alpha(x.lp :=_{DA} e) & \quad \widehat{=} \quad A
 \end{aligned}$$

Hiding and filtering variables from programs: In Section 2.3.4 we discussed why hiding was used in the modelling of the $\mathcal{O}$ -calculus. Here, variables were temporarily hidden from the scope of a subprogram (or procedure call) and then restored following its completion.

Aside 6.5: This was achieved via the use of existential quantification, which essentially stored a *copy* of the value of each variable to be hidden, prior to that variable being deleted (i.e. ending the variable's scope).

□

Providing a facility to hide variables can, in principle, lead to a situation where some heap locations are only referenced by these hidden variables. Therefore, as our abstract heap model automatically removes unreferenced entries, we need to ensure that the hidden variables are stored in the abstract heap. Further, as the hiding operations are supplied with an unordered list of variables to hide, we support the creation and deletion of hidden local variable frames, where a frame is a mapping between variables and their values.

We now introduce another two first-level nodes – $\{\text{hide}\}$ and $\{\text{hTop}\}$ – to store the *stack* of hidden variable frames, where the node:

- $\{\text{hide}.i\}$ contains the i^{th} frame on the stack; and
- $\{\text{hTop}\}$ contains the index of the top stack entry (or zero if no such entry exists).

Having specified the structure of the hidden variable frame stack, it is useful to provide some operations to manage it. First, we provide the hTop macro that extracts the index on the top of the Π_{ALTG} hidden local variable frame stack.

$$\text{hTop} \quad \equiv \quad *_G(\Pi_{\text{ALTG}}, \text{hTop})$$

Second, we provide two UTP commands for incrementing and decrementing this counter.

$$\begin{aligned}
 \text{incrFr} & \quad \widehat{=} \quad \{\Pi_{\text{ALTG}}\} : (\text{true} \vdash \Pi_{\text{ALTG}}' = (\Pi_{\text{ALTG}}, \text{hTop}) :=_G \text{hTop} + 1) \\
 \text{decrFr} & \quad \widehat{=} \quad \{\Pi_{\text{ALTG}}\} : (\text{hTop} > 0 \vdash \Pi_{\text{ALTG}}' = (\Pi_{\text{ALTG}}, \text{hTop}) :=_G \text{hTop} - 1)
 \end{aligned}$$

These commands are to be used along side the commands for adding and deleting a frame, where the adding of a frame creates an empty variable-value map, which then needs to be populated.

$$\begin{aligned}
 \text{addFrame} & \quad \widehat{=} \quad \text{HC23b}_D(\\
 & \quad \text{true} \\
 & \quad \vdash \\
 & \quad \Pi_{\text{ALTG}}' = (\Pi_{\text{ALTG}}, \text{hide.hTop}) :=_G \{\} \\
 & \quad) \\
 \text{delFrame} & \quad \widehat{=} \quad \text{HC23b}_D(\\
 & \quad \text{true} \\
 & \quad \vdash \\
 & \quad \Pi_{\text{ALTG}}' = \text{unlink}_G(\Pi_{\text{ALTG}}, \text{hide.hTop}) \\
 & \quad)
 \end{aligned}$$

Once a hidden variable frame has been created (added) it is possible to populate it, by copying a current local variable (and its associated value) into it.

$$\text{initVar}_{Ax} \hat{=} \text{HC23b}_D(\text{true} \vdash \Pi_{\text{ALTG}}' = \text{copy}_G(\Pi_{\text{ALTG}}, \text{var.}'x, \text{hide.hTop.}'x))$$

where x is in the alphabet A

The last operation on the hidden variable frames is to recover the data stored in the frame. This is achieved by copying the hidden variable into a variable of the same name in the current local variable block (frame).

$$\text{restoreVar}_{Ax} \hat{=} \text{HC23b}_D(\text{true} \vdash \Pi_{\text{ALTG}}' = \text{copy}_G(\Pi_{\text{ALTG}}, \text{hide.hTop.}'x, \text{var.}'x))$$

where x is in the alphabet A

Hide variables command: The hide variables command $\text{hide } x_{i=1}^k$ from P removes the variables $x_{i=1}^k$ from the alphabet of P , but following P 's completion restores these variables with their previous values.

$$\begin{aligned} \text{hide } _ \text{ from } P &\hat{=} P \\ \text{hide } x_{i=1}^k \text{ from } P &\hat{=} \text{incrFr} \circ \text{addFrame} \circ \\ &(\circ_{i=1}^k \text{initVar } x_i) \circ (\circ_{i=1}^k \text{end } x_i) \circ \\ &P \circ \\ &(\circ_{i=1}^k \text{var } x_i) \circ (\circ_{i=1}^k \text{restoreVar } x_i) \circ \\ &\text{delFrame} \circ \text{decrFr} \end{aligned}$$

Aside 6.6: Limiting a subprogram to a specific set of variables can be achieved in an almost identical manner to that described in [Section 2.3.4](#); the only difference is that this model's variant of the hide command is used in place of the original hide command.

□

6.2 Location Graph Model of the Object Calculus

We now present an outline of a way in which the combined theory of locations and designs can be used to model the $\mathcal{O}$ -calculus. First, we discuss how the $\mathcal{O}$ -calculus's concrete operand values can be represented in an ALTG. Second, we discuss how the UTP constant-map-design ($\mathcal{U}_M$ -design), presented in [Section 3.4](#), can be adapted to use such ALTG-based values. Third, and last, we discuss the feasibility of completing this model.

6.2.1 Value representation

For the purposes of the following discussion we consider the $\mathcal{O}$ -calculus labels and methods to be values; they were often treated as values in our previous UTP models of the $\mathcal{O}$ -calculus ([Chapter 3](#)). Literal values, such as the integers, the booleans, the null pointer, and the explicit

unset value, can be represented in the ALTG by their texts. Here, the node that represents the location of the literal value has a single outgoing edge (to a leaf node) that is labelled with the literal value's text. An object method's label can be similarly represented by its own text.

An object essentially contains a set of distinctly labelled methods. Therefore, it can be represented by a compound value node where each outbound edge connects to a node representing the location of a method value; and is labelled with the method's name. To avoid potential confusion with labelled values within the ALTG, the object method labels can be prefixed with the \$ symbol.

An $\mathcal{O}$ -calculus method is represented by a pair, a *self parameter* and an *expression*, where the expression (which represents the method's body) provides the method's definition. An obvious way of representing such a method is by a compound value node that has two outbound edges labelled `_self` and `_body`, where: the `_self` edge leads to a node with a single outbound edge (to a leaf node) that is labelled with the self parameter's name; and the `_body` edge leads to a node with a single outbound edge that is labelled with the program text that represents the method's definition. However, this representation turns out to be problematic, due to a combination of the following four issues:

1. the body of a method can contain inner methods that refer to outer methods' parameters;
2. the method parameters are replaced by their argument value on invocation;
3. the argument provided to a method invocation can be a pointer value; and
4. the ALTG pointer values are unstable – they change during program execution.

The first two issues require the ability to replace a parameter within the program text (and inner program texts) with a value. This is feasible so long as there is a textual representation of any value that is to be substituted; such a substitution approach was used for the first two object models in [Chapter 3](#). The third issue explicitly mentions that pointers can be the values for substitution, thus these values would need a textual representation. The ALTG represents a pointer value by a potentially infinite path-set; such path-sets can be textually denoted by a finite regular expression. The fourth issue says that the path-sets denoting a pointer's value can change during the execution of a program (e.g. assigning the pointer to a new variable adds some additional paths to the path-set that denotes its value). Therefore, any changes in the abstract pointer's values would invalidate the textual representation of the pointer within a method's body. A method of preventing such an invalidation, would be to introduce a unique tag for each pointer, which could be stored textually in the method. However, this would prevent our model of pointers from being fully abstract, as we would have re-introduced pointer labelling.

An alternative approach to the modelling of a method is to split it into a triple, a *self parameter*, a *constant map*, and an *expression*, where the *constant map* is used to store the values of parameters as they are bound. Note that the reason it is called a constant map, is that by the time it is used for method invocation, all its parameters have been bound to their constant values. Having said this, the map is also used to provide unique names to, the otherwise anonymous, inner methods, functions, and objects (during $\mathcal{O}$ -calculus compilation as discussed in [Section 3.4.3](#)). It is important that the anonymous inner methods and functions are extracted in this manner, because they also can contain outer (pointer) parameters.

Aside 6.7: Extracting the inner-objects is not strictly required, but it does mean that their values are expanded into the ALTG during compilation, rather than being contained within the program-text label representing the method's body.

□

Hence, we model a method by a compound value node that has three outbound edges labelled `_self`, `_map`, and `_body`, where: the `_self` and `_body` edges are modelled as before; and the

`_map` edge leads to a compound value node that contains an entry (labelled outbound edge) for each parameter, inner method, inner function and inner object. These outbound edges lead to the location of general, method, function, and object value nodes respectively, and are labelled with their distinct (locally unique) names.

An $\mathcal{O}$ -calculus function is represented in the ALTG in an identical manner to that of a method. Function and method definitions differ in their usage, rather than their data representation (e.g. they differ in how they are invoked).

6.2.2 Model definition

The UTP location-graph-design ($\mathcal{U}_L$ -design) is similar to that of the $\mathcal{U}_M$ -design (in Section 3.4), in the sense that it adapts this UTP model of objects to work with the combined theory of locations and designs. The significant difference is that of the heap representation, which has gone from a straightforward partial map between named abstract locations and operand values, to an abstract graph (i.e. an ALTG) that represents the value of any in scope (or hidden) variable capable of containing a reference value (pointer). Hence, the heap is only one aspect of the more general model of storage provided by the combined theory of locations and designs.

The $\mathcal{U}_L$ -design uses the following three pairs of observational variables to define a UTP model of the $\mathcal{O}$ -calculus.

- Π_{OK} and Π_{OK}' to represent the start and termination of a program;
- Π_{ALTG} and Π_{ALTG}' to represent the input and output states of a program;
- Π_{RES} and Π_{RES}' to represent the current (working) and final result of a program.

A significant amount of the $\mathcal{U}_M$ -design can be directly inherited by the $\mathcal{U}_L$ -design, such as its compilation of literal values and its definition of evaluation function, because:

1. the combined theory of locations and designs provides support for most of the operations that the original design supported (including variable declaration and hiding);
2. the UTP commands used to support the object modelling are frequently written in terms of these design operators; and
3. the $\mathcal{O}$ -calculus compilation schemes mainly target these supported commands.

Note that these conditions are not sufficient (in all cases) to enable an operation to be inherited. For example, consider the function transformation command $\text{trans}_M(f, k)$, which is inherited from the UTP result-value-design ($\mathcal{U}_R$ -design), and defined in Section 3.3.2; it applies function f to the values of the local variables that represent its arguments (i.e. $x_{i=1}^k$). If any argument x_i contains a reference value (pointer), which the function f uses, then it is possible that the result will store the current representation of the pointer value, which can become invalid (as previously discussed in Section 6.2.1). Therefore, all functions that previously stored the state of a reference value, or accessed the heap (based on the pointer value), need to be updated or replaced so that they perform actions that are consistent with the ALTG's model of pointers.

Aside 6.8: The arithmetic and logical transformation functions worked only on numeric and boolean values respectively, thus they can be inherited from the $\mathcal{U}_M$ -design. Further, their enclosing, command expression based, compilation schemes can also be inherited, as these are written in terms of operators that are supported by both the $\mathcal{U}_M$ -design and $\mathcal{U}_L$ -design.

□

We now consider how a sample of $\mathcal{U}_M$ -design operations can be modelled – specifically those operations that manipulate pointers (references) and methods.

Heap operations: The three $\mathcal{O}$ -calculus operations that manipulate the heap can be modelled in a straightforward manner in the $\mathcal{U}_L$ -design. First, the fresh location operation **fresh** is compiled to the **fresh_L** UTP command. This command deliberately clears the result variable before assigning the explicitly unset value ($\mathfrak{!}$) to a fresh shareable location. Note that if the clear operation had not been performed, then the assignment could set the value of an existing location to $\mathfrak{!}$.

$$\text{fresh}_L \hat{=} \text{clear}(\text{res}) \mathbin{\text{\&}} (\Pi_{\text{RES}}.\$:=_L \mathfrak{!})$$

where

$$\begin{aligned} \text{clear}(p) \hat{=} & \text{HC23b}_D(\\ & \text{true} \\ & \vdash \\ & \Pi_{\text{ALTG}}' = \text{unlinkChildren}_G(\Pi_{\text{ALTG}}, p) \\ &) \end{aligned}$$

The dereferencing operation $*e$ is compiled to the UTP program $\llbracket e \rrbracket_L \mathbin{\text{\&}} \text{deref}_L$. Here, the program representing the expression e results in a value, which the **deref_L** command uses to dereference a pointer; if the value is not a pointer then the command is specified to behave chaotically, as this is the UTP semantics for a *stuck* $\mathcal{O}$ -calculus program.

$$\begin{aligned} \text{deref}_L \hat{=} & \Pi_{\text{RES}} :=_L \Pi_{\text{RES}}.\$ \\ & \triangleleft \text{labels}_G(\Pi_{\text{ALTG}}, \text{res}) = \{\$ \} \triangleright \\ & \text{chaos} \end{aligned}$$

The pointer assignment (update) operation is slightly more complex to define, as it requires two expressions to be evaluated: one to represent the pointer to the location to be updated; and the other the value of the update. In order to simplify its definition we introduce the **pair_L** utility command, which stores the results of two programs in a pair.

$$\begin{aligned} \text{pair}_L(P_1, P_2) \hat{=} & \text{decl } x_1, x_2 \text{ in} \\ & (\text{hide } x_1, x_2 \text{ from } P_1) \mathbin{\text{\&}} x_1 :=_L \Pi_{\text{RES}} \mathbin{\text{\&}} \\ & (\text{hide } x_1, x_2 \text{ from } P_2) \mathbin{\text{\&}} x_2 :=_L \Pi_{\text{RES}} \mathbin{\text{\&}} \\ & \Pi_{\text{RES}} :=_L \{ \mathbf{1} \mapsto \mathfrak{!}, \mathbf{2} \mapsto \mathfrak{!} \} \mathbin{\text{\&}} \\ & \Pi_{\text{RES}}.\mathbf{1} :=_L x_1 \mathbin{\text{\&}} \\ & \Pi_{\text{RES}}.\mathbf{2} :=_L x_2 \end{aligned}$$

We are now in a situation to define the pointer assignment compilation scheme, which pre-evaluates the arguments of the pointer assignment operation before applying the **update_L** command.

$$e_1 * e_2 \hat{=} \text{pair}_L(\llbracket e_1 \rrbracket_L, \llbracket e_2 \rrbracket_L) \mathbin{\text{\&}} \text{update}_L$$

The **update_L** command itself is now straightforward to define. First, it checks whether the first element of the pair is actually a pointer value. If the check fails, then the program is invalid and thus behaves chaotically. If the check passes, then the value stored in the second part of the pair is assigned to this pointer. Having done this, the result variable is then updated to equal that of the pointer value (to mirror the $\mathcal{O}$ -calculus definition).

$$\begin{aligned} \text{update}_L \hat{=} & \Pi_{\text{RES}}.\mathbf{1}.\$:=_L \Pi_{\text{RES}}.\mathbf{2} \mathbin{\text{\&}} \\ & \Pi_{\text{RES}} :=_L \Pi_{\text{RES}}.\mathbf{1} \\ & \triangleleft \text{labels}_G(\Pi_{\text{ALTG}}, \text{res}.\mathbf{1}) = \{\$ \} \triangleright \\ & \text{chaos} \end{aligned}$$

Note that the order of processing is slightly different to that of the $\mathcal{O}$ -calculus being modelled. Here, the second expression gets evaluated before the check on the first argument's value is performed. This is not a significant problem, because the overall program semantics are preserved. First, we observe that the check does not alter the state of the program, therefore delaying its application does not affect the evaluation of the second expression. Second, we observe that the evaluation of the expression does not affect the outcome of the check, so long as it terminates. If it does not terminate or the check fails, then the $\mathcal{O}$ -calculus program never returns or becomes stuck respectively; in either case the result is modelled by the chaotic UTP program, as our UTP semantics does not distinguish between types of failure (i.e. ways in which a program can be invalid).

Method definitions and updates: The method definition follows $\mathcal{U}_M$ -design's alternative method definition scheme, presented in Section 3.4.3. This scheme enables the function values and method definitions to be represented by an ALTG, as it removes anonymous object values and program texts.

A method is defined as a triple of compiled program texts that represent the method's *self-variable*, *constant-map* and *body*, where the constant map provides definitions for the self-variable and each of the free-variables within the method's body.

$$\langle\langle \varsigma(x) e \rangle\rangle_M \hat{=} \left\{ \begin{array}{l} \text{\textcolor{violet}{_self}} \mapsto x, \\ \text{\textcolor{violet}{_map}} \mapsto M_1 \cup M_2, \\ \text{\textcolor{violet}{_body}} \mapsto \langle\langle e \{ t \leftarrow y \mid y \mapsto t \in M_2 \} \rangle\rangle_M \end{array} \right\} \quad \begin{array}{l} \text{Self variable} \\ \text{Constant map} \\ \text{Body} \end{array}$$

where

$$\begin{array}{ll} Y = \text{FV}(e) \cup \{x\} & \text{The method body's free variables} \\ M = \text{freshMap}(Y, \text{objsAndProcs}(e)) & \text{The named inner procedure and object map} \\ M_1 = \{y \mapsto \text{\textcolor{violet}{i}} \mid y \in Y\} & \text{The map of, un-set (free), external variables} \\ M_2 = \{y \mapsto \langle\langle e \rangle\rangle_M \mid y \mapsto e \in M\} & \text{The map of simulated external variables.} \end{array}$$

This compilation scheme ensures that no terms representing either a procedure definition or an object value are contained within a program text (which represents the definition of a method or function). Such procedures and objects have been replaced by local variables that refer to the definitions within a method's constant map. Therefore the $\mathcal{O}$ -calculus method update operation is provided with a method identifier (variable), rather than the method itself. The compilation of the transformed method update operation now follows, where: the expression e representing the object to be updated is evaluated, then checked to see if it contains the method label l ; if so it replaces the existing method indexed by label l with that stored in the method identifier variable x .

$$\langle\langle e.l \leftarrow x \rangle\rangle_L \hat{=} \langle\langle e \rangle\rangle_L ; \left(\begin{array}{l} \Pi_{\text{RES}}.l :=_L \Pi_{\text{RES}}.l.\text{\textcolor{violet}{_map}}.\text{\textcolor{violet}{x}} \\ \triangleleft l \in \text{labels}_G(\Pi_{\text{ALTG}}, \text{\textcolor{violet}{res}}) \triangleright \\ \text{chaos} \end{array} \right)$$

Field assignment is actually slightly more complex to define, because it has an extra expression to evaluate, namely, the one that results in the value to be assigned. This is similar to the issue that we encountered for pointer assignment. Therefore, we follow a similar approach of pre-evaluating the expressions and storing them in a pair. Having done this, the next step is to check whether the value in the first argument is an object that contains the member to be updated (the one with label l). If the check fails, then the program is invalid and thus behaves

chaotically. If the check passes, then the program: constructs a method that returns the value to be assigned; and updates the relevant object member with this constructed method value.

$$\begin{aligned} \llbracket e_1.l := e_2 \rrbracket_L &\hat{=} \text{pair}_L(\llbracket e_1 \rrbracket_L, \llbracket e_2 \rrbracket_L) \circ ; \\ & (\Pi_{\text{RES.1}}.l :=_L \{ \text{\textcolor{violet}{_self}} \mapsto _, \text{\textcolor{violet}{_map}} \mapsto \{ \mathbf{x} \mapsto \mathbf{i} \}, \text{\textcolor{violet}{_body}} \mapsto \mathbf{x} \} \circ ; \\ & \quad \Pi_{\text{RES.1}}.l. \text{\textcolor{violet}{_map}}.\mathbf{x} :=_L \Pi_{\text{RES.2}} \circ ; \\ & \quad \Pi_{\text{RES}} :=_L \Pi_{\text{RES.2}} \\ & \quad \triangleleft l \in \text{labels}_G(\Pi_{\text{ALTG}}, \text{\textcolor{violet}{res.1}}) \triangleright \\ & \quad \text{chaos} \\ &) \end{aligned}$$

Note that the order of processing is slightly different to that of the $\mathcal{O}$ -calculus being modelled, in precisely the same manner as that of the pointer assignment.

Method invocation: Method invocation in the $\mathcal{O}$ -calculus is compiled in three parts. First, an object-member pair is constructed from the invocation arguments: an expression (e) representing the target object (o) and a label (l). This is achieved by retrieving the method with label l from object o . Second, all the program (non-observational) variables are hidden from view, in order to create the right conditions for method invocation. And third, the actual method invocation is performed using a generic procedure call command (call_L). This executes the body of the method, where the method's self variable has been instantiated with the calling object's value.

$$\begin{aligned} \llbracket e.l \rrbracket_L &\hat{=} \text{pair}_L(\llbracket e \rrbracket_L, \text{getMethod}_L(\llbracket l \rrbracket_L^E)) \\ &\circ \Pi_{\text{OK}}, \Pi_{\text{ALTG}}, \Pi_{\text{RES}} \upharpoonright \\ &\text{call}_L \end{aligned}$$

where

$$\begin{aligned} \text{getMethod}_L l &\hat{=} \Pi_{\text{RES}} :=_L \Pi_{\text{RES.}l} \\ &\triangleleft l \in \text{labels}_G(\Pi_{\text{ALTG}}, \text{\textcolor{violet}{res}}) \triangleright \\ &\text{chaos} \end{aligned}$$

Note that getMethod_L makes use of the result of the evaluation of the previous program within the pair, which produces the object that the method is going to be extracted from.

Function invocation is compiled in a similar manner. The key difference is that the function's argument is a general expression, rather than a label.

$$\begin{aligned} \llbracket e_1 \ e_2 \rrbracket_R &\hat{=} (\text{pair}_L(\llbracket e_1 \rrbracket_L, \llbracket e_2 \rrbracket_L) \circ \text{swap}_L(\Pi_{\text{RES.1}}, \Pi_{\text{RES.2}})) \\ &\circ \Pi_{\text{OK}}, \Pi_{\text{ALTG}}, \Pi_{\text{RES}} \upharpoonright \\ &\text{call}_L \\ \text{swap}_L(x, y) &\hat{=} \text{decl } z \text{ in} \\ &\quad z :=_L x \circ x :=_L y \circ y :=_L z \end{aligned}$$

Note that the arguments within the pair are swapped so that they are in the same order as that for method invocation.

The generic procedure call is defined as in the earlier models as the least fix-point of the apply_L command, which: checks to ensure that the result variable contains a value-program-text pair; assigns the value to the procedure's – and any inner-procedure's – parameter; initialises the procedure's local call block (the parameters in its constant map); and then executes the

program contained within the procedure's program text.

$$\begin{aligned}
\text{call}_L &\hat{=} \mu z \bullet \text{apply}_L(z) \\
\text{apply}_L(z) &\hat{=} (\text{setParam} \ ; \\
&\quad \text{getProc} \ ; \\
&\quad \text{substCall}(z) \ ; \\
&\quad \text{addCallFrame} \ ; \\
&\quad \text{extProgText} \ ; \\
&\quad \text{delCallFrame} \\
&\quad) \\
&\triangleleft \Pi_{\text{RES.1}} \in \text{Value} \wedge \Pi_{\text{RES.2}} \in \text{ProgramText} \triangleright \\
&\text{chaos}
\end{aligned}$$

Here, the subprograms

<i>setParam</i>	Replaces all free occurrences of the procedure's parameter within the constant-map (texts) defined within this procedure. Note that for the purposes of this variable replacement the procedure being invoked is not considered to bind the variable being replaced.
<i>getProc</i>	Select the updated procedure (i.e. $\Pi_{\text{RES}} :=_L \Pi_{\text{RES.2}}$).
<i>substCall</i> (<i>z</i>)	Replaces all occurrences of the program-text call_L which are defined within the procedure (but not inner procedures), with the contents of the fix-point meta-variable <i>z</i> . Note this can always be represented by a program-text of the form $\text{apply}_L^i(\text{chaos})$, where <i>i</i> is a natural number, $\text{apply}_L^0(P) = \text{chaos}$ and $\text{apply}_L^{i+1}(P) = \text{apply}_L(\text{apply}_L^i(P))$.
<i>addCallFrame</i>	Moves the constant map representing the invoked procedure's call frame from the res._var node to the var node and adds these variables to the programs alphabet.
<i>extProgText</i>	Extracts the definition of the procedure's subprogram to execute from the res._body node, unsets the res node, and then starts executing the subprogram.
<i>delCallFrame</i>	Deletes the call frame variables introduced by the <i>addCallFrame</i> subprogram.

Each of these informally specified subprograms performs a significant role in the procedure invocation. Arguably the most awkward of these subprograms to define are the *setParams* and *substCall*(*z*), as both modify several locations within the ALTG, which correspond to the locations of the procedure's (and inner procedures') constant-map definitions and the procedure's recursive call program-texts respectively. Here, pattern matching can be used to determine the locations of all inner procedure nodes, which will always have three outbound edges with the reserved labels **_self**, **_map**, and **_body**. Note that care needs to be taken to ensure that the pattern match searching is stopped when a pointer edge (i.e. one with the label **\$**) is encountered; otherwise program-texts outside the procedure's definition could be incorrectly updated.

6.2.3 Cell – Method Invocation Example

We now illustrate the $\mathcal{U}_L$ -design method invocation process via an example – specifically a cut down version of the restorable cell example, in Section 2.2.5, which only contains the **data** and **get** methods. This cut down version of the restorable cell object is denoted by the following *ExampleCell* $\mathcal{O}$ -calculus object.

$$\text{ExampleCell} \equiv [\text{data} = \varsigma(x) 42, \text{get} = \varsigma(x) x.\text{data}]$$

The $\mathcal{U}_L$ -design version of *ExampleCell* object can be systematically constructed by applying the compilation schemes previously presented in this chapter (and those inherited from Chapter 3). The result of this compilation is now denoted by the following *ExCel* UTP object.

$$ExCel \equiv \{ \text{data} \mapsto \{ _self \mapsto x, _map \mapsto \{x \mapsto i\}, _body \mapsto \mathcal{E}_L 42 \}, \\ \text{get} \mapsto \{ _self \mapsto x, _map \mapsto \{x \mapsto i\}, _body \mapsto G_PT \} \}$$

where

$$G_PT \equiv \text{pair}_L(\mathcal{E}_L x, \text{getMethod}_L \text{data}) \\ \circ \Pi_{OK}, \Pi_{ALTG}, \Pi_{RES} \mid \\ \text{call}_L$$

The *ExCel* object can be represented by a tree, within an ALTG, as illustrated in Figure 6.1. Having said this, future ALTG illustrations that contain this tree, typically represent the whole

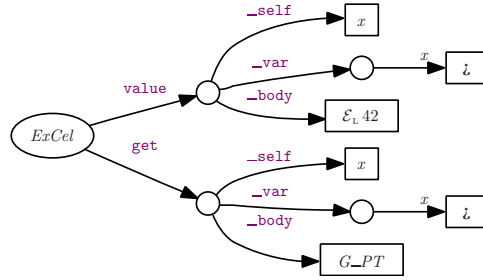


Figure 6.1: Example cell object – tree representation

tree by an *ExCel* labelled node, where multiple instances of the *ExCel* node represent distinct copies of the tree.

Representing the method invocation: We now consider the *ExampleCell.get* $\mathcal{O}$ -calculus program, which invokes the *ExampleCell*'s *get* method. This is systematically compiled to the following $\mathcal{U}_L$ -design program as follows.

$$\text{pair}_L(\mathcal{E}_L ExCel, \text{getMethod}_L \text{get}) \\ \circ \Pi_{OK}, \Pi_{ALTG}, \Pi_{RES} \mid \\ \text{call}_L$$

Note that this variant of the ternary filter operation ($P \circ \Pi_{OK}, \Pi_{ALTG}, \Pi_{RES} \mid Q$) sequentially composes the result of the first sub-program (P) with that of the second sub-program (Q), where the non-observational variables within the first subprogram are hidden from the scope of the second sub-program (i.e. all variables but Π_{OK} , Π_{RES} and Π_{ALTG} are hidden). In this case, the first sub-program ($\text{pair}_L(\mathcal{E}_L ExCel, \text{getMethod}_L \text{get})$) is only defined in terms of the observational variables, so the empty set of variables is removed from the second sub-program.

$$\text{pair}_L(\mathcal{E}_L ExCel, \text{getMethod}_L \text{get}) ; \\ (\text{hide } _ \text{ from} \\ \text{call}_L \\)$$

This is equivalent to not performing any hiding at all. So, the resultant program for our method invocation can be simplified to:

$$\text{pair}_L(\mathcal{E}_L ExCel, \text{getMethod}_L \text{get}) ; \\ \text{call}_L$$

Here, the outer level sequential composition operator provides a convenient break point in our example presentation, as it forms an imperative step between two parts of the example calculation. In general, we aim to reduce our UTP program into a sequence of imperative steps, as this illustrates the sequence of steps that an execution of the program would essentially take. Having said this, we do not reduce the variable declaration and hiding to this form, as it is useful to keep their block scoped form. Further, we unwind the method invocation as each method is invoked, in order to manage the amount of program that is being viewed at any point in time.

Evaluating the object-method pair: For the moment we focus on the first part of the method invocation process, which evaluates the arguments, and stores them – as a pair – in the observational result variable (Π_{RES}).

$\text{pair}_L(\mathcal{E}_L \text{ ExCel}, \text{getMethod}_L \text{ get})$

This subprogram can be rewritten to a block structured imperative style by recursively applying the definitions of the $\mathcal{U}_L$ -design commands (such as pair_L). This results in the following program definition, where each step of the execution is labelled by its step number.

```

( decl  $x_1, x_2$  in                (1)
  ( hide  $x_1, x_2$  from          (2)
     $\mathcal{E}_L \text{ ExCel}$           (3)
  ) ;                             (4)
   $x_1 :=_L \Pi_{\text{RES}}$  ;             (5)
  ( hide  $x_1, x_2$  from          (6)
     $\text{getMethod}_L \text{ get}$         (7)
  ) ;                             (8)
   $x_2 :=_L \Pi_{\text{RES}}$               (9)
   $\Pi_{\text{RES}} :=_L (\mathbf{i}, \mathbf{i})$         (10)
   $\Pi_{\text{RES}.1} :=_L x_1$            (11)
   $\Pi_{\text{RES}.2} :=_L x_2$            (12)
)                                  (13)
```

Here, step 1 introduces local variables x_1 and x_2 to store the result of evaluating the first and second arguments respectively. Steps 2–5 perform the evaluation of the first argument; i.e. the subprogram that evaluates the object that is the target of the method invocation. Note that the local variables are hidden from the evaluation of this sub-program ($\mathcal{E}_L \text{ ExCel}$), as their only purpose is to temporarily store the result of executing the sub-program.

Steps 6–9 perform the evaluation of the second argument in a similar manner to that of the first argument. The difference is that the ($\text{getMethod}_L \text{ get}$) sub-program atypically makes use of the value stored within the Π_{RES} variable, which is that of the target object (i.e. the same value as stored in x_1).

Steps 10–12 perform the creation of the object-method pair, and the storing of the object and method values within the first and second elements of that pair respectively. Step 13 removes the local variables from scope.

The overall affect of these thirteen steps is summarised in Figure 6.2, which illustrates both the initial state of the UTP program's Π_{ALTG} and the state of the program's Π_{ALTG} after the thirteenth step of execution. Recall that *ExCel* and *G_PT* denote the tree representing the example cell object and the program text of the *get* method respectively. Appendix D contains a step-by-step illustration of the thirteen changes to the UTP program's Π_{ALTG} .

Expanding the first method call: The second part of the method invocation is the generic procedure call_L program, which is defined to be the weakest fix-point of the apply_L program (i.e.

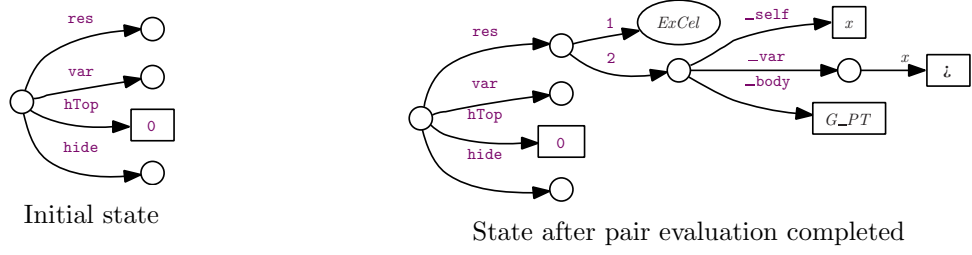


Figure 6.2: Evaluation pair summary

$\mu z \bullet \text{apply}_L(z)$). In the case of our example the weakest fix-point is the program $\text{apply}_L^2(\text{chaos})$, which can be rewritten with the definitions of the apply_L^2 and apply_L to generate the following UTP program.

```
( setParam ;
  getProc ;
  substCall( $\text{apply}_L(\text{chaos})$ ) ;
  addCallFrame ;
  extProgText ;
  delCallFrame
)
 $\triangleleft \Pi_{\text{RES.1}} \in \text{Value} \wedge \Pi_{\text{RES.2}} \in \text{ProgramText} \triangleright$ 
chaos
```

The above program can be further simplified by observing that following the evaluation of the object-method pair, the Π_{RES} variable contains a *Value* and a *ProgramText*, thus the **true** branch of the conditional can be taken. This produces the following program, where the twenty three steps (steps 18–40) involved with executing the body of the method itself are denoted by the *extProgText* program.

<i>setParam</i> ;	(14)
<i>getProc</i> ;	(15)
<i>substCall</i> ($\text{apply}_L(\text{chaos})$) ;	(16)
<i>addCallFrame</i> ;	(17)
<i>extProgText</i> ;	(18 – 40)
<i>delCallFrame</i>	(41)

For the moment we focus on the first four steps, which are step 14–17 of our example’s execution. These are illustrated in Figure 6.3, where G_PT' denotes the program-text that is created by substituting all the call_L texts within the method’s program-text with the text $\text{apply}_L(\text{chaos})$, and:

$$G_PT' \hat{=} \text{pair}_L(\mathcal{E}_L x, \text{getMethod}_L \text{data}) \\ \circ \Pi_{\text{OK}}, \Pi_{\text{ALTG}}, \Pi_{\text{RES}} \upharpoonright \\ \text{apply}_L(\text{chaos})$$

Step 18 is more involved as it extracts the above program text, which represents a method invocation that is similar to that of the initial program. The three main differences between this second invocation and the original are that:

1. the target object (*ExCel*) is stored in the local variable x , but is otherwise a duplicate of the original object;
2. the method to be selected is labelled with **data**, rather than **get**; and

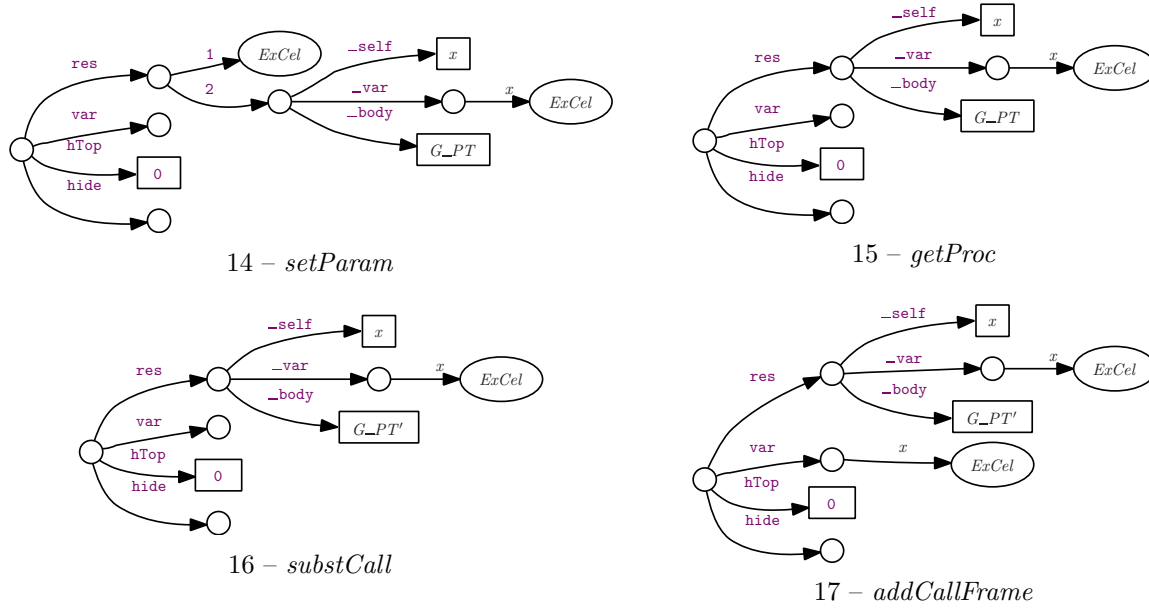


Figure 6.3: Steps 14 to 17 – Preparing the call frame

- the call_L program has been replaced by the $\text{apply}_L(\text{chaos})$ (as part of the original method's weakest fix-point calculation).

The effect of this extraction is illustrated in Figure 6.4. Following the extraction, the program

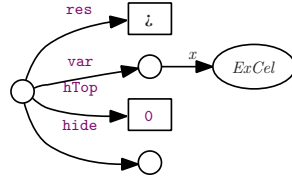


Figure 6.4: Inner program text extraction

contained within the inner method is essentially inserted into the execution path at this point in the program's execution; it takes twenty two steps (steps 19-40) to execute before returning control back to the original outer method.

Expanding the second method invocation: The second, inner, method invocation follows the same pattern as the original method invocation. First, the ternary filter is rewritten to produce the sequential composition of two program texts: the first performing the construction of an object-method pair; and the second performing the actual method invocation, once the original method's local variable x is hidden from scope.

```

pairL( $\mathcal{E}_L x, \text{getMethod}_L \text{data}$ ) ;    (19 – 31)
(  hide  $x$  from                        (32)
  applyL(chaos)                       (33 – 39)
)                                       (40)

```

The thirteen steps (steps 19–31) to perform pair evaluation are almost identical to those of the original object-method pair evaluation. The overall effect of these steps on the Π_{ALTG} is summarised in Figure 6.5.

Having got the inner object-method pair, the next step (step 32) is to hide all the non-observational variables from the scope of the inner-method's execution, as illustrated in Figure 6.6, so that the program's alphabet at the point of a procedure call contains only the

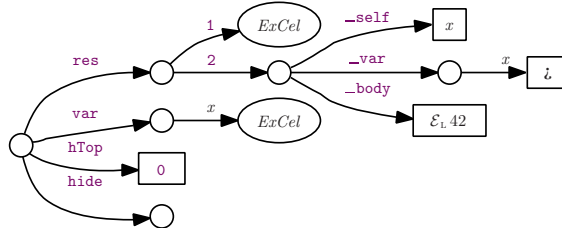


Figure 6.5: Steps 19–31 – Result of the inner method pair evaluation

$\mathcal{U}_L$ -design observational variables. Recall that this is a requirement of the way our model handles potentially mutually recursive procedure invocation as discussed in [Section 3.1.2](#).

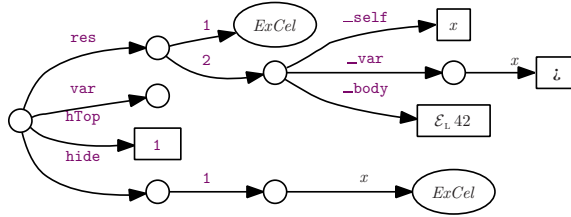


Figure 6.6: Step 32 – Hide the local variables

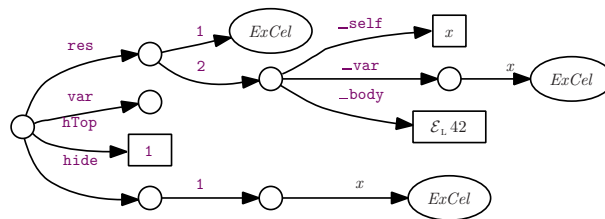
The apply_L program can now be rewritten in a similar manner to before, to produce the following program.

```

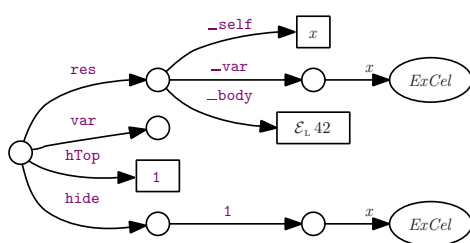
setParam ;           (33)
getProc ;            (34)
substCall(chaos) ;   (35)
addCallFrame ;       (36)
extProgText ;        (37 – 38)
delCallFrame         (39)

```

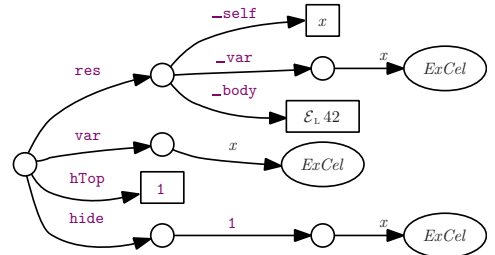
The first four steps (steps 33–36) prepare the inner method call frame as before, except that this time there is no call_L text to substitute, and thus no change in state following the execution of the *substCall* command. The affects of these steps on the Π_{ALTG} are illustrated in [Figure 6.7](#).



Step 33 – SetParam



Step 34 – GetProc (and Step 35 – SubstCall)



Step 36 – AddCallFrame

Figure 6.7: Steps 33–36 – Prepare the inner call frame

Following the preparation of the call frame the inner method is extracted and executed, as illustrated in Figure 6.8. Here the inner program consists of a single evaluation command ($\mathcal{E}_L 42$), which is step 38.

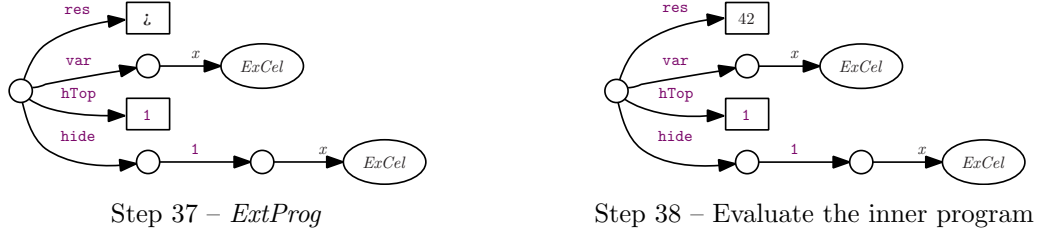


Figure 6.8: Steps 37–38 Extract and execute inner method

Complete the method invocations: Now all that remains is to complete the two method invocations. This amounts to deleting the inner method’s call frame, restoring the outer method’s local variable, and then deleting the outer method’s call frame. Figure 6.9 illustrates the affect these last three steps (steps 39–41) have on the Π_{ALTG} . Note that the final

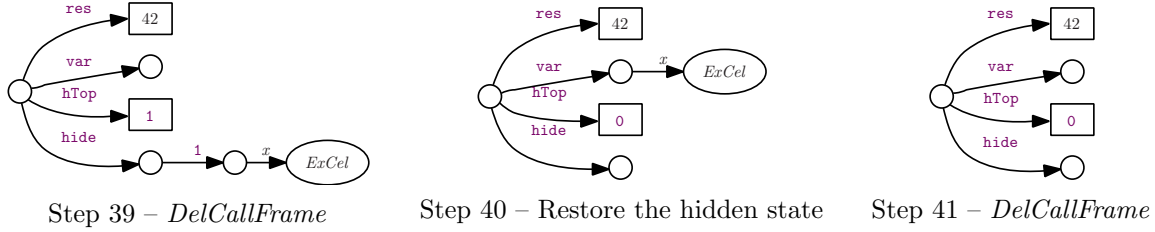


Figure 6.9: Steps 39–40 – Complete both invocations

Π_{ALTG} contains the value 42 in its result variable, which is the expected result of executing the example program.

6.2.4 Feasibility

So far we have illustrated an approach to adapting the $\mathcal{U}_M$ -design so that it can use the combined model of locations and designs. In particular, we have discussed how the ALTG can be used to store the various sorts of data used within the $\mathcal{O}$ -calculus, and how these graphs can be used to represent the evolving state of a $\mathcal{U}_L$ -design program.

There is still a significant amount of work required to complete the $\mathcal{U}_L$ -design. Having said this, we have illustrated how some of the more awkward UTP commands could be modelled in terms of an ALTG. We argue that this is sufficient to show that it is possible to use an ALTG to represent the evolving state of a program. Hence, it is feasible to construct a $\mathcal{U}_L$ -design from the combined model of locations and designs to represent the $\mathcal{O}$ -calculus.

A more difficult question is whether the $\mathcal{U}_L$ -design is better than the preceding models. It inherits a fully abstract model of heap storage from the UTP model of locations (Chapter 5), which is an improvement on the previous models. However, a cost of this improvement is that the abstract *path-set* representation of a pointer’s address (i.e. the pointer’s value rather than its contents) can change during the execution of a program, such as when a variable (x) containing a pointer is assigned to null. In this case the path-set representing the pointer is updated so that it no longer contains paths that originate from the given variable (i.e. those that are prefixed by $\text{var}.x$). Such changes in value reflect the updates to the ALTG that is being used to model the logical structure of a program’s storage. Therefore, we cannot use existential quantification to represent the value of a pointer, whenever the underpinning ALTG that it is associated with might change.

Example 6.9: We used existential quantification to hide variables from the execution of a subprogram; when the subprogram had completed the variables were then reinstated with the existentially quantified values. Our solution to this issue was to add some extra nodes to the graph to store such previously existentially quantified values; the pointers within these values would then be maintained by the ALTG operations.

□

Whether the complications added by the introduction of the extra nodes to store variables that are temporarily outside the scope of a subprogram is worth the benefit of a fully abstract heap is a subjective matter of judgement. It is likely to depend on the purpose for which the model is being used, as this will determine how frequently two heaps will be compared to determine whether they are equal (or one contains the other). If this is happening frequently, then a fully abstract model of memory is likely to be beneficial, whereas if this is a rare occurrence then it may be easier to compare concrete heap representations. Note that it is straightforward to construct an abstract heap representation from a concrete heap (say for the purposes of such a comparison).

Overall, we argue that it is feasible to construct a $\mathcal{U}_L$ -design from the combined model of locations and designs to represent the $\mathcal{O}$ -calculus. Thus, it is possible to provide an object-based object orientation to the UTP, which has value- and reference-based objects, and a fully abstract model of references; this was our thesis ([Section 1.5](#)).

6.3 Summary

We have outlined how the UTP model of locations presented in [Chapter 5](#) can be extended to the standard UTP theory of designs and our UTP theory of objects ([Chapter 3](#)). In particular, we have shown that it is feasible to use the UTP model locations to provide a fully-abstract model of pointers for our UTP theory of objects. This completes the demonstration of our thesis, which stated that ‘it is possible to provide an object-based object orientation to the UTP with value- and reference-based objects, and a fully abstract model of references’ [Section 1.5](#).

Chapter 7

Discussion

We have extended the Unifying Theories of Programming (UTP) [HH98, WC04] so that it can model object-based object-oriented (OO) programs, which are written in an Abadi–Cardelli-style object calculus [AC96] – specifically our $\mathcal{O}$ -calculus [SG07] (as described in Section 2.2). This provides the $\mathcal{O}$ -calculus with a relational predicate denotational semantics. Further, as the UTP already has several models of concurrency, this encoding provides the potential for adding one or more of these concurrency models to the untyped object calculus (ζ -calculus).

This chapter provides a summary of the work done (Section 7.1), a discussion of *closely* related work (Section 7.2 and Section 7.3), some future possibilities (Section 7.4), some personal reflections (Section 7.5), and the conclusions (Section 7.6). By ‘closely related’ we mean any work that considers using the UTP to model objects or pointers, or any work that we have been inspired by. Work that is not closely related has already been discussed in the introduction (Chapter 1) and Appendix A; this includes a general discussion of OO-concepts and alternative approaches to the formal modelling of object orientation (Section 1.3).

7.1 Summary

We presented four UTP models of an Abadi–Cardelli-style object calculus (the $\mathcal{O}$ -calculus). This $\mathcal{O}$ -calculus directly supports value-based objects, eagerly evaluated functions, references, literal values and their operations, conditional evaluation, local declaration blocks, and sequential composition – but not types, classes, or other class-based notions such as inheritance.

All four models extend the notion of a UTP design. The first three introduce a concrete model of a heap, which mirrors that of the $\mathcal{O}$ -calculus; i.e. a partial map from abstract locations to values. The last model provides an alternative – fully abstract – model of the heap, which is based on a directed graph that has anonymous nodes and distinctly labelled edges. Here, the abstract locations are represented by the, potentially infinite, set of finite paths to a given location node within the Abstract Location Trace Graph (ALTG) from the root node.

Operand stack design: The UTP operand–stack–design ($\mathcal{U}_s$ -design) model of the $\mathcal{O}$ -calculus does not make use of the UTP’s local variable introduction (and completion) feature; instead all temporary values are stored on the operand stack. This simplifies the handling of method and function invocation, as the alphabet of all subprograms is the same; i.e. there are no alphabet issues to consider when determining the weakest fix-point of a procedure (method or function) call. Most of the other operations within the $\mathcal{O}$ -calculus are straightforward expressions, in the sense that they apply a simple function (e.g. addition) to an evaluation of the expressions representing its arguments. Such operations (expressions) are transformed into UTP programs via the command expression (cmdExp_s) constructor, which stores the results of executing the programs representing the operation’s arguments on the operand stack, before applying the appropriate underlying function.

Result variable design: The UTP result–value–design ($\mathcal{U}_R$ -design) replaces the operand stack with a combination of a special result variable (which is similar to an accumulator) and some local variables for storing intermediate results in a calculation (of a command expression). Note, the alphabet at the point of a procedure invocation call is guaranteed to contain only the observational variables; this follows from the observation that the local variables for storing the results of evaluating a command expression’s arguments are hidden from the subprograms that evaluate its arguments.

Constant map design: The UTP constant–map–design ($\mathcal{U}_M$ -design) extends the $\mathcal{U}_R$ -design by using a local variable to represent a procedure’s formal parameter. A nested procedure also *inherits* the parameters of its enclosing parent (grandparent, etc.) procedures, which are similarly represented by UTP local variables. Now as the UTP local variables that represent a procedure’s parameters are visible throughout the definition of its body, they will be visible at the time of a procedure’s invocation. A consequence of this is that the alphabet just prior to the procedure invocation can vary, as distinct procedures may have different parameter sets. This is at odds with our use of the weakest fix-point to model procedure invocation, which requires the alphabets at the point of the procedure invocation call to be the same. The solution we adopted was:

1. To hide the current procedure’s local variables (parameters) as part of the procedure invocation, just prior to the actual procedure call; these variables are then reinstated once the called procedure has terminated.
2. To change the procedure definition from a pair (representing the explicit parameter and the body) to a triple (representing the explicit parameter, the body, and a constant parameter map); here the map is from variable names to values, which once set cannot be changed.

Note that the constant map is only changed as part of a procedure invocation, where the procedure’s explicit parameter is bound in both its own map and the maps of any inner procedure (that does not redeclare that parameter).

Location graph design: The UTP location–graph–design ($\mathcal{U}_L$ -design) uses an ALTG to represent complex relationships between abstract locations and their data. Here, both shareable and containable locations are modelled by a path-set. They differ in that only shareable locations can be dereferenced by a pointer, whose value is the shareable location’s path-set itself. The key point is that containable locations actually contain rather than reference their contents; thus, when they are copied their contents are duplicated rather than shared. This mirrors situations where the whole of a compound value, such as a Pascal **record** or a C++ **struct**, is duplicated on assignment. In general, being able to control the amount of data that gets duplicated on assignment provides a means for directly supporting different levels of containment within a data structure.

One consequence of modelling a pointer’s value as the path-set that defines its location is that it is possible to directly represent the concept of a handle (i.e. a pointer to a pointer). The combination of having direct support for contained locations and pointer values mirrors the features of the $\mathcal{O}$ -calculus.

Garbage collection: The $\mathcal{U}_L$ -design’s model of a heap can only represent reachable locations. Therefore, any completely faithful implementation of this heap model must free memory as soon as it was no longer reachable. This places a significant obligation on the implementation, which may not be desirable, say for efficiency or applicability reasons; e.g. we may wish to use this model of a heap to describe a system that uses a less eager garbage collection scheme. In such cases, we can relax the acceptable implementation criterion, to say that all reachable locations

in the specification have a corresponding location in the implementation (which stores the same value).

Consistency: The $\mathcal{O}$ -calculus’s semantics is defined by a collection of small-step reduction rules, which are deterministic in the sense that at most one rule can be applied to any given $\mathcal{O}$ -calculus term (in its reduction context). The semantics of the corresponding UTP designs are each defined as a relational predicate, which contains a four-member alphabet: two members for representing the input and output values of a program’s heap storage (i.e. Π_{HEAP} and Π_{HEAP}'); and two members representing the input and output values of the program’s result variable, which is either Π_{STK} and Π_{STK}' or Π_{RES} and Π_{RES}' depending on the UTP model being used. Note that there are no input parameters, as these are required to be bound to the program, as we only demonstrate the consistency of closed programs; i.e. programs without free variables.

We consider the $\mathcal{O}$ -calculus and UTP semantics to be consistent so long as the UTP semantics of the $\mathcal{O}$ -calculus term prior to an operational step *equals* that of the UTP semantics of the $\mathcal{O}$ -calculus term after the operational step. This consistency was proven for the UTP operand-stack, result-variable, and constant-map designs; it has not been proven for the location-graph design, which was created to illustrate how our UTP object and location models could be combined.

7.2 Predicative Objects

Within this section we discuss the work on representing objects as (relational) predicates. Having done this, we identify those works that already support a refinement calculus, as this is an important step in enabling the theories to support practical programming techniques. We finish this section with a comparison between the work we have just discussed and that being presented in this dissertation.

This is not the first time object-oriented ideas have been added to (or modelled in) the UTP. Two teams have made significant progress with adding class-based OO-support to the UTP; for convenience we refer to them as the rCOS and Circus teams. They have both used higher-order programming concepts, in the sense that object values essentially contain the methods that can be invoked on them.

rCOS: The rCOS team’s work [HLL02a, HLL02b, HLXS04, HLL06, CHH⁺07] has resulted in the development of the Refinement Calculus of Object Systems (rCOS). Here, an object value is denoted by a unique object-reference. Such references are bound to:

- the name of the class they are implementing (which is unique);
- the name of the class that they have been cast to (which is either the class’s own name or the name of one of its ancestor classes); and
- the values for the class’s – direct or inherited – attributes (i.e. fields).

Each class directly contains those fields and methods that it declares and defines respectively; it also indirectly contains any inherited methods and fields, from the class that it optionally extends. An rCOS OO-program is denoted by a pair, where the first element is a list of class declarations, and the second is the ‘main’ algorithm of a program that may make use of these classes. The semantics of an rCOS program is also split into two parts, which correspond to the two elements within the program pair. First, the class declaration list is used to construct a ‘static’ class-relationship environment. Second, this environment is used to determine the ‘dynamic’ evaluation of the program’s main algorithm.

Circus: The Circus team’s work [CSW03, CSW05, SCS06, CHW06, HCW08] is being used to extend the Circus language [WC02, OCW07], which is a concurrent language for refinement that unifies Z [Spi92, WD96], CSP (Communicating Sequential Processes) [Hoa85, Ros98], and the refinement calculus [Mor94, BvW98] in the UTP. Here, the work on objects and pointers has been done independently, but with a view to merging the results. We discuss the pointer work in Section 7.3.

The object extension to the Circus language introduces a mechanism for declaring classes and their associated methods, prior to the creation of the objects that use them. A significant difference between this work and the rCOS work is that the notion of object references is not bound in with the modelling of the objects themselves. Another significant difference is that methods are not directly contained by their classes. Instead, each method is an entity in its own right, which specifies the class it is directly bound to. This has the affect of allowing the definition of recursive (or mutually recursive) method invocations to be defined directly in terms of a fix-point over methods, rather than a fix-point over classes and their methods.

TCOZ: In addition to these teams, Qin, Dong, and Chin have adapted [HLL02a] to support the Timed Communicating Object-Z (TCOZ) class-based OO-specification language [QDC03]. This work integrates trace-based notions of concurrency, in a similar manner to that in [HH98]. Method invocation of passive methods – i.e. methods that update the local object state – is performed in a similar manner to that of the rCOS team. Active methods are used to interact with other threads of communication (which are called processes); here, the communication event contains the arguments that are being used in the active method invocation.

Predicative objects: Within the more general field of predicative programming [Heh93], another notion of class-based object orientation has been modelled [Kas04, Kas06]. Here, both objects and classes are special forms of a more general predicative construct. In [Kas04] these predicative constructs are called ‘object-specifications’, which are essentially a subset (though technically a subbunch) of the object-values (a named bunch) that conform to a given predicate. Such predicates, are used to define the relationships (e.g. axioms) between the object-values’s methods. Objects and classes place further constraints on the object-specifications. Here, an object belongs to any class whose defining predicate is implied by the object’s defining predicate.

In [Kas06] the object specifications are replaced by ‘theories’ each of which is a named collection of identifiers (called the vocabulary) and a boolean expression (called the axiom of the theory). This is similar to the UTP concept of a theory, where the UTP’s observational variables and healthiness conditions correspond to the vocabulary and the axiom of [Kas06] respectively.

Aside 7.1: The axiom of the theory may use identifiers that are not in its vocabulary; such identifiers are called the external names of the theory, and are used to parametrise a theory. Theories can also be combined (by conjunction), so long as their defining predicates are consistent.

□

Objects and classes are defined in a [Kas06] programming theory – i.e. ‘a theory that describes computations’. Here, an object consists of a pair: a type, which represents the class(es) of object it belongs to; and a record that binds distinctly named attributes to their mutually disjoint addresses. A class is essentially a set of objects that share the same type and attribute names. One consequence of this is that methods are not strictly part of a class, but rather are procedures whose first argument is the ‘self’ object (of a given type of class).

Refinement Calculus: The objective of a refinement calculus [Mor94, BvW98] is to provide a collection of laws that support an engineer when they perform the stepwise transformation of

a specification into an implementation, and the verification that an implementation meets its specification. Here, the usefulness of the refinement calculus can be subjectively measured, by the amount of support it provides to the engineer (e.g. does it provide simple transformations at the source code, rather than semantic model, level?). Technically, what we have done is a refinement calculus – it has a refinement ordering (i.e. implication) and some laws – but it is not useful as such because the laws do not support refinement at the source code level (i.e. $\mathcal{O}$ -calculus level). Both [HLL06] and [Kas06] provide refinement calculi, for their variants of a predicative programming language.

Comparison: Both the class-based UTP and predicative programming models differ from our approach, as each object is considered to be an instance of a class, rather than a class being a special sort of object; i.e. the related approaches are class-based, whereas our approach is object-based. One consequence of this is that we can model classless objects, whereas the other approaches cannot. This opens the possibility of considering prototype-based languages, such as Self [US87, SU95] and JavaScript [Fla06].

Aside 7.2: Within our $\mathcal{O}$ -calculus classes can be modelled as objects that provide certain facilities [AC96], such as a method which creates a new object instance and a repository of methods (or functions) that can be shared with each of its created objects. Different techniques for sharing methods between objects are presented in Appendix B. Note that it is not as straightforward to model an object as a special sort of class (rather than an instance of a class).

□

Having said that our work is not class-based, it is similar to the class-based UTP work in the sense that we also use higher-order concepts to represent methods as designs, whose semantics is determined by an appropriate weakest fix-point construction. Further, like rCOS each of our object models presents a conservative extension to the unifying theory of designs [HLL06]. Here, our extensions are designs that can access the additional observational variables for modelling heap storage and the accumulator (e.g. an operand stack or result variable). Therefore, all the existing laws concerning designs still hold.

Aside 7.3: Class-based languages frequently use the class (and interface) hierarchy to form a type hierarchy. Here, the idea is that an object of a child class can be used wherever a value of a parent class is required. This provides a form of structural typing, where code reuse through inheritance implies sub typing. This is not always what is wanted.

Object-based languages can also be typed (and even statically typed [AC96]). Here, the typing mechanisms tend to be based on the signature of the object (i.e. the collection of members that it contains and their signatures). Our $\mathcal{O}$ -calculus is not typed, as the standard UTP theory of designs that we are embedding it into is not typed. The benefits of either object or class typing could be added in the future, along with their associated constraints on what can be legitimately expressed.

□

Our model of objects is not designed to support a refinement calculus, though it has the standard UTP notion of refinement, namely logical implication. Instead, we were focusing on the task of accurately capturing the operational semantics of $\mathcal{O}$ -calculus in the UTP. Having done this, a future step is to consider how the UTP models can be adapted to support a refinement calculus (as discussed in Section 7.4).

7.3 Abstract Location Model

Within this section we relate our model of locations to those of: the trace-based models inspired by Hoare and He [HH99]; the UTP pointer models based on the notions of an *entity group*

[CHW06] and *automata* [HCW08]; and the rCOS object references [HLL06].

Trace-based modelling: Our model of locations was inspired by Hoare and He’s trace-based model of pointers [HH99]. It introduces the notion of a containable location. This significantly complicates the — already non-trivial — notion of assignment, which in [HH99] is defined in terms of *swinging* the pointer of the assigned location to its new contents. In our model, several contained pointers can be swung at once by an assignment operation, as the contents of:

- a containable location are duplicated on assignment;
- a shareable location are referenced (shared) on assignment;
- a location can include many shareable and containable locations.

The benefit of this extra complexity is that our location model enables the atomicity of assignment to be directly specified (or supported). For example, the copying of a **struct** in C⁺⁺ or a **record** in Pascal can be captured.

Schieder has also adapted Hoare and He’s work on trace-based pointers to provide a weakest precondition semantics for pointers [Sch04]. Here, the object maps have been totalised in order to avoid undefinedness; this leads to the null pointer being modelled by a node that has outbound edges (all of which point to itself). However, like [HH99] it does not support the notion of a containable location.

Entity group modelling: Cavalcanti, Harwood, and Woodcock have an *entity group* [PO04] inspired model of pointers and records [CHW06]. Here, an *entity group* contains the set (equivalence class) of path names that can be used to access the same value, where a path name is either:

1. a *simple name* of a UTP user variable; or
2. a *rooted field name*, which is a simple name extended by a dot-separated sequence of record field labels.

This notion of a path is similar to the one we use, in the sense that both use a sequence of labels to define a route from a given starting point to a location. The main difference is that every location in [CHW06] is potentially shareable, whereas only some of the locations within our model are shareable. Specifically, [CHW06] does not support the notion of a containable location.

Automata modelling: Harwood, Cavalcanti, and Woodcock have an alternative UTP model of pointers, which models the heap in terms of an *automaton* [HCW08]. Their UTP model of pointers uses three observational variables to represent the automaton: one to represent the set of addresses (i.e. paths in [CHW06]) the automaton will accept; one to represent a partial mapping from *terminal* addresses to primitive values; and one to represent an equivalence relation on addresses. A feature of this model is that the value of a non-terminal address – i.e. one that does not contain a primitive value – is a partial function that maps addresses to values (that can be reached from this point). Having said this, there are no operations for assigning such values to variables.

This work also presents a theory of ‘conjunctive healthiness conditions’ that is used to simplify the formal reasoning. They demonstrate that this theory can be combined with the UTP theory of designs, via the existence of a Galois connection, provided that the conjunctive healthiness conditions do not mention the design’s observational variables. We took a different approach, which was to adapt our location model’s healthiness conditions so that they were written in terms of UTP designs (Section 6.1). Another difference is that every location in

[HCW08] is potentially shareable, whereas only some of the locations within our model are shareable. Therefore, like [CHW06], [HCW08] does not support the notion of a containable location.

rCOS object references: In rCOS [HLL06] a non-null object value is modelled as a member of an abstract set of *object identities*, which is the set of references, without the special null-reference (for the null object). The underpinning object itself is denoted by a data-structure that contains its identity (object reference), its constructed and current class-type index (class name), and a map from attributes of its constructed class-type to their values. Further, the ‘dynamic’ observational variable Π is used to store the set of currently allocated objects. The set of objects acts much like our simple heap model – a partial map from abstract locations to values. Here, each object within Π has a unique object identity, which corresponds to our abstract location, and the object itself can be considered the value. The main difference is that our heap can contain any type of value, whereas the rCOS equivalent (i.e. the Π set) can only contain object values.

First class pointer values: A further difference between our model and those of [HH99], [Sch04], [CHW06], [HCW08], and [HLL06] is that we have an explicit notion of a pointer value (i.e. shareable location), as represented by a path-set of the form ‘ $Lp.\$$ ’. One consequence of this is that our model directly supports the notion of a *handle*, which is a pointer to a pointer, whereas the other approaches do not. Here, the second pointer value (path-set) includes a path of the form ‘ $lp.\$. \$$ ’.

7.4 Future Directions

Within this section we consider three ways of progressing our work. First, we discuss how the existing models might be improved or extended to handle additional OO-features (e.g. exception handling). Second, we discuss some possible extensions to the abstract location model. Third, we discuss the possibility of adapting our models so that they can be used to support another programming language’s refinement calculus.

Object modelling: An Abadi–Cardelli-style of object calculus [AC96] can be modelled in the UTP in a variety of different ways. One issue with the models we have presented is that the encoding (compilation) from the $\mathcal{O}$ -calculus to the UTP is essentially one way. It would be interesting to consider whether one can provide a UTP encoding that was more bidirectional, in the sense that it is possible write a UTP program using a straightforward subset of the UTP predicate language, and have this represented as an $\mathcal{O}$ -calculus program via a reversal (inverse) of the compilation function.

Another issue is that the $\mathcal{O}$ -calculus model presented in this dissertation does not include several modern programming language features, such as exception handling and multi-threading (concurrency). Here, both the UTP and the object/lambda calculi have existing models of both exception handling, such as [He08, Pie02], and concurrency, such as [HH98, GH98]. A way forward would be to survey the approaches to modelling these features in both frameworks, and identify which of these models correspond (or almost correspond) to each other.

Location modelling: We have presented both a concrete and an abstract model of locations. What we have not done is to prove that the two models are consistent, so a potential future direction is to prove the consistency of these models; having done this the next step would be to prove the consistency of the $\mathcal{U}_L$ -design against the $\mathcal{O}$ -calculus, as the $\mathcal{U}_L$ -design is underpinned by this abstract location model.

We have also not considered the issues of:

- location ownership and encapsulation (e.g. as presented in [Cla02, NCP99, Rey02]);
- location typing (e.g. augment a location with the type of its contents);
- location visibility (e.g. augment a location with read-only or scope modifiers).

Augmenting the UTP model of locations to handle any of these issues is a potential way forward.

Refinement calculus: Each of our UTP models of the $\mathcal{O}$ -calculus have an implication refinement ordering. This implication ordering could be used to prove laws that define refinement relationships between $\mathcal{O}$ -calculus expressions, which in turn could be used to prove refinement laws for any programming language that defines its semantics in terms of the $\mathcal{O}$ -calculus. Hence, one line of future work could be to provide the $\mathcal{O}$ -calculus with such a refinement calculus. Alternatively, it may be desirable to use one (or more) of our UTP object models to directly prove refinement laws for a programming language whose semantics is defined in terms of the $\mathcal{O}$ -calculus. There is already a significant body of work on OO-refinement calculi; Shield’s thesis [Shi04] provides a good starting point, as the first three chapters provide a survey of this area. Further, his own work provides a refinement calculus for an Abadi–Cardelli-style object calculus that does not contain references.

7.5 Reflection

We now reflect on some of the decisions we took regarding the embedding of the $\mathcal{O}$ -calculus in the UTP. Within our work context the UTP had been proposed as a framework for combining the results of various industrial analysis tools. At that time object orientation had only just started to be addressed within the UTP community, and only for class-based OO-languages. This left object-based object orientation within the UTP unaddressed. Further, as class-based OO-languages were (and still are) dominating the popular languages it appeared that object-based object orientation might be overlooked by other UTP researchers. Therefore, it appeared that this was a challenging area, which was relevant to our job and had the potential to be successfully explored on a part-time basis (where it would not be feasible to compete with a full-time research group). This eventually led to the decision to identify a research issue within this area of study.

Abadi and Cardelli’s work on the ζ -calculus as defined in their book [AC96] had generated a significant amount of interest within the object-based research community. It was being used and adapted by several researchers to meet their own needs. Further, at its core was a small calculus whose core concept was that of an abstract representation of an object. Therefore, this calculus was adopted as the starting point for our modelling of object-orientation within the UTP.

We spent some time considering how the ζ -calculus could be used, and the way in which it had been extended to directly support a range of programming issues. Part of this work is documented in [Appendix B](#), which focuses on the modelling of classes via the sharing of methods. In hindsight, too much time was spent examining what the ζ -calculus could do and support, rather than focusing on how it could be related to the UTP.

One significant concern was how to demonstrate the validity of any UTP model of objects. We already had a deterministic operational semantics for our variant of the ζ -calculus (which we called the $\mathcal{O}$ -calculus); it had been validated by running examples through a term-rewriting machine (which implemented the operational semantics). Therefore, we decided that the UTP semantics of objects should be consistent with this already validated operational semantics. This was one of the contributing factors that led to our decision to embed a variant of the $\mathcal{O}$ -calculus within the UTP, rather than use it to inspire a UTP specific model of objects. Another

significant factor was that a model as fundamental as the ζ -calculus ought to be representable in the UTP, if the UTP enabled different styles of programming to be unified (by relational predicates). In hindsight this led to the UTP being used to define an alternative denotational semantics for the $\mathcal{O}$ -calculus, rather than having a model of objects that was inspired by the $\mathcal{O}$ -calculus. Nevertheless, such an encoding is useful in its own right, as a denotational semantics is often useful for the proof of laws concerning its source language(s).

The precise choice of features that were included in the $\mathcal{O}$ -calculus came from two conflicting objectives. First, to keep the language small, so that the proofs concerning its consistency were feasible. And second, to provide a selection of programming language operators to simplify examples and illustrate that it could be of practical use. In hindsight, it is arguable whether functions should have been included as first class values. We were able to generalise our encoding of methods in a straightforward, and possibly obvious, way to handle both function values and methods, but the cost in time and effort was significant.

7.6 Conclusions

Hoare and He's UTP [HH98, WC04] provides a rich model of programs as relational predicates. Our first contribution is to extend this work with a notion of object-based object orientation, in contrast with the existing class-based models. Some might argue that class-based object orientation dominates the mainstream OO-languages to such an extent that the alternative object-based OO-languages do not need to be addressed. However, the newer features of some of these languages, such as Java's dynamic proxies [FF04, D'A04], enable classes to be generated (compiled) and then used (loaded and linked) at run-time. These sorts of features blur the boundaries between classes and objects. Further, some of the prominent scripting languages, such as JavaScript [Fla06] are incorporating object-based object orientation. Therefore, we would argue that the UTP should support object-based OO-languages, and that our work provides a step in this direction.

Cavalcanti, Harwood, and Woodcock have provided the UTP model with a notion of pointers [CHW06, HCW08]. He, Liu, and Li have also provided a UTP model of pointers, which is limited to object references, where an object's value is its reference [HLL06]. Our second contribution is to provide an alternative model of pointers for the UTP that supports both value-based compound values (e.g. objects) and references, in contrast to existing UTP models with pointers that support only reference-based compound values. We argue that a general UTP model of pointers ought to consider both shareable and containable locations. Such models will provide support for languages like $C^\#$ and the $\mathcal{O}$ -calculus, which have language constructs for building containable locations and handles (i.e. pointers to pointers).

Abadi and Cardelli's ζ -calculi [AC96] provide a prominent formalism of object-based OO-programs, which models programs as objects. In general, we believe that the UTP ought to be able to model all significant computer programming language formalisms, in order for it to be considered a unifying theory. Hence, our third contribution is to model this Abadi-Cardelli notion of an object in the UTP, and thus demonstrate that it can unify this style of object formalism.

Finally, we observe that our work could be used to provide a denotational semantics for any programming language whose semantics was defined in terms of the $\mathcal{O}$ -calculus. Here, the designer of the programming language could have the benefit of using an operational OO-semantics to define his language, and get our four corresponding denotational UTP models for 'free'. This would provide the designer with a choice of formal models for proving properties (and laws) concerning his programming language; which may well be easier in a denotational rather than operational semantic setting.

References

- [AC96] Martin Abadi and Luca Cardelli. *A Theory of Objects*. Springer, 1996.
- [AL98] Martin Abadi and K. Rustan M. Leino. A logic of object-oriented programs. Research Report 161, Systems Research Center, Digital, September 1998.
- [AZ96] Davide Ancona and Elena Zucca. An algebraic approach to mixins and modularity. In *5th International Conference on Algebraic and Logic Programming*, pages 179–193. Springer Verlag, 1996. 1139 in Lecture Notes in Computer Science.
- [AZ98] Davide Ancona and Elena Zucca. A theory of mixin modules: Basic and derived operators. *Mathematical Structures in Computer Science*, 8(4):401–446, 1998. Cambridge University Press.
- [AZ02] D. Ancona and E. Zucca. A calculus of module systems. *Journal of Functional Programming*, 11(2):91–132, March 2002.
- [BC90] Gilad Bracha and William Cook. Mixin-based inheritance. In Norman Meyrowitz, editor, *ECOOP 1990 – European Conference on Object-Oriented Programming*, pages 303–311, Ottawa, Canada, 1990. ACM Press.
- [BC99] Michele Bugliesi and Silvia Crafa. Object calculi for dynamic messages. In *FOOL 6 – Sixth International Workshop on the Foundations of Object-Oriented Languages*, 1999. Available via the FOOL web-site <http://www.cs.hmc.edu/~stone/FOOL/-index.html>.
- [BCG00] Roland Backhouse, Roy Crole, and Jeremy Gibbons, editors. *Algebraic and Coalgebraic Methods in the Mathematics of Program Construction*, volume 2297 of *Lecture Notes in Computer Science*. Springer-Verlag, April 2000. International Summer School and Workshop, Oxford, Revised Lectures.
- [BCP99] Kim B. Bruce, Luca Cardelli, and Benjamin C. Pierce. Comparing object encodings. *Information and Computation*, 155(1/2):108–133, 1999.
- [BDMN79] Graham M. Birtwistle, Ole-Johan Dahl, Bjorn Myhrhaug, and Kristen Nygaard. *Simula Begin*. Studentlitteratur (Lund, Sweden), Bratt Institute Fuer Neues Lerne (Goch, FRG), Chartwell-Bratt Ltd. (Kent, England), 2nd edition, 1979.
- [BL05] Christie Bolton and Gavin Lowe. A hierarchy of failures-based models: theory and application. *Theoretical Computer Science*, 330(3):407–438, 2005.
- [BMvW00] Ralph-Johan Back, Anna Mikhajlova, and Joakim von Wright. Class refinement as semantics of correct object substitutability. Technical Report 333, Turku Centre for Computer Science, 2000.
- [Bol02] Christie Bolton. *On the refinement of state-based and event-based models*. PhD thesis, Oxford University, 2002.

- [BPG00] Michele Bugliesi and Santiago Pericas-Geersten. Depth subtyping and type inference for object calculi. In *FOOL 7 – Seventh International Workshop on the Foundations of Object-Oriented Languages*, 2000. Available via the FOOL web-site <http://www.cs.hmc.edu/~stone/FOOL/index.html>.
- [BPS99] Viviana Bono, Amit Patel, and Vitaly Shmatikov. A core calculus of classes and mixins. In *ECOOP 1999 – 13th European Conference on Object-Oriented Programming*, volume 1628 of *Lecture Notes in Computer Science*, pages 43–66. Springer, 1999.
- [Bra92] Gilad Bracha. *The Programming Language Jigsaw: Mixins, Modularity and Multiple Inheritance*. PhD thesis, University of Utah, 1992.
- [Bru96] Kim Bruce. Typing in object-oriented languages: Achieving expressiveness and safety. author’s web site, September 1996. Report.
- [Bru02] Kim B. Bruce. *Foundations of Object-Oriented Languages: Types and Semantics*. MIT Press, 2002.
- [BvW98] Ralph-Johan Back and Joakim von Wright. *Refinement Calculus: A Systematic Introduction*. Springer-Verlag, 1998. Graduate Texts in Computer Science.
- [Car88] Luca Cardelli. A semantics of multiple inheritance. *Information and Computation*, 76(2/3):138–164, 1988.
- [Cas97] Giuseppe Castagna. *Object-Oriented Programming: A Unified Foundation*. Birkhauser Verlag AG, 1997.
- [CD02] Daev Clarke and Sophia Drossopoulou. Ownership, encapsulation and the disjointness of type and effect. In *OOPSLA 2002 – Object-Oriented Programming Systems, Languages and Applications*, pages 292–310. ACM, 2002.
- [CG98] Luca Cardelli and Andrew Gordon. Mobile ambients. In Maurice Nivat, editor, *FoSSaCS 1998 – First International Conference Foundations of Software Science and Computation Structure*, volume 1378 of *Lecture Notes in Computer Science*, pages 140–155, 1998.
- [CHH⁺07] Zhenbang Chen, Abdel Hakim Hannousse, Dang Van Hung, Istvan Knoll, Xiaoshan Li, Zhiming Liu, Yang Liu, Qu Nan, Joseph C. Okika, Anders P. Ravn, Volker Stolz, Lu Yang, and Naijun Zhan. Modelling with relational calculus of object and component systems - rCOS. In Andreas Rausch, Ralf Reussner, Raffaella Mirandola, and Frantisek Plasil, editors, *CoCoME – The Common Component Modeling Example: Comparing Software Component Models*, volume 5153 of *Lecture Notes in Computer Science*, pages 116–145, 2007.
- [Chu36] Alonzo Church. An unsolvable problem of elementary number theory. *American Journal of Mathematics*, 58:345–363, 1936.
- [Chu40] Alonzo Church. A formulation of the simple theory of types. *Journal of Symbolic Logic*, 5:56–68, 1940.
- [Chu85] Alonzo Church. *The Calculi of Lambda Conversion. (AM-6) (Annals of Mathematics Studies)*. Princeton University Press, Princeton, NJ, USA, 1985.
- [CHW06] Ana Cavalcanti, Will Harwood, and Jim Woodcock. Pointers and records in the Unifying Theories of Programming. In Steve Dunne and Bill Stoddart, editors, *First International Symposium on Unifying Theories of Programming*, number 4010 in *Lecture Notes in Computer Science*. Springer-Verlag, 2006.

- [Cla01] David Clarke. An object calculus with ownership and containment (extended abstract). In *FOOL 8 – Eighth International Workshop on the Foundations of Object-Oriented Languages*, 2001. Available via the FOOL web-site <http://www.cs.hmc.edu/~stone/FOOL/index.html>.
- [Cla02] David Clarke. *Object Ownership and Containment*. PhD thesis, University of New South Wales, October 2002.
- [CM92] R.C.H. Connor and R. Morrison. “Subtyping without tears”. In *15th Australian Computer Science Conference*, pages 209–225, Hobart, Tasmania, 1992.
- [CN00] Ana Cavalcanti and David A. Naumann. Simulation and class refinement for Java. In S. Drossopoulou, S. Eisenbach, B. Jacobs, G. T. Leavens, P. Müller, and A. Poetzsch-Heffter, editors, *ECOOP Workshop on Formal Techniques for Java Programs*. Technical Report 269, Fernuniversität Hagen, 2000. Available from <http://www.informatik.fernuni-hagen.de/pi5/publications.html>.
- [CN02] Ana Cavalcanti and David A. Naumann. Forward simulation for data refinement of classes. In L. Eriksson and P. A. Lindsay, editors, *FME 2002 – Formal Methods Europe*, volume 2391 of *Lecture Notes in Computer Science*, pages 471–490. Springer-Verlag, 2002.
- [Coo89] William R. Cook. *A Denotational Semantics of Inheritance*. PhD thesis, Brown University, May 1989. Tech Report CS-89-33.
- [CSW03] A. L. C. Cavalcanti, A. C. A. Sampaio, and J. C. P. Woodcock. A unified language of classes and processes. In *St Eve: State-Oriented vs. Event-Oriented Thinking in Requirements Analysis, Formal Specification and Software Engineering*, Satellite Workshop at FM’03, 2003.
- [CSW05] A. L. C. Cavalcanti, A. C. A. Sampaio, and J. C. P. Woodcock. Unifying Classes and Processes. *Software and System Modelling*, 4(3):277–296, 2005.
- [CW85] Luca Cardelli and Peter Wegner. On understanding types, data abstraction, and polymorphism. *Computing Surveys*, 17(4):471–522, December 1985.
- [D’A04] Lara D’Abreo. Java dynamic proxies: One step from aspect-oriented programming. Technical Report 21463, DevX, July 2004. <http://www.devx.com/Java/Article/21463>.
- [DG87] Linda G. DeMichiel and Richard P. Gabriel. The Common Lisp Object System: An overview. In *ECOOP 1987 - European Conference for Object-Oriented Programming*, pages 151–170. Springer-Verlag, June 1987. LNCS, vol. 276.
- [DP02] B.A. Davey and H.A. Priestley. *Introduction to Lattices and Order*. Cambridge University Press, 2nd edition, 2002.
- [DS90] E.W. Dijkstra and C.S. Scholten. *Predicate Calculus and Program Semantics*. Texts and Monographs in Computer Science. Springer-Verlag, 1990.
- [FF98] Robert Bruce Findler and Mathew Flatt. Modular object-oriented programming with units and mixins. In *3rd ACM International Conference on Functional Programming*, pages 94–104. ACM, September 1998.
- [FF04] Ira R. Forman and Nate Forman. *Java Reflection in Action*. Manning Publications Co., 2004.

- [FKF98] Matthew Flatt, Shriram Krishnamurthi, and Matthias Felleisen. Classes and mixins. In *POPL 1998 – 25th ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages*, pages 171–183. ACM, 1998.
- [Fla06] David Flanagan. *JavaScript: The Definitive Guide*. O'Reilly & Associates, 5th edition, August 2006.
- [FR00] Kathleen Fisher and John Reppy. Inheritance-based subtyping. In *FOOL 7 – Seventh International Workshop on the Foundations of Object-Oriented Languages*, 2000. Available via the FOOL web-site <http://www.cs.hmc.edu/~stone/FOOL/-index.html>.
- [Fse05] Formal Systems (Europe) Ltd. *Failures-Divergence and Refinement: FDR2 User Manual*, 6 edition, June 2005. Available from www.fsel.com/fdr_2_manual.html.
- [GH98] Andrew D. Gordon and Paul D. Hankin. A concurrent object calculus: Reduction types. *Electronic Notes in Theoretical Computer Science*, 16(3):248–264, 1998.
- [GHJV94] Erich Gamma, Richard Helm, Ralph E. Johnson, and John Vlissides. *Design Patterns*. Addison-Wesley, 1994.
- [GHL99] Andrew D. Gordon, Paul D. Hankin, and Søren B. Lassen. Compilation and equivalence of imperative objects. *Journal of Functional Programming*, 9(4):373–426, 1999.
- [GJS96] James Gosling, Bill Joy, and Guy Steele. *Java Language Specification*. The Java Series. Addison Wesley, August 1996.
- [GM94] Carl A. Gunter and John C. Mitchell, editors. *Theoretical Aspects of Object-Oriented Programming: Types, Semantics, and Language Design*. MIT Press, 1994.
- [Gol03] Michael Goldsmith. Personal correspondence concerning the semantics of CSP and its application to the development of FDR, April 2003.
- [GR89] Adele Goldberg and David Robson. *Smalltalk 80: The Language*. Computer Science. Addison-Wesley, 1989.
- [GTWW77] J. A. Goguen, J. W. Thatcher, E. G. Wagner, and J. B. Wright. Initial algebra semantics and continuous algebras. *Journal of the ACM*, 24(1):68–95, January 1977.
- [HCW08] Will Harwood, Ana Cavalcanti, and Jim Woodcock. A theory of pointers for the UTP. In John S. Fitzgerald, Anne Elisabeth Haxthausen, and Hüsni Yenigün, editors, *ICTAC 2008 – 5th International Colloquium on Theoretical Aspects of Computing*, volume 5160 of *Lecture Notes in Computer Science*, pages 141–155. Springer, 2008.
- [He08] Jifeng He. Transaction Calculus. In Andrew Butterfield, editor, *UTP 2008 – Second International Unifying Theories of Programming Symposium*, 2008. To appear in volume 5713 of Springer-Verlag's *Lecture Notes in Computer Science*.
- [Heh93] Eric C.R. Hehner. *A Practical Theory of Programming*. Springer-Verlag, 1993. Electronic edition freely available on line from: www.cs.utoronto.ca/~hehner/aPToP.
- [HH98] C.A.R. Hoare and J. He. *Unifying Theories of Programming*. Computer Science. Prentice Hall, 1998.

- [HH99] C.A.R. Hoare and Jifeng He. A trace model for pointers and objects. In *ECOOP 1999 – 13th European Conference on Object-Oriented Programming*, pages 1–17, 1999.
- [HHJT98] Ulrich Hensel, Marieke Huisman, Bart Jacobs, and Hendrik Tews. Reasoning about classes in object-oriented languages: Logical models and tools. In Chris Hankin, editor, *European Symposium on Programming*, pages 105–121. LNCS 1381, 1998.
- [HL02] Tom Hirschowitz and Xavier Leroy. Mixin modules in a call-by-value setting. In Daniel Le Métayer, editor, *European Symposium on Programming*, volume 2305 of *Lecture Notes in Computer Science*, pages 6–20. Springer, 2002.
- [HLL02a] Jifeng He, Zhiming Liu, and Xiaoshan Li. A relational model for specification of object-oriented systems. Technical Report 262, UNU/IIST, P.O. Box 3058, Macau, October 2002.
- [HLL02b] Jifeng He, Zhiming Liu, and Xiaoshan Li. Towards a refinement calculus for object systems. In *ICCI 2002 – 1st IEEE International Conference on Cognitive Informatics*, pages 69–76. IEEE Computer Society, 2002.
- [HLL06] Jifeng He, Xiaoshan Li, and Zhiming Liu. rcos: A refinement calculus of object systems. *Theoretical Computer Science*, 365(1-2):109–142, 2006.
- [HLXS04] Jifeng He, Zhiming Liu, Li Xiaoshan, and Qin Shengchao. A relational model for object-oriented designs. In Wei-Ngan Chin, editor, *APLAS 2004 – Second Asian Symposium on Programming Languages and Systems*, volume 3302 of *Lecture Notes in Computer Science*, pages 415–436. Springer, 2004.
- [Hoa85] C.A.R. Hoare. *Communicating Sequential Processes*. Computer Science. Prentice Hall, 1985.
- [Hoa99] C.A.R. Hoare. A theory of programming: Denotational, algebraic and operational semantics. Technical report (in proceedings), Microsoft Research, November 1999.
- [Hog91] John Hogg. Islands: Aliasing protection in object-oriented languages. In *OOPSLA 1991 – 6th Conference on Object Oriented Programming Systems Languages and Applications*, volume 26 of *SIGPLAN Notices*, pages 271–285, 1991.
- [ISO02] ISO. *Information Technology – Z Formal Specification Notation – Syntax, Type System and Semantics*, 2002. ISO/IEC 13568.
- [Jac96] Bart Jacobs. Objects and classes, co-algebraically. In B. Freitag, C.B. Jones, C. Lengauer, and H.-J. Schek, editors, *Object-Oriented Programming with Parallelism and Persistence*, chapter 6, pages 83–103. Kluwer Academic Publishers, 1996.
- [Jac98] Bart Jacobs. Coalgebraic reasoning about classes in object-oriented languages. In Horst Reichel Bart Jacobs, Larry Moss and Jan Rutten, editors, *International Workshop on Coalgebraic Methods in Computer Science*, volume 11 of *Electronic Notes in Theoretical Computer Science*. Elsevier Science Publishers, 1998.
- [Jon93] C. B. Jones. Process-algebraic foundations for an object-based design notation. Technical Report UMCS-93-10-1, University of Manchester, October 1993.
- [Kas04] Ioannis T. Kassios. Objects as predicates. Technical report, Computer Systems Research Group, University of Toronto, 2004.

- [Kas06] Ioannis T. Kassios. *A Theory of Object Oriented Refinement*. PhD thesis, University of Toronto, 2006.
- [KR88] Brian W. Kernighan and Dennis M. Ritchie. *The C Programming Language*. Prentice Hall, 2nd edition, 1988.
- [LA03] Mark Lutz and David Ascher. *Learning Python*. O'Reilly & Associates, 2nd edition, December 2003.
- [LBE00] K. Lano, J. Bicarregui, and A. Evans. Structured axiomatic semantics for UML models. In *Rigorous Object-Oriented Methods*, Workshops in Computing. BCS, 2000.
- [Lem00] Lemma 1 Ltd. *ProofPower Tutorial Manual*, 2000. Available from www.lemma-one.com/ProofPower/doc/doc.html.
- [Lew95] Simon Lewis. *The Art and Science of Smalltalk*. Prentice-Hall, 1995.
- [Lib05] Jesse Liberty. *Programming C#*. O'Reilly Media, 4th edition, 2005.
- [Lum99] Mark Lumpe. *A Pi-Calculus Based Approach for Software Composition*. PhD thesis, Universitat Bern, 1999.
- [LW01] Barbara H. Liskov and Jeannette M. Wing. Behavioral subtyping using invariants and constraints. In Howard Bowman and John Derrick, editors, *FMDP 2001 – Formal Methods for Distributed Processing, an Object Oriented Approach*, pages 254–280. Cambridge University Press, 2001. Also available as a Carnegie Mellon University technical report CMU-CS-99-156.
- [Mer00] Massimo Merro. *Locality in the Pi calculus and applications to distributed objects*. PhD thesis, l'Ecole des Mines de Paris, 2000.
- [Mey96] Bertrand Meyer. The many faces of inheritance: a taxonomy of taxonomy. *Computer*, 29(5):105–108, May 1996.
- [Mey97] Bertrand Meyer. *Object-Oriented Software Construction*. Prentice Hall, 2nd edition, 1997.
- [Mil80] Robin Milner. *A Calculus of Communicating Systems*. Springer-Verlag, 1980.
- [Mil99] Robin Milner. *Communicating and Mobile Systems: the π -Calculus*. Cambridge University Press, 1999.
- [Mor94] Carroll Morgan. *Programming from Specifications*. Computer Science. Prentice Hall, 2nd edition, 1994.
- [MT92] K. Meinke and J.V. Tucker. Universal algebra. In S. Abramsky, D. Gabbay, and T. Maibaum, editors, *Handbook of Logic in Computer Science*, volume I: Mathematical Structures, pages 189–411. Oxford University Press, 1992.
- [NCP99] James Noble, Daev Clarke, and John Potter. Object ownership for dynamic alias protection. In *TOOLS 1999 – 32nd International Conference on Technology of Object-Oriented Languages and Systems*, pages 176–187. IEEE Computer Society, 1999.
- [Nor92] Else Nordhagen. Pi-calculus semantics of a Smalltalk-like language. In *Norsk Informatikk Konferanse*, 1992. Available from <http://folk.uio.no/lc/lit-LC.html>.

- [Nor98] Else K. Nordhagen. *A Computational Framework for Verifying Object Component Substitutability*. PhD thesis, University of Oslo, November 1998.
- [NVP98] James Noble, Jan Vitek, and John Potter. Flexible alias protection. In *ECOOP 1998 – 12th European Conference on Object Oriented Programming*, pages 158–185, 1998. volume 1445 of *Lecture Notes in Computer Science*.
- [OCRZ03] Martin Odersky, Vincent Cremet, Christine Rockl, and Matthias Zenger. A nominal theory of objects with dependent types. In *ECOOP 2003 – 17th European Conference on Object Oriented Programming*, volume 2743 of *Lecture Notes in Computer Science*, pages 201–224, 2003.
- [OCW07] M. V. M. Oliveira, A. L. C. Cavalcanti, and J. C. P. Woodcock. A UTP Semantics for *Circus*. *Formal Aspects of Computing*, 21(1):3–32, 2007.
- [OSV08] Martin Odersky, Lex Spoon, and Bill Venners. *Programming in Scala: A comprehensive step-by-step guide*. Artima, 2008.
- [Par05] Matthew Parkinson. *Local Reasoning for Java*. PhD thesis, University of Cambridge, 2005.
- [PB08] Matthew J. Parkinson and Gavin M. Bierman. Separation logic, abstraction and inheritance. In George C. Necula and Philip Wadler, editors, *POPL 2008 – 35th ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages*, pages 75–86. ACM, 2008.
- [Pie02] Benjamin C. Pierce. *Types and Programming Languages*. MIT Press, 2002.
- [PNCB06] Alex Potanin, James Noble, Dave Clarke, and Robert Biddle. Generic ownership for Generic Java. In Peri L. Tarr and William R. Cook, editors, *Proceedings of the 21st Annual ACM SIGPLAN Conference on Object-Oriented Programming, Systems, Languages, and Applications*, pages 311–324. ACM, October 2006.
- [PO04] R.F. Paige and J.S. Ostroff. ERC: an object-oriented refinement calculus for Eiffel. *Formal Aspects of Computing*, 16(1):51–79, April 2004.
- [Pol01] Erik Poll. A coalgebraic semantics of subtyping. *Theoretical Informatics and Applications*, 35(1):61–82, 2001.
- [PT97] Benjamin Pierce and David Turner. Pict: A programming language based on the pi-calculus. Technical report, Indiana University, 1997.
- [QDC03] Shengchao Qin, Jin Song Dong, and Wei-Ngan Chin. A semantic foundation for TCOZ in unifying theories of programming. In *FME 2003: Formal Methods*, volume 2805 of *Lecture Notes in Computer Science*, pages 321–340. Springer Berlin / Heidelberg, 2003.
- [Qui40] Willard Van Orman Quine. *Mathematical Logic*. Harvard University Press, first edition, 1940. Revised edition, first printed 1951, corrected inconsistencies in the axioms of class theory.
- [Ree00] Dafydd L. L. Rees. *An Algebraic Theory of Software Interfaces*. PhD thesis, Swansea University, September 2000. Revised March 2001.
- [Rei95] Horst Reichel. An approach to object semantics based on terminal co-algebras. *Mathematical Structures in Computer Science*, 5(2):129–152, 1995.

- [Rey74] John C. Reynolds. Towards a theory of type structure. In *Colloque sur la Programmation*, volume 19 of *Lecture Notes in Computer Science*, pages 408–425. Springer-Verlag, 1974.
- [Rey75] John C. Reynolds. User-defined types and procedural data structures as complementary approaches to data abstraction. In Stephen A. Schuman, editor, *New Directions in Algorithmic Languages*, pages 157–186, 1975. Reprinted in: David Gries, editor, *Programming Methodology, A Collection of Papers by Members of IFIP Working Group 2.3*, 1978, Springer-Verlag, pages 309–317; and also in Carl A. Gunter and John C. Mitchel, editors, *Theoretical Aspects of Object-Oriented Programming*, pages 13–24, MIT Press, 1994.
- [Rey02] John C. Reynolds. Separation logic: A logic for shared mutable data structures. In *17th Annual IEEE Symposium on Logic in Computer Science*, pages 55–74. IEEE Computer Society, 2002.
- [Ros98] A.W. Roscoe. *The Theory and Practice of Concurrency*. Computer Science. Prentice Hall, 1998.
- [SB98] Yannis Smaragdakis and Don Batory. Implementing layered designs with mixin layers. In *12th European Conference on Object-Oriented Programming (ECOOP 98)*, number 1445 in *Lecture Notes in Computer Science*. Springer, 1998.
- [SB02] Yannis Smaragdakis and Don Batory. Mixin layers: An object-oriented implementation technique for refinements and collaboration-based designs. *ACM Transactions on Software Engineering and Methodology*, 11(2):215–255, April 2002.
- [Sca98] Bryan Scattergood. *The Semantics and Implementation of Machine-Readable CSP*. PhD thesis, University of Oxford, 1998.
- [Sch04] Birgit Schieder. Pointer theory and weakest preconditions without addresses and heap. In Dexter Kozen and Carron Shankland, editors, *MPC 2004 – 7th International Conference on the Mathematics of Program Construction*, volume 3125 of *Lecture Notes in Computer Science*, pages 357–380. Springer, 2004.
- [SCS06] T. Santos, A. L. C. Cavalcanti, and A. C. A. Sampaio. Object-Orientation in the UTP. In S. Dunne and B. Stoddart, editors, *UTP 2006 – First International Symposium on Unifying Theories of Programming*, volume 4010 of *LNCS*, pages 20–38. Springer-Verlag, 2006.
- [Sek96] E. Sekerinski. A type-theoretic basis for an object-oriented refinement calculus. In S. J. Goldack and S. J. H. Kent (Eds.), editors, *Formal Methods and Object Technology*. Springer-Verlag, 1996.
- [SG07] Michael Anthony Smith and Jeremy Gibbons. Unifying Theories of Objects. In Jim Davies and Jeremy Gibbons, editors, *IFM 2007 – Integrated Formal Methods Conference*, volume 4591 of *Lecture Notes in Computer Science*, pages 599–618. Springer-Verlag, July 2007.
- [SG08] Michael Anthony Smith and Jeremy Gibbons. Unifying Theories of Locations. In Andrew Butterfield, editor, *UTP 2008 – Second International Unifying Theories of Programming Symposium*, 2008. To appear in volume 5713 of Springer-Verlag’s *Lecture Notes in Computer Science*.
- [Shi04] Jamie Shield. *Towards an Object-Oriented Refinement Calculus*. Doctor of philosophy, University of Queensland, 2004.

- [Sim95] Anthony James Howard Simons. *A Language with Class: The Theory of Classification Exemplified in an Object-Oriented Programming Language*. PhD thesis, University of Sheffield, December 1995.
- [Smi95] Graeme Smith. A logic for Object-Z (additional rules). Technical report 95-26, Software Verification Research Centre, School of Information Technology, The University of Queensland, Brisbane 4072. Australia, September 1995.
- [Smi99] Michael Anthony Smith. Type-based assessment of Java's bytecode. Master's thesis, Oxford University, February 1999.
- [SP94] Martin Steffen and Benjamin Pierce. Higher-order subtyping: Revised version with fun subtyping. Technical Report ECS-LFCS-94-280, University of Edinburgh, February 1994.
- [Spi92] Mike Spivey. *The Z Notation: A Reference Manual*. Computer Science. Prentice Hall, 2nd edition, 1992.
- [Str91] Bjarne Stroustrup. *The C++ Programming language*. Addison Wesley, 2nd edition, August 1991.
- [SU95] Randall B. Smith and David Ungar. Programming as an experience: The inspiration for Self. In *ECOOP 1995 – 9th European Conference on Object-Oriented Programming*, volume 952 of *Lecture Notes in Computer Science*, pages 303–330, 1995.
- [SW01] Davide Sangiorgi and David Walker. *The π -calculus: A Theory of Mobile Processes*. Cambridge University Press, 2001.
- [Tar55] Alfred Tarski. A lattice-theoretical fixpoint theorem and its applications. *Pacific Journal of Mathematics*, 5(2):285–309, 1955.
- [Tur95] David Turner. *The Polymorphic Pi-calculus: Theory and Implementation*. PhD thesis, University of Edinburgh, 1995.
- [US87] David Ungar and Randall B. Smith. Self: The power of simplicity. In *OOPSLA 1987 – Object-Oriented Programming Proceedings of the 2nd Annual ACM SIGPLAN Conference on Object-Oriented Programming, Systems, Languages, and Applications*, pages 227–241, 1987. Published as SIGPLAN Notices 22(12), December, 1987. Also published in *Lisp and Symbolic Computation* 4(3), Kluwer Academic Publishers, June, 1991.
- [WC02] Jim Woodcock and Ana Cavalcanti. The semantics of Circus. In *2nd Z and B Conference: Formal Specification and Development in Z and B*, volume 2272 of *Lecture Notes in Computer Science*, pages 184–208. Springer-Verlag, 2002.
- [WC04] J. C. P. Woodcock and A. L. C. Cavalcanti. A Tutorial Introduction to Designs in Unifying Theories of Programming. In *IFM 2004 – Integrated Formal Methods*, volume 2999 of *Lecture Notes in Computer Science*, pages 40–66. Springer-Verlag, 2004. Invited tutorial.
- [WD96] Jim Woodcock and Jim Davies. *Using Z: Specification, Refinement and Proof*. Computer Science. Prentice Hall, 1996.
- [Wei01] Zhuangjian Wei. *A framework for studying the semantics of object-oriented programs*. PhD thesis, Temple University, 2001.

References

- [WO99] Jeannette M. Wing and John Ockerbloom. Respectful type converters for mutable types. Technical Report CMU-CS-99-142, Carnegie Mellon University, Pittsburgh PA 15213-389, June 1999.

Appendix A

Concepts and Terminology

This appendix provides a slightly more detailed presentation of our understanding of object-oriented (OO) concepts, semantics, and aliasing. The concept map in [Figure A.1](#) illustrates our view of how the typical OO-concepts are related. Essentially, we group the OO-concepts into four groups that contain illustrative sub-concepts and have links to their nearest *neighbouring* concept groups.

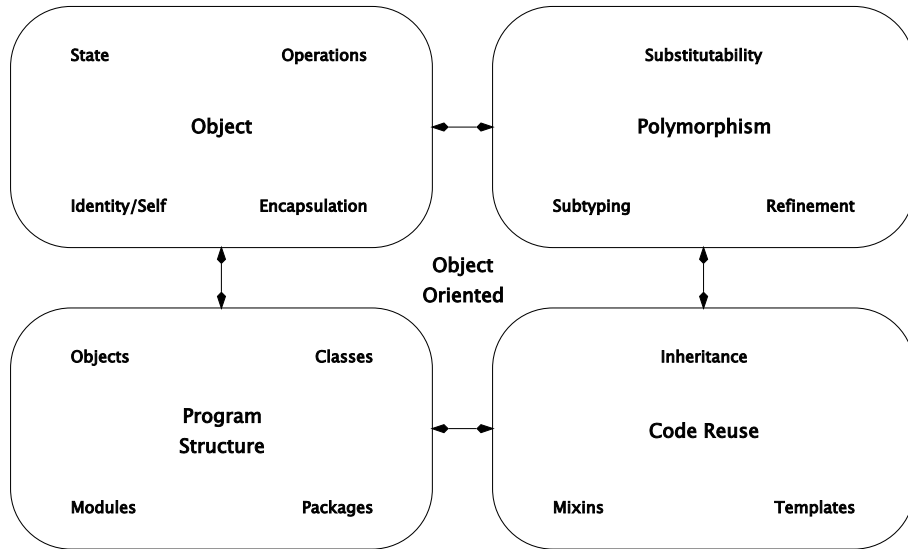


Figure A.1: Object-oriented concept map

These OO-concepts fit into a semantic framework, which may have a variety of characteristics, such as the *single threaded* assumption. In order to discuss these semantic issues, it is useful to organise them as illustrated in [Figure A.2](#), where some of the entries in the diagram contain a brief context setting comment as represented by *slanted* text. This time, the ring presentation is more for convenience of explanation rather than demonstrating a genuine closeness of conceptual relationship. Having said this, the semantics of *concurrency* builds on that of the *state values* and *observable behaviour*.

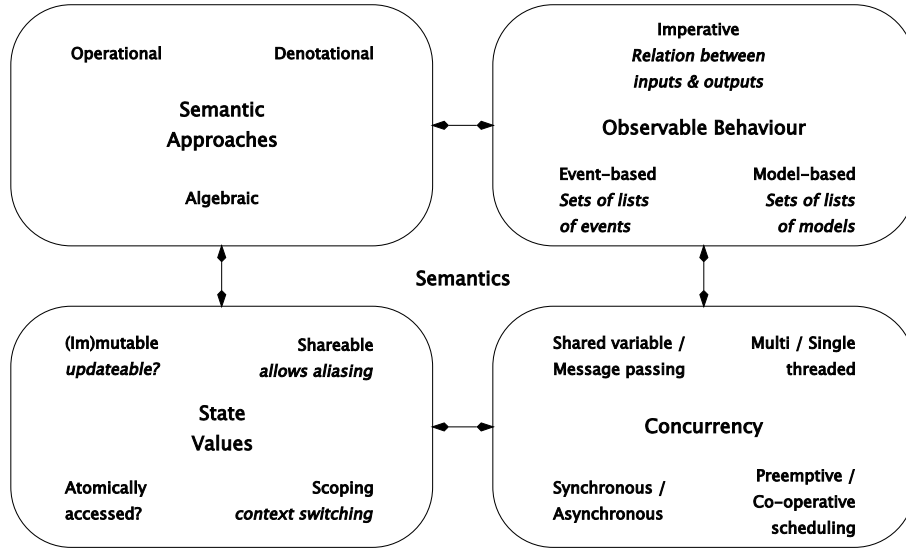


Figure A.2: Semantic concept map

The remainder of this appendix explores these ideas in more detail, and also briefly discusses aliasing.

A.1 Semantics

A language's *semantics* provides its constructs (*terms*), both individually and collectively, with a meaning; this can be anything from an informal description to a fully formal mathematical definition.

Example A.1: The JVM (Java virtual machine) has a semi-formal definition, where an informal description is augmented with some formal notation. Here, a simple list-based model of a method's operand stack is used to aid in the clarity of defining the effect of executing a bytecode instruction.

□

Frequently, the semantics of a language is given in terms of another language or model that is already understood, such as the untyped lambda calculus (λ -calculus). However, this need not be the case, as it is sometimes appropriate to describe the semantics of a language in terms of an abstract mathematical model. Such models can then be validated, experimentally, against the world that they are supposed to represent.

Example A.2: The semantics of the λ -calculus can be defined in terms of a set of operational rewrite (*reduction*) rules [Pie02]. This is augmented by what is known as an *evaluation scheme* which states when each rule can be applied.

□

Each of the concept-groups presented in Figure A.2 are now discusses in more detail.

A.1.1 Semantic approaches

There are three basic approaches to the formalisation of semantics [Pie02, Hoa99, Ros98]:

Operational semantics – Here the semantics is given in terms of how each operation affects the state of the system, which is sometimes referred to as the system's environment. Two common approaches are the *virtual machine*¹ and *reduction-rule* approaches. The first of

¹ Virtual machines are sometimes referred to as abstract (state) machines and vice-versa.

these approaches models the environment (or state) as a *virtual machine* and defines each operation as a partial mapping from current environment to the next (possibly updated) environment.² The second approach models the state of the system as a single – often compound – term, and its operations are defined in terms of a collection of *reduction rules* along with an evaluation scheme for indicating when each rule can be applied.

Denotational semantics – Here the semantics of the source language is given in terms of existing semantic models (domains), which may be captured by a target language. Hence, compilers essentially provide a denotational semantics for their source languages, though it may not be what is expected (or wanted). This is why denotational semantics are often targeted at well understood abstract machines or calculi, such as the λ -calculus. However, in general any compatible formal collection of semantic domains may be used. Note that the chosen target domains may immediately indicate whether certain source language properties are (or are not) available.

Algebraic semantics – Here the semantics is given by a collection of *laws* (equations) that specify when two constructs in the language are equivalent. Such semantics are very useful in the high level transformation of a design, whether this be for intuitive understanding, efficiency, modularisation, etc.

Before continuing, it is worth considering the following example, which illustrates the distinction between semantic models and approaches.

Example A.3: The CSP (Communicating Sequential Processes) language [Hoa85, Ros98], has three standard semantic interpretations (models), namely the *trace*, the *stable-failures*, and the *failure-divergence* models, which represent different granularities of observable behaviour. These models form a partial order, where valid refinements³ in either the stable-failures or failure-divergence models are also valid refinements of the traces model.

The CSP *failure-divergence* model has three semantically equivalent presentations, which use the operational, denotational, and algebraic approaches respectively. Note that being able to prove (demonstrate) results in one approach and apply them in another has been invaluable in the efficient construction of the CSP FDR (Failures-Divergence and Refinement) model checking tool [Fse05, Sca98, Ros98, Gol03].

□

A.1.2 State values

Main stream OO-programming languages, such as Smalltalk [GR89], Eiffel [Mey97], Python [LA03], C++ [Str91], C# [Lib05], and Java [GJS96], are essentially defined in terms of, what is commonly known as, a *reference semantics*; that is a semantics which is based on a pair of mappings, one from labels to locations, and the other from locations to values. This differs from several formal treatments of OO, such as Abadi and Cardelli's untyped object calculus (ζ -calculus) [AC96], System-F (a variant of the λ -calculus) [Pie02], and various object refinement calculi [CN02, HLL02b, Sek96], which are defined in terms of a *value semantics*; that is a single map from labels to values.

The semantic gap between the mainstream OO-programming languages and the value-based formal models is problematic, as value-semantics essentially prevents the sharing of an object, such as used in a typical implementation of a doubly-linked list.⁴ Having said this, Chapter

² Partial mappings are important, as they define which operations are available only in a given environmental state.

³ The notion of refinement in this case is essentially reverse containment; that is the behaviours of the implementation are contained within the behaviours of the specification.

⁴ Suppose $A - B - C$ represents a three element doubly linked list, then both A and C share object B .

13 of [Pie02] presents a mechanism for adding references to the simply typed λ -calculus ($\lambda_{\rightarrow}$ -calculus) with *Unit* (a type containing a single value *unit*). In principle, similar extensions can be made to: other varieties of the λ -calculus, such as System-F; and various object calculi, such as the ζ -calculus. An example of an object calculus with this type of extension is given in Section 2.2.

Clarifying Terms: Unfortunately, the term *reference semantics* is in itself ambiguous, as it can either be interpreted as a semantics of references, or a base-line semantics for discussing other semantics. Therefore, within this document the term *location semantics* will be used instead, as this is less likely to cause confusion. It also pattern matches into the usual abstract way of representing references that cannot be arbitrarily calculated via pointer arithmetic [Pie02]. Further, the term *value semantics*, could be similarly replaced by the term *copy semantics*, to more accurately capture its nature; however, it is felt that the term *value semantics* is sufficiently clear, and provides a more natural conceptual correspondence to location semantics, so it is retained.

A.1.3 Observable behaviour

The concept of observable behaviour is key to the semantics of software, as two software entities are considered equivalent if in *all contexts* no distinction can be made between their patterns of observable behaviour. Ironically, the notion of *all contexts* depends upon its context, i.e. its underlying modelling assumptions.

Example A.4: The assumptions underpinning the CSP modelling language include the constraint that observable events are atomic and instantaneous. These assumptions (constraints) ensure that it is impossible to have two distinct observable events happening at the same time. Hence, the notion of *all contexts* that is associated with the CSP language does not include the possibility of either independent simultaneous events, or a single event happening over a period of time. Having said this, it is straightforward to model events with durations, by having named start and stop events. In the author's experience the highlighted CSP assumptions have not caused significant problems when modelling real systems.

□

All programming languages have an associated model of expected behaviour, which is broken down, at least conceptually, into a collection of related sub-models, such as the arithmetic model, the (sub)routine invocation model, the memory model, and the exception handling model. These sub-models describe (or define) various aspects of the programming language at varying degrees of depth and completeness. In particular, some aspects of a programming language may be deliberately undefined, in order to enable a variety of context dependent (e.g. operating system specific) interpretations.

Example A.5: In the 'C' programming language [KR88] the size of both the integer and pointer types may vary, where the pointer type's size is set to reflect that of the architecture of the target computer's Central Processing Unit (CPU).

□

What impact does such undefinedness have on the observable behaviour of a program? At best, it has the potential of adding well defined context dependent aspects to the values and types of behaviour that can be observed. At worst, it may introduce subtle non-deterministic ambiguity into the behaviour of a language, whose affects may be very hard to predict and reason about. Having said this, it does allow for optimisation of code for a particular computer's hardware architecture and operating systems infrastructure.

Some of the undefinedness in a programming language can be mitigated by adopting programming practices that are designed to ensure that however a semantic ambiguity is resolved, the resulting program semantics is stable; that is given a suitably chosen granularity of observation, all potential implementations of the program are still guaranteed to meet their specification. For example, one could ensure that regardless of whether preemptive or cooperative multi-threading is used by the implementation, the program would still not deadlock.

The notion of granularity of observation is important. If it is too fine (detailed) then it distinguishes software components that are actually equivalent to each other in all relevant contexts; this is analogous to a specification stating ‘how something is to be implemented’, rather than ‘what is to be implemented’. Conversely, if the granularity of observation is too coarse, then software components that should be distinguished as having significantly different behaviours are not distinguished; this is analogous to an ambiguous specification. The correct granularity of observation can be dependent on the property that is being checked.

Example A.6: Being able to observe the internal state of an encapsulated entity, such as the private members (internal state and actions) of an object, when considering how that entity interacts with its peers, provides too fine a level of granularity. In general, it prevents the replacement of the entity with an entity that has the same behaviour from the peers’ perspective but a different internal behaviour (e.g. different sequence of internal actions). However, being able to observe the internal behaviour of an entity may be useful, when the implementation of that entity is updated, so as to treat that entity’s internal behaviour as the new system context. Alternatively, it may be useful to observe the internal state of an entity for checking properties such as type-correctness of the program; though in this case, it could be argued that such type correctness checks should be done in a layered manner, thus in any individual layer there is no need to check the internal type correctness of its components.

□

In summary, establishing the correct level of visibility of observable behaviour is critical when determining the semantic effects of a program. In particular, it governs when one software entity can be replaced by another software entity without changing the semantics of the program as a whole.

A.1.4 Concurrency

Modern operating systems, such as Microsoft Windows and Linux, provide their users with an inherently concurrent environment; in the sense that it is possible for the computer to be performing several independent jobs for their user at the *same time*. For example, a user can be playing a CD, whilst browsing the Internet and background printing some interesting articles. From the user’s perspective, all these events appear to be happening simultaneously (e.g. at the *same time*), but in practise these jobs are likely to be sharing the same CPU core. Hence, some form of job management is being provided – behind the scenes – by the operating system⁵ in order to maintain this illusion. The clever, or tricky, bit is ensuring that the independent jobs do not adversely interfere with each other.

In general, avoiding adverse interference between various processes (jobs), or sub-processes, is a non-trivial task, particularly when these processes are supposed to interact with each other. This leads to the study of the semantics of concurrency; with such semantics it is possible to demonstrate that an application’s design and implementation will be well behaved. For example, the application is both deadlock and livelock free; that is all processes are eventually able to make progress and no process can infinitely loop without user (environmental) intervention respectively.

⁵ The operating system may be supported by various run-time environments, such as a Java virtual machine.

Essentially, the semantics behind concurrency can be summarised by the semantics of interaction; that is when and how are processes allowed to interact. This suggests that a natural model of concurrent behaviour is a monitoring of these interactions or communication events. It is possible that such observations led to the design of process algebraic languages, such as CSP [Hoa85, Ros98] and CCS (Calculus of Communicating Systems) [Mil80]. More recently these ideas have been enhanced to include more modern notions of portability, or reconfigurability, of the channels of communication, such as in the pi-calculus (π -calculus) [Mil99, SW01] and the ambient calculus [CG98]. The latter ideas have been suggested as a natural basis for modelling OO-languages [SW01, PT97, Lum99, Mer00, Nor92, Nor98]. However, we believe that the notions of concurrency and object orientation, though related, are sufficiently independent to be studied in their own rights, and then combined. In particular, studying them in a concurrent context may mean that the notion of an object is not as general as it could be.

A.2 Object-Oriented Infrastructure

For convenience we reprint the OO concept map in Figure A.3. The concepts in this map are now discussed in the following four subsections on objects, polymorphism, program structure and code reuse.

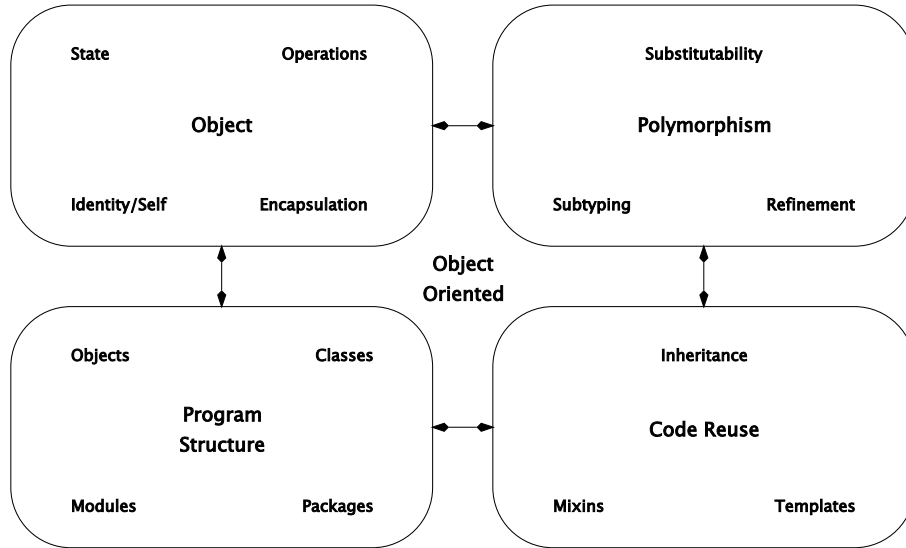


Figure A.3: Object-oriented concept map (reprint)

A.2.1 Objects

For the purposes of this dissertation the key concepts that an object possesses are:

State – An object contains data, which is typically represented by a collection of named fields, much like a record (or form). This data may be hidden from external view by the object's, possibly context dependent, *encapsulation* mechanism(s). For example, in Java [GJS96] an object's field could be marked as package visible; that is visible to operations (methods) that are defined in the same package as the field.

Operations – An object contains named operations for accessing and modifying its state. These operations are typically referred to as the methods of an object. Similarly to fields, methods can be hidden from external view, by the object's, possibly context dependent, *encapsulation* mechanism(s).

Note that in some OO-languages methods and fields are treated uniformly, as updateable members. Conceptually, this can be presented from two distinct points of view. First, a method can be treated as a special sort of data value, which defines an operation; this is sometimes referred to as *higher-order* programming. Second, fields can be treated as constant operations which return the field's value; thus writing to a field is represented by the updating (replacement) of the constant operation. This latter view is adopted in this dissertation.

Identity/Self – An object has a unique identity that can be used to distinguish it from other objects that otherwise share the same state (i.e. have the same value). In practise, this is often simply the value of the location at which the object is stored in the *heap* (i.e. the memory address). However several significant formal models of objects ignore this point. For example the simplest of the Abadi–Cardelli ζ -calculi does not have this property [AC96], as two objects that differ only in the *self* variable are considered to be equivalent, in all contexts. This is an example, of a value-based semantics of objects, where objects are: passed by value, updated by replacement, and considered equal whenever they have the *same* state. In general, the unique identity of an object is less important in a value-based semantics as it is impossible for two variables to share the same object, the best they can do is have a duplicate (copy) of each other's object. Having said this, the notion of *self* (limited ownership) is still important and does not require unique identity, and this is supported by the Abadi–Cardelli ζ -calculi.

Encapsulation – An object provides a container that ought to be used to group related data and operations. Such a grouping is sometimes considered to be a form of encapsulation. However, object encapsulation usually refers to a combination of hiding and access control, where strictly internal members are completely hidden and visible members are given context dependent access control markings, such as *read-only* and *write-once*. The intention is that an object can only be accessed via operations on its published interfaces.

Note that the visibility of these interfaces may be context dependent (as in Java). Further, aspects of an object's internal representation may be compromised, by *reference exposure*; for example, where a reference to an internal object is passed to the object's environment. In this case, the object's environment can update the inner-object directly. At first sight this may appear to be a poor idea (as proponents of a value-based object semantics might suggest), but in practise this feature provides an efficient way of manipulating structured data. Some commentators believe that sharing objects via references (*aliasing*) is “fundamental to good object-oriented design” [Cla01]; it is certainly required for the implementation of a doubly linked list, or any other cyclic data structure. In such contexts the problem is how to bound the range and scope of object references (i.e. control reference exposure).

A.2.2 Polymorphism

Before the topic of polymorphism can be discussed, the standard notion of a *type* needs to be introduced. Essentially operations such as addition and multiplication only make sense, when applied to certain *types* of values, such as the integers, but not when applied to other types of values such as the booleans. At the lowest level a type describes the interpretation of bit patterns. For example, both the ASCII character ‘A’ and the 8-bit representation of the number 65 have the same bit pattern, namely 01000001. At more abstract levels of reasoning, a type can be used to characterise a class of functions or objects, by defining their interface (or signature). Further, types can be placed in a *subtype* hierarchy,⁶ where “*A* is a subtype of *B*”

⁶ The subtype relationship forms a pre-order over object types, i.e. it is reflexive and transitive. Note that the reason it is not necessarily a partial order, is that two distinct object types may be subtypes of each other.

informally means that a term of type A can be used anywhere a term of type B is expected; this is known as subsumption.

“The term polymorphism refers to a range of language mechanisms that allow a single part of a program to be used with different types in different contexts” [Pie02, Section 22.7]. One example, is operation (or method) overloading, which enables the same operation to be applied to different types of data. For example, a single addition or multiplication operation that can be applied to different sizes of integer. Within the OO-context, the ability to use one object in place of another in a *controlled manner* is a form of polymorphism. The precise definition of what is an adequate measure of control is dependent on the context, as discussed below.

From a compiler writer’s point of view, the rules governing when one object can validly replace another, would ideally be statically determinable, by a *straightforward and efficient algorithm*. This leads to the standard type-based sanity checking algorithms, which focus on an entity’s syntactic description (*signature*), rather than its *behaviour*. In this context, a type correct program is guaranteed to be able to apply each of the operations to the supplied operands, but there is no guarantee that the subtyped operators actually behave appropriately. In general this amounts to the use of subsumption.

From a program verifier’s (semantic) point of view, it is not sufficient for an object to be replaced by an object with a compatible signature. What is required is that the replacing object is *semantically equivalent*. Precisely what is meant by semantically equivalent is a matter for interpretation, but essentially amounts to ensuring that the specified behaviour of the original object is implemented by the replacement object. Hence, this style of subtyping is commonly known as *behavioural subtyping*. The notions of behavioural equivalence and refinement are typical in this area of research.

From now on the notion of subtyping refers to standard subtyping, rather than behavioural subtyping, unless explicitly indicated otherwise.

In untyped languages the concept of polymorphism is not particularly useful, as there is no mechanism for classifying different families of values. Here, any value can be passed as an argument to any operation. In general, being able to statically check that every operation within a program will always receive sensible operands requires a full program analysis. For example, ensuring that the divisor of every division operation is non-zero can require a full program analysis.

Aside A.7: The concepts of polymorphism and subtyping are distinct from those of inheritance, mixins and other code reuse mechanisms. Some real OO-programming languages, such as Java, merge these concepts; in this case inheritance implies subtyping (i.e. A inherits-from $B \Rightarrow A$ subtypes B). This strictly limits what inheritance can be *safely* used for.
□

Polymorphism and subtyping is, and has been, a popular field of academic research [CW85, CM92, SP94, Sim95, Bru96, WO99, LW01, BMvW00, BPG00, FR00, Pol01]. Pierce’s text book on “Types and Programming Languages” [Pie02] summarises a significant portion of this work, though it is essentially limited to the compatibility, rather than behavioural, style of subtyping. Behavioural subtyping is normally discussed in the context of program equivalence, program refinement and module composition [Jac96, BMvW00, Pol01, CN00].

A.2.3 Program structure (composability)

Most programming languages include some *structural* mechanisms, such as modules, packages, classes, procedures, and objects, which enable the code to be organised into a collection of components and *links* between components. Such components ought to provide access to their resources via, and only via, a published (e.g. documented) interface.

Interfaces are the key to the *structural* mechanisms, as they provide the ability to describe services (or resources) at a higher level of abstraction than that of the implementing code. In

particular, they enable the hiding of implementation details, which can then be changed to improve efficiency, fault tolerance, security, etc.

The ability to update or replace an implementation of a component without performing a complete program analysis, assumes: that the published component interface accurately describes what the component provides; and that the clients of this component only use published features of the interface (i.e. the client respects the components interface).

Aside A.8: There are many techniques for demonstrating that a component conforms to its interface, such as bisimulation or refinement. Hoare and He's Unifying Theories of Programming (UTP) is constructed in terms of the logical implication ordering, which can also be used as a refinement order. Hence, refinement is the natural way of checking interface descriptions within the UTP.

□

Aside A.9: A new *sort* of object can be introduced into a multi-method program that affects the resolution of existing objects, even if no object of that *sort* is ever instantiated. Essentially, the new object definition introduces a method which *better fits* (pattern matches) an existing method call within the program, and hence is used instead of the existing method call.

□

A.2.4 Code reuse

There are several motivating reasons for supporting the ability to reuse code, both on global and local scales. On the global scale, the idea is to build a resource of library components that can be used in the construction of an application. Modern programming languages, such as Java and C#, come with a large library of utility routines, which provide facilities for things like network access (e.g. sockets and remote method calls), data storage (e.g. linked lists, trees, and hash tables), and user interface design (e.g. windows, buttons and text boxes). In general, however, there are some significant difficulties to overcome, such as: how to find suitable components for a given context; who owns the *rights* to those components and the limitations on their use; who is responsible for failures of a component to meet its design; etc.

On the local scale, the idea is to factor out common aspects of a design, and only write the code for implementing a common aspect once. This has the benefit of reducing the size of the source code, and aiding in the construction of a logical program design. However, there are sometimes significant run-time overheads with such a design methodology, as it may introduce some level of indirection, such as additional method calls, and method lookups. These can be worked around to some extent, by *in-lining* the code during compilation process (e.g. macros rather than procedures or methods). Further, the very mechanisms that introduce the ability to factor out common aspects of a design can complicate the formal model of a language.

Example A.10: Within the class-based OO-context a new class can typically declare that it *extends* (inherits from) a selection of existing classes (and interfaces), which are called its parent classes, or simply its parents. In such a context, the new *child* class inherits the visible members of its parent's class, which are promoted to the child class's interface. These visible members can then be augmented (and sometimes overridden) by the child class definition, in order to provide the necessary functionality.

There are many subtleties in the precise semantics of inheritance, particularly when the effects of scoping and aliasing are taken into account, see [Coo89, CW85, Cla02], [Smi99, Appendix D] and [GM94, Part 5] for more details.

□

Inheritance mechanisms vary significantly from language to language; but they essentially amount to some form of entity specialisation or extension. These ideas can be generalised by

the notion of a *mixin*, as illustrated by Bracha and Cook in [BC90] and extended by Bracha's PhD. thesis [Bra92]. Here *mixin* is informally defined to be an abstract subclass that may be used to specialise the behaviour of parent classes [BC90, Section 3.2]. In other words, a mixin can be thought of as a form of class generator function, which takes a class and generates a new class. Mixins have since, been used by other researchers in their module and inheritance systems [Sim95, AZ96, AZ98, AZ02, BPS99, FF98, FKF98, HL02, SB98, SB02, OCRZ03, OSV08].

A.3 Aliasing

Sharing (*aliasing*) objects is a natural concept in OO-software design because of the location semantics in mainstream OO-languages, such as C++ and Java. In practise, the notion of object sharing is linked with that of object identity. "Object identity is the foundation of object-oriented programming" [NVP98]. Once created an object retains its unique identity for duration of its existence; in other words the object's identity is immutable.

Aside A.11: The existence of such identities provides a platform, whereby the operation performed on the object contained in a typed-variable can be dependent on the run-time value of the variable's object, rather than compile- or load-time type of the variable itself. This is one of the ways in which an object is distinguished from the concept of a *record*; a record is a compound data-type that essentially maps labels, known as fields or attributes, onto values (or typed-values, for a typed-language).

□

Aliasing is a well known software engineering issue [Hog91, NVP98]. It can be usefully examined at various levels of abstraction, which may depend on *strong* assumptions about the nature of the context.

Aside A.12: The assumption that a location in memory can only be obtained via *creation* or *duplication* requires a check that no *pointer arithmetic* is used; that is memory locations are not obtained by direct *numeric address* conversion, whether this be absolute or relative to an existing memory location (i.e. an address). This assumption, though *strong*, is not sufficient to guarantee that a memory location is valid, as it is still possible to *free* the contents of that location before its last use.

□

In general, it is useful to make several *strong* assumptions about the way in which memory locations, and their contents, can be used. These assumptions then need to be checked, via appropriate analyses for the context in question, where the context includes all issues introduced by languages and *design patterns* [GHJV94] that are used.

Example A.13: Some languages, such as C [KR88], provide almost no memory protection. Hence, analysis concerning the *valid* use of memory is potentially very difficult, as it is hard to provide *safe* abstractions for high level reasoning. Having said this, the use of well understood *design patterns* and *safe* coding practises can significantly simplify matters.

Other languages, such as *pure* Java [GJS96]⁷, provide certain guarantees about the way in which memory locations and their contents can be used. In the case of *pure* Java these guarantees, ensure that memory access is always *valid*; that is the contents of an available memory location always exist and are either *type correct*, or will throw a run-time exception, such as the array storage exception⁸.

□

⁷ *Pure* Java refers to Java which does not embed (interface to) other languages.

⁸ Java's `ArrayStorageException` is *thrown* whenever an attempt is made to store an element of an incompatible type in an array.

Aliasing can provide a means for exposing the internal representation of an object and dynamically updating it, as illustrated by the following example.

Example A.14: Consider a simple object that has a private inner-object O of type T , and a public method M that returns an object reference to an object of type T . If M returns a reference to O , it could be argued that this has broken the encapsulation of the simple object. However, it could also be argued, that the simple object's private data was actually only a reference to another external object, and thus the encapsulation was not broken. Either way, if M had returned a reference to a deep copy of object O there would not have been an issue; note that a shallow copy may expose the contained elements within object o .

□

The problem of bounding or controlling the exposure of object references is well known. A significant amount of work has been done in this area, such as [NVP98, NCP99, Cla01, Cla02, CD02].

Appendix B

Classes in the Object Calculus

A class essentially provides a means for generating objects of the same form (e.g. type). Here, objects of the same class would typically have the same members. Further, the *genuine* methods of these objects (rather than methods that are used to represent fields) typically never change, so can legitimately be shared.

This appendix presents seven different ways of sharing methods, from simple *static* copy based approaches, to *dynamic* approaches that make use of object and function references. Here, the point is to demonstrate that the flexible primitives we have defined (e.g. method update and heap access) allow us to explicitly model and compare a number of different implementation techniques for sharing of methods.

The *static* approaches to modelling the sharing of methods constrain the use of method labels. Specifically, the *label* used to identify which of an object's methods is to be invoked or update cannot be calculated. In other words, within the *static* approach labels are not considered to be values, therefore cannot be passed into or returned out of a method invocation or function call.

The *dynamic* approaches to modelling the sharing of methods allow the value of a method's label to be calculated. Having said this, such label calculations may provide no means for modifying a literal label value, which essentially limits the dynamism to a selection process between existing label values.

The remainder of this chapter is split into two sections, one describing the *static* approaches and the other the *dynamic*.

B.1 Static Approaches to Sharing Methods

B.1.1 Prototype approach

The prototype approach simply shares method definitions, by copying an existing object. This is illustrated by the following simple cell example.

Example B.1: A simple cell has a **value** field, a **get** method, and a **set** method. It also has a **new** method for generating new instances of itself; this is strictly unnecessary, but helps mirror the class-definitions that follow later. Further, the **new** operation does more than simply copy the current cell; it resets the value of the copy.

$$\begin{aligned} \text{Cell} \triangleq [& \\ & \text{value} = \text{InitVal}, \\ & \text{new} = \text{InitVal}, \\ & \text{get} = \varsigma(s) \text{ GetCell}, \\ & \text{set} = \varsigma(s) \text{ SetCell} \\ &] \end{aligned}$$

where:

$$InitVal \equiv \varsigma(s) (s.\text{value} \Leftarrow \varsigma(z) z.\text{value})$$

$$GetCell \equiv s.\text{value}$$

$$SetCell \equiv \lambda(v) s.\text{value} := v$$

Note that the definitions of the initial-value and the methods to get and set the cell value are abstracted, so that they can be used to highlight the differences between this approach and the class-based approaches that follow. Further, note that the values in the cells are initialised to a method that infinitely invokes itself. An alternative would have been to initialise to the unset value $\perp$. It is unclear which definition is best: the former does not require a special value; whereas the latter is simpler (at least in an untyped environment). $\square$

B.1.2 Naive class-based approach

One approach to modelling classes, with objects, is to have a special *class* object, which can be used to generate instance objects of that class. The naive model of a class contains a function to model each method within one of its instances, where the function's argument is set to the instance's self parameter. The class also contains a special member for generating new instance objects, denoted by **new**, which creates members for both the instance's fields and methods.

Example B.2: The naive model of a simple cell class has three members, a **new** method, a **get** field, and a **set** field. The two fields contain functions for getting and setting the **value** field of an instance object. The **new** method generates an instance object, with a circular **value** field, a **get** method, and a **set** method. Note that an instance object's **get** and **set** methods both contain a complete copy of the underlying class, along with an appropriate field selection label.

$$\begin{aligned} Cell \hat{=} [& \\ & \text{new} = \varsigma(c) [\\ & \quad \text{value} = InitVal, \\ & \quad \text{get} = \varsigma(s) c.\text{get}(s), \\ & \quad \text{set} = \varsigma(s) c.\text{set}(s) \\ &], \\ & \text{get} = \lambda(s) GetCell, \\ & \text{set} = \lambda(s) SetCell \\ &] \end{aligned}$$

$\square$

In general, the instance objects created in this approach contain multiple copies of the creating class, one for each instance method. This duplication can be avoided by pre-processing the class field selection during object instance generation, as we discuss next.

B.1.3 Pre-processed class-based approach

In order to avoid unnecessary duplication of the class, it is possible to pre-process the class-field selection, as illustrated by the following example.

Example B.3: The pre-processing of the class-field selections is achieved via the use of the

let-in construct.

$$\begin{aligned} \text{Cell} \triangleq [& \\ & \text{new} = \varsigma(c) \text{ let } f_1 \triangleq c.\text{get}, f_2 \triangleq c.\text{set} \text{ in } [\\ & \quad \text{value} = \text{InitVal}, \\ & \quad \text{get} = \varsigma(s) f_1(s), \\ & \quad \text{set} = \varsigma(s) f_2(s) \\ &], \\ & \text{get} = \lambda(s) \text{ GetCell}, \\ & \text{set} = \lambda(s) \text{ SetCell} \\ &] \end{aligned}$$

□

In this approach, each instance object has its own copy of the instance methods. This duplication can be avoided via the use of references to objects or functions (or both).

B.1.4 Sharing method instances

So far the sharing of methods has been achieved via duplication. The introduction of object references allows method instances to be shared. A class is no longer represented directly by an object; instead, it is represented by a reference to a class-object. Here, each object-method selection dereferences its base class, then selects and invokes the appropriate class method, which takes the object's self parameter as its argument.

Example B.4: This variant of the simple cell is defined as follows, where r represents the unique heap location of the class, which is freshly generated on its first usage.

$$\begin{aligned} \text{Cell} \triangleq \text{let } r \triangleq \text{fresh in } r * = [& \\ & \text{new} = [\\ & \quad \text{value} = \text{InitVal}, \\ & \quad \text{get} = \varsigma(s) *r.\text{get}(s), \\ & \quad \text{set} = \varsigma(s) *r.\text{set}(s) \\ &], \\ & \text{get} = \lambda(s) \text{ GetCell}, \\ & \text{set} = \lambda(s) \text{ SetCell} \\ &] \end{aligned}$$

□

One issue with this representation of a class is that its instances dereference the whole of the class on each method invocation (selection). Having done that the appropriate class method needs to be selected, and then invoked with the instance's self parameter.

B.1.5 Caching shared methods

One way of reducing the number of steps required to access a shared method is to use a cache. The idea here is that the class contains a collection of references to methods, which are then directly used by its instances.

Before looking at a full example, it is worth considering the format of caching the location of function f in a member n of an object o , where o is stored in heap location ℓ_i .

$$o \equiv [\mathbf{n} = \varsigma(c) \ell_i * = c.n := \text{fresh} * = f]$$

Given that all update operations are of equal precedence and right associative, the above member definition is equivalent to:

$$o \equiv [\mathbf{n} = \varsigma(c) \ell_i * = (c.n := (\text{fresh} * = f))]$$

The stepwise evaluation of $o.n$ is:

1. create a fresh (new) heap location ℓ_j ;
2. store function f in location ℓ_j ;
3. update member n of object o with location value ℓ_j ;
4. store the updated object in location ℓ_i .

$$o.n \equiv [\mathbf{n} = \ell_j]$$

Example B.5: The caching method instances variant of the simple cell is defined as follows, where r represents the unique heap location of the class, and f_i represents a pointer to an appropriate function.

$$\begin{aligned} \text{Cell} \hat{=} & \text{let } r \hat{=} \text{fresh in } r * = [\\ & \text{new} = \varsigma(c) * (\text{let } f_1 \hat{=} c.\text{get} \text{ in let } f_2 \hat{=} c.\text{set} \text{ in } r * = c.\text{new} := [\\ & \quad \text{value} = \text{InitVal}, \\ & \quad \text{get} = \varsigma(s) * f_1(s), \\ & \quad \text{set} = \varsigma(s) * f_2(s) \\ &]).\text{new}, \\ & \text{get} = \varsigma(c) * (r * = c.\text{get} := (\text{fresh} * = (\lambda(s) \text{GetCell}))).\text{get}, \\ & \text{set} = \varsigma(c) * (r * = c.\text{set} := (\text{fresh} * = (\lambda(s) \text{SetCell}))).\text{set} \\ &] \end{aligned}$$

□

B.2 Dynamic Approaches to Sharing Methods

B.2.1 Simulating dynamic method sharing

So far all the models of sharing methods between a *class* of objects have required each object to contain an explicit entry for every shared method. One way of avoiding this duplication is to enable dynamic method resolution, where the *name* of the method is passed as an argument to the invocation process. However, as labels are not values within the object calculi presented so far, they cannot be passed as arguments to an invocation process. This leaves us with the problem of how to represent the name of a method. One technique would be to represent each method by an enumerated value, which could be modelled as either an integer or a string literal. Such enumerated names could then be used to construct a dynamic dispatch lookup table (function) that mapped the name onto the corresponding method. Here the result of a lookup is a function that expects the *self* object as its argument.

These lookup tables can then be used to set the generic apply method of an object: $\text{apply} = \varsigma(s) \lambda(n) \text{lookup}(n)(s)$, where n represents the name of the method to be applied.

Example B.6: In the case of the simple cell, the lookup table that contains the two shared methods, get and set, can be modelled as follows:

$$\begin{aligned} \text{GetSetLookup} \hat{=} & \text{fresh} * = \lambda(n) \\ & \text{if } \mathbf{n} = \text{"get"} \text{ then} \\ & \quad \lambda(s) \text{GetCell} \\ & \text{else if } \mathbf{n} = \text{"set"} \text{ then} \\ & \quad \lambda(s) \text{SetCell} \\ & \text{else} \\ & \quad \text{let } \mathbf{v} = [\mathbf{l} = \varsigma(z) z.l].l \text{ in } v \end{aligned}$$

where the *GetSetLookup* table is stored in a fresh location on the heap. Note that the issue of handling invalid method names is resolved in this case by ensuring that the evaluation infinitely loops. An alternative would have been to set this to the undeclared value, so that the rule interpreter could detect the error.

It is now possible to define the simple cell object in terms of the *GetSetLookup* table, where the get and set methods are replaced by a generic apply method, which takes the names "get" or "set" as its arguments.

$$\begin{aligned} \text{Cell} \hat{=} [& \\ & \text{new} = [\\ & \quad \text{value} = \text{InitValue}, \\ & \quad \text{apply} = \varsigma(s) \lambda(n) * \text{GetSetLookup}(n)(s) \\ &] \\ &] \end{aligned}$$

Note that the lookup function could have been defined within the simple Cell class object. There is some merit in this approach as it would co-locate data with functionality (increasing the cohesiveness of the modelling of a class). $\square$

The lookup scheme presented in [Example B.6](#) does not provide a dynamic means for determining what *method names* are available. In this case, it is up to the *programmer* or *compiler* to ensure that only valid names are requested. If dynamic determination of what method names are available is required, such as for modelling some reflective language features, then this functionality could be provided by a companion *what names are defined* lookup function.

B.2.2 Directly supporting dynamic method sharing

It is also possible for dynamic method sharing to be directly supported by the object calculus. One approach would be to model an object's label as a string literal, as method names are represented by identifier strings. A side effect of this approach, is that a label's value could be calculated (generated) by the application of general string operations, such as string concatenation. An alternative approach is to classify the existing method labels as literal values, thus enabling them to be stored in an object's state and passed as the argument to a function. Semantically, this is all that is required as parameters are substituted for their values as part of the method and function invocation process. Hence there is no need to update the existing $\mathcal{O}$ -calculus reduction rules.

Aside B.7: Having said that there is no need to update the existing rules, it might be worth extending them so that they treat the label arguments as evaluable expressions, rather than literal values (or identifiers that are substituted with literal values). Without the suggested rule extensions such expressions would have to be pre-evaluated, in a local definition block, before use within the context of either a method invocation or update operation.

$$\begin{array}{ll} \frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1.e_2 \longrightarrow \Gamma' \bullet e'_1.e_2} \text{ INVM-1} & \frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet o.e \longrightarrow \Gamma' \bullet o.e'} \text{ INVM-2} \\ \frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1.e_2 \Leftarrow m \longrightarrow \Gamma' \bullet e'_1.e_2 \Leftarrow m} \text{ UPDM-1} & \frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet o.e \Leftarrow m \longrightarrow \Gamma' \bullet o.e' \Leftarrow m} \text{ UPDM-2} \\ \frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1.e_2 := e_3 \longrightarrow \Gamma' \bullet e'_1.e_2 := e_3} \text{ FLDA-1} & \frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet o.e_1 := e_2 \longrightarrow \Gamma' \bullet o.e_1 := e_2} \text{ FLDA-2} \\ \frac{\Gamma \bullet e_1 \longrightarrow \Gamma' \bullet e'_1}{\Gamma \bullet e_1?e_2 \longrightarrow \Gamma' \bullet e'_1?e_2} \text{ HASM-1} & \frac{\Gamma \bullet e \longrightarrow \Gamma' \bullet e'}{\Gamma \bullet o?e \longrightarrow \Gamma' \bullet o'?e} \text{ HASM-2} \end{array}$$

$\square$

Example B.8: Direct support for dynamic method binding enables the simple *Cell* class to define objects that have both: a single application function; and a single reference to their defining class.

$$\begin{aligned} \text{Cell} \hat{=} [\\ & \text{new} = [\\ & \quad \text{value} = \text{InitValue}, \\ & \quad \text{apply} = \varsigma(s) \lambda(n) * \text{GetSetLookup}(n)(s) \\ &] \\] \end{aligned}$$

Note that this direct model of the dynamic method application provides no support for handling invalid method names. If an invalid name was passed to the apply function, none of the existing rules would pattern match with a non-existent method selection, so the term evaluation as a whole would become blocked. Fundamentally this is not a new type of failure, because it has previously been possible to ask for a non-existent method of an object; it is just that doing this dynamically has highlighted the issue.

□

The previous example highlights the issue that the non-existent method invocation on an object causes all previously discussed models of the object calculus to block, and thus fail to produce a result. This is not necessarily a problem, as the *code* is invalid.

Aside B.9: When operations are being called outside their domains, it might be convenient to produce an *invalid* result, instead of blocking, in order to support error reporting and/or exception handling. In this case, providing additional rules for addressing method selection and update when invalid member names are encountered. The topic of error and exception handling is beyond the scope of the work being presented here.

□

Instead of waiting for an error to occur, it might be better to prevent an error occurring in the first place. This can be achieved in a number of ways, from arguments concerning correctness by construction, through static analysis to ensure that each operation is always supplied with arguments of the correct form (type), to dynamic testing of values *just* prior to their use. The last of these approaches is already supported by a combination of our $\mathcal{O}$ -calculus's membership test and conditional evaluation rules.

Example B.10: The previous dynamic dispatch example has been updated to return an unset ($\mathfrak{!}$) value whenever the class c does not contain a method with name n .

$$\begin{aligned} \text{Cell} \hat{=} \text{let } c = \text{fresh in } c * = [\\ & \text{new} = [\\ & \quad \text{value} = \text{InitValue}, \\ & \quad \text{apply} = \varsigma(s) \lambda(n) (\text{if } c?n \text{ then } *c.n(s) \text{ else } \mathfrak{!}) \\ &], \\ & \text{get} = \lambda(s) \text{GetCell}, \\ & \text{set} = \lambda(s) \text{SetCell} \\ &] \end{aligned}$$

□

Appendix C

UTP Design Laws

This appendix provides some basic laws concerning designs within Unifying Theories of Programming (UTP). It is split into three sections. The first section demonstrates that our theory of extended UTP designs is closed; that is each of our UTP design commands generates a design, whenever their subprogram arguments are designs (or constructed from a UTP commands that generates a design). The second section presents some general laws for manipulating UTP these designs. The third section presents a specific law for manipulating our UTP models' of objects.

C.1 Closed Design Laws

In this section we demonstrate that the conditional evaluation, sequential composition, variable introduction, variable elimination, variable hiding, and block scoped variable declaration commands return UTP designs whenever their subprogram arguments are designs. Note that the iteration command also has a subprogram argument. This awkward case has already been considered within the literature [HH98, WC04] and has been shown to generate a design so long as the other UTP commands generate designs (via Tarski weakest fix-point theorem on lattices [Tar55]). We use this result. Informally, the iteration command produces a design because it either: generates a UTP program which infinitely loops, in which case it is equivalent to the chaos design; or it sequentially composes a finite number of – non-iterative – designs, which is a design.

C.1.1 Conditional binding

The following law states that a conditional choice between two designs is a design.

Law C.1 – Conditional design:

$$(p \vdash P) \triangleleft b \triangleright (q \vdash Q) = (p \triangleleft b \triangleright q) \vdash (P \triangleleft b \triangleright Q)$$

□

Proof

$$\begin{aligned} & (p \vdash P) \triangleleft b \triangleright (q \vdash Q) \\ = & \dots\dots\dots \text{Defn. of conditional} \\ & b \wedge (p \vdash P) \vee \neg b \wedge (q \vdash Q) \\ = & \dots\dots\dots \text{Defn. of design} \\ & b \wedge (\Pi_{\text{OK}} \wedge p \Rightarrow \Pi_{\text{OK}}' \wedge P) \\ & \vee \\ & \neg b \wedge (\Pi_{\text{OK}} \wedge q \Rightarrow \Pi_{\text{OK}}' \wedge Q) \end{aligned}$$

$$\begin{aligned}
&= \dots\dots\dots \text{Disjunctive normal form} \\
&\quad \neg \Pi_{\text{OK}} \vee b \wedge \neg p \vee \neg b \wedge \neg q \\
&\quad \vee \\
&\quad \Pi_{\text{OK}}' \wedge b \wedge P \vee \Pi_{\text{OK}}' \wedge \neg b \wedge Q \\
&= \dots\dots\dots \text{Introduce implication} \\
&\quad \Pi_{\text{OK}} \wedge (\neg b \vee p) \wedge (b \vee q) \\
&\quad \Rightarrow \\
&\quad \Pi_{\text{OK}}' \wedge (b \wedge P \vee \neg b \wedge Q) \\
&= \dots\dots\dots \text{Propositional Logic} \\
&\quad \Pi_{\text{OK}} \wedge (b \wedge p \vee \neg b \wedge q) \\
&\quad \Rightarrow \\
&\quad \Pi_{\text{OK}}' \wedge (b \wedge P \vee \neg b \wedge Q) \\
&= \dots\dots\dots \text{Defn. of conditional} \\
&\quad \Pi_{\text{OK}} \wedge (p \triangleleft b \triangleright q) \\
&\quad \Rightarrow \\
&\quad \Pi_{\text{OK}}' \wedge (P \triangleleft b \triangleright Q) \\
&= \dots\dots\dots \text{Defn. of design} \\
&\quad (p \triangleleft b \triangleright q) \vdash (P \triangleleft b \triangleright Q) \\
&\square
\end{aligned}$$

C.1.2 Sequential composition

The following law states that the sequential composition of two designs is a design.

Law C.2 – Composition design:

$$p \vdash P \mathbin{;} q \vdash Q \quad = \quad p \wedge \neg(P \mathbin{;} \neg q) \vdash P \mathbin{;} Q$$

□

Proof

$$\begin{aligned}
&p \vdash P \mathbin{;} q \vdash Q \\
&= \dots\dots\dots \text{Defn. of design} \\
&\quad (\Pi_{\text{OK}} \wedge p \Rightarrow \Pi_{\text{OK}}' \wedge P) \mathbin{;} (\Pi_{\text{OK}} \wedge q \Rightarrow \Pi_{\text{OK}}' \wedge Q) \\
&= \dots\dots\dots \text{Defn. of } \mathbin{;} \\
&\quad \exists \Pi_{\text{OK}_0}, w_0 \bullet (\Pi_{\text{OK}} \wedge p \Rightarrow \Pi_{\text{OK}_0} \wedge P[w_0/w']) \\
&\quad \quad \wedge \\
&\quad \quad (\Pi_{\text{OK}_0} \wedge q[w_0/w] \Rightarrow \Pi_{\text{OK}}' \wedge Q[w_0/w]) \\
&= \dots\dots\dots \text{Let } P_0 = P[w_0/w'] \\
&\quad \exists \Pi_{\text{OK}_0}, w_0 \bullet (\Pi_{\text{OK}} \wedge p \Rightarrow \Pi_{\text{OK}_0} \wedge P_0) \quad \quad \quad q_0 = q[w_0/w] \\
&\quad \quad \wedge \quad \quad \quad Q_0 = Q[w_0/w] \\
&\quad \quad (\Pi_{\text{OK}_0} \wedge q_0 \Rightarrow \Pi_{\text{OK}}' \wedge Q_0) \\
&= \dots\dots\dots \text{Case analysis of } \Pi_{\text{OK}_0} \\
&\quad \exists w_0 \bullet (\Pi_{\text{OK}} \wedge p \Rightarrow \text{false} \wedge P_0) \wedge (\text{false} \wedge q_0 \Rightarrow \Pi_{\text{OK}}' \wedge Q_0) \\
&\quad \quad \vee \\
&\quad \quad (\Pi_{\text{OK}} \wedge p \Rightarrow \text{true} \wedge P_0) \wedge (\text{true} \wedge q_0 \Rightarrow \Pi_{\text{OK}}' \wedge Q_0) \\
&= \dots\dots\dots \text{Simplification} \\
&\quad \exists w_0 \bullet ((\Pi_{\text{OK}} \wedge p) \Rightarrow \text{false}) \\
&\quad \quad \vee \\
&\quad \quad (\Pi_{\text{OK}} \wedge p \Rightarrow P_0) \wedge (q_0 \Rightarrow \Pi_{\text{OK}}' \wedge Q_0) \\
&= \dots\dots\dots \text{Removal of implication} \\
&\quad \exists w_0 \bullet \neg(\Pi_{\text{OK}} \wedge p) \\
&\quad \quad \vee \\
&\quad \quad (\neg(\Pi_{\text{OK}} \wedge p) \vee P_0) \wedge (\neg q_0 \vee \Pi_{\text{OK}}' \wedge Q_0)
\end{aligned}$$

$$\begin{aligned}
&= \dots\dots\dots \text{Disjunctive normal form} \\
&\quad \exists w_0 \bullet \neg \Pi_{\text{OK}} \vee \neg p \vee (P_0 \wedge \neg q_0) \\
&\quad \quad \vee \\
&\quad \quad \Pi_{\text{OK}}' \wedge P_0 \wedge Q_0 \\
&= \dots\dots\dots \text{Predicate logic} \\
&\quad \neg \Pi_{\text{OK}} \vee \neg p \vee (\exists w_0 \bullet P_0 \wedge \neg q_0) \\
&\quad \quad \vee \\
&\quad \quad \Pi_{\text{OK}}' \wedge (\exists w_0 \bullet P_0 \wedge Q_0) \\
&= \dots\dots\dots \text{Defn. of } \circledast, P_0, q_0, \& Q_0 \\
&\quad \neg \Pi_{\text{OK}} \vee \neg p \vee (P \circledast \neg q) \vee \Pi_{\text{OK}}' \wedge (P \circledast Q) \\
&= \dots\dots\dots \text{Introduce implication} \\
&\quad \Pi_{\text{OK}} \wedge p \wedge \neg(P \circledast \neg q) \Rightarrow \Pi_{\text{OK}}' \wedge P \circledast Q \\
&= \dots\dots\dots \text{Defn. of design} \\
&\quad p \wedge \neg(P \circledast \neg q) \vdash P \circledast Q \\
&\square
\end{aligned}$$

There are at least two useful corollaries, for the special cases where: the second design has a true precondition; and the first design's post condition explicitly sets the output value of each variable in the alphabet.

Corollary C.3 – Pre-true composition:

$$p \vdash P \circledast \text{true} \vdash Q = p \vdash P \circledast Q$$

□

Proof

$$\begin{aligned}
&p \vdash P \circledast \text{true} \vdash Q \\
&= \dots\dots\dots \text{Design comp. law} \\
&\quad p \wedge \neg(P \circledast \neg \text{true}) \vdash P \circledast Q \\
&= \dots\dots\dots \text{false – zero of } \circledast \\
&\quad p \wedge \neg(\text{false}) \vdash P \circledast Q \\
&= \dots\dots\dots \text{Propositional logic} \\
&\quad p \vdash P \circledast Q \\
&\square
\end{aligned}$$

Corollary C.4 – Known value composition:

$$p \vdash w' = V \circledast q \vdash Q = p \wedge q[V/w] \vdash Q[V/w]$$

□

Proof

$$\begin{aligned}
&p \vdash w' = V \circledast q \vdash Q \\
&= \dots\dots\dots \text{Design comp. law} \\
&\quad p \wedge \neg(w' = V \circledast \neg q) \vdash w' = V \circledast Q \\
&= \dots\dots\dots \text{Defn. of } \circledast \\
&\quad p \wedge \neg(\exists w_0 \bullet w_0 = V \circledast \neg q[w_0/w]) \\
&\quad \vdash \\
&\quad \exists w_0 \bullet w_0 = V \circledast Q[w_0/w] \\
&= \dots\dots\dots \text{Predicate Logic} \\
&\quad p \wedge \neg(\neg q[V/w]) \vdash Q[V/w] \\
&= \dots\dots\dots \text{Propositional Logic} \\
&\quad p \wedge q[V/w] \vdash Q[V/w] \\
&\square
\end{aligned}$$

C.1.3 Variable introduction and elimination

The following law says that the variable introduction command can be rewritten as a design.

Law C.5:

$$\text{var } x \quad = \quad \text{true} \vdash \exists x \bullet w' = w \wedge x' = x$$

$$\text{where } w \equiv x_{i=1}^k, x \notin \{w\} \text{ and } \alpha(\text{var } x) = \{w, w', x'\}.$$

□

Proof

$$\begin{aligned}
 & \text{var } x \\
 = & \dots\dots\dots \text{Defn. of var } _ \\
 & \exists x \bullet \text{skip}^D \\
 = & \dots\dots\dots \text{Defn. of skip}^D \\
 & \exists x \bullet \text{true} \vdash w' = w \wedge x' = x \\
 = & \dots\dots\dots \text{Defn. of } _ \vdash _ \\
 & \exists x \bullet \Pi_{\text{OK}} \Rightarrow (\Pi_{\text{OK}}' \wedge w' = w \wedge x' = x) \\
 = & \dots\dots\dots \text{Predicate Logic} \\
 & \Pi_{\text{OK}} \Rightarrow (\Pi_{\text{OK}}' \wedge \exists x \bullet w' = w \wedge x' = x) \\
 = & \dots\dots\dots \text{Defn. of } _ \vdash _ \\
 & \text{true} \vdash \exists x \bullet w' = w \wedge x' = x
 \end{aligned}$$

□

Note that the last step can be further simplified to remove the existential quantification altogether, as we can use the one-point rule to set the variable x equal to x' . In this case, we end up with the design $\text{true} \vdash w' = w$.

The following law says that the variable elimination command can be rewritten as a design.

Law C.6:

$$\text{end } x \quad = \quad \text{true} \vdash \exists x' \bullet w' = w \wedge x' = x$$

$$\text{where } w \equiv x_{i=1}^k, x \notin \{w\} \text{ and } \alpha(\text{end } x) = \{w, w', x\}.$$

□

Proof

$$\begin{aligned}
 & \text{end } x \\
 = & \dots\dots\dots \text{Defn. of var } _ \\
 & \exists x' \bullet \text{skip}^D \\
 = & \dots\dots\dots \text{Defn. of skip}^D \\
 & \exists x' \bullet \text{true} \vdash w' = w \wedge x' = x \\
 = & \dots\dots\dots \text{Defn. of } _ \vdash _ \\
 & \exists x' \bullet \Pi_{\text{OK}} \Rightarrow (\Pi_{\text{OK}}' \wedge w' = w \wedge x' = x) \\
 = & \dots\dots\dots \text{Predicate Logic} \\
 & \Pi_{\text{OK}} \Rightarrow (\Pi_{\text{OK}}' \wedge \exists x' \bullet w' = w \wedge x' = x) \\
 = & \dots\dots\dots \text{Defn. of } _ \vdash _ \\
 & \text{true} \vdash \exists x' \bullet w' = w \wedge x' = x
 \end{aligned}$$

□

C.1.4 Hiding

The following law states that hiding some variables from a design returns a design.

Law C.7 – Hide Design:

$$\text{hide } u \text{ from } p[w] \vdash P[w, w'] = p[w] \vdash P[w, w'] \wedge u' = u$$

□

Proof

Let P_1 , P_2 , and P_3 denote $(\exists u' \bullet u' = u \wedge w' = w)$, $(\exists u \bullet u' = u \wedge w' = w)$, and $(u' = u_0 \wedge w' = w)$ respectively.

$$\begin{aligned}
& \text{hide } u \text{ from } p[w] \vdash P[w, w'] \\
= & \dots \text{Defn. of hide } _ \text{ from } _ \\
& \exists u_0 \mid u_0 = u \bullet \\
& \quad (\text{end } u) \circ (p[w] \vdash P[w, w']) \circ (\text{var } u) \circ (u := u_0) \\
= & \dots \text{Design variants of} \\
& \quad \text{end } _, \text{ var } _, \text{ and } _ := _ \\
& \exists u_0 \mid u_0 = u \bullet \\
& \quad (\text{true} \vdash P_1) \circ (p[w] \vdash P[w, w']) \circ \\
& \quad (\text{true} \vdash P_2) \circ (\text{true} \vdash P_3) \\
= & \dots \text{Law C.2 comp. design} \\
& \exists u_0 \mid u_0 = u \bullet \\
& \quad \text{true} \wedge \neg(P_1 \circ \neg p[w]) \wedge \\
& \quad \neg(P_1 \circ P[w, w'] \circ \neg \text{true}) \wedge \neg(P_1 \circ P[w, w'] \circ P_2 \circ \neg \text{true}) \\
& \quad \vdash \\
& \quad P_1 \circ P \circ P_2 \circ P_3 \\
= & \dots \text{false is the right zero of } _ \circ _ \\
& \quad \vdash \\
& \quad \neg(P_1 \circ \neg p[w]) \\
& \quad \vdash \\
& \quad P_1 \circ P[w, w'] \circ P_2 \circ P_3 \\
= & \dots \text{Defn. of } _ \circ _ \\
& \exists u_0 \mid u_0 = u \bullet \\
& \quad \neg(\exists w_0 \bullet P_1[w_0/w'] \wedge \neg p[w_0/w]) \\
& \quad \vdash \\
& \quad P_1 \circ P[w, w'] \circ P_2 \circ P_3 \\
= & \dots \text{Defn. of } P_1 \\
& \exists u_0 \mid u_0 = u \bullet \\
& \quad \neg(\exists w_0 \bullet (\exists u' \bullet u' = u \wedge w_0 = w) \wedge \neg p[w_0/w]) \\
& \quad \vdash \\
& \quad P_1 \circ P \circ P_2 \circ P_3 \\
= & \dots u' = u, w_0 = w, \text{ and} \\
& \quad \text{simplify} \\
& \exists u_0 \mid u_0 = u \bullet \\
& \quad p[w] \vdash P_1 \circ P[w, w'] \circ P_2 \circ P_3 \\
= & \dots \text{Defn. of } _ \circ _, \\
& \quad \text{used 3 times} \\
& \exists u_0 \mid u_0 = u \bullet \\
& \quad p[w] \vdash \exists w_0, w_1, w_2 \bullet \\
& \quad \quad P_1[w_0/w'] \wedge P[w_0, w_1/w, w'] \wedge \\
& \quad \quad P_2[w_1, w_2/w, w'] \wedge P_3[w_2/w]
\end{aligned}$$

$$\begin{aligned}
&= \dots\dots\dots \text{Defn. of } P_1, P_2, \text{ and } P_3, \\
&\quad \exists u_0 \mid u_0 = u \bullet \text{ where bound variables are} \\
&\quad p[w] \vdash \exists w_0, w_1, w_2 \bullet \text{ renamed} \\
&\quad (\exists u_1 \bullet u_1 = u \wedge w_0 = w) \wedge P[w_0, w_1/w, w'] \wedge \\
&\quad (\exists u_2 \bullet u' = u_2 \wedge w_2 = w_1) \wedge \\
&\quad (u' = u_0 \wedge w' = w_2)[w_2/w] \\
&= \dots\dots\dots u_1 = u, u_2 = u, \text{ and} \\
&\quad \exists u_0 \mid u_0 = u \bullet \text{ simplify} \\
&\quad p[w] \vdash \exists w_0, w_1, w_2 \bullet \\
&\quad w_0 = w \wedge P[w_0, w_1/w, w'] \wedge \\
&\quad w_2 = w_1 \wedge (u' = u_0 \wedge w' = w_2)[w_2/w] \\
&= \dots\dots\dots w_0 = w, w_1 = w', w_2 = w', \\
&\quad \exists u_0 \mid u_0 = u \bullet \text{ and simplify} \\
&\quad p[w] \vdash P[w, w'] \wedge u' = u_0 \\
&= \dots\dots\dots \text{Defn. of } \vdash \text{ and} \\
&\quad \exists u_0 \bullet u_0 = u \wedge (\text{Defn. of } \exists \vdash \bullet \text{ } \\
&\quad \Pi_{\text{OK}} \wedge p[w] \Rightarrow \Pi_{\text{OK}'} \wedge P[w, w'] \wedge u' = u_0 \\
&\quad) \\
&= \dots\dots\dots \text{Predicate Logic} \\
&\quad \exists u_0 \bullet \\
&\quad (\Pi_{\text{OK}} \wedge p[w]) \vee u_0 \neq u \\
&\quad \Rightarrow \\
&\quad \Pi_{\text{OK}'} \wedge (P[w, w'] \wedge u' = u_0 \wedge u_0 = u) \\
&= \dots\dots\dots \text{Predicate Logic} \\
&\quad (\forall u_0 \bullet (\Pi_{\text{OK}} \wedge p[w]) \vee u_0 \neq u) \\
&\quad \Rightarrow \\
&\quad (\exists u_0 \bullet (\Pi_{\text{OK}'} \wedge P[w, w'] \wedge u' = u_0 \wedge u_0 = u)) \\
&= \dots\dots\dots u_0 \text{ not free in } p[w] \text{ or} \\
&\quad (\Pi_{\text{OK}} \wedge p[w]) \vee (\forall u_0 \bullet u_0 \neq u) \text{ } P[w, w'] \\
&\quad \Rightarrow \\
&\quad (\Pi_{\text{OK}'} \wedge P[w, w'] \wedge (\exists u_0 \bullet u' = u_0 \wedge u_0 = u)) \\
&= \dots\dots\dots \text{Predicate Logic} \\
&\quad \Pi_{\text{OK}} \wedge p[w] \Rightarrow \Pi_{\text{OK}'} \wedge P[w, w'] \wedge u' = u \\
&= \dots\dots\dots \text{Defn. of } \vdash \text{ } \\
&\quad p[w] \vdash P[w, w'] \wedge u' = u \\
&\square
\end{aligned}$$

C.1.5 Block scoped variable declaration

The following law states that when a block scoped declaration command, which introduces an uninitialised vector of variables for use within a subprogram, is provided with a design it returns a design.

Law C.8 – Declaration Design:

$$\text{decl } u \text{ in } p[u] \vdash P[u, u'] = (\forall u \bullet p[u]) \vdash (\exists u, u' \bullet P[u, u'])$$

□

Proof

$$\begin{aligned}
&\text{decl } u \text{ in } p[u] \vdash P[u, u'] \\
&= \dots\dots\dots \text{Defn. of decl } _ \text{ in } _ \\
&\quad (\text{var } u) \circ (p[u] \vdash P[u, u']) \circ (\text{end } u)
\end{aligned}$$

$$\begin{aligned}
&= \dots\dots\dots \text{Law C.11 var scoping} \\
&\quad \exists u \bullet (p[u] \vdash P[u, u']) \S (\text{end } u) \\
&= \dots\dots\dots \text{Law C.12 end scoping} \\
&\quad \exists u, u' \bullet p[u] \vdash P[u, u'] \\
&= \dots\dots\dots \text{Defn. of } \vdash \vdash \\
&\quad \exists u, u' \bullet \Pi_{\text{OK}} \wedge p[u] \Rightarrow \Pi_{\text{OK}'} \wedge P[u, u'] \\
&= \dots\dots\dots \text{Predicate Logic} \\
&\quad (\forall u, u' \bullet \Pi_{\text{OK}} \wedge p[u]) \Rightarrow (\exists u, u' \bullet \Pi_{\text{OK}'} \wedge P[u, u']) \\
&= \dots\dots\dots u' \text{ not free in } \Pi_{\text{OK}} \wedge p[u]. \\
&\quad (\forall u \bullet \Pi_{\text{OK}} \wedge p[u]) \Rightarrow (\exists u, u' \bullet \Pi_{\text{OK}'} \wedge P[u, u']) \\
&= \dots\dots\dots \text{Predicate Logic} \\
&\quad \Pi_{\text{OK}} \wedge (\forall u \bullet p[u]) \Rightarrow \Pi_{\text{OK}'} \wedge (\exists u, u' \bullet P[u, u']) \\
&= \dots\dots\dots \text{Defn. of } \vdash \vdash \\
&\quad (\forall u \bullet p[u]) \vdash (\exists u, u' \bullet P[u, u'])
\end{aligned}$$

□

The following law states that when a block scoped declaration command, which uses a partial map of variables to values to introduce an initialised collection of variables to a sub-program, is provided with a design it returns a design.

Law C.9 – Declaration Map Design: Let u denote the list of distinct undashed variables $x_{i=1}^k$. Let V denote the list of values $v_{i=1}^j$.

$$\text{decl}\{\{x_{i=1}^k \mapsto v_i\} \text{ in } p \vdash P \quad = \quad p[V/u] \vdash (\exists u' \bullet P[V/u])$$

□

Proof

$$\begin{aligned}
&\text{decl}\{\{x_{i=1}^k \mapsto v_i\} \text{ in } p \vdash P \\
&= \dots\dots\dots \text{Defn. of decl } \vdash \text{ in } \vdash \\
&\quad \text{decl } u \text{ in } (u := V) \S (p \vdash P) \\
&= \dots\dots\dots \text{Defn. of } \vdash := \vdash \\
&\quad \text{decl } u \text{ in } (\text{true} \vdash u' = V) \S (p \vdash P) \\
&= \dots\dots\dots \text{Law C.2 comp. design} \\
&\quad \text{decl } u \text{ in} \\
&\quad \quad \text{true} \wedge \neg((u' = V) \S \neg p) \\
&\quad \quad \vdash \\
&\quad \quad (u' = V) \S P \\
&= \dots\dots\dots \text{Defn. of } \vdash \S \vdash \\
&\quad \text{decl } u \text{ in} \\
&\quad \quad \neg(\exists u_0 \bullet u_0 = V \wedge \neg p[u_0/u]) \\
&\quad \quad \vdash \\
&\quad \quad (\exists u_0 \bullet (u_0 = V) \wedge P[u_0/u]) \\
&= \dots\dots\dots \text{One point rule} \\
&\quad \text{decl } u \text{ in } \neg(\neg p[V/u]) \vdash P[V/u] \\
&= \dots\dots\dots \text{Law C.8 decl. design} \\
&\quad (\forall u \bullet p[V/u]) \vdash (\exists u, u' \bullet P[V/u]) \\
&= \dots\dots\dots u \text{ not free in } \neg[V/u] \\
&\quad p[V/u] \vdash (\exists u' \bullet P[V/u])
\end{aligned}$$

□

C.2 General Design Laws

In this section we present (and prove) a collection of basic laws for the theory of UTP designs. Several of these are variants on existing work [HH98, WC04]; the differences in the proofs result from our use of [HLL02a]’s updated variant of the design definition.

C.2.1 Skip the unit of sequential composition

We have two distinct notions of skip, one for the basic relational unifying theory (i.e. `skip`) and one for the design-based unifying theory (i.e. `skipD`). They are both considered to be the unit of sequential composition. Note the design-based definition of skip is not strictly necessary; as it can be shown to be equal to that of the basic one, when applied to designs, however, it has the convenient form of a design.

We now state and prove the law that the design-based version of skip is the unit of composition for designs. The proof that the basic relational version of skip is also a unit of composition can be proved in a similar manner (it is actually easier to do).

Law C.10 – Design-skip composition unit:

$$\text{skip}^D \circ (p \vdash P) = p \vdash P = (p \vdash P) \circ \text{skip}^D$$

□

This is proved in two parts, one where skip is on the left hand side (LHS) of the sequential composition, and the other where it is on the right hand side (RHS).

Proof of LHS

$$\begin{aligned} & \text{skip}^D \circ (p \vdash P) \\ = & \dots\dots\dots \text{Defn. of skip}^D \\ & (\text{true} \vdash w' = w) \circ (p \vdash P) \\ = & \dots\dots\dots \text{Known val. Corol. C.4} \\ & \text{true} \wedge p[w/w] \vdash P[w/w] \\ = & \dots\dots\dots \text{Simplification} \\ & p \vdash P \end{aligned}$$

□

Proof of RHS

$$\begin{aligned} & (p \vdash P) \circ \text{skip}^D \\ = & \dots\dots\dots \text{Defn. of skip}^D \\ & (p \vdash P) \circ (\text{true} \vdash w' = w) \\ = & \dots\dots\dots \text{Pre-true Corol. C.3} \\ & p \vdash P \circ w' = w \\ = & \dots\dots\dots \text{Defn. of } \circ \\ & p \vdash \exists w_0 \bullet P[w_0/w'] \wedge w' = w_0 \\ = & \dots\dots\dots \text{1-point rule } (w_0 = w') \\ & p \vdash P \end{aligned}$$

□

C.2.2 Scoping variables

The following law states that the scope of an introduced variable extends to the RHS of a sequential composition of designs.

Law C.11 – Variable scoping law:

$$\text{var } x \circ (q \vdash Q) = \exists x \bullet (q \vdash Q)$$

□

Proof

$$\begin{aligned}
& \text{var } x \circ (q \vdash Q) \\
= & \dots \text{Defn. of var} \\
& (\exists x \bullet \text{skip}^D) \circ (q \vdash Q) \\
= & \dots \text{Defn. of } \circ \\
& \exists w_0 \bullet (\exists x \bullet \text{skip}^D)[w_0/w'] \wedge (q \vdash Q)[w_0/w] \\
= & \dots \text{Pred. calc. } (x \notin w_0, w') \\
& \exists w_0 \bullet (\exists x \bullet \text{skip}^D[w_0/w'] \wedge (q \vdash Q)[w_0/w]) \\
= & \dots \text{Pred. calc.} \\
& \exists x \bullet (\exists w_0 \bullet \text{skip}^D[w_0/w'] \wedge (q \vdash Q)[w_0/w]) \\
= & \dots \text{Defn. of } \circ \\
& \exists x \bullet \text{skip}^D \circ (q \vdash Q) \\
= & \dots \text{skip}^D \text{ is a unit } \circ \\
& \exists x \bullet (q \vdash Q)
\end{aligned}$$

□

The following law states that the scope of an eliminated variable extends to the LHS of a sequential composition of designs.

Law C.12 – Variable completion law:

$$(q \vdash Q) \circ \text{end } x = \exists x' \bullet (q \vdash Q)$$

□

Proof

$$\begin{aligned}
& (q \vdash Q) \circ \text{end } x \\
= & \dots \text{Defn. of end} \\
& (q \vdash Q) \circ (\exists x' \bullet \text{skip}^D) \\
= & \dots \text{Defn. of } \circ \\
& \exists w_0 \bullet (q \vdash Q)[w_0/w'] \wedge (\exists x' \bullet \text{skip}^D)[w_0/w] \\
= & \dots \text{Pred. calc. } (x' \notin w_0, w) \\
& \exists w_0 \bullet (\exists x' \bullet (q \vdash Q)[w_0/w'] \wedge \text{skip}^D[w_0/w]) \\
= & \dots \text{Pred. calc.} \\
& \exists x' \bullet (\exists w_0 \bullet (q \vdash Q)[w_0/w'] \wedge \text{skip}^D[w_0/w]) \\
= & \dots \text{Defn. of } \circ \\
& \exists x' \bullet (q \vdash Q) \circ \text{skip}^D \\
= & \dots \text{skip}^D \text{ is a unit } \circ \\
& \exists x' \bullet (q \vdash Q)
\end{aligned}$$

□

C.2.3 Assignment composition

We now introduce four laws associated with applying assignment within a sequential composition context. The first of these laws is the one that is most used, as it states how to apply the affects of an assignment operation to the next sequentially composed subprogram.

Left assignment composition: The following law states that the affect of applying an assignment of the value v to the variable x to the next sequentially composed subprogram, is to perform a semantic replacement of the composed subprogram's variable x with the value v . Note that this generalises to multiple assignment case in a straightforward manner.

Law C.13 – Left Assignment Composition:

$$x := v \circ (p \vdash P) = (p \vdash P)[v/x]$$

□

Proof

$$\begin{aligned}
& x := v \mathbin{\circ} (p \vdash P) \\
= & \dots \text{Defn. } _ := _ \\
& (\text{true} \vdash x' = v \wedge w' = w) \mathbin{\circ} (p \vdash P) \\
= & \dots \text{Law C.2 comp. design} \\
& \neg((x' = v \wedge w' = w) \mathbin{\circ} \neg p) \\
& \vdash \\
& (x' = v \wedge w' = w) \mathbin{\circ} P \\
= & \dots \text{Defn. } _ \mathbin{\circ} _ \\
& \neg(\exists w_0, x_0 \bullet (x_0 = v \wedge w_0 = w) \wedge \neg p[w_0, x_0/w, x]) \\
& \vdash \\
& \exists w_0, x_0 \bullet x_0 = v \wedge w_0 = w \wedge P[w_0, x_0/w, x] \\
= & \dots \text{1-point rule} \\
& \neg(\neg p[v/x]) \vdash P[v/x] \\
= & \dots \text{Defn. } _[-/-] \\
& (p \vdash P)[v/x]
\end{aligned}$$

□

Double assignment composition: The following law states that assigning two values to the same variable in succession, is the same as assigning the second value to the variable. Note this relies on the UTP constraint that expressions cannot have side-effects (or at least side-effects that affect the program's state).

Law C.14 – Double Assignment Composition:

$$x := e_1 \mathbin{\circ} x := e_2 = x := e_2$$

□

Proof

$$\begin{aligned}
& x := e_1 \mathbin{\circ} x := e_2 \\
= & \dots \text{Defn. } _ := _ \\
& (\text{true} \vdash x' = e_1 \wedge w' = w) \mathbin{\circ} (\text{true} \vdash x' = e_2 \wedge w' = w) \\
= & \dots \text{Corol. C.3 comp. pre-true} \\
& \text{true} \vdash (x' = e_1 \wedge w' = w) \mathbin{\circ} (x' = e_2 \wedge w' = w) \\
= & \dots \text{Defn. } _ \mathbin{\circ} _ \\
& \text{true} \vdash \exists x_0, w_0 \bullet x_0 = e_1 \wedge w_0 = w \wedge x' = e_2 \wedge w' = w_0 \\
= & \dots \text{1 point rule} \\
& \text{true} \vdash x' = e_2 \wedge w' = w \\
= & \dots \text{Defn. } _ := _ \\
& x := e_2
\end{aligned}$$

□

Right assignment elimination: The following two laws demonstrate how to eliminate an assignment that does not change the state of a program. The first law states that the affect of assigning a value to a variable can be pushed back into the preceding sequentially composed subprogram, which does not otherwise alter the variable, by updating this subprogram so that it also sets the variable to the assigned value.

Law C.15 – Right Assignment Composition:

$$(p \vdash P) \mathbin{\circ} x := v = p \vdash P[x/x'] \wedge x' = v \quad \text{where } P = (P \wedge x' = x)$$

□

Proof

$$\begin{aligned}
& (p \vdash P) \circledast x := v \\
= & \dots\dots\dots \text{Defn. } _ := _ \\
& (p \vdash P) \circledast (\text{true} \vdash x' = v \wedge w' = w) \\
= & \dots\dots\dots \text{Corol. C.3 comp. pre-true} \\
& p \vdash P \circledast (x' = v \wedge w' = w) \\
= & \dots\dots\dots \text{Defn. } _ \circledast _ \\
& p \vdash \exists w_0, x_0 \bullet P[w_0, x_0/w', x'] \wedge (x' = v \wedge w' = w_0) \\
= & \dots\dots\dots 1 \text{ point rule} \\
& p \vdash \exists x_0 \bullet P[x_0/x'] \wedge x' = v \\
= & \dots\dots\dots \text{A1 i.e. } P = (P \wedge x' = x) \\
& p \vdash \exists x_0 \bullet (P \wedge x' = x)[x_0/x'] \wedge x' = v \\
= & \dots\dots\dots \text{Defn. } _[-/_] \text{ and Pred. L.} \\
& p \vdash (\exists x_0 \bullet P[x_0/x'] \wedge x_0 = x) \wedge x' = v \\
= & \dots\dots\dots 1 \text{ point rule} \\
& p \vdash P[x/x'] \wedge x' = v
\end{aligned}$$

□

The second law states that a program that the reassigns of the same value to a variable can be eliminated.

Law C.16 – Special Assignment Composition:

$$x := v \circledast (p \vdash P[x/x'] \wedge x' = v) = x := v \circledast (p \vdash P) \quad \text{where } P = (P \wedge x' = x)$$

□

Proof

$$\begin{aligned}
& x := v \circledast (p \vdash P[x/x'] \wedge x' = v) \\
= & \dots\dots\dots \text{Law C.13 assign comp} \\
& (p \vdash P[x/x'] \wedge x' = v)[v/x] \\
= & \dots\dots\dots \text{Defn. } _[-/_] \\
& (p[v/x] \vdash P[x/x'] [v/x] \wedge x' = v) \\
= & \dots\dots\dots \text{Defn. } _[-/_] \\
& (p[v/x] \vdash P[v/x'] [v/x] \wedge x' = v) \\
= & \dots\dots\dots \text{Defn. } _[-/_] \\
& (p \vdash P[v/x'] \wedge x' = v)[v/x] \\
= & \dots\dots\dots v = x' \\
& (p \vdash P \wedge x' = v)[v/x] \\
= & \dots\dots\dots \text{A1 i.e. } P = (P \wedge x = x') \\
& (p \vdash P \wedge x' = x \wedge x' = v)[v/x] \\
= & \dots\dots\dots x = x' \\
& (p \vdash P \wedge x' = x \wedge x = v)[v/x] \\
= & \dots\dots\dots \text{A1 i.e. } P = (P \wedge x = x') \\
& (p \vdash P \wedge x = v)[v/x] \\
= & \dots\dots\dots \text{Similar rewriting} \\
& (p \vdash P)[v/x] \\
= & \dots\dots\dots \text{Law C.13 assign comp} \\
& x := v \circledast (p \vdash P)
\end{aligned}$$

□

C.2.4 Conjunctive implication

The following law demonstrates that the a conjunctive predicate involving a design can be weakened, by apply the other argument to the post-condition within the design, so long as this other argument does not involve the design specific observational variables.

Law C.17:

$$(p \vdash P) \wedge R \Rightarrow (p \vdash P \wedge R) \quad \text{provided } \Pi_{\text{OK}}, \Pi_{\text{OK}}' \notin R.$$

□

Proof

$$\begin{aligned} & (p \vdash P) \wedge R \\ = & \dots\dots\dots \text{Defn. of } _ \vdash _ \\ & (\Pi_{\text{OK}} \wedge p \Rightarrow \Pi_{\text{OK}}' \wedge P) \wedge R \\ = & \dots\dots\dots \text{Propositional Logic} \\ & (\Pi_{\text{OK}} \wedge p) \vee \neg R \Rightarrow (\Pi_{\text{OK}}' \wedge P) \wedge R \\ \Rightarrow & \dots\dots\dots \text{Propositional Logic} \\ & (\Pi_{\text{OK}} \wedge p) \Rightarrow (\Pi_{\text{OK}}' \wedge P) \wedge R \\ = & \dots\dots\dots \text{Defn. of } _ \vdash _ \\ & p \vdash P \wedge R \end{aligned}$$

□

C.3 Object Model Design Laws and Example Proofs

In this section we present (and prove) a law that is specific to the UTP models of the $\mathcal{O}$ -calculus, and proofs of the ‘factorial example’ lemmas of [Section 4.6.3](#). Most of the proofs concerning the UTP object models are contained in [Chapter 4](#).

C.3.1 Extract identity

The extract identity law states that substituting the UTP program call_j for the call_j program text within the body of a method has no effect on that method. This special case occurs when one iteration of the fix-point definition of procedure invocation is unwound.

Law C.18 – Extract identity:

$$\text{ext}_j(pt, \text{call}_j) = \llbracket pt \rrbracket_j$$

□

Proof

$$\begin{aligned} & \text{ext}_j(pt, \text{call}_j) \\ = & \dots\dots\dots \text{Defn. ext}_j \\ & \llbracket \text{fix}_j(pt, \text{text}_j(\text{call}_j)) \rrbracket \\ = & \dots\dots\dots \text{Defn. text}_j \\ & \llbracket \text{fix}_j(pt, \text{call}_j) \rrbracket \\ = & \dots\dots\dots \text{Defn. fix}_j \\ & \llbracket pt \rrbracket \end{aligned}$$

□

C.3.2 Factorial example UTP correspondence proofs

We now present the correspondence proof for each of the four steps in the factorial example ([Section 4.6.3](#)).

Step 1 – Method invocation to function invocation:

$$\llbracket sFact.\mathbf{fact}(n) \rrbracket_s = \llbracket sFactFn(n) \rrbracket_s$$

Proof

$$\begin{aligned}
& \llbracket sFact.\mathbf{fact}(n) \rrbracket_s \\
= & \dots\dots\dots \text{Compile } sFact.\mathbf{fact}(n) \\
& \llbracket SF(OF(O), n) \rrbracket \\
= & \dots\dots\dots \text{Defn. of } SF(OF(O), n) \\
& \llbracket vfPair_s^{CE}(\mathcal{E}_s n, OF(O), \mathcal{E}_s n) \rrbracket_s \\
= & \dots\dots\dots \text{Defn. of } OF(O) \text{ and } \llbracket _ \rrbracket \\
& vfPair_s^{CE}(\mathcal{E}_s n, omPair_s^{CE}(\mathcal{E}_s x, \mathcal{E}_s \mathbf{fact}) \rrbracket_s, \mathcal{E}_s n) \\
= & \dots\dots\dots \text{Defn. of } omPair_s^{CE} \\
& vfPair_s^{CE}(\mathcal{E}_s n, cmdExp_s(\mathcal{E}_s n, trans_s(omPair, 2), (\mathcal{E}_s O, \mathcal{E}_s \mathbf{fact})) \rrbracket_s, \mathcal{E}_s n) \\
= & \dots\dots\dots \text{Lemma 4.1 cmd. exp., where } (O, \mathbf{fact}) \in omPair \\
& vfPair_s^{CE}(\mathcal{E}_s n, \mathcal{E}_s omPair(O, \mathbf{fact}) \rrbracket_s, \mathcal{E}_s n) \\
= & \dots\dots\dots \text{Defn. of } omPair_s^{CE} \\
& vfPair_s^{CE}(\mathcal{E}_s n, \mathcal{E}_s(O, \llbracket x, F(x) \rrbracket) \rrbracket_s, \mathcal{E}_s n) \\
= & \dots\dots\dots \text{Lemma 4.3 proc. call} \\
& vfPair_s^{CE}(\mathcal{E}_s n, ext_s(F(x) \llbracket x \leftarrow O \rrbracket_s, call_s), \mathcal{E}_s n) \\
= & \dots\dots\dots \text{Defn. of } \llbracket _ \leftarrow _ \rrbracket_s \\
& vfPair_s^{CE}(\mathcal{E}_s n, ext_s(F(O), call_s), \mathcal{E}_s n) \\
= & \dots\dots\dots \text{Law C.18 extract id.} \\
& vfPair_s^{CE}(\mathcal{E}_s n, \llbracket F(O) \rrbracket, \mathcal{E}_s n)
\end{aligned}$$

$$\begin{aligned}
&= \dots \dots \dots \text{Defn. of } \llbracket - \rrbracket \\
&\quad \llbracket \text{vfPair}_s^{\text{CE}}(F(O), \mathcal{E}_s n) \mathbin{\text{\texttt{;}}} \text{call}_s \rrbracket \\
&= \dots \dots \dots \text{Compile } sFactFn(n) \\
&\quad \llbracket sFactFn(n) \rrbracket_s
\end{aligned}$$

□

Step 2 – Function invocation to conditional expression:

$$\llbracket sFactFn(n) \rrbracket_s = \llbracket sFactCond1 \rrbracket_s$$

Proof

$$\begin{aligned}
&\llbracket sFactFn(n) \rrbracket_s \\
&= \dots \dots \dots \text{Compile } sFactFn(n) \\
&\quad \llbracket SF(F(O), n) \rrbracket \\
&= \dots \dots \dots \text{Defn. of } SF(F(O), n) \\
&\quad \llbracket \text{vfPair}_s^{\text{CE}}(\\
&\quad \quad F(O), \mathcal{E}_s n \\
&\quad) \mathbin{\text{\texttt{;}}} \text{call}_s \rrbracket \\
&= \dots \dots \dots \text{Defn. of } F(O) \text{ and } \llbracket - \rrbracket \\
&\quad \text{vfPair}_s^{\text{CE}}(\\
&\quad \quad \mathcal{E}_s(\llbracket i, E1(O, i) \rrbracket), \mathcal{E}_s n \\
&\quad) \mathbin{\text{\texttt{;}}} \text{call}_s \\
&= \dots \dots \dots \text{Defn. of } \text{vfPair}_s^{\text{CE}} \\
&\quad \text{cmdExp}_s(\\
&\quad \quad \text{trans}_s(\text{vfPair}, 2), \\
&\quad \quad (\mathcal{E}_s(\llbracket i, E1(O, i) \rrbracket), \mathcal{E}_s n) \\
&\quad) \mathbin{\text{\texttt{;}}} \text{call}_s \\
&= \dots \dots \dots \text{Lemma 4.1 cmd. exp., where} \\
&\quad \mathcal{E}_s \text{vfPair}(\llbracket i, E1(O, i) \rrbracket, n) \mathbin{\text{\texttt{;}}} \quad (\llbracket i, E1(O, i) \rrbracket, n) \in \text{vfPair} \\
&\quad \text{call}_s \\
&= \dots \dots \dots \text{Defn. of } \text{vfPair} \\
&\quad \mathcal{E}_s(n, \llbracket i, E1(O, i) \rrbracket) \mathbin{\text{\texttt{;}}} \\
&\quad \text{call}_s \\
&= \dots \dots \dots \text{Lemma 4.3 proc. call} \\
&\quad \text{ext}_s(E1(O, i) \llbracket i \leftarrow n \rrbracket_s, \text{call}_s) \\
&= \dots \dots \dots \text{Defn. of } \llbracket - \leftarrow - \rrbracket_s \\
&\quad \text{ext}_s(E1(O, n), \text{call}_s) \\
&= \dots \dots \dots \text{Law C.18 extract id.} \\
&\quad \llbracket E1(O, n) \rrbracket \\
&= \dots \dots \dots \text{Compile } sFactCond1 \\
&\quad \llbracket sFactCond1 \rrbracket_s
\end{aligned}$$

□

Step 3 – Conditional expression decision evaluation:

$$\llbracket sFactCond1 \rrbracket_s = \llbracket sFactCond2 \rrbracket_s$$

Proof

$$\begin{aligned}
&\llbracket sFactCond1 \rrbracket_s \\
&= \dots \dots \dots \text{Compile } sFactCond1 \\
&\quad \llbracket E1(O, n) \rrbracket
\end{aligned}$$

$$\begin{aligned}
&= \dots\dots\dots \text{Defn. of } E1(O, n) \text{ and } \llbracket - \rrbracket \\
&\quad (\mathcal{E}_s n <_s^{\text{CE}} \mathcal{E}_s 2) \mathbin{\text{\textcircled{;}}} \llbracket SW(O, n) \rrbracket \\
&= \dots\dots\dots \text{Defn. of } (- <_s^{\text{CE}} -) \\
&\quad \text{cmdExp}_s(\text{trans}_s(- < -, 2), (\mathcal{E}_s n, \mathcal{E}_s 2)) \mathbin{\text{\textcircled{;}}} \\
&\quad \llbracket SW(O, n) \rrbracket \\
&= \dots\dots\dots \text{Lemma 4.1 cmd. exp.,} \\
&\quad \mathcal{E}_s(- < -)(n, 2) \mathbin{\text{\textcircled{;}}} \text{ where } (n, 2) \in (- < -) \\
&\quad \llbracket SW(O, n) \rrbracket \\
&= \dots\dots\dots \text{Assumption that } n > 1 \\
&\quad \mathcal{E}_s \text{false} \mathbin{\text{\textcircled{;}}} \llbracket SW(O, n) \rrbracket \\
&= \dots\dots\dots \text{Defn. of } \llbracket - \rrbracket \\
&\quad \llbracket \mathcal{E}_s \text{false} \mathbin{\text{\textcircled{;}}} SW(O, n) \rrbracket \\
&= \dots\dots\dots \text{Compile } sFactCond2 \\
&\quad \llbracket sFactCond2 \rrbracket_s \\
&\square
\end{aligned}$$

Step 4 – Selection of else conditional expression branch:

$$\llbracket sFactCond2 \rrbracket_s = \llbracket sFactIter \rrbracket_s$$

Proof

$$\begin{aligned}
&\llbracket sFactCond2 \rrbracket_s \\
&= \dots\dots\dots \text{Compile } sFactCond2 \\
&\quad \llbracket E2(O, n, \text{false}) \rrbracket \\
&= \dots\dots\dots \text{Defn. of } E2(O, n, \text{false}) \\
&\quad \llbracket \mathcal{E}_s \text{false} \mathbin{\text{\textcircled{;}}} SW(O, n) \rrbracket \\
&= \dots\dots\dots \text{Defn. of } SW(O, n) \\
&\quad \llbracket \mathcal{E}_s \text{false} \mathbin{\text{\textcircled{;}}} \text{switch}_s(\mathcal{E}_s 1, ITER(O, n)) \rrbracket \\
&= \dots\dots\dots \text{Defn. of } \llbracket - \rrbracket \\
&\quad \mathcal{E}_s \text{false} \mathbin{\text{\textcircled{;}}} \text{switch}_s(\mathcal{E}_s 1, \llbracket ITER(O, n) \rrbracket) \\
&= \dots\dots\dots \text{Lemma 4.9 cond. false} \\
&\quad \llbracket ITER(O, n) \rrbracket \\
&= \dots\dots\dots \text{Compile } sFactIter \\
&\quad \llbracket sFactIter \rrbracket_s \\
&\square
\end{aligned}$$

Appendix D

Combined model extended example

This appendix contains the detailed working of the evaluation of the first object-method pair in [Section 6.2.3](#). Here the *Excel* method invocation starts by constructing an object-method pair. This takes thirteen steps from the initial state, which is illustrated in [Figure D.1](#).

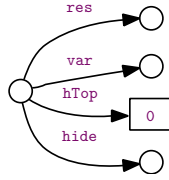


Figure D.1: Evaluate pair – Step 00 – Initial state

For convenience we now present another copy of the thirteen step program:

```

( decl  $x_1, x_2$  in           (1)
  ( hide  $x_1, x_2$  from      (2)
     $\mathcal{E}_L \text{Excel}$     (3)
  ) §                       (4)
   $x_1 :=_L \Pi_{\text{RES}}$  §       (5)
  ( hide  $x_1, x_2$  from      (6)
     $\text{getMethod}_L \text{get}$     (7)
  ) §                       (8)
   $x_2 :=_L \Pi_{\text{RES}}$          (9)
   $\Pi_{\text{RES}} :=_L (\mathfrak{!}, \mathfrak{!})$     (10)
   $\Pi_{\text{RES}.1} :=_L x_1$       (11)
   $\Pi_{\text{RES}.2} :=_L x_2$       (12)
)                             (13)

```

The first step declares two local variables (x_1 and x_2), which are used to store the contents of the pair. This is illustrated in [Figure D.2](#).

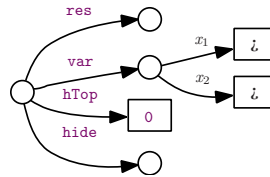


Figure D.2: Evaluate pair – Step 01 – Declare local variables

The second step hides the introduced local variables from the calculation of the first element within the pair. This is illustrated in [Figure D.3](#).

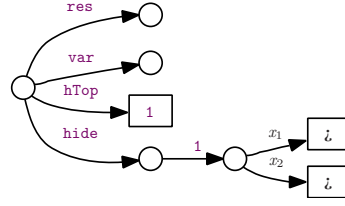


Figure D.3: Evaluate pair – Step 02 – Hide local variables

The third step evaluates the first element of the pair, which is the target object of the method invocation. This is illustrated in Figure D.4.

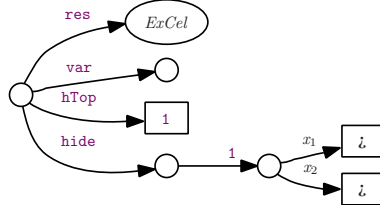


Figure D.4: Evaluate pair – Step 03 – Evaluated object

The fourth step restores the local variables following the evaluation of the first element of the pair. This is illustrated in Figure D.5.

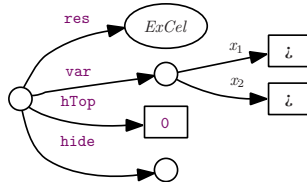


Figure D.5: Evaluate pair – Step 04 – Restore local variables

The fifth step stores the first element of the pair (i.e. the target object) in local variable x_1 . This is illustrated in Figure D.6.

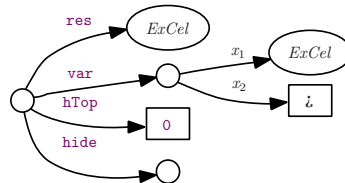


Figure D.6: Evaluate pair – Step 05 – Store object in first local variable

The sixth step hides the introduced local variables from the calculation of the second element within the pair. This is illustrated in Figure D.7.

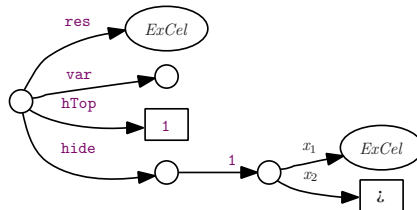


Figure D.7: Evaluate pair – Step 06 – Hide local variables

The seventh step evaluates the second element of the pair by looking up the method that is associated with the given label (**get**) within the target object. This is illustrated in Figure D.8.

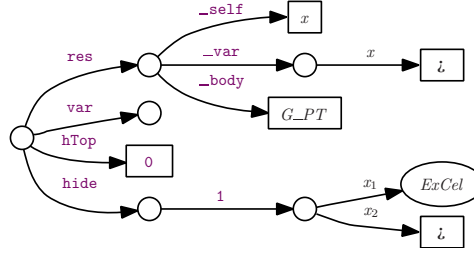


Figure D.8: Evaluate pair – Step 07 – Get method with label **get**

The eighth step restores the local variables following the evaluation of the first element of the pair. This is illustrated in Figure D.9.

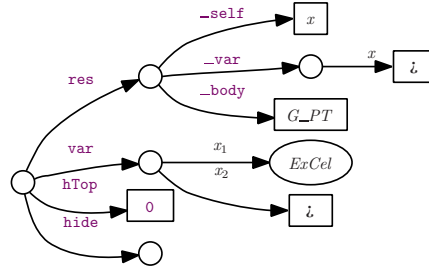


Figure D.9: Evaluate pair – Step 08 – Restore local variables

The ninth step stores the second element of the pair (i.e. the method to be invoked) in local variable x_2 . This is illustrated in Figure D.10.

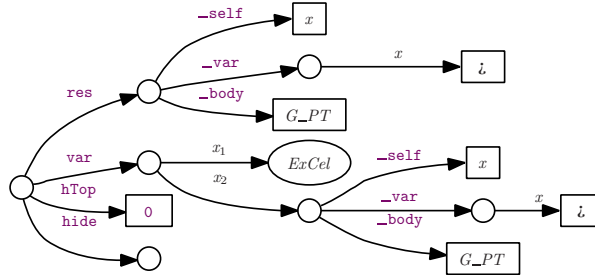


Figure D.10: Evaluate pair – Step 09 – Store method in second local variable

The tenth step creates an empty output pair, which will eventually contain the result of the object-method pair calculation that is being performed. This is illustrated in Figure D.11.

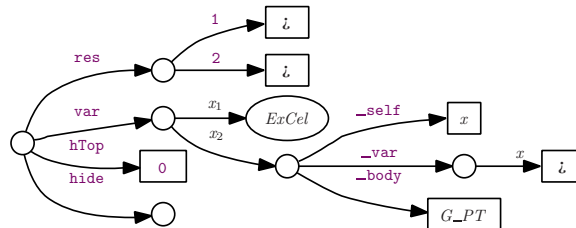


Figure D.11: Evaluate pair – Step 10 – Create an empty output pair

The eleventh step copies the target object to the first element in the output pair. This is illustrated in Figure D.12.

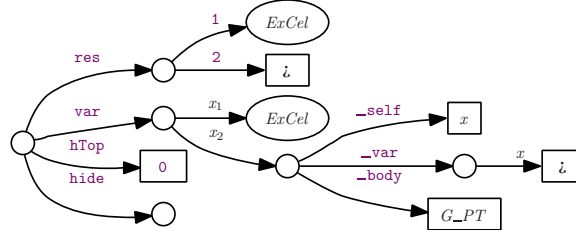


Figure D.12: Evaluate pair – Step 11 – Copy object to output pair

The twelfth step copies the method to be invoked to the second element in the output pair. This is illustrated in Figure D.13.

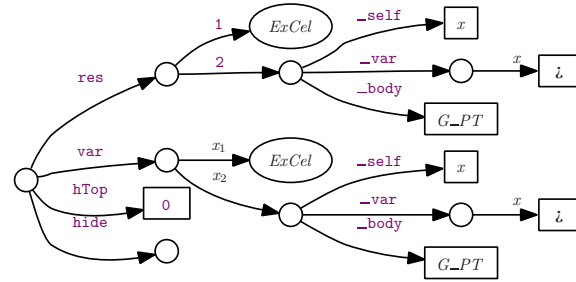


Figure D.13: Evaluate pair – Step 12 – Copy method to output pair

The thirteenth step removes the local variables from scope. This is illustrated in Figure D.14.

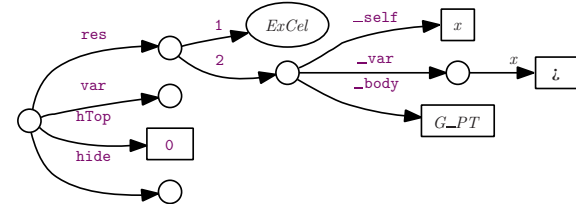


Figure D.14: Evaluate pair – Step 13 – Remove local variables from scope

List of Abbreviations

The entries in this list of abbreviations are either *general* or *product-specific*, where: the text for a *general* abbreviation is in lower-case words; and the text for a *product-specific* abbreviation is in capitalised words. The only exception to this rule is when an abbreviation contains another abbreviation, in which case the inner abbreviation is formatted as an upper case word.

ALTG	Abstract Location Trace Graph
CCS	Calculus of Communicating Systems
CLOS	Common Lisp Object System
CORBA	Common Object Request Broker Architecture
CPU	Central Processing Unit
CSP	Communicating Sequential Processes
FDR	Failures-Divergence and Refinement
JVM	Java virtual machine
λ-calculus	untyped lambda calculus
$\lambda_{\rightarrow}$-calculus	simply typed λ -calculus
LHS	left hand side
OO	object-oriented
π-calculus	pi-calculus
rCOS	Refinement Calculus of Object Systems
RHS	right hand side
RMI	Remote Method Invocation
SA	Systems Assurance
ς-calculus	untyped object calculus
TCOZ	Timed Communicating Object-Z
$\mathcal{U}_M$-design	UTP constant-map-design

List of Abbreviations

$\mathcal{U}_R$-design	UTP result–value-design
$\mathcal{U}_S$-design	UTP operand–stack-design
$\mathcal{U}_L$-design	UTP location–graph-design
UML	Unified Modelling Language
UTP	Unifying Theories of Programming