

A Logical Approach to Type Soundness

AMIN TIMANY, Aarhus University, Denmark

ROBBERT KREBBERS, Radboud University Nijmegen, The Netherlands

DEREK DREYER, MPI-SWS, Germany

LARS BIRKEDAL, Aarhus University, Denmark

Type soundness, which asserts that “well-typed programs cannot go wrong”, is widely viewed as the canonical theorem one must prove to establish that a type system is doing its job. It is commonly proved using the so-called *syntactic approach* (aka *progress and preservation*), which has had a huge impact on the study and teaching of programming language foundations. Unfortunately, syntactic type soundness is a rather weak theorem. It only applies to programs that are well-typed in their entirety, and thus tells us nothing about the many programs written in “safe” languages that make use of “unsafe” language features. Even worse, it tells us nothing about whether type systems achieve one of their main goals: enforcement of data abstraction. One can easily define a language that enjoys syntactic type soundness and yet fails to support even the most basic modular reasoning principles for abstraction mechanisms like closures, objects, and abstract data types.

Given these concerns, we argue that programming languages researchers should no longer be satisfied with proving syntactic type soundness, and should instead start proving *semantic type soundness*, a more useful theorem which captures more accurately what type systems are actually good for. Semantic type soundness is an old idea—Milner’s original account of type soundness from 1978 was semantic—but it fell out of favor in the 1990s due to limitations and complexities of denotational models. In the succeeding decades, thanks to a series of technical advances—notably, *step-indexed Kripke logical relations* constructed over operational semantics, and *higher-order concurrent separation logic* as consolidated in the *Iris* framework in Coq—we can now build (machine-checked) semantic soundness proofs at a much higher level of abstraction than was previously possible.

The resulting “logical” approach to semantic type soundness has already been employed to great effect in a number of recent papers, but those papers typically (a) concern advanced problem scenarios that complicate the presentation, (b) assume significant prior knowledge of the reader, and (c) suppress many details of the proofs. Here, we aim to provide a gentler, more pedagogically motivated introduction to logical type soundness, targeted at a broader audience that may or may not be familiar with logical relations and Iris. As a bonus, we also show how logical type soundness proofs can easily be generalized to establish an even stronger *relational* property—*representation independence*—for realistic type systems.

CCS Concepts: • **Theory of computation** → **Abstraction; Separation logic; Type theory; Logic and verification**.

Additional Key Words and Phrases: Type soundness, data abstraction, logical relations, step-indexing, concurrent separation logic, Iris, Coq

ACM Reference Format:

Amin Timany, Robbert Krebbers, Derek Dreyer, and Lars Birkedal. 2024. A Logical Approach to Type Soundness. *J. ACM* 1, 1, Article 1 (June 2024), 74 pages. <https://doi.org/xx/xx>

Authors’ addresses: [Amin Timany](#), Aarhus University, Denmark, timany@cs.au.dk; [Robbert Krebbers](#), Radboud University Nijmegen, The Netherlands, mail@robbertkrebbers.nl; [Derek Dreyer](#), MPI-SWS, Saarland Informatics Campus, Germany, dreyer@mpi-sws.org; [Lars Birkedal](#), Aarhus University, Denmark, birkedal@cs.au.dk.

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

© 2024 Copyright held by the owner/author(s).

ACM 0004-5411/2024/6-ART1

<https://doi.org/xx/xx>

Type structure is a syntactic discipline for enforcing levels of abstraction.

– Reynolds [1983]

Although types and assertions may be semantically similar, the actual development of type systems for programming languages has been quite separate from the development of approaches to specification such as Hoare logic... the real question is whether the dividing line between types and assertions can be erased.

– Reynolds [2002]

This paper is dedicated to the memory of John C. Reynolds.

1 INTRODUCTION

The *type soundness* (or *type safety*) theorem for a programming language states that if a program in that language passes the type checker, then it is guaranteed to have well-defined behavior when executed. Introduced over 40 years ago by Milner [1978], type soundness has become the canonical property that type systems for “safe” programming languages are expected to satisfy.

In Milner’s original formulation for a λ -calculus with ML-style polymorphism, type soundness was characterized using denotational semantics. Ill-behaved programs were assigned a special denotation “wrong”, and the type soundness theorem stated that well-typed programs could not “go wrong” (*i.e.*, have “wrong” as their denotation). However, it turned out to be difficult to scale this methodology to richer type systems with features such as general recursive types, higher-order mutable state, control operators, and concurrency.

Syntactic type soundness. Today, the most common formulation of type soundness is the so-called “syntactic approach”, pioneered by Wright and Felleisen [1994] and subsequently simplified by Robert Harper into the two theorems known as “progress and preservation”.¹ Instead of employing a “wrong” denotation, the syntactic approach characterizes undefined behavior *operationally*: a program has undefined behavior if its execution under a small-step operational semantics “gets stuck” (*i.e.*, reaches a non-terminal state where there is no next step of execution to take). The preservation theorem states that a program remains well-typed as it executes, and the progress theorem states that a well-typed program is either in a terminal state or its next step of execution is well-defined.² Together, these theorems imply that the execution of a well-typed program is well-defined in the sense that it never gets stuck.

The syntactic approach to type soundness via progress and preservation is arguably one of the “greatest hits” of programming languages research of the past three decades. In addition to being conceptually simple, the approach scales easily to handle a wide range of programming language features, and it has been popularized effectively through the central organizing role it plays in the textbooks of Pierce [2002] and Harper [2016]. As a result, it has become one of the most widely known, widely taught, and widely applied formal methods in the entire area of programming language foundations, with countless research papers on type systems concluding triumphantly with a statement of progress and preservation.

The limitations of syntactic type soundness. Unfortunately, syntactic type soundness also suffers from two significant limitations that are not (in our experience) widely recognized.

¹According to Felleisen and Morrisett [personal communication, Nov. 2017], Harper is responsible for suggesting the reorganization of syntactic type soundness using a progress theorem (and the associated “canonical forms” lemma), which superseded Wright and Felleisen’s more complex analysis of “faulty” expressions. To our knowledge, this revised proof structure was deployed for the first time in Morrisett et al. [1995]. See Harper [2016] for a modern presentation.

²In order to state these theorems, one must first generalize the notion of well-typed programs to a notion of well-typed machine states, but this is typically straightforward.

The first limitation pertains to *data abstraction*. One of the primary functions of the type systems of many languages is to give programmers a way of enforcing data abstraction boundaries, so that one can place invariants on the private data representations of modules, objects, abstract data types, *etc.* and be sure that client code will not violate those invariants. However, syntactic type soundness offers no guarantees about whether a programming language’s data abstraction facility actually works—it is easy to prove syntactic soundness of a type system whose data abstraction mechanism is completely broken.

The second limitation pertains to *unsafe features*. In practice, most “safe” languages provide unsafe escape hatches—*e.g.*, `Obj.magic` in OCaml, `unsafePerformIO` in Haskell, `unsafe blocks` in Rust—which enable programmers to perform potentially unsafe operations (such as unchecked type casts or array accesses) that the safe fragment of the language disallows. These unsafe escape hatches have proven indispensable, both for functionality—when the language’s safe type system is more restrictive than necessary—and for efficiency—when performance concerns demand the use of a lower-level unsafe abstraction. However, syntactic type soundness has nothing to say about programs that use unsafe features: it simply declares such programs out of scope.

These two limitations are in fact closely connected, in that it is common to justify the “safe” use of unsafe features by appeal to data abstraction. Specifically, programmers often argue informally that their use of unsafe operations is harmless because said operations have been *encapsulated* behind a “safe API”. That is, they argue that, thanks to the abstraction boundary of the API, the implementation of the API can enforce invariants on its private data representation which ensure that its use of unsafe features does not lead to any undefined behavior. But of course, to make this reasoning formal, one needs to know whether the language is enforcing data abstraction properly, precisely one of the issues on which syntactic type soundness is silent.

Together, these limitations suggest that syntactic type soundness does not provide a sufficient foundation for judging whether a type system is really doing its job. One may then rightly wonder: can we do any better? And we are here to say: yes, we can!

Logical type soundness. We propose an alternative to syntactic type soundness that overcomes the aforementioned limitations, offering a flexible foundation for reasoning about data abstraction, as well as the safe use of unsafe features. We call our approach *logical type soundness*. The essence of logical type soundness is not new: it is the age-old idea of *semantic type soundness*, as exemplified by the formulation in the original paper of Milner [1978]. Under the semantic soundness approach, one defines a semantic model of types, which offers an extensional view of typing rather than an intensional one. In other words, unlike syntactic typing, which dictates the syntactic structure of well-typed terms, semantic typing merely places restrictions on their observable behavior. Accordingly, it enables us to explain when a term behaves safely at a given type, even if the term employs unsafe or low-level operations internally.

Although Milner built his semantic model over a denotational semantics, we follow more recent approaches [Birkedal et al. 2011; Schwinghammer et al. 2013] and build ours over an operational semantics. In particular, we are inspired by the work by Appel, Ahmed, and their collaborators on the Foundational Proof-Carrying Code project [Appel and Felty 2000; Appel 2001; Appel and McAllester 2001; Ahmed et al. 2002; Ahmed 2004; Ahmed et al. 2010], which demonstrated how to scale semantic soundness to account for a wide variety of programming language features using the powerful technique of *step-indexed models*.

The key point of difference between our approach and theirs is the level of abstraction at which the semantic soundness proof is conducted. As we explain in detail in §4.3 (and as previously noted by Appel et al. [2007] and Dreyer et al. [2011]), prior work that built semantic soundness proofs directly using step-indexed models involved a great deal of explicit reasoning about step-indexing

and about the quasi-circular constructions that step-indexing serves to disentangle. Such reasoning quickly became very tedious, to the point that the high-level structure of a proof would become obscured if one were to write out all the low-level details.

In contrast, we show how to lift semantic soundness proofs to a much higher level of abstraction by employing recent advances in *higher-order concurrent separation logic* [Svendsen et al. 2013; Svendsen and Birkedal 2014]—hence the name “logical type soundness”. Specifically, we show how by using the separation-logic framework *Iris* [Jung et al. 2015, 2016; Krebbers et al. 2017a; Jung et al. 2018b], we can formulate semantic soundness proofs—for *feature-rich, realistic languages*—in a clear and concise manner, uncluttered by the low-level details of prior accounts. As a major added bonus, *Iris* is implemented in the Coq proof assistant and provides effective tactic support for constructing machine-checked logical type soundness proofs with relative ease [Krebbers et al. 2017b, 2018].

Type soundness expresses a property of a single program and hence it is sometimes referred to as a *unary* property. It turns out that our logical approach to type soundness can be easily adapted to support *relational reasoning* as well. Here, relational reasoning refers to properties about *pairs* of programs, sometimes referred to as *binary* properties. A particularly important example of such a binary property is *representation independence* [Reynolds 1974; Mitchell 1986]. Representation independence is a strong guarantee on the effectiveness of a language’s data abstraction facility, even stronger than semantic soundness—it ensures that one can change the internal data representation of an *abstract data type* (ADT)³ without affecting the behavior of its clients. Yet, similarly to semantic type soundness, prior state-of-the-art semantic models for proving representation independence have typically been expressed directly in set theory using explicit step-indexing, see e.g., Neis et al. [2009], Ahmed et al. [2009], Dreyer et al. [2010, 2012], Thamsborg and Birkedal [2011], and Birkedal et al. [2012, 2013]. We demonstrate by example how the advanced features of *Iris* can be used to formalize (machine-checked) proofs of representation independence at a higher level of abstraction than was previously possible.

Goal of this paper. Over the last five years, many papers have demonstrated that the logical approach to type soundness in *Iris* is eminently practical and scalable: among other things, it has been used for a machine-checked proof of type soundness of a significant subset of the Rust programming language [Jung et al. 2018a, 2021; Jung 2020; Dang et al. 2020], an extension of Scala’s core type system DOT [Giarrusso et al. 2020], session types [Hinrichsen et al. 2021; Jacobs et al. 2024], and refinement types for the C programming language [Sammler et al. 2021]. Aside from type soundness, the logical approach has also been used to prove robust safety [Swasey et al. 2017; Sammler et al. 2020; Rao et al. 2023], various forms of representation independence and program refinement [Krogh-Jespersen et al. 2017; Tassarotti et al. 2017; Timany et al. 2018; Timany and Birkedal 2019; Frumin et al. 2018, 2021b; Jacobs et al. 2021], and various security properties [Frumin et al. 2021a; Gregersen et al. 2021; Georges et al. 2021].

The aforementioned papers are driven by particular applications and thus use the logical approach in sophisticated ways, typically in the context of a complicated programming language, type system, or program property. As a consequence, those papers typically omit many details and presuppose expert knowledge. Our goal in this paper is instead pedagogical: to make the general technique of logical type soundness better known to a wider audience. Thus, we present it in the context of a simple programming language with a pedestrian set of features and without assuming that the reader is already well-versed in separation logic and step-indexing. Our intention is that this paper should be accessible to researchers and students who are familiar with basic textbooks in programming language theory such as Pierce [2002] or Harper [2016].

³The acronym ADT is sometimes used to mean “*algebraic data type*”, but in this paper we always mean “*abstract data type*”.

A note about the proofs in this paper: Most of our proofs are carried out *within* the Iris logic. As Iris is a modal and substructural logic, proofs in Iris are of a rather different (and likely unfamiliar) nature compared with proofs in ordinary (higher-order) logic. Hence, we use proof trees to spell out our Iris proofs in great detail, showing exactly which proof rules are applied where. However, we hasten to note that this is merely for formal clarity; in practice, when you are developing such Iris proofs *in Coq* (as nearly all Iris users do), much of this explicit detail is kept implicit, since the Iris Proof Mode [Krebbers et al. 2017b, 2018] keeps track of the Iris proof context and performs many “boring” proof steps *automatically*. Though a presentation of the Iris Proof Mode is beyond the scope of this paper, we refer the interested reader to the above-cited papers and accompanying online tutorials (see §10) for further details.

Outline. In §2, we define a small but rich programming language—with higher-order state, recursive types, abstract types, and concurrency—which we will use throughout the rest of the paper, and we sketch the syntactic type soundness result for it. In §3, we explain the limitations of syntactic type soundness in more detail. In §4, we give a high-level description of the logical approach to type soundness and provide an extensive comparison of our approach to prior work on semantic type soundness. In §5, we present the definition of a *logical relation*—the core ingredient for proving logical type soundness in Iris—and, in §6, we present the corresponding proofs. In §7, we show how the logical approach allows us to reason about safe encapsulation of unsafe features, and in §8, we extend the logical approach to support relational reasoning about representation independence. The relevant features and proof rules of Iris are introduced along the way in §5–§8. Finally, in §9, we discuss related work, and in §10 we conclude with a brief discussion of recent work that has employed our logical approach to type soundness and relational reasoning.

Origin of this paper. The technical content of this paper is based in part on Krebbers et al. [2017b, §6] and Timany [2018, Chapter 5]. Krebbers et al. provide a (2-page) case study showing that Iris and the Iris Proof Mode can be used for the mechanization of both semantic type soundness and representation independence proofs, and Chapter 5 of Timany [2018]’s PhD thesis provides a more extensive description of this case study. The present paper can be seen as a significant expansion of the above, explaining semantic type soundness from first principles in a more didactic fashion, without requiring prior knowledge of Iris, and with motivating examples drawn from Dreyer’s keynote talk at the POPL 2018 conference [Dreyer 2018].

2 THE LANGUAGE MYLANG AND ITS SYNTACTIC TYPE SOUNDNESS

We present the syntax and the semantics of our subject of study: the language **MyLang**—a call-by-value λ -calculus with impredicative polymorphism, iso-recursive types, higher-order state, and fine-grained concurrency. We start by describing the syntax (§2.1), typing (§2.2), and operational semantics (§2.3) of **MyLang**. Finally, we make the notion of type soundness formal (§2.4) and show how it is proved using the standard syntactic approach (§2.5).

2.1 Syntax

We present the syntax of **MyLang** in two variations: *static* expressions $\hat{e} \in \widehat{Expr}$, and *dynamic* expressions $e \in Expr$. The idea behind this distinction is that the static syntax is used for writing surface programs, but in order for these programs to be executed, they must first be transformed by an *erasure* function $|_| : \widehat{Expr} \rightarrow Expr$ into dynamic programs. Since ultimately we would like to prove the safety of programs that are not syntactically well-typed, throughout most of this paper we will work with dynamic syntax. In particular, we will define the operational semantics (§2.3)

$A, B \in \text{Type} ::= \alpha \in \text{Tvar} \mid 1 \mid 2 \mid \mathbb{Z} \mid A \times A \mid A + A \mid A \rightarrow A \mid \forall \alpha. A \mid \exists \alpha. A \mid \mu \alpha. A \mid \text{ref } A$	
$\odot \in \text{BinOp} ::= + \mid * \mid - \mid < \mid =$	
$\hat{e} \in \widehat{\text{Expr}} ::= x \in \text{Var} \mid \text{rec } f(x) = \hat{e} \mid \hat{e} \hat{e} \mid \Lambda \alpha. \hat{e} \mid \hat{e} \langle A \rangle \mid$	(Polymorphic λ -calculus)
$() \mid n \in \mathbb{Z} \mid \hat{e} \odot \hat{e} \mid$	(Unit type and arithmetic)
$\text{true} \mid \text{false} \mid \text{if } \hat{e} \text{ then } \hat{e} \text{ else } \hat{e} \mid$	(Booleans)
$(\hat{e}, \hat{e}) \mid \pi_1 \hat{e} \mid \pi_2 \hat{e} \mid$	(Products)
$\text{inj}_1 \hat{e} \mid \text{inj}_2 \hat{e} \mid (\text{match } \hat{e} \text{ with } \text{inj}_1 x \Rightarrow \hat{e}_1 \mid \text{inj}_2 x \Rightarrow \hat{e}_2 \text{ end}) \mid$	(Sums)
$\text{pack } \langle B, \hat{e} \rangle \text{ as } \exists \alpha. A \mid \text{match } \hat{e} \text{ with pack } \langle \alpha, x \rangle \Rightarrow \hat{e} \text{ end} \mid$	(Existentials)
$\text{fold } \hat{e} \mid \text{unfold } \hat{e} \mid$	(Iso-recursive types)
$\text{ref } \hat{e} \mid !\hat{e} \mid \hat{e} \leftarrow \hat{e} \mid \text{CAS}(\hat{e}, \hat{e}, \hat{e}) \mid \text{FAA}(\hat{e}, \hat{e}) \mid$	(References)
$\text{fork } \{\hat{e}\}$	(Concurrency)
$e \in \text{Expr} ::= x \in \text{Var} \mid \text{rec } f(x) = e \mid e \mid e \mid \Lambda. e \mid e \langle \rangle \mid$	(Polymorphic λ -calculus)
$() \mid n \in \mathbb{Z} \mid e \odot e \mid$	(Unit type and arithmetic)
$\text{true} \mid \text{false} \mid \text{if } e \text{ then } e \text{ else } e \mid$	(Booleans)
$(e, e) \mid \pi_1 e \mid \pi_2 e \mid$	(Products)
$\text{inj}_1 e \mid \text{inj}_2 e \mid \text{match } e \text{ with } \text{inj}_i x \Rightarrow e_i \text{ end} \mid$	(Sums)
$\text{pack } \langle e \rangle \mid \text{match } e \text{ with pack } \langle x \rangle \Rightarrow e \text{ end} \mid$	(Existentials)
$\text{fold } e \mid \text{unfold } e \mid$	(Iso-recursive types)
$\ell \in \text{Loc} \mid \text{ref } e \mid !e \mid e \leftarrow e \mid \text{CAS}(e, e, e) \mid \text{FAA}(e, e) \mid$	(References)
$\text{fork } \{e\}$	(Concurrency)

Fig. 1. Syntax of **MyLang**: Types A, B , binary operators $\odot$, static expressions $\hat{e}$, and dynamic expressions e .

and our semantic type system (§4–§6) on the dynamic syntax. However, when presenting example programs, we use the static syntax, since these programs are written by users of **MyLang**.⁴

The syntax of types, static expressions, and dynamic expressions is shown in Figure 1. We let α range over Tvar , a countably infinite set of type variables, and let x and f range over Var , a countably infinite set of term variables.

The static expressions of **MyLang** are in Church style, *i.e.*, they include type annotations (marked in red). The dynamic, Curry-style syntax of **MyLang** is obtained by simply erasing all type annotations and by adding an additional literal $\ell \in \text{Loc}$ for memory locations. Locations only appear in the dynamic syntax because the programmer is not permitted to write them directly in the source program—they only get created dynamically during execution.

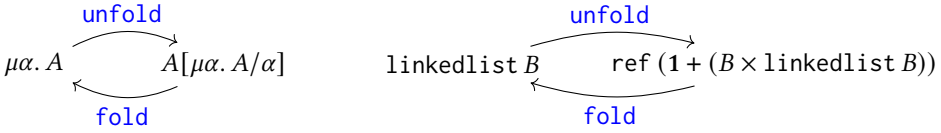
The ground types of **MyLang** are: the unit type 1, the type of Booleans 2, and the type of integers $\mathbb{Z}$. Basic type formers include products ($A \times B$), sums ($A + B$), and function types ($A \rightarrow B$). Types also include recursive types ($\mu \alpha. A$), polymorphic types ($\forall \alpha. A$), and existential types ($\exists \alpha. A$), which

⁴In our dynamic expression language, the reader may notice that we leave in the “markers” of type abstractions ($\Lambda. e$) and type instantiations ($e \langle \rangle$), and likewise for the existentially-typed constructs, even though the type arguments have been erased. This is following the approach taken by Ahmed [2006]. It has the benefit, *e.g.*, that the erasure of a type abstraction remains a value—if instead the erasure of $\Lambda \alpha. \hat{e}$ were simply the erasure of e , this would not be the case. An alternative approach would be to impose a *value restriction* [Wright 1995] on type abstractions—see Pitts [2005] for example.

classify *abstract data types* (ADTs). The type $\text{ref } A$ is the type of memory locations that store values of type A .

Following Mitchell and Plotkin [1988], the expression $\text{pack } \langle B, \hat{e} \rangle \text{ as } \exists \alpha. A$ represents an abstract data type (ADT), *i.e.*, an expression of existential type $(\exists \alpha. A)$, which “packs” the type “witness” B (representing the abstract type α) together with the term $\hat{e}$ (representing the operations on the ADT of type A). The type witness is “abstract” in the sense that there is no way for clients of the ADT to observe the implementation of α as B . ADTs can be unpacked using a syntax similar to ML-style pattern matching: $\text{match } \hat{e}_1 \text{ with pack } \langle \alpha, x \rangle \Rightarrow \hat{e}_2 \text{ end}$. Here, the term component of the ADT $\hat{e}_1$ can be referred to as x (and its type witness as α) within the scope of the expression $\hat{e}_2$.

Recursive types in **MyLang** are *iso-recursive*, meaning that explicit **fold** and **unfold** operations are used to coerce an expression between a recursive type $(\mu \alpha. A)$ and its expansion $(A[\mu \alpha. A/\alpha])$. For example, linked lists with elements of type B are given by $\text{linkedlist } B \triangleq \mu \alpha. \text{ref } (1 + (B \times \alpha))$, where the sum $+$ indicates that the list is either “nil” (*i.e.*, 1) or a “cons” (*i.e.*, $B \times \alpha$).



An alternative would be to support *equi-recursive* types, whereby a recursive type is equivalent to its unrolling. We employ iso-recursive types in this paper merely for simplicity, to avoid a non-trivial *syntactic* type equivalence relation. Concerning the logical approach to type soundness presented later in the paper, the technical development could be adapted easily to handle equi-recursive types. See Jung et al. [2018a] and Hinrichsen et al. [2021] for examples of such developments.

References can be allocated, read from, and written to using the **ref** $\hat{e}$, **!** $\hat{e}$, and $\hat{e}_1 \leftarrow \hat{e}_2$ expressions, respectively. The compare-and-set (**CAS**) and fetch-and-add (**FAA**) operations are **MyLang** primitives for fine-grained concurrency. The expression **CAS** $(\hat{e}_1, \hat{e}_2, \hat{e}_3)$ evaluates the three subexpressions to values v_1 , v_2 , and v_3 , where v_1 must be a memory location ℓ ; it then *atomically* checks if the value stored in memory at ℓ is equal to v_2 , and, if so, updates ℓ to store v_3 instead; otherwise, it does nothing. The expression **FAA** $(\hat{e}_1, \hat{e}_2)$ *atomically* increments the value stored in the location described by $\hat{e}_1$ by the result of $\hat{e}_2$. The expression **fork** $\{\hat{e}\}$ forks a new thread to execute $\hat{e}$ and then immediately returns $()$ to the current thread.

Syntactic sugar. Non-recursive functions $\lambda x. e$ are defined as **rec** $_ (x) = e$, let-bindings **let** $x = e_1$ **in** e_2 are defined as $(\lambda x. e_2) e_1$, and sequential composition $e_1; e_2$ is defined as **let** $_ = e_1$ **in** e_2 . Here, we use the underscore $_$ to denote an anonymous binder, which is not used in the body of the binding expression.

2.2 Typing

We write $\Gamma \vdash \hat{e} : A$ for the syntactic typing judgment, which expresses that the expression $\hat{e}$ has the type A under the typing context Γ . The typing context Γ is a list of the form $x_1 : A_1, \dots, x_n : A_n$, which associates free variables (that may appear in $\hat{e}$) to their types. The empty typing context is denoted by $\emptyset$.

The syntactic typing rules of **MyLang** are displayed in Figure 2. We adopt Barendregt’s variable convention [Barendregt 1985], which means that in typing rules we assume that bound variables in expressions or types are “fresh”—*i.e.*, they do not conflict with any other variables in scope. Accordingly, our typing judgment $\Gamma \vdash \hat{e} : A$ does not keep track of the free type variables in Γ and A (*e.g.*, through an additional context $\Delta \subseteq \text{Tvar}$) and leaves conditions on freshness of variables

$\frac{\text{T-VAR} \quad x : A \in \Gamma}{\Gamma \vdash x : A}$	$\text{T-UNIT} \quad \Gamma \vdash () : 1$	$\frac{\text{T-BOOL} \quad b \in \{\text{true}, \text{false}\}}{\Gamma \vdash b : 2}$	$\frac{\text{T-INT} \quad n \in \mathbb{Z}}{\Gamma \vdash n : \mathbb{Z}}$
$\frac{\text{T-BINOP} \quad \Gamma \vdash \hat{e}_1 : A_1 \quad \Gamma \vdash \hat{e}_2 : A_2 \quad \odot : A_1 \times A_2 \Rightarrow B}{\Gamma \vdash \hat{e}_1 \odot \hat{e}_2 : B}$	$\frac{\text{T-REC} \quad \Gamma, x : A, f : A \rightarrow B \vdash \hat{e} : B}{\Gamma \vdash \text{rec } f(x) = \hat{e} : A \rightarrow B}$		
$\frac{\text{T-APP} \quad \Gamma \vdash \hat{e}_1 : A \rightarrow B \quad \Gamma \vdash \hat{e}_2 : A}{\Gamma \vdash \hat{e}_1 \hat{e}_2 : B}$	$\frac{\text{T-TLAM} \quad \Gamma \vdash \hat{e} : A}{\Gamma \vdash \Lambda \alpha. \hat{e} : \forall \alpha. A}$	$\frac{\text{T-TAPP} \quad \Gamma \vdash \hat{e} : \forall \alpha. A}{\Gamma \vdash \hat{e} \langle B \rangle : A[B/\alpha]}$	
$\frac{\text{T-IF} \quad \Gamma \vdash \hat{e} : 2 \quad \Gamma \vdash \hat{e}_1 : B \quad \Gamma \vdash \hat{e}_2 : B}{\Gamma \vdash \text{if } \hat{e} \text{ then } \hat{e}_1 \text{ else } \hat{e}_2 : B}$	$\frac{\text{T-PAIR} \quad \Gamma \vdash \hat{e}_1 : A_1 \quad \Gamma \vdash \hat{e}_2 : A_2}{\Gamma \vdash (\hat{e}_1, \hat{e}_2) : A_1 \times A_2}$		
$\frac{\text{T-PROJ} \quad \Gamma \vdash \hat{e} : A_1 \times A_2 \quad i \in \{1, 2\}}{\Gamma \vdash \pi_i \hat{e} : A_i}$	$\frac{\text{T-INJ} \quad \Gamma \vdash \hat{e} : A_i \quad i \in \{1, 2\}}{\Gamma \vdash \text{inj}_i \hat{e} : A_1 + A_2}$		
$\frac{\text{T-MATCH-SUM} \quad \Gamma \vdash \hat{e} : A_1 + A_2 \quad \forall i \in \{1, 2\}. \Gamma, x : A_i \vdash \hat{e}_i : B}{\Gamma \vdash \text{match } \hat{e} \text{ with } \text{inj}_1 x \Rightarrow \hat{e}_1 \mid \text{inj}_2 x \Rightarrow \hat{e}_2 \text{ end} : B}$	$\frac{\text{T-PACK} \quad \Gamma \vdash \hat{e} : A[B/\alpha]}{\Gamma \vdash \text{pack } \langle B, \hat{e} \rangle \text{ as } \exists \alpha. A : \exists \alpha. A}$		
$\frac{\text{T-MATCH-EX} \quad \Gamma \vdash \hat{e} : \exists \alpha. A \quad \Gamma, x : A \vdash \hat{e}_2 : B}{\Gamma \vdash \text{match } \hat{e} \text{ with pack } \langle \alpha, x \rangle \Rightarrow \hat{e}_2 \text{ end} : B}$	$\frac{\text{T-FOLD} \quad \Gamma \vdash \hat{e} : A[\mu \alpha. A/\alpha]}{\Gamma \vdash \text{fold } \hat{e} : \mu \alpha. A}$		
$\frac{\text{T-UNFOLD} \quad \Gamma \vdash \hat{e} : \mu \alpha. A}{\Gamma \vdash \text{unfold } \hat{e} : A[\mu \alpha. A/\alpha]}$	$\frac{\text{T-ALLOC} \quad \Gamma \vdash \hat{e} : A}{\Gamma \vdash \text{ref } \hat{e} : \text{ref } A}$	$\frac{\text{T-LOAD} \quad \Gamma \vdash \hat{e} : \text{ref } A}{\Gamma \vdash !\hat{e} : A}$	
$\frac{\text{T-STORE} \quad \Gamma \vdash \hat{e}_1 : \text{ref } A \quad \Gamma \vdash \hat{e}_2 : A}{\Gamma \vdash \hat{e}_1 \leftarrow \hat{e}_2 : 1}$	$\frac{\text{T-CAS} \quad \Gamma \vdash \hat{e}_1 : \text{ref } A \quad \Gamma \vdash \hat{e}_2 : A \quad \Gamma \vdash \hat{e}_3 : A \quad \text{EqType}(A)}{\Gamma \vdash \text{CAS}(\hat{e}_1, \hat{e}_2, \hat{e}_3) : 2}$		
$\frac{\text{T-FAA} \quad \Gamma \vdash \hat{e}_1 : \text{ref } \mathbb{Z} \quad \Gamma \vdash \hat{e}_2 : \mathbb{Z}}{\Gamma \vdash \text{FAA}(\hat{e}_1, \hat{e}_2) : \mathbb{Z}}$	$\frac{\text{T-FORK} \quad \Gamma \vdash \hat{e} : A}{\Gamma \vdash \text{fork } \{\hat{e}\} : 1}$		
$\text{EQTYP-UNIT} \quad \text{EqType}(1)$	$\text{EQTYP-BOOL} \quad \text{EqType}(2)$	$\text{EQTYP-INT} \quad \text{EqType}(\mathbb{Z})$	$\text{EQTYP-REF} \quad \text{EqType}(\text{ref } A)$

Fig. 2. Typing rules of **MyLang**.

implicit (e.g., in **T-TLAM**). Such conditions are also absent in our Coq mechanization because there we use de Bruijn indices [de Bruijn 1972] to handle variable binding.

The typing rule **T-CAS** for the **CAS** operation has the side-condition $\text{EqType}(A)$, which ensures that a **CAS** can only be performed on word-sized data types. The rules for the EqType predicate are also displayed in Figure 2. The typing rule **T-BINOP** for a binary operator $\odot$ has the side-condition $\odot : A_1 \times A_2 \Rightarrow B$, which expresses that the operator, when supplied with arguments of type A_1 and A_2 , produces a result of type B . The rules are $\odot : Z \times Z \Rightarrow Z$ for $\odot \in \{+, *, -\}$, and $(<) : Z \times Z \Rightarrow 2$, and $(=) : A \times A \Rightarrow 2$ for $\text{EqType}(A)$. (The $\text{EqType}(A)$ side-condition is not necessary for syntactic or semantic type soundness, but we include it in the typing rules since it is necessary for relational reasoning (Remark 8.3). For a semantics that is closer to a machine implementation, one would need to adjust the reduction rules of **CAS**; see Jung et al. [2018a] for an example of how this can be done for a language with an Iris-based semantic model of type soundness.)

We also define a typing judgment $\Gamma \vdash e : A$ on dynamic expressions analogously to the typing judgment for static expressions. We omit the definition here for brevity; it can be derived from the typing judgment for static expressions by simply removing all the red text from Figure 2 and replacing all the $\hat{e}$'s with e 's. It is then straightforward to show that $\Gamma \vdash e : A$ if and only if there exists a static expression $\hat{e}$ such that $|\hat{e}| = e$ and $\Gamma \vdash \hat{e} : A$.

2.3 Operational Semantics

To define the operational semantics of **MyLang**, we first define *values*, *states*, and *evaluation contexts* as shown in Figure 3. These definitions are mostly standard. The states $\sigma \in \text{State}$ of **MyLang** are heaps, which we model as partial functions with finite support from memory locations to values. Evaluation contexts $K \in \text{Ctx}$ are used to define a left-to-right call-by-value (CBV) evaluation strategy for **MyLang**.

With these notions in hand, we define the small-step operational semantics of **MyLang** in three stages:

- (1) We first define a *base reduction* relation, $(\sigma, e) \rightarrow_b (\sigma', e')$, which describes how e reduces under initial state σ to a new e' and (possibly) updated state σ' . This definition of the base reduction relation makes use of the auxiliary *pure reduction* relation $e \rightarrow_{\text{pure}} e'$ to handle state-independent reductions. The rules are shown in Figure 3. The function $\llbracket \odot \rrbracket : \text{Val} \times \text{Val} \rightarrow \text{Val}$ assigns a denotation to each binary operator $\odot$. This function is partial to account for operators that are applied wrongly, e.g., $10 \llbracket + \rrbracket \text{true}$ is undefined. We write $\uplus$ for the disjoint union operation on heaps.
- (2) Following Felleisen and Hieb [1992], we use evaluation contexts to lift the base reduction relation to a *thread-local reduction* relation $(\sigma, e) \rightarrow_t (\sigma', e')$.
- (3) Finally, the *thread-pool reduction* relation $(\sigma, \vec{e}) \rightarrow_{\text{tp}} (\sigma', \vec{e}')$ for our programs is a relation defined on machine states, i.e., pairs of a state and a thread pool (represented as a sequence of expressions $\vec{e}$ executing in different threads). The thread-pool reduction relation expresses that a machine state reduces by picking an arbitrary thread and either making a thread-local reduction step in that expression or else executing a **fork** and spawning a new thread.

2.4 Type Soundness

A programming language is *type-sound* (or *type-safe*) if every closed well-typed expression is safe (i.e., has well-defined behavior). To define this formally, we first give some auxiliary definitions:

- We say that a machine state $(\sigma, \vec{e})$ is *progressive*, written $\text{progressive}(\sigma, \vec{e})$, if any thread in that state is either a value (i.e., it has finished executing), or it is *reducible* (i.e., it can make at

Values, states, and evaluation contexts:

$$v \in \text{Val} ::= \text{rec } f(x) = e \mid \Lambda. e \mid () \mid n \mid \text{true} \mid \text{false} \mid \\ (v, v) \mid \text{inj}_1 v \mid \text{inj}_2 v \mid \text{pack}\langle v \rangle \mid \text{fold } v \mid \ell \in \text{Loc}$$

$$\sigma \in \text{State} \triangleq \text{Loc} \rightarrow_{\text{fin}} \text{Val}$$

$$K \in \text{Ctx} ::= [] \mid K e \mid v K \mid K \langle \rangle \mid K \odot e \mid v \odot K \mid \text{if } K \text{ then } e \text{ else } e \\ (K, e) \mid (v, K) \mid \pi_1 K \mid \pi_2 K \mid \\ \text{inj}_1 K \mid \text{inj}_2 K \mid (\text{match } K \text{ with } \text{inj}_1 x \Rightarrow e_1 \mid \text{inj}_2 x \Rightarrow e_2 \text{ end}) \mid \\ \text{pack}\langle K \rangle \mid \text{match } K \text{ with } \text{pack}\langle x \rangle \Rightarrow e \text{ end} \mid \\ \text{fold } K \mid \text{unfold } K \mid \\ \text{ref } K \mid !K \mid K \leftarrow e \mid v \leftarrow K \mid \\ \text{CAS}(K, e, e) \mid \text{CAS}(v, K, e) \mid \text{CAS}(v, v, K) \mid \\ \text{FAA}(K, e) \mid \text{FAA}(v, K)$$
Pure reduction:

$$(\text{rec } f(x) = e) v \rightarrow_{\text{pure}} e[v/x][\text{rec } f(x) = e/f] \\ (\Lambda. e) \langle \rangle \rightarrow_{\text{pure}} e \\ v_1 \odot v_2 \rightarrow_{\text{pure}} v_1 [\odot] v_2 \\ \text{if true then } e_1 \text{ else } e_2 \rightarrow_{\text{pure}} e_1 \\ \text{if false then } e_1 \text{ else } e_2 \rightarrow_{\text{pure}} e_2 \\ \pi_i(v_1, v_2) \rightarrow_{\text{pure}} v_i \quad (\text{if } i \in \{1, 2\}) \\ \text{match inj}_i v \text{ with inj}_1 x \Rightarrow e_1 \mid \text{inj}_2 x \Rightarrow e_2 \text{ end} \rightarrow_{\text{pure}} e_i[v/x] \quad (\text{if } i \in \{1, 2\}) \\ \text{match pack}\langle v \rangle \text{ with pack}\langle x \rangle \Rightarrow e \text{ end} \rightarrow_{\text{pure}} e[v/x] \\ \text{unfold}(\text{fold } v) \rightarrow_{\text{pure}} v$$
Base reduction:

$$(\sigma, e_1) \rightarrow_b (\sigma, e_2) \quad (\text{if } e_1 \rightarrow_{\text{pure}} e_2) \\ (\sigma, \text{ref } v) \rightarrow_b (\sigma \uplus \{(\ell, v)\}, \ell) \quad (\text{if } \ell \notin \text{dom}(\sigma)) \\ (\sigma, !\ell) \rightarrow_b (\sigma, v) \quad (\text{if } (\ell, v) \in \sigma) \\ (\sigma \uplus \{(\ell, v)\}, \ell \leftarrow w) \rightarrow_b (\sigma \uplus \{(\ell, w)\}, ()) \\ (\sigma \uplus \{(\ell, v)\}, \text{CAS}(\ell, v, u)) \rightarrow_b (\sigma \uplus \{(\ell, u)\}, \text{true}) \\ (\sigma \uplus \{(\ell, v)\}, \text{CAS}(\ell, w, u)) \rightarrow_b (\sigma \uplus \{(\ell, v)\}, \text{false}) \quad (\text{if } v \neq w) \\ (\sigma \uplus \{(\ell, n)\}, \text{FAA}(\ell, m)) \rightarrow_b (\sigma \uplus \{(\ell, n+m)\}, n)$$
Thread-local and thread-pool reduction:

$$\frac{(\sigma, e) \rightarrow_b (\sigma', e')}{(\sigma, K[e]) \rightarrow_t (\sigma', K[e'])} \quad \frac{(\sigma, e) \rightarrow_t (\sigma', e')}{(\sigma, (\vec{e}_1; e; \vec{e}_2)) \rightarrow_{\text{tp}} (\sigma', (\vec{e}_1; e'; \vec{e}_2))} \\ (\sigma, (\vec{e}_1; K[\text{fork } \{e\}]; \vec{e}_2)) \rightarrow_{\text{tp}} (\sigma, (\vec{e}_1; K[()]; \vec{e}_2; e))$$
Fig. 3. Operational semantics of **MyLang**.

least one further step of computation):

$$\begin{aligned} \text{progressive}(\sigma, (e_1; \dots; e_n)) &\triangleq \forall i \in \{1, \dots, n\}. (e_i \in \text{Val} \vee \text{red}(\sigma, e_i)) \\ \text{red}(\sigma, e) &\triangleq (\exists \sigma', e'. (\sigma, e) \rightarrow_t (\sigma', e')) \vee (\exists K, e'. e = K[\text{fork } \{e'\}]) \end{aligned}$$

- We then say that a closed expression e , representing a complete program, is *safe*, written $\text{safe}(e)$, if any machine state reachable by evaluating e for any number of steps is progressive:

$$\text{safe}(e) \triangleq \forall \sigma_2, \vec{e}_2. (\emptyset, e) \rightarrow_{\text{tp}}^* (\sigma_2, \vec{e}_2) \Rightarrow \text{progressive}(\sigma_2, \vec{e}_2)$$

Now we can formally define type soundness as:

$$(\emptyset \vdash e : A) \text{ implies } \text{safe}(e)$$

2.5 Syntactic Type Soundness via Progress and Preservation

The syntactic approach to proving type soundness involves two key theorems:

- (1) **Progress (Theorem 2.1)**. Well-typed machine states are progressive.
- (2) **Preservation (Theorem 2.2)**. Reduction preserves well-typedness of machine states.

These theorems rely on a notion of a *well-typed machine state* $(\sigma, \vec{e})$, which intuitively expresses that each value in the heap σ is well-typed and each expression in the thread-pool $\vec{e}$ is well-typed. To formalize this notion, we need to account for location literals ℓ . While location literals do not appear in static expressions, they may appear in runtime expressions and values during reduction (when memory cells are allocated), and their types need to match up with the types of values in the heap σ . We thus define a generalized typing judgment, written $\Sigma; \Gamma \vdash e : A$, which extends the typing judgment $\Gamma \vdash e : A$ with a *heap typing* $\Sigma : \text{Loc} \rightarrow_{\text{fin}} \text{Type}$. A heap typing is a partial function with finite support that assigns a closed type to each location. The essential rule of the generalized typing judgment is the one for location literals:

$$\frac{\text{DT-LOC} \quad \Sigma(\ell) = A}{\Sigma; \Gamma \vdash \ell : \text{ref } A}$$

In all rules but **DT-LOC** above, the heap typing is simply threaded through. For example, the rules for function application and allocation become:

$$\begin{array}{c} \text{DT-APP} \\ \hline \Sigma; \Gamma \vdash e_1 : A \rightarrow B \quad \Sigma; \Gamma \vdash e_2 : A \\ \hline \Sigma; \Gamma \vdash e_1 e_2 : B \end{array} \qquad \begin{array}{c} \text{DT-ALLOC} \\ \hline \Sigma; \Gamma \vdash \hat{e} : A \\ \hline \Sigma; \Gamma \vdash \text{ref } \hat{e} : \text{ref } A \end{array}$$

We can now define the notion of well-typed machine states with the judgment $\Sigma \vdash_{\text{MS}} (\sigma, \vec{e}) : \vec{A}$:

$$\frac{\begin{array}{l} \text{dom}(\Sigma) = \text{dom}(\sigma) \quad \text{length}(\vec{e}) = \text{length}(\vec{A}) \\ \forall \ell \in \text{dom}(\Sigma). \Sigma; \emptyset \vdash \sigma(\ell) : \Sigma(\ell) \quad \forall i < \text{length}(\vec{e}). \Sigma; \emptyset \vdash e_i : A_i \end{array}}{\Sigma \vdash_{\text{MS}} (\sigma, \vec{e}) : \vec{A}}$$

This judgment says that each value $\sigma(\ell)$ in σ has the corresponding type $\Sigma(\ell)$ from the heap typing Σ , and that, under Σ , each expression e_i in $\vec{e}$ has the corresponding type A_i from the list $\vec{A}$.

THEOREM 2.1 (PROGRESS). *Every machine state $(\sigma, \vec{e})$ that is typed for some heap environment Σ is safe. Formally, if $\Sigma \vdash_{\text{MS}} (\sigma, \vec{e}) : \vec{A}$, then $\text{progressive}(\sigma, \vec{e})$.*

THEOREM 2.2 (PRESERVATION). *Typing of machine states is preserved by the reduction relation $\rightarrow_{\text{tp}}$. Formally, if $\Sigma \vdash_{\text{MS}} (\sigma, \vec{e}) : \vec{A}$ and $(\sigma, \vec{e}) \rightarrow_{\text{tp}} (\sigma', \vec{e}')$, then there exists an extended heap typing $\Sigma' \supseteq \Sigma$ and an extended list of thread types $\vec{A}' \supseteq_{\text{prefix}} \vec{A}$ such that $\Sigma' \vdash_{\text{MS}} (\sigma', \vec{e}') : \vec{A}'$.*

Progress and preservation are typically proven by induction and straightforward case analysis on the given typing and reduction derivations. The proofs involve some easy helper lemmas related to substitution and weakening w.r.t. extension of the heap typing, as well as (in some variations) lemmas for decomposing a well-typed term into a well-typed evaluation context and a well-typed redex. The interested reader can find a detailed account of the proofs in the course lecture notes of [Dreyer et al. \[2022\]](#). For present purposes, what is important is that progress and preservation give rise to the following result:

COROLLARY 2.3 (SYNTACTIC TYPE SOUNDNESS). *Every closed well-typed expression e is safe. Formally, if $(\emptyset \vdash e : A)$, then $\text{safe}(e)$.*

PROOF. Assume that we have $(\emptyset \vdash e : A)$ and that we are given a reduction $(\emptyset, e) \rightarrow_{\text{tp}}^* (\sigma_2, \vec{e}_2)$. Our goal is to prove $\text{progressive}(\sigma_2, \vec{e}_2)$.

By definition of the judgment for well-typed machine states, we obtain $\emptyset \vdash_{\text{MS}} (\emptyset, e) : A$ from the assumption $(\emptyset \vdash e : A)$. By repeatedly using preservation ([Theorem 2.2](#)) for each reduction step in $(\emptyset, e) \rightarrow_{\text{tp}}^* (\sigma_2, \vec{e}_2)$, we obtain $\Sigma \vdash_{\text{MS}} (\sigma_2, \vec{e}_2) : \vec{A}'$ for some Σ and $\vec{A}'$. By progress ([Theorem 2.1](#)), we obtain $\text{progressive}(\sigma_2, \vec{e}_2)$, which concludes the proof. $\square$

3 LIMITATIONS OF SYNTACTIC TYPE SOUNDNESS

Now that we have reviewed a typical formulation of syntactic type soundness, we are in a position to expound our criticisms of it:

- (1) It says nothing about whether a programming language enforces *data abstraction* (§3.1).
- (2) It says nothing about programs that use *unsafe features* in a *safely encapsulated* way (§3.2).

3.1 Data Abstraction

Consider the following type `symbol_type`, which describes an (extremely simplified) interface of an abstract data type (ADT) of *symbols*:

$$\text{symbol_type} \triangleq \exists \alpha. (1 \rightarrow \alpha) \times (\alpha \rightarrow 2)$$

Following [Mitchell and Plotkin \[1988\]](#), we model the type of an ADT using an *existential type*.⁵ Here, the existential type describes the interface of an ADT that exports an abstract type α representing “symbols”, along with two operations: a `gensym` function of type $1 \rightarrow \alpha$ with which one can generate fresh symbols, and a `check` function of type $\alpha \rightarrow 2$ which one can use to check if a symbol is valid. This interface is obviously not very useful in the stripped-down form presented here, but it gives the flavor of the interface one often sees for symbol tables in compilers and will suffice to get across our point about syntactic type soundness.

Now, consider the implementation `symbol` of the `symbol_type` interface:

```
symbol  $\triangleq$  let  $c = \text{ref } 0$  in
  pack  $\left\langle \text{Z}, \left( \begin{array}{l} \lambda (). \text{FAA}(c, 1), \\ \lambda s. s < ! c \end{array} \right) \right\rangle$  as symbol_type
```

This implementation employs a private integer counter c , which is allocated when the expression defining `symbol` is evaluated. The counter c is used as a perpetual source of fresh symbols. When the `gensym` function (the first closure returned by `symbol`) is called, it uses `fetch-and-add (FAA)` to atomically increment the value of c and return the previous value. Thus, when called repeatedly, `gensym` will return 0, 1, 2, etc. The `check` function (the second closure returned by `symbol`) checks validity of its symbol argument by checking that it is less than the current value of the counter.

⁵See [Rossberg et al. \[2014\]](#) for an explanation of how existential types also provide a foundation for understanding and formally modeling much more complex data abstraction facilities, such as those of the ML module system [\[MacQueen 1984\]](#).

It is easy to argue by appeal to intuitive reasoning that the check function returned by `symbol` must always return `true`. To see this, we first observe that the only values of the abstract type `symbol` (i.e., α) that can ever be generated are those returned by the `gensym` function. However, whenever such a value is returned by `gensym`, it is at that instant one less than the current value of the counter c . Furthermore, the value of the counter only increases over time. Put together, these imply that, at all times, all values of type `symbol` are always less than the current value of the counter, so `check` applied to a value of type `symbol` should always return `true`.

This is the kind of informal reasoning about program correctness that programmers employ all the time. Crucially, though, it relies on the assumption that the programming language properly enforces two forms of data abstraction: *private state* (via *closures*) and *abstract types*. To enforce the invariant that the value of the counter c only increases over time, it is essential that the only way to modify c is by applying one of the closures returned by `symbol`—i.e., that c is maintained as private state of those closures. Otherwise, clients could update c willy-nilly, and thus break the invariant. To enforce the invariant that the only values of the `symbol` type are the ones produced by calls to the `gensym` function, it is essential that the representation of the `symbol` type as Z be held abstract from clients. Otherwise, clients could take an arbitrary integer and “forge” a value of type `symbol` from it, which could cause the `check` function to return `false`.

Fortunately, the language **MyLang** *does* properly enforce data abstraction, so the above informal reasoning is in fact valid. Unfortunately, proper treatment of data abstraction is not in any way a consequence of syntactic type soundness. To demonstrate this point, we extend **MyLang** with a new and rather devious feature that we call, for lack of a better name, `gremlin`. This new feature has the property that on the one hand it is “harmless” in that it preserves the syntactic type soundness of **MyLang**, but on the other hand it completely breaks the language’s support for data abstraction and hence the ability to reason modularly about **MyLang** code.

The static and dynamic semantics of `gremlin` are as follows:

$$\Gamma \vdash \text{gremlin} : 1$$

$$(\sigma, \text{gremlin}) \rightarrow_b (\sigma, ()) \quad (\sigma \uplus \{(\ell, n)\}, \text{gremlin}) \rightarrow_b (\sigma \uplus \{(\ell, 0)\}, ())$$

In short, `gremlin` is an expression of type 1, and its execution non-deterministically proceeds in one of two ways. Either it is simply a no-op, or else it non-deterministically selects some memory location ℓ currently storing an integer value n , and it updates ℓ to store 0.

As far as syntactic type soundness is concerned, `gremlin` is almost trivially a “safe” operator. First, the no-op evaluation rule ensures that `gremlin` can always make progress. Second, if `gremlin` does have an effect, it is merely to replace the integer value stored at some memory location with another integer value (namely, 0), thus preserving syntactic well-typedness of the heap. Consequently, it is easy to extend the syntactic soundness result from §2.5 to account for `gremlin`.

However, as should be intuitively clear, `gremlin` is a *terrible* feature because it destroys the programmer’s ability to place invariants on the private state of their ADTs. To make this point perfectly concrete, consider the following client of the `symbol` ADT:

```
evil_client  $\triangleq$  match symbol with pack  $\langle \alpha, x \rangle \Rightarrow$ 
  let gensym =  $\pi_1 x$  in
  let check =  $\pi_2 x$  in
  let s = gensym () in
  gremlin; check s
end
```

After unpacking the existential representing the ADT, the client first calls the `gensym` function to create a fresh symbol value s . It then invokes `gremlin`, and finally calls `check` on s .

When this client is executed, the call to `gensym` will have the effect of updating the ADT's private counter c to 1, and returning the value 0 for s . When `gremlin` is invoked, one possible behavior is that the counter c will be set back to 0. If that happens, the subsequent application of `check` to s (i.e., to 0) will return `false`! The problem here, of course, is that with `gremlin` in the language, the private state of the symbol ADT is no longer truly private since `gremlin` can modify it. As a result, the programmer cannot depend on any invariants on the private counter being maintained.

Now `gremlin` may seem like a rather contrived operator, but it is actually just an absurdly pointless variation of what is already supported, for example, by the Reflection API in Java. Using reflection, one can inspect the private fields and methods of an arbitrary object and freely modify its private state [Nasi 2011], thus achieving the same devastating effect on data abstraction as `gremlin` does.

Fortunately, since version 9, Java has given programmers the choice of whether the private state of their ADTs should be accessible via reflection from other ADTs. (Prior to Java 9, there was no way to limit reflective access.) Consequently, we would really like to be able to prove *some* theorem about data abstraction in Java 9 that would not hold of prior versions of the language. But seeing as reflection—like `gremlin`—is perfectly “type-safe” in the syntactic sense, syntactic type soundness is not that theorem.

In summary, syntactic type soundness of a programming language tells us essentially nothing about whether the language is sensible to program in.

3.2 Safe Encapsulation of Unsafe Features

As noted in the introduction, the price that “safe” programming languages pay for safety is that they do not always allow the programmer to write the code they want to write. Consequently, most such languages provide *unsafe escape hatches* by which programmers can circumvent the restrictions of their type systems. For example, OCaml provides `Obj.magic`, an unchecked type cast operator. Haskell provides `unsafeCoerce` (similar to `Obj.magic`), `unsafePerformIO` (for escaping the IO monad), and more. Rust provides a whole host of low-level C-style operations, although uses of them must be confined to blocks of code explicitly marked `unsafe`. These unsafe escape hatches are widely used and depended upon in real-world applications.

Of course, given that these unsafe features are blatantly dangerous, programmers are advised to only make use of them if “you know what you are doing”, and a major aspect of “knowing what you are doing” is knowing how to use such features in a *safely encapsulated* way. That is, when a programmer uses unsafe features in the implementation of some ADT (or module or object) M , they typically rely on the data abstraction mechanisms of the programming language to enforce invariants on the private state or data representation of M —invariants which imply that the local uses of unsafe features within M will never lead to any undefined behavior for M 's clients. In this way, one can see the data abstraction mechanisms of a language as their own saving grace: though the restrictions they place on programmers sometimes necessitate the use of unsafe workarounds, it is the data abstraction afforded by these mechanisms that make it possible to use those workarounds “locally”—i.e., without breaking the safety guarantees of the language as a whole.

Hence, the ability to safely encapsulate uses of unsafe features is inextricably linked with data abstraction: understanding whether a programming language supports one is tantamount to understanding whether it supports the other. To drive this point home, let us explain how we can recast the example of the symbol ADT from §3.1 in terms of safe encapsulation of unsafe features.

First, let us extend (the static and dynamic) syntax of **MyLang** with a new feature: *assertions*, written `assert e`. The dynamic semantics of assertion expressions is given as follows:

$$K ::= \dots \mid \text{assert } K \qquad \text{assert true} \rightarrow_{\text{pure}} \text{true}$$

In words, `assert e` will first evaluate e to a value v , and then return `true` iff $v = \text{true}$. If v evaluates to anything else, `assert e` will get stuck.⁶ Thus, in order for `assert e` to be a safe expression, e must never evaluate to any value other than `true`. However, seeing as there is no type in **MyLang** which would ensure that e satisfies this property, `assert` is an example of an *unsafe* feature.

Now consider the following, slightly revised implementation of the `symbol` ADT from §3.1 (the changed code is underlined):

```
symbol  $\triangleq$  let  $c = \text{ref } 0$  in
  pack  $\left\langle \text{Z}, \left( \begin{array}{l} \lambda (). \text{FAA}(c, 1), \\ \lambda s. \text{assert } (s < ! c) \end{array} \right) \right\rangle$  as symbol_type
```

Rather than simply returning the result of $s < ! c$, the check function now *asserts* it. Consequently, check will now only be safe to execute (i.e., not get stuck) if the expression $s < ! c$ indeed evaluates to `true`. As we have already mentioned in §3.1, **MyLang**'s support for data abstraction *should* ensure that $s < ! c$ *does* always evaluate to `true` in all well-typed contexts—or equivalently, that the new `symbol` ADT has safely encapsulated the potentially unsafe behavior of the `assert` expression in its body, so that no well-typed client of `symbol` will ever encounter undefined (stuck) behavior. But syntactic type soundness—by restricting attention to syntactically well-typed programs—does not offer us a means to prove this! In the next section, we will see a more powerful approach to formalizing type soundness that does.

4 SEMANTIC TYPE SOUNDNESS

In this section, we explain how to overcome the limitations of syntactic type soundness described in §3 via the alternative approach of *semantic type soundness*. We begin with a brief high-level explanation of how a semantic type soundness proof is structured (§4.1). We then discuss prior approaches to semantic type soundness and the problems they suffer, which might explain why such approaches have not been more widely adopted (§4.2 and §4.3). We conclude by explaining how Iris addresses these problems, thus providing an ideal logical framework in which semantic type soundness can be more effectively formalized (§4.4). The sections that follow (§5–§7) will then present our *logical approach* to proving semantic type soundness in great detail.

4.1 High-Level Overview of Semantic Type Soundness

The central problem with syntactic type soundness is that it identifies “safe” with “syntactically well-typed”. As a result, it is unable to account for the safety of code that uses potentially unsafe features in a well-encapsulated way, such as the `symbol` example at the end of §3.2.

To overcome this problem, semantic type soundness models safety instead using a more liberal view of well-typedness, which we call *semantic typing* and write as $\Gamma \models e : A$. The difference is that, whereas syntactic typing is *intensional* (it dictates a fixed set of syntactic rules by which safe terms can be constructed), semantic typing is *extensional* (it merely requires that terms behave in a safe way when executed). For example, the `symbol` ADT from §3.2, though not syntactically well-typed due to its use of the unsafe `assert` expression, will be shown to be semantically well-typed at the type `symbol_type`, thus establishing that `symbol` is in fact safe to use at that type. Of course, the price paid for this extensionality is that semantic typing is in general not a property that can

⁶One can in fact encode a primitive extremely similar to `assert e` in **MyLang** without any extension, via the following syntactic sugar: `assert e  $\triangleq$  if e then true else 42(42)`.

be checked algorithmically. Rather, proving that a term is semantically well-typed may require arbitrarily interesting verification effort. But this is to be expected, given that the goal of semantic soundness is to help establish that ADTs are properly maintaining their internal invariants, a task which often amounts to proving full functional correctness of the code.

The high-level structure of a semantic type soundness proof is simple:

- **Adequacy.** First, we prove an *adequacy* theorem, which establishes that closed, semantically well-typed terms are indeed safe to execute. Formally, this means that $\emptyset \models e : A$ implies $\text{safe}(e)$. This theorem is usually almost trivial to prove because, as explained above, it is typically more or less baked into the extensional definition of semantic typing.
- **Semantic typing rules.** Second, and more interestingly, we prove semantic versions of all the syntactic typing rules of the language, where the semantic version of a typing rule simply replaces all the syntactic $\vdash$'s in it with semantic $\models$'s. For instance, concerning function applications in **MyLang**, we prove the following *semantic typing rule* as a lemma stating that the premises imply the conclusion:

$$\frac{\Gamma \models e_1 : A \rightarrow B \quad \Gamma \models e_2 : A}{\Gamma \models e_1 e_2 : B}$$

These semantic typing rules serve to demonstrate that semantic typing is compositional in the same way that syntactic typing is. Semantic typing rules are sometimes also referred to as “compatibility lemmas” [Pitts 2005].

One immediate consequence of the semantic typing rules is that syntactic typing implies semantic typing, *i.e.*, $\Gamma \vdash e : A$ implies $\Gamma \models e : A$. Historically, this property is often referred to as the *fundamental theorem* or *fundamental property*. (It is provable by a straightforward induction on syntactic typing derivations.) As a result, closed syntactically well-typed programs are also semantically well-typed and thus, by adequacy, safe to execute. In other words, *once we have proven semantic type soundness, syntactic type soundness falls out as a simple corollary*.⁷

But the main reason we care about semantic type soundness is that it is a more useful result than syntactic type soundness: it shows that we can safely compose syntactically *well*-typed pieces of a program with other pieces that may be syntactically *ill*-typed (*e.g.*, use unsafe features) so long as those other pieces are *semantically well-typed*. For instance, once we prove that the `symbol` ADT is semantically well-typed at the type `symbol_type`, we will be able to deduce that if `symbol` is used within any syntactically (or semantically) well-typed program context C , then the resulting whole program $C[\text{symbol}]$ will be safe to execute—implying that the assertion inside `symbol`'s check function will always succeed.

4.2 Prior Work on Semantic Type Soundness

As noted in the introduction, there is a great deal of prior work on semantic type soundness, dating back to the original paper of Milner [1978], in which the idea of type soundness was introduced. However, the approach has never really “taken off” as a method for proving type soundness of more realistic languages in the same way that syntactic type soundness has. We now explain why

⁷Technically speaking, although semantic type soundness implies syntactic type soundness (Corollary 2.3)—*i.e.*, that syntactically well-typed programs are safe to execute—it does not imply Preservation (Theorem 2.2)—*i.e.*, that syntactically well-typed programs remain syntactically well-typed throughout execution. However, we would argue that, for realistic languages like **MyLang**, the Preservation property is not independently that useful, as it relies on an essentially *ad hoc* notion of typing on machine states.

we believe this is so, as it helps to provide a clearer motivation for the “logical” formulation of type soundness that is our main contribution.⁸

Milner’s original semantic soundness proof for a core ML-like calculus, as well as subsequent semantic soundness proofs for more expressive type systems with higher-order polymorphism, recursive types, and subtyping—e.g., [MacQueen et al. 1986; Bruce and Mitchell 1992]—were formulated using “realizability” models, in which types were interpreted as certain kinds of subsets or partial equivalence relations over a domain-theoretic (*i.e.*, denotational) model of untyped computation. Such models were also used to study parametricity and data abstraction, both for functional languages with polymorphic types—e.g., [Bainbridge et al. 1990; Abadi and Plotkin 1990]—and for imperative languages with local variables—e.g., [O’Hearn and Tennent 1992]. However, developing denotational semantics for programming languages with higher-order state (*i.e.*, general mutable references to values of arbitrary type) turned out to be quite challenging. Indeed, despite the ubiquity of higher-order state in programming languages for the past several decades, it was only in the work of Birkedal et al. [2010b] that the realizability approach over domains was finally extended to handle this feature.

In much of this work, it was understood implicitly that, due to the inherent compositionality of denotational models, they could serve to establish the safety of combining syntactically well-typed programs with syntactically unsafe, but semantically well-typed, programs (as we described in §4.1). But this capability was not commonly exploited or even remarked upon, perhaps because the focus was on building semantic models of higher-order, richly-typed λ -calculi, not on modeling realistic languages with low-level unsafe primitives (such as the language studied in the RustBelt project [Jung et al. 2018a, 2021; Jung 2020; Dang et al. 2020]).

Unfortunately, this state of affairs has meant that, for realistic languages, Milner-style semantic soundness based on denotational semantics has not offered a viable solution to the problem we posed in the introduction. And for the more modest goal of simply proving well-defined behavior for syntactically well-typed programs, progress and preservation has offered a more elementary and broadly applicable technique.

In the 1980s and 1990s, there arose a related but distinct line of work on building semantic models of typed languages over an *operational* semantics rather than a denotational one. In particular, partial equivalence relations over operational semantics were used early on in seminal work on the NuPRL type theory [Constable et al. 1986; Allen 1987]. This approach was further developed to account for recursive types [Birkedal and Harper 1999], local state [Pitts and Stark 1998], and the combination of recursive types and polymorphism [Crary and Harper 2007]. The approach to recursive types in [Birkedal and Harper 1999; Crary and Harper 2007] employed a syntactic adaptation of the denotational idea of *minimal invariance* [Pitts 1996], but this was quite technically involved and, as with denotational methods, it was for a long time not clear how to generalize the approach to handle higher-order state.

An important breakthrough came in 2001 when Appel and McAllester [2001] developed their *step-indexed* model of recursive types. The basic idea of step-indexing is to stratify the quasi-circular definition of semantic typing for recursive types by the number of steps for which the term in question is allowed to execute.⁹ One immediate benefit of step-indexing was that it supplied a much more elementary model of recursive types than the previous approaches based on minimal invariance. But the more important benefit was that the basic idea scaled to account for more complex features that were beyond the scope of denotational models. In particular, Ahmed in her

⁸The literature on semantics of type systems is voluminous, so this section should not be viewed as a comprehensive survey, but rather a brief dive into the literature in order to suggest where existing approaches to semantic type soundness come up short. See §9 for additional discussion.

⁹See Ahmed [2004] for a more detailed exposition of step-indexing.

PhD thesis [Ahmed 2004] (building on prior work with Appel and Virga [Ahmed et al. 2002]) showed how to apply step-indexing in a more sophisticated fashion in order to construct a semantic model of higher-order state.

Ahmed’s thesis led in turn to a flood of follow-on work, including [Appel et al. 2007; Ahmed et al. 2009; Neis et al. 2009; Benton and Hur 2009; Dreyer et al. 2010; Birkedal et al. 2011; Krishnaswami and Benton 2011; Schwinghammer et al. 2013; Dreyer et al. 2012; Thamsborg and Birkedal 2011; Birkedal et al. 2012; Turon et al. 2013b; Birkedal et al. 2013]. This line of work has demonstrated that step-indexing—in conjunction with various other techniques, notably *biorthogonality* [Krivine 1994; Pitts and Stark 1998] and *Kripke logical relations* [Jung and Tiuryn 1993]—could be used to construct operational-semantics-based models of much more realistic languages, featuring (among other things) control effects, substructural types, intensional polymorphism, and concurrency. What is more, some of these models (e.g., [Ahmed et al. 2009; Dreyer et al. 2010, 2012; Schwinghammer et al. 2013])¹⁰ were developed for the express purpose of verifying the kind of invariants on the private state of ADTs that we saw in the symbol example from §3.1. (In fact, that example is adapted from one proven by Ahmed et al. [2009].) Other models used step-indexing and semantic soundness to reason about low-level code, e.g., to capture what it means for a piece of low-level code to implement a high-level function and to prove correctness of a simple compiler [Benton and Hur 2009; Hur and Dreyer 2011]. These more sophisticated semantic models—constructed using step-indexing and companion techniques—are often referred to as *step-indexed Kripke logical-relations* (or *SKLR*) models.

At this point, the reader may rightly wonder: if SKLR models already address the limitations of syntactic type soundness from §3, why are we writing this article? And if they are so powerful, why have they not gained widespread adoption? Why does syntactic type soundness continue to be much more commonly known and used?

4.3 The Problems with SKLR Models

Based on our personal experience with SKLR models, we believe the reason they have not been more widely adopted is that, if one works *directly* with these models, one’s proofs become painfully tedious, low-level, and difficult to maintain. Specifically:

- (1) **Explicit step-index arithmetic:** When working directly in a step-indexed model of any kind, one ends up performing a great deal of tedious “step-index arithmetic”—i.e., counting of how many computation steps different operations take—even though it seems for the most part completely irrelevant to what one is proving.¹¹
- (2) **Explicit reasoning about global state:** When working directly in an SKLR model for a stateful language, one ends up reasoning explicitly about the global state of memory, even though the operations one is reasoning about only affect local pieces of that memory (e.g., a single location).¹²
- (3) **Explicit reasoning about possible worlds:** When working directly in one of the more advanced SKLR models, one ends up performing a lot of tedious manipulation of and quantification over “possible worlds”, which describe the set of invariants that have been established on the program state.¹³

¹⁰The cited works formalized invariants on private state *relationally*—i.e., by proving certain kinds of *contextual refinements*—rather than in the setting of semantic soundness. We find it useful to start first in this section by formalizing such invariants in the simpler “unary” setting of semantic soundness, before showing in §8 how our approach generalizes to the more complex “binary” relational setting of prior work.

¹¹See for instance the proofs in [Ahmed 2004] or [Ahmed et al. 2009].

¹²See for instance the proofs in [Ahmed et al. 2009] or [Schwinghammer et al. 2013].

¹³See for instance the proofs in [Ahmed et al. 2009], [Schwinghammer et al. 2013], or [Turon et al. 2013b].

For a representative example of all these points, we refer the reader to Amal Ahmed’s PhD thesis [Ahmed 2004] and the technical appendices accompanying several of her papers, e.g., [Ahmed 2006]. Her formal developments are unusual (and commendable) in that they spell out step-indexing-based proofs in full, and with great attention to detail.¹⁴ The end result, however, is that her proofs are cluttered with seemingly unnecessary low-level technical details about step-indexing, the global state, and possible worlds.¹⁵ For example, see Ahmed’s proof of the rule mentioned earlier in this section—the semantic typing rule for function applications—which appears as Theorem 3.21 in her thesis. The proof involves explicit step-index arithmetic throughout (e.g., “let $k^* = k - j - i - 1$ ”), as well as manipulation of three global states and four possible worlds—and that is all for a semantic typing rule that has nothing to do with mutable state! As a result, compared with the cases of a progress-and-preservation proof concerning function applications, the semantic soundness proof appears significantly more low-level and complex, and for no clear reason.

This problem was noted fairly early on in the development of SKLR models [Appel et al. 2007], and led to a fruitful line of work on *program logics* for encoding SKLR models at a much higher level of abstraction [Dreyer et al. 2011, 2010; Turon et al. 2013a]. In combination with a line of work on higher-order concurrent separation logic [Svendsen et al. 2013; Svendsen and Birkedal 2014], this line of work culminated in the development of Iris [Jung et al. 2015, 2016; Krebbers et al. 2017a; Jung et al. 2018b]: a unifying, language-generic framework for *higher-order concurrent separation logic*, implemented in the Coq proof assistant [Krebbers et al. 2017b, 2018], into which a variety of SKLR models can be (and have already been) encoded. We give an overview of SKLR models that have been encoded in Iris in §10.

4.4 How Iris Solves the Problems of SKLR Models

Using Iris, the pain points of working with SKLR models—the global, “tedious” reasoning—can largely be made to disappear—replaced by local, “interesting” reasoning. In particular, concerning the complications of working directly with SKLR models that we mentioned in §4.3, Iris addresses them head-on:

- (1) **Eliminating tedious step-indexed reasoning:** In Iris, the tedious details of step-index arithmetic are (to a large extent) hidden within the soundness proof of Iris itself, so that proofs done on top of Iris are not cluttered with them. To the limited extent that step-indexed reasoning is necessary (to avoid circular paradoxes), it is handled abstractly—following prior work by Appel et al. [2007]—using the so-called “later” ($\triangleright$) modality [Nakano 2000], which enables one to assert that a proposition should hold “one step of computation later”.
- (2) **Local reasoning about state:** In contrast to direct proofs with SKLR models, which involve manipulation of global state, Iris builds on separation logic [O’Hearn et al. 2001; Reynolds 2002] to support *local reasoning* about state. Local reasoning makes proofs about stateful code much more pleasant: when proving semantic soundness of an expression e , we need only reason about what happens to the piece of state that e itself manipulates. Moreover, separation logic is a good fit for formalizing semantic soundness because semantic soundness is a compositional property and separation-logic proofs are compositional by construction.
- (3) **High-level reasoning about stateful invariants:** Iris extends vanilla separation logic with two logical mechanisms—*impredicative invariants* (a higher-order generalization, first developed in [Svendsen and Birkedal 2014], of the shared resource invariants from O’Hearn’s

¹⁴The technical appendix accompanying [Dreyer et al. 2012] is similarly detail-oriented in this respect.

¹⁵In contrast, some subsequent papers employing step-indexed proofs, such as [Krishnaswami et al. 2012; Turon et al. 2013b], may seem marginally less cluttered, but that is only because they systematically elide “boring details” related to step-indexing that are quite easy to get wrong.

$$\begin{aligned}
\tau &::= 0 \mid 1 \mid \mathbb{B} \mid \mathbb{N} \mid \mathbb{Z} \mid \text{Val} \mid \text{Expr} \mid \text{iProp} \mid \tau \times \tau \mid \tau + \tau \mid \tau \rightarrow \tau \mid \dots & (\text{Types}) \\
t, u, P, Q &::= x \mid \lambda x : \tau. t \mid t(u) \mid \text{True} \mid \text{False} \mid P \wedge Q \mid P \vee Q \mid P \Rightarrow Q \mid & (\text{Propositional logic}) \\
&\quad \forall x : \tau. P \mid \exists x : \tau. P \mid t = u \mid & (\text{Higher-order logic}) \\
&\quad P * Q \mid P \multimap Q \mid \ell \mapsto v \mid \text{wp}_{\mathcal{E}} e \{ \Phi \} \mid & (\text{Separation logic}) \\
&\quad \Box P \mid \triangleright P \mid \mu x : \tau. t \mid \models_{\mathcal{E}} P \mid \boxed{P}^N \mid \dots & (\text{Iris-specific connectives})
\end{aligned}$$

Fig. 4. Syntax of Iris. (We use P, Q to represent terms of type iProp , and Φ, Ψ to represent (persistent) Iris predicates (i.e., functions to iProp or $\text{iProp}_{\Box}$), and t, u to represent arbitrary terms.)

original concurrent separation logic [O’Hearn 2007; Brookes 2007]) and *user-defined ghost state* (the ability to define custom, domain-specific notions of logical resource). Used in tandem, these two mechanisms enable one to express complex invariants on the private state of modules, ADTs, *etc.*, and to reason about those invariants at a much higher level of abstraction than is afforded by the possible worlds of SKLRs.

In short, Iris provides an ideal framework for formalizing *logical type soundness*—i.e., semantic type soundness proofs for richly-typed programming languages encoded in higher-order separation logic. We now proceed to concretely demonstrate the abovementioned benefits of logical type soundness in the context of our example language **MyLang**.

5 DEFINING A LOGICAL RELATION IN IRIS

Figure 4 shows the syntax of Iris, a higher-order logic extended with connectives from separation logic as well as a few other custom modalities that we will present in due course. In this and the next section, we show, step-by-step, how indeed Iris provides a natural logical language in which semantic type soundness for realistic languages can be formalized.

Before we begin, we note that there are multiple ways to formalize higher-order logic. Iris is formalized as a two-sorted system, with a sort of types and sort of terms: there are typing rules formalizing when a term is well-typed (omitted here), and equality rules formalizing when two terms are equal, including standard β - and η -rules for product, coproduct, and function types (also omitted here). Types include those of the simply-typed lambda calculus, and also the special type iProp of Iris propositions. To a first approximation, one can think of Iris propositions as being like the assertions of standard separation logic: predicates over some underlying types of resources (e.g., pieces of the heap), which implicitly describe ownership of those resources.

It is also important to note that Iris is a language-generic framework, meaning that it can be instantiated and used to reason about any language defined by a relatively common form of operational semantics. We provide more information about Iris’s language parameterization in §6.1. For space reasons, we will not attempt to present Iris in its full generality. Rather, to make things concrete, we will instantiate Iris specifically with **MyLang** and show how one can build a semantic soundness proof for this representative language.

Furthermore, instead of first explaining the design of Iris and then showing how to use it, we will present the features of Iris on-demand, as they arise in defining a *logical relation* (§5.1)—the key ingredient to encoding our semantic typing judgment for **MyLang** (§5.8). Then, in §6, we will show how to use this logical relation to prove semantic type soundness, and in §7 we will show how to formalize the idea of “safe encapsulation of unsafe features” that we presented informally in §3.2.

5.1 The Value and Expression Interpretations of Types

As explained in §4.1, proving semantic type soundness for the language **MyLang** involves defining a semantic version of the **MyLang** typing judgment, $\Gamma \models e : A$. Before defining the general semantic typing relation, we will first define a *logical relation*, which represents semantic typing for *closed* terms. It will then be straightforward in §5.8 to lift the logical relation on closed terms to a semantic typing relation on open terms using closing substitutions.

The logical relation for **MyLang**, shown in Figure 5, consists of two semantic interpretations of types A —a *value interpretation* $\llbracket A \rrbracket_\delta$ and an *expression interpretation* $\llbracket A \rrbracket_\delta^e$ —which are defined by structural recursion on A . These interpretations describe which values and expressions *behave like* valid inhabitants of A . Here, δ is a semantic *environment*, mapping type variables to their semantic value interpretations—hence, $\llbracket \alpha \rrbracket_\delta = \delta(\alpha)$. We will explain the need for this semantic environment when we come to the cases of the logical relation for types that bind variables, namely universal types, existential types, and recursive types. Until then, the reader can simply ignore the δ parameter, since it is otherwise merely threaded through the definition.

Crucially, note that $\llbracket A \rrbracket_\delta$ and $\llbracket A \rrbracket_\delta^e$ are interpretations of **MyLang** types *in Iris*—i.e., they are simply Iris predicates of type $Val \rightarrow iProp$ and $Expr \rightarrow iProp$.¹⁶ We now proceed to explain the definition of these predicates, step by step.

The expression interpretation. The first line of Figure 5 line shows how the expression interpretation is defined in terms of the value interpretation. Intuitively, a closed expression is in the expression interpretation of a type A if it *computes* a result that is in the value interpretation of A . This intuitive idea can be concisely captured using Iris’s *weakest precondition* connective:¹⁷

$$\llbracket A \rrbracket_\delta^e \triangleq \lambda e. \text{wp } e \{ \llbracket A \rrbracket_\delta \}$$

Given a postcondition $\Phi : Val \rightarrow iProp$, the connective $\text{wp } e \{ \Phi \}$ represents the weakest precondition ensuring that (1) e is safe to execute, and (2) any result value that e computes will satisfy Φ . Accordingly, $\llbracket A \rrbracket_\delta^e(e)$ expresses that e is safe to execute (i.e., it will not get stuck), and whatever value it evaluates to will be in the value interpretation of the type A . For administrative reasons, the weakest precondition $\text{wp}_\mathcal{E} e \{ \Phi \}$ is equipped with an *invariant mask* $\mathcal{E}$. We discuss the purpose of the invariant mask in §6.9 and omit it until then.

The remainder of Figure 5 defines the value interpretation of types, which we now explain.

5.2 Ground Types

The value interpretations of ground types are exactly what one would expect:

$$\llbracket 1 \rrbracket_\delta \triangleq \lambda v. v = () \quad \llbracket 2 \rrbracket_\delta \triangleq \lambda v. v \in \{\text{true}, \text{false}\} \quad \llbracket \mathbb{Z} \rrbracket_\delta \triangleq \lambda v. v \in \mathbb{Z}$$

The only value of the unit type **1** is the unit value $()$, the values of the Boolean type **2** are **true** and **false**, and the values of the integer type **Z** are the integers $\mathbb{Z}$.

5.3 Product and Sum Types

The value interpretations of product and sum types are similarly straightforward:

$$\begin{aligned} \llbracket A_1 \times A_2 \rrbracket_\delta &\triangleq \lambda v. \exists v_1, v_2. (v = (v_1, v_2)) * \llbracket A_1 \rrbracket_\delta(v_1) * \llbracket A_2 \rrbracket_\delta(v_2) \\ \llbracket A_1 + A_2 \rrbracket_\delta &\triangleq \lambda v. \bigvee_{i \in \{1, 2\}} \exists w. (v = \text{inj}_i w) * \llbracket A_i \rrbracket_\delta(w) \end{aligned}$$

¹⁶More specifically, since the type system of **MyLang** is an intuitionistic type system (where variables can be used repeatedly), the value interpretation uses *persistent* Iris predicates—i.e., their return type is really $iProp_\square$ —which intuitively means that they are predicates that can be freely duplicated. See §6.2 for more on this point.

¹⁷The particular weakest precondition predicate we use here is specific to the **MyLang** instantiation of Iris (see §6.1).

$$\begin{aligned}
\llbracket A \rrbracket_\delta^e &\triangleq \lambda e. \text{wp } e \{ \llbracket A \rrbracket_\delta \} \\
\llbracket \alpha \rrbracket_\delta &\triangleq \delta(\alpha) \\
\llbracket 1 \rrbracket_\delta &\triangleq \lambda v. v = () \\
\llbracket 2 \rrbracket_\delta &\triangleq \lambda v. v \in \{\text{true}, \text{false}\} \\
\llbracket Z \rrbracket_\delta &\triangleq \lambda v. v \in \mathbb{Z} \\
\llbracket A_1 \times A_2 \rrbracket_\delta &\triangleq \lambda v. \exists v_1, v_2. (v = (v_1, v_2)) * \llbracket A_1 \rrbracket_\delta(v_1) * \llbracket A_2 \rrbracket_\delta(v_2) \\
\llbracket A_1 + A_2 \rrbracket_\delta &\triangleq \lambda v. \bigvee_{i \in \{1, 2\}} \exists w. (v = \text{inj}_i w) * \llbracket A_i \rrbracket_\delta(w) \\
\llbracket A \rightarrow B \rrbracket_\delta &\triangleq \lambda v. \Box (\forall w. \llbracket A \rrbracket_\delta(w) \multimap \llbracket B \rrbracket_\delta^e(v w)) \\
\llbracket \forall \alpha. A \rrbracket_\delta &\triangleq \lambda v. \Box \left(\forall (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(v \langle \rangle) \right) \\
\llbracket \exists \alpha. A \rrbracket_\delta &\triangleq \lambda v. \exists (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \exists w. (v = \text{pack}(w)) * \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w) \\
\llbracket \mu \alpha. A \rrbracket_\delta &\triangleq \mu (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \lambda v. \exists w. (v = \text{fold } w) * \triangleright \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w) \\
\llbracket \text{ref } A \rrbracket_\delta &\triangleq \lambda v. \exists (\ell : \text{Loc}). (v = \ell) * \boxed{\exists w. \ell \mapsto w * \llbracket A \rrbracket_\delta(w)}^{\mathcal{N}_\ell} \\
\llbracket \emptyset \rrbracket_\delta^c(\emptyset) &\triangleq \text{True} \\
\llbracket \Gamma, x : A \rrbracket_\delta^c(\gamma, x \mapsto w) &\triangleq \llbracket \Gamma \rrbracket_\delta^c(\gamma) * \llbracket A \rrbracket_\delta(w) \\
\Gamma \models e : A &\triangleq \Box (\forall \delta, \gamma. \llbracket \Gamma \rrbracket_\delta^c(\gamma) \multimap \llbracket A \rrbracket_\delta^e(\gamma(e)))
\end{aligned}$$

Fig. 5. The expression interpretation $\llbracket _ \rrbracket^e$, value interpretation $\llbracket _ \rrbracket$, typing context interpretation $\llbracket _ \rrbracket^c$, and semantic typing judgment for **MyLang**.

Values of type $A_1 \times A_2$ are tuples (v_1, v_2) , where v_1 and v_2 are in the interpretation of A_1 and A_2 , respectively. Values of type $A_1 + A_2$ are either $\text{inj}_1 w$ or $\text{inj}_2 w$, where w is in the interpretation of A_1 or A_2 , respectively.

The reader may wonder why we use separating conjunction ($P * Q$) in the definition of $\llbracket A_1 \times A_2 \rrbracket_\delta$ rather than ordinary conjunction ($P \wedge Q$)—especially because, as we explain in §6.2, one *can* in fact replace all occurrences of $*$ in Figure 5 by $\wedge$ without changing the meaning of the logical relation. The short answer is that, in Iris proofs, particularly when mechanized in Coq, separating conjunction is the “default” and most commonly used form of conjunction. In the general case, where $*$ and $\wedge$ are not interchangeable, the appropriate connective to use is almost always $*$, not $\wedge$. (There are exceptions to this rule, but we will not encounter any in this paper.) So in cases where $*$ and $\wedge$ turn out to be interchangeable, we use $*$ as well for uniformity of notation.

We will return to this point in more detail in §6.2, but for the moment, the reader can simply think of $*$ as being synonymous with $\wedge$.

5.4 Function Types

The value interpretation of the function type $A \rightarrow B$ is perhaps the most iconic and familiar case, as some slight variation of it appears in any proof that calls itself a “logical relations” proof:

$$\llbracket A \rightarrow B \rrbracket_\delta \triangleq \lambda v. \Box (\forall w. \llbracket A \rrbracket_\delta(w) \multimap \llbracket B \rrbracket_\delta^e(v w))$$

It expresses that a value v inhabits the type $A \rightarrow B$ if v maps arguments in $\llbracket A \rrbracket_\delta$ to results in $\llbracket B \rrbracket_\delta^e$. Note that this definition imposes *no syntactic restriction* on v —it merely insists that, when v is

used like a function of type $A \rightarrow B$ (i.e., when applied to an argument of type A), it *behaves* like a function of type $A \rightarrow B$ (i.e., it is safe to execute and returns a result of type B).

There are three technical points of note here.

First, note that we use the *separating implication* connective ($P \multimap Q$), aka *magic wand*, instead of ordinary implication ($P \Rightarrow Q$). The reason is simple: since magic wand is the adjoint connective to separating conjunction—i.e., $P \multimap Q \vdash R$ iff $P \vdash Q \ast R$ (where $\vdash$ is the entailment relation of Iris)—and since in Iris (as noted above) we work mostly with $\ast$ rather than $\wedge$, we correspondingly work mostly with $\multimap$ rather than $\Rightarrow$. However, just as with $\ast$ vs. $\wedge$, and as we detail in §6.2, the distinction between $\multimap$ and $\Rightarrow$ is not important in the context of our logical relation, and the reader can comfortably gloss over it.

Second, note that $\llbracket A \rrbracket_\delta$ appears in a *negative* (contravariant) position in the definition of $\llbracket A \rightarrow B \rrbracket_\delta$. If one attempted to define the logical relation directly as an inductive predicate, this negative occurrence would cause a problem because it would render the inductive generating function non-monotone, so the definition would not be well-founded. However, as mentioned above, the logical relation is in fact defined by structural recursion on its type parameter, so since A is smaller than $A \rightarrow B$, the definition is in fact well-founded. Indeed, it is precisely this function case that necessitates defining $\llbracket A \rrbracket_\delta$ by structural recursion on A .

Lastly, note that the definition of $\llbracket A \rightarrow B \rrbracket_\delta$ is wrapped in Iris’s *persistence modality* ($\Box$); we defer explanation of this modality until §6.2.

5.5 Universal and Existential Types

The cases we have seen so far exhibit a common pattern. Types are interpreted semantically using the logical connective to which they are associated via the Curry-Howard correspondence: product types by conjunction, sum types by disjunction, and function types by implication. This pattern explains what is “logical” about a logical relation.

The next two cases continue this pattern, interpreting universal and existential types using logical propositions that are universally and existentially quantified, respectively. For readers familiar with prior work on logical relations—the “reducibility candidates” of Girard [1972], parametricity à la Reynolds [1983], or the logical characterization of parametricity due to Plotkin and Abadi [1993]—these cases should look very familiar. For other readers, some explanation is in order.

Naively, one might expect that, since the type variable α in $\forall \alpha. A$ (resp. $\exists \alpha. A$) represents an unknown syntactic type B , the definition of the logical relation for these types should universally (resp. existentially) quantify over a syntactic type B and then recurse on $A[B/\alpha]$, as follows:

Ill-founded attempt to define logical relation for $\forall \alpha. A$ and $\exists \alpha. A$:

$$\begin{aligned} \llbracket \forall \alpha. A \rrbracket_\delta &\triangleq \lambda v. \forall (B : \text{Type}). \llbracket A[B/\alpha] \rrbracket_\delta^e(v\langle \rangle) \\ \llbracket \exists \alpha. A \rrbracket_\delta &\triangleq \lambda v. \exists (B : \text{Type}). \exists w. (v = \text{pack}\langle w \rangle) \ast \llbracket A[B/\alpha] \rrbracket_\delta(w) \end{aligned}$$

But we have seen that the logical relation is crucially defined by structural recursion on its type parameter, and $A[B/\alpha]$ is not structurally smaller than $\forall \alpha. A$ and $\exists \alpha. A$. Even if we were to define the logical relation by recursion on the *size* of the type, we would run into the same problem because, thanks to the *impredicativity* of polymorphism in **MyLang**, the type B could be of arbitrary size.¹⁸ Hence, this naive definition is not well-founded.

¹⁸It is not accidental that polymorphism in **MyLang** is impredicative: impredicativity enables the programmer to represent an abstract data type internally using whatever type they want.

The solution, due originally to Girard [1972], is instead to define these cases of the logical relation by quantifying over a *semantic type* Ψ , rather than a syntactic type B :

$$\begin{aligned} \llbracket \forall \alpha. A \rrbracket_\delta &\triangleq \lambda v. \Box \left(\forall (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(v\langle \rangle) \right) \\ \llbracket \exists \alpha. A \rrbracket_\delta &\triangleq \lambda v. \exists (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \exists w. (v = \text{pack}\langle w \rangle) * \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w) \end{aligned}$$

By “semantic type”, we mean an arbitrary element Ψ drawn from the same space to which the value interpretation of types belongs—that is, Ψ is any (persistent) Iris predicate on values, of type $\text{Val} \rightarrow i\text{Prop}_\Box$. (Again we defer discussion of persistence until §6.2.) Once we have quantified over Ψ , we can then recurse over A , interpreting (free) occurrences of α in A using Ψ . This is achieved by extending the semantic environment δ to map α to Ψ . On a purely technical level, it is easy to see that this solves the problem with well-foundedness, since A is structurally smaller than $\forall \alpha. A$ and $\exists \alpha. A$. It also goes to show why we needed the semantic environment δ in the first place.

However, if one has not seen Girard’s method before, one may well wonder how this could possibly work and what ramifications it has. In particular, the space of semantic types includes many value predicates that are not the value interpretation of any syntactic type, so by quantifying over semantic types, does the definition of $\llbracket \forall \alpha. A \rrbracket_\delta$ not become too strong, and the definition of $\llbracket \exists \alpha. A \rrbracket_\delta$ too weak? The short answer why this works is *parametricity*: in **MyLang**, abstract types α are “really abstract”, in the sense that the language provides no way for the client of α to syntactically analyze the type B by which α is implemented (*i.e.*, the type with which α ultimately gets instantiated at runtime). Hence, there is no need to require that α be modeled as a syntactic **MyLang** type; it is fine to instead model α as belonging to the larger space of semantic types. Moreover, Girard’s method has the major side benefit that it will enable us to establish invariants on the private data representations of existentially-typed ADTs. We will see how this works when we formalize “safe encapsulation” in §7.

Finally, note that, when we quantify over semantic types, we are fundamentally relying on Iris’s support for higher-order (in this case, second-order) impredicative quantification.

5.6 Recursive Types

The value interpretation of recursive types poses yet another challenge. In principle, we would like to say that $\text{fold } w$ inhabits the type $\mu \alpha. A$ if w inhabits the type $A[\mu \alpha. A/\alpha]$, just as syntactic typing dictates. However, this would mean defining the value interpretation of $\mu \alpha. A$ in terms of the value interpretation of the *larger* type $A[\mu \alpha. A/\alpha]$, which is not well-founded.

Enter *guarded recursive predicates*, a distinctive feature of Iris which offers a way out of our predicament. In Iris, the *guarded fixed-point* operator $\mu x. t$ can be used to define recursive predicates without a restriction on the variance of the recursive occurrences of x in t . In return for this flexibility, all recursive occurrences of x must be *guarded*, meaning that they must appear below a *later modality* ($\triangleright$)—*i.e.*, within a term of the form $\triangleright P$. Subject to this restriction, $(\mu x. t) = t[\mu x. t/x]$. Using guarded recursion, the interpretation of recursive types becomes:

$$\llbracket \mu \alpha. A \rrbracket_\delta \triangleq \mu (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \lambda v. \exists w. (v = \text{fold } w) * \triangleright \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w)$$

By unrolling the μ we obtain the following:¹⁹

$$\llbracket \mu \alpha. A \rrbracket_\delta(v) = (\exists w. (v = \text{fold } w) * \triangleright \llbracket A[\mu \alpha. A/\alpha] \rrbracket_\delta(w))$$

This is *almost* exactly what we wanted. The only wrinkle here is the $\triangleright$ modality, which ensures well-foundedness of the guarded fixed-point. Roughly speaking, $\triangleright P$ means that “ P will hold in the

¹⁹The proof of this equivalence relies on the μ -equation, together with the fact that $\llbracket A[B/\alpha] \rrbracket_\delta = \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket B \rrbracket_\delta}$, which is proved by straightforward induction on the structure of A (see Lemma 6.3).

future after one step of computation”. That means that, if we have $\triangleright P$ in our context when proving a weakest precondition $\text{wp } e \{ \Phi \}$, we cannot make use of P right away; but, after we have verified that e can safely take a step of reduction to e' and the goal reduces to showing $\text{wp } e' \{ \Phi \}$, we are allowed to strip the $\triangleright$ off of $\triangleright P$ and use P in the rest of the proof. As we shall see in §6.8, this “under a later” knowledge about the semantic well-typedness of w is strong enough for us to be able to establish semantic type soundness.

5.7 Reference Types

Intuitively, values of the reference type $\text{ref } A$ are memory locations ℓ at which the value w stored may change over time but must always be of type A . In order to formalize this intuition in Figure 5, we make use of two features of Iris:

- The *points-to connective* $\ell \mapsto v$ (from vanilla separation logic) asserts exclusive *ownership* of location ℓ , along with the knowledge that it currently stores value v . (We will return to the concept of ownership in §6.2.)
- The *invariant assertion* $\boxed{P}^{\mathcal{N}}$ expresses the knowledge that a proposition P holds *invariantly*—i.e., at all times. Here, $\mathcal{N}$ denotes the *namespace* of the invariant, which is needed to ensure that invariants are not accessed repeatedly in an unsound fashion. See §6.9 for details.

With these connectives in hand, the interpretation of references types becomes:

$$\llbracket \text{ref } A \rrbracket_{\delta} \triangleq \lambda v. \exists (\ell : \text{Loc}). (v = \ell) * \boxed{\exists w. \ell \mapsto w * \llbracket A \rrbracket_{\delta}(w)}^{\mathcal{N}_{\ell}}$$

It says that v inhabits the type $\text{ref } A$ if v is a location ℓ and there is an invariant enforcing that ℓ always points to some value w that inhabits the type A . Here, $\mathcal{N}_{\ell}$ is the unique namespace that we use to designate the invariant on location ℓ .

5.8 The Semantic Typing Judgment

We now proceed to define the semantic typing judgment $\Gamma \models e : A$, which lifts the expression interpretation to *open* expressions e using a *closing substitution* $\gamma \in \text{Subst}$. Here, γ is a list of the form $x_1 \mapsto v_1, \dots, x_n \mapsto v_n$, which associates variables to (closed) values; we write $\gamma(e)$ to denote the result of applying the substitution γ to e . We will quantify over closing substitutions γ belonging to the *context interpretation* $\llbracket \Gamma \rrbracket_{\delta}^c : \text{Subst} \rightarrow i\text{Prop}_{\square}$, which says that γ maps every $x : B \in \Gamma$ to a value v that is in the value interpretation $\llbracket B \rrbracket_{\delta}$ of x 's type B . The formal definition is shown in Figure 5.

The *semantic typing judgment* $\Gamma \models e : A$ is defined as a relation in the Iris logic as follows:

$$\Gamma \models e : A \triangleq \square (\forall \delta, \gamma. \llbracket \Gamma \rrbracket_{\delta}^c(\gamma) * \llbracket A \rrbracket_{\delta}^c(\gamma(e)))$$

It says that e should inhabit the expression interpretation of A under any semantic environment δ and any closing substitution γ that satisfies the context interpretation $\llbracket \Gamma \rrbracket_{\delta}^c$. One can see this definition as essentially consisting of iterated applications of the function and universal type cases of the logical relation, which serve to abstract e over its typing context. The definition uses the persistence modality ($\square$) to ensure that semantic typing is a freely duplicable proposition.

6 LOGICAL TYPE SOUNDNESS: PROVING SEMANTIC TYPE SOUNDNESS IN IRIS

In this section, we demonstrate the method of *logical type soundness*—i.e., proving semantic type soundness within the separation logic framework of Iris. In particular, this involves proving (in Iris) semantic versions of the typing rules of **MyLang**. For instance, we will prove the following semantic typing rule for function application:

$$\frac{\Gamma \models e_1 : A \rightarrow B \quad \Gamma \models e_2 : A}{\Gamma \models e_1 \ e_2 : B}$$

Rules for separating conjunction and separating implication:

$$\begin{array}{c}
 \text{True} * P \vdash P \\
 P * Q \vdash Q * P \\
 (P * Q) * R \vdash P * (Q * R)
 \end{array}
 \quad
 \begin{array}{c}
 \text{*--MONO} \\
 \frac{P_1 \vdash Q_1 \quad P_2 \vdash Q_2}{P_1 * P_2 \vdash Q_1 * Q_2}
 \end{array}
 \quad
 \begin{array}{c}
 \text{*--INTRO} \\
 \frac{P * Q \vdash R}{P \vdash Q \multimap R}
 \end{array}
 \quad
 \begin{array}{c}
 \text{*--ELIM} \\
 \frac{P \vdash Q \multimap R \quad P * Q \vdash R}{P * Q \vdash R}
 \end{array}$$

Rules for the persistence modality:

$$\begin{array}{c}
 \text{□--MONO} \\
 \frac{P \vdash Q}{\Box P \vdash \Box Q}
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--DUP} \\
 \Box P \vdash (\Box P * \Box P)
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--ELIM} \\
 \Box P \vdash P
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--IDEMP} \\
 \Box P \vdash \Box \Box P
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--AND-SEP} \\
 (\Box P \wedge \Box Q) \vdash (\Box P * \Box Q)
 \end{array}$$

$$\begin{array}{c}
 \text{□--IMPL-WAND} \\
 (\Box P \Rightarrow Q) \vdash (\Box P \multimap Q)
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--SEP} \\
 \Box(P * Q) \vdash (\Box P * \Box Q)
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--TRUE} \\
 \text{True} \vdash \Box \text{True}
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--FALSE} \\
 \text{False} \vdash \Box \text{False}
 \end{array}$$

$$\begin{array}{c}
 \text{□--EQUAL} \\
 t = u \vdash \Box(t = u)
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--AND} \\
 \Box(P \wedge Q) \vdash (\Box P \wedge \Box Q)
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--OR} \\
 \Box(P \vee Q) \vdash (\Box P \vee \Box Q)
 \end{array}$$

$$\begin{array}{c}
 \text{□--FORALL} \\
 \Box(\forall x. P) \vdash (\forall x. \Box P)
 \end{array}
 \quad
 \begin{array}{c}
 \text{□--EXISTS} \\
 \Box(\exists x. P) \vdash (\exists x. \Box P)
 \end{array}$$

Rules for guarded recursion and the later modality:

$$\begin{array}{c}
 \mu\text{-UNFOLD} \\
 \vdash (\mu x. t) = t[\mu x. t/x]
 \end{array}
 \quad
 \begin{array}{c}
 \triangleright\text{-MONO} \\
 \frac{P \vdash Q}{\triangleright P \vdash \triangleright Q}
 \end{array}
 \quad
 \begin{array}{c}
 \triangleright\text{-INTRO} \\
 P \vdash \triangleright P
 \end{array}
 \quad
 \begin{array}{c}
 \text{LöB} \\
 \frac{\triangleright P \vdash P}{\vdash P}
 \end{array}
 \quad
 \begin{array}{c}
 \triangleright\text{-SEP} \\
 \triangleright(P * Q) \vdash (\triangleright P * \triangleright Q)
 \end{array}$$

$$\begin{array}{c}
 \triangleright\text{-AND} \\
 \triangleright(P \wedge Q) \vdash (\triangleright P \wedge \triangleright Q)
 \end{array}
 \quad
 \begin{array}{c}
 \triangleright\text{-OR} \\
 \triangleright(P \vee Q) \vdash (\triangleright P \vee \triangleright Q)
 \end{array}
 \quad
 \begin{array}{c}
 \triangleright\text{-FORALL} \\
 \triangleright(\forall x. P) \vdash (\forall x. \triangleright P)
 \end{array}$$

$$\begin{array}{c}
 \triangleright\text{-EXISTS} \\
 \text{inhabited}(\tau) \\
 \hline
 \triangleright(\exists x : \tau. P) \vdash (\exists x : \tau. \triangleright P)
 \end{array}
 \quad
 \begin{array}{c}
 \text{EXISTS-}\triangleright \\
 (\exists x. \triangleright P) \vdash \triangleright(\exists x. P)
 \end{array}
 \quad
 \begin{array}{c}
 \triangleright\text{-PERS} \\
 \triangleright(\Box P) \vdash \Box(\triangleright P)
 \end{array}$$

Rules for weakest preconditions (specialized to MyLang):

$$\begin{array}{c}
 \Phi(v) \vdash \text{wp}_{\mathcal{E}} v \{ \Phi \} \quad (\text{WP-VAL}) \\
 \text{if } e_1 \rightarrow_{\text{pure}} e_2 \text{ then } (\triangleright(\text{wp}_{\mathcal{E}} e_2 \{ \Phi \}) \vdash \text{wp}_{\mathcal{E}} e_1 \{ \Phi \}) \quad (\text{WP-PURE}) \\
 \text{wp}_{\mathcal{E}} e \{ w. \text{wp}_{\mathcal{E}} K[w] \{ \Phi \} \} \vdash \text{wp}_{\mathcal{E}} K[e] \{ \Phi \} \quad (\text{WP-BIND}) \\
 (\forall w. \Phi(w) \multimap \Psi(w)) * \text{wp}_{\mathcal{E}} e \{ \Phi \} \vdash \text{wp}_{\mathcal{E}} e \{ \Psi \} \quad (\text{WP-WAND}) \\
 \triangleright(\forall \ell. \ell \mapsto v \multimap \text{wp}_{\mathcal{E}} \ell \{ \Phi \}) \vdash \text{wp}_{\mathcal{E}} \text{ref } v \{ \Phi \} \quad (\text{WP-ALLOC}) \\
 \triangleright(\ell \mapsto v * (\ell \mapsto v \multimap \text{wp}_{\mathcal{E}} v \{ \Phi \})) \vdash \text{wp}_{\mathcal{E}} !\ell \{ \Phi \} \quad (\text{WP-LOAD}) \\
 \triangleright(\ell \mapsto v * (\ell \mapsto w \multimap \text{wp}_{\mathcal{E}} () \{ \Phi \})) \vdash \text{wp}_{\mathcal{E}} (\ell \leftarrow w) \{ \Phi \} \quad (\text{WP-STORE}) \\
 \triangleright(\ell \mapsto v * (\ell \mapsto w \multimap \text{wp}_{\mathcal{E}} \text{true} \{ \Phi \})) \vdash \text{wp}_{\mathcal{E}} \text{CAS}(\ell, v, w) \{ \Phi \} \quad (\text{WP-CAS-SUC}) \\
 \triangleright((v \neq v') * \ell \mapsto v * (\ell \mapsto v \multimap \text{wp}_{\mathcal{E}} \text{false} \{ \Phi \})) \vdash \text{wp}_{\mathcal{E}} \text{CAS}(\ell, v', w) \{ \Phi \} \quad (\text{WP-CAS-FAIL}) \\
 \triangleright(\ell \mapsto n * (\ell \mapsto (n + m) \multimap \text{wp}_{\mathcal{E}} n \{ \Phi \})) \vdash \text{wp}_{\mathcal{E}} \text{FAA}(\ell, m) \{ \Phi \} \quad (\text{WP-FAA}) \\
 \triangleright(\text{wp}_{\mathcal{E}} () \{ \Phi \} * \text{wp}_{\top} e \{ v. \text{True} \}) \vdash \text{wp}_{\mathcal{E}} \text{fork } \{ e \} \{ \Phi \} \quad (\text{WP-FORK})
 \end{array}$$

Fig. 6. Selected rules of the Iris logic.

Rules for invariants:

$$\begin{array}{c}
\text{INV-ALLOC} \quad \triangleright P \vdash \models_{\mathcal{E}} \boxed{P}^N \quad \text{INV-PERSIST} \quad \boxed{P}^N \vdash \Box \boxed{P}^N \quad \text{INV-OPEN-UPD} \quad \frac{\mathcal{N}^\uparrow \subseteq \mathcal{E}}{\boxed{P}^N * (\triangleright P \multimap \models_{\mathcal{E} \setminus \mathcal{N}^\uparrow} (\triangleright P * Q)) \vdash \models_{\mathcal{E}} Q} \\
\text{INV-OPEN-WP} \quad \frac{\text{atomic}(e) \quad \mathcal{N}^\uparrow \subseteq \mathcal{E}}{\boxed{P}^N * (\triangleright P \multimap \text{wp}_{\mathcal{E} \setminus \mathcal{N}^\uparrow} e \{v. \triangleright P * \Phi(v)\}) \vdash \text{wp}_{\mathcal{E}} e \{\Phi\}}
\end{array}$$

Rules for the update modality:

$$\begin{array}{c}
\text{Iris-MONO} \quad \frac{P \vdash Q}{\models_{\mathcal{E}} P \vdash \models_{\mathcal{E}} Q} \quad \text{Iris-INTRO} \quad P \vdash \models_{\mathcal{E}} P \quad \text{Iris-IDEMP} \quad \models_{\mathcal{E}} \models_{\mathcal{E}} P \vdash \models_{\mathcal{E}} P \quad \text{Iris-FRAME} \quad Q * \models_{\mathcal{E}} P \vdash \models_{\mathcal{E}} (Q * P) \\
\text{Iris-TIMELESS} \quad \frac{\text{timeless}(P)}{\triangleright P \vdash \models_{\mathcal{E}} P} \quad \text{Iris-WP} \quad \models_{\mathcal{E}} \text{wp}_{\mathcal{E}} e \{\Phi\} \vdash \text{wp}_{\mathcal{E}} e \{\Phi\} \quad \text{Iris-WP-Frame} \quad \text{wp}_{\mathcal{E}} e \{w. \models_{\mathcal{E}} \Phi(w)\} \vdash \text{wp}_{\mathcal{E}} e \{\Phi\}
\end{array}$$

Fig. 7. Selected rules for invariants and the update modality of the Iris logic.

Since the semantic typing judgment is an Iris definition, semantic typing rules are simply implications in Iris. For instance, the above inference rule should be read as:

$$(\Gamma \models e_1 : A \rightarrow B \quad * \quad \Gamma \models e_2 : A) \multimap \Gamma \models e_1 e_2 : B$$

Semantic typing rules are proven by unfolding the definition of the semantic typing judgment, the expression interpretation, and the value interpretation. Before carrying out these proofs in detail, we first discuss how Iris can be instantiated with a concrete programming language (§6.1), and return to a key technical issue that we have thus far glossed over—persistent propositions and the persistence modality $\Box$ (§6.2). We then explain the proof rules for weakest preconditions (§6.3), and how they are used to derive higher-level reasoning principles for the logical relation (§6.4). At that point, we are well equipped to prove the semantic typing rules (§6.5–§6.9). We then prove the *fundamental theorem*, which states that syntactic typing implies semantic typing, and the *adequacy theorem*, which states that closed semantically well-typed terms are indeed safe to execute (§6.10). In §7, we then demonstrate the key benefit of semantic typing—the ability to reason about “safe encapsulation of unsafe features”.

6.1 Language Parameterization and Basics of Iris

To make Iris applicable to a variety of programming languages, it is defined to be parametric in the types of expressions, values and states, and in a reduction relation [Jung et al. 2018b, §7.3]. The choice of programming language influences the semantics of Iris’s connective $\text{wp } e \{\Phi\}$ for weakest preconditions. Consequently, Iris’s proof rules for weakest preconditions are specific to the choice of language, while all other rules are language independent. For the purpose of this paper we instantiate Iris with the expressions, states, and reduction relation of **MyLang** (§2.3). This instantiation of Iris satisfies the proof rules given in Figure 6 and Figure 7, which we explain throughout this section.

Iris’s proof rules are not axioms. Their soundness is justified by a step-indexed model that is detailed in Jung et al. [2018b]. To use Iris, it is not necessary to understand the Iris model. Rather,

the key purpose of the model is to establish the adequacy theorem of weakest preconditions, which says that a closed proof of a weakest precondition implies safety with respect to the operational semantics.

THEOREM 6.1 (ADEQUACY OF WEAKEST PRECONDITIONS). *If $\text{True} \vdash wp\ e\ \{\Phi\}$, then $\text{safe}(e)$.*

Adequacy of weakest preconditions is essential in proving the adequacy theorem for our logical relation ([Theorem 6.6](#)). But it is important to emphasize that neither Iris’s connectives, nor its proof rules, nor its weakest preconditions, nor its adequacy theorem are specific to semantic soundness. These features and meta-theorems are also exploited by many Iris developments which have a completely different motivation, such as functional verification of low-level systems code.

Since we develop our semantic soundness proofs *within* Iris, these proofs are rather different from proofs in ordinary logic—they involve reasoning in separation logic and use the various Iris connectives. To allow the reader to get accustomed to the way such proofs are carried out, we visualize many proofs using proof trees. These proof trees make explicit many low-level details, so as to make clear exactly how the Iris rules are used. However, we should stress that, when Iris proofs are carried out in practice, they are done in Coq using the Iris Proof Mode [[Krebbers et al. 2017b, 2018](#)], which takes care of many of the low-level proof steps automatically.

The entailment relation $P \vdash Q$ says that P entails Q . For brevity, we use the following notations:

- $P \dashv\vdash Q$ means $P \vdash Q$ and $Q \vdash P$.
- $\vdash Q$ means $\text{True} \vdash Q$.
- If we write “proof of Q ” or “ Q is true” or “ Q holds”, we mean $\vdash Q$.

6.2 Persistent vs. Ephemeral Propositions and the $\Box$ Modality

In separation logic, propositions may express *exclusive ownership of resources*. We have already seen the prototypical example of such a proposition, namely the *points-to connective*, $\ell \mapsto v$, which denotes exclusive ownership of a location ℓ storing value v . Along with exclusive ownership of a resource typically comes the right to mutate the resource, which may have the effect of invalidating previously valid assertions about the resource. For example, if we can assert $\ell \mapsto 3$, we have the right to update ℓ to 5, after which we can assert $\ell \mapsto 5$, but at that point $\ell \mapsto 3$ no longer holds. Thus, propositions P expressing exclusive ownership are (to use Iris’s terminology) *ephemeral*: although P may hold at one point in a program proof, it may cease to hold later on.

There are many propositions, however, that do not assert exclusive ownership of resources; these propositions are (again following Iris’s terminology) *persistent*: once they hold, they hold forever. Examples of persistent propositions include pure facts like equality ($t = u$), as well as invariant assertions $\Box P$. Although persistent propositions do not offer any exclusive capabilities to their asserters, they do have an advantage over ephemeral propositions, namely that they are *duplicable*. That is, if P is persistent, then $P \dashv\vdash P * P$. Being duplicable is very useful because it means that once P is proven, it represents *freely shareable knowledge*: P can be used repeatedly, as often as needed, in the rest of the proof. For instance, if we need to prove $P \vdash Q * R$, we can reduce this to proving $P \vdash Q$ and $P \vdash R$, which we could not do if P were ephemeral.

Due to Iris’s support for ghost state, there are many other examples of ephemeral and persistent propositions besides the ones mentioned above. For example, in §7, we will see the connectives $\gamma \hookrightarrow_{=n}$ and $\gamma \hookrightarrow_{>n}$ of *ghost counters*, with the former being ephemeral and the latter persistent.

The notion of being persistent is expressed in Iris by means of the *persistence modality* $\Box$. The purpose of $\Box P$ is to say that P holds without depending on any ephemeral propositions. The most important rules for the persistence modality are $\Box P \dashv\vdash \Box P * \Box P$ (rule $\Box\text{-DUP}$) and $\Box P \vdash P$ (rule $\Box\text{-ELIM}$), which allow one to freely duplicate $\Box P$, and use it to obtain P when desired. Using the

persistence modality, we can formally define the class $iProp_{\Box}$ of persistent propositions, and what it means for a proposition P to be persistent (denoted $\text{persistent}(P)$):²⁰

$$\begin{aligned} iProp_{\Box} &\triangleq \{P : iProp \mid \text{persistent}(P)\} \\ \text{persistent}(P) &\triangleq P \vdash \Box P \end{aligned}$$

As usual for $\Box$ in modal logic, we have $\Box P \vdash \Box \Box P$ ($\Box$ -IDEMP), which implies that $\Box P$ is persistent regardless of what P is. Furthermore, using the fact that the $\Box$ modality commutes with most logical connectives (see Figure 6), we can show that the class of persistent propositions is closed under separating conjunction, conjunction, disjunction, universal quantification, and existential quantification. Lastly, using the fact that $\Box$ is monotone ($\Box$ -MONO), we can derive the following introduction rule, which says that a $\Box$ modality can be introduced if the context is persistent:

$$\frac{\Box\text{-INTRO} \quad P \vdash Q \quad \text{persistent}(P)}{P \vdash \Box Q}$$

Persistent propositions play an important role when defining a logical relation in separation logic. In particular, in the syntactic typing derivation $x : A \vdash e : B$, the assumption that x has type A may be used repeatedly. In establishing the corresponding semantic typing relation $x : A \models e : B$, we quantify over a value v , and must then prove that $\llbracket A \rrbracket^e(v) * \llbracket B \rrbracket^e(e[v/x])$. To prove that implication, we will need to use the assumption $\llbracket A \rrbracket^e(v)$ repeatedly, namely for each occurrence of x in e , which we can only do if it is persistent.

Consequently, we have set up the logical relation in Figure 5 so that $\llbracket A \rrbracket_{\delta}(v)$ is persistent by definition.²¹ In particular, the value and expression interpretations have the following types:

$$\begin{aligned} \llbracket _ \rrbracket_{(_)}^e &: \text{Type} \rightarrow (\text{Tvar} \rightarrow (\text{Val} \rightarrow iProp_{\Box})) \rightarrow \text{Expr} \rightarrow iProp \\ \llbracket _ \rrbracket_{(_)} &: \text{Type} \rightarrow (\text{Tvar} \rightarrow (\text{Val} \rightarrow iProp_{\Box})) \rightarrow \text{Val} \rightarrow iProp_{\Box} \end{aligned}$$

Here, we require the semantic environments δ (over which the interpretations are parameterized) to map type variables to persistent value predicates (i.e., functions from Val to $iProp_{\Box}$), and we require the value interpretation to return a persistent value predicate as well. Most cases of the value interpretation are persistent by construction. The only two that require some “intervention” in order to ensure persistence are the function and universal type cases. In these cases, the definition involves $\llbracket _ \rrbracket_{(_)}^e$, which is not persistent in general; so to make the definition persistent, we place its entire right-hand side under a $\Box$ modality.

Another important property of persistent propositions is $(\Box P \wedge Q) \dashv\vdash (\Box P * Q)$ (rule $\Box$ -AND-SEP), which says that ordinary conjunction and separating conjunction coincide when one conjunct is persistent, i.e., we have $P \wedge Q \dashv\vdash P * Q$ if P or Q is persistent. Similarly, ordinary implication $P \Rightarrow Q$ and magic wand $P * Q$ coincide when P is persistent, which follows from the fact that $(\Box P \Rightarrow Q) \dashv\vdash (\Box P * Q)$ (rule $\Box$ -IMPL-WAND). As a result, in the definition of the logical relation in Figure 5, the choice between $\wedge$ vs. $*$, and $\Rightarrow$ vs. $*$, is actually irrelevant. Nevertheless, as explained in §5.3–§5.4, we prefer to stick to $*$ and $\dashv\vdash$ in this paper for uniformity of notation (and thus not having to worry about how to associate $*$ and $\wedge$) and because that is what we actually do in Coq.

²⁰Iris’s step-indexed model, which is based on the category of OFEs (Ordered Families of Equivalences) [Birkedal et al. 2010a], does not support subset types in general—i.e., $\{x : \tau \mid \Phi x\}$ is not well-defined for every Iris type τ and predicate Φ . We omit the technical conditions on τ and Φ , but note that $iProp_{\Box}$ is indeed a well-defined subset type.

²¹Iris can of course also be used to model substructural type systems, in which the value interpretation $\llbracket A \rrbracket^e(v)$ will no longer be persistent, although persistence is useful in Iris for a number of other reasons as well. Examples of substructural type systems modeled in Iris include the Rust programming language [Jung et al. 2018a, 2021; Jung 2020; Dang et al. 2020] and session types [Tassarotti et al. 2017; Hinrichsen et al. 2021; Jacobs et al. 2024].

6.3 Weakest Preconditions

At the heart of the logical relation—in the definition of the expression interpretation $\llbracket e \rrbracket_\delta^c$, as shown in Figure 5—we use Iris’s connective $\text{wp } e \{ \Phi \}$ for weakest preconditions. Recall from §5.1 that, given a postcondition $\Phi : \text{Val} \rightarrow iProp$, the connective $\text{wp } e \{ \Phi \}$ represents the weakest precondition ensuring that (1) e is safe to execute, and (2) any result value e computes will satisfy Φ . To improve readability, we often write $\text{wp } e \{ w. Q \}$ instead of $\text{wp } e \{ \lambda w. Q \}$, and we completely omit the binder w in the postcondition if we do not say anything about the return value.

We will now go over the proof rules for weakest preconditions from Figure 6. The occurrences of *invariant masks* ($\mathcal{E}$) in this figure can be ignored for now; we will come back to those in §6.9.

Pure expressions. The simplest proof rules for weakest preconditions are **WP-VAL** and **WP-PURE**. The rule **WP-VAL** expresses that if an expression is a value v , then proving $\text{wp } v \{ \Phi \}$ can be reduced to proving the postcondition $\Phi(v)$. After all, a value v is vacuously safe, and values are results themselves. The rule **WP-PURE** expresses that if e_1 reduces to e_2 by a pure step (see Figure 3 for the definition of $\rightarrow_{\text{pure}}$), then proving $\text{wp } e_1 \{ \Phi \}$ can be reduced to proving $\triangleright(\text{wp } e_2 \{ \Phi \})$. The later modality ($\triangleright$) makes the new goal $\triangleright(\text{wp } e_2 \{ \Phi \})$ weaker (*i.e.*, easier to prove) than $\text{wp } e_2 \{ \Phi \}$, since we have $P \vdash \triangleright P$ (**P-INTRO**). Towards the end of this subsection, we explain the use of the later modality, but much of the time we can ignore it by applying **P-INTRO**.

For example, we can prove $\text{wp } (\text{if true then } () \text{ else } 42(42)) \{ v. v = () \}$ using **WP-PURE**, **P-INTRO**, and **WP-VAL** consecutively:

$$\frac{\frac{\frac{}{\vdash () = ()} \text{=refl}}{\vdash \text{wp } () \{ v. v = () \}} \text{WP-VAL}}{\vdash \triangleright \text{wp } () \{ v. v = () \}} \text{P-INTRO}}{\vdash \text{wp } (\text{if true then } () \text{ else } 42(42)) \{ v. v = () \}} \text{WP-PURE}$$

There are a couple of important things we should point out.

First, here and in the following text, we explain the proof trees starting with the conclusion and applying inference rules *bottom-up* to reduce it to simpler hypotheses, as one does generally when mechanizing these proofs in Coq.

Second, to compose proof rules, we implicitly use transitivity of the entailment relation ($\vdash$). For example, the bottom part of the proof above is actually (recall that $\vdash P$ means $\text{True} \vdash P$):

$$\frac{\dots \quad \text{True} \vdash \triangleright \text{wp } () \{ v. v = () \} \quad \triangleright \text{wp } () \{ v. v = () \} \vdash \text{wp } (\text{if true then } () \text{ else } 42(42)) \{ v. v = () \}}{\text{True} \vdash \text{wp } (\text{if true then } () \text{ else } 42(42)) \{ v. v = () \}} \text{WP-PURE} \quad \vdash\text{-trans}$$

Finally, we note that the weakest precondition in this example is logically equivalent to the semantic typing judgment $\models (\text{if true then } () \text{ else } 42(42)) : 1$. We show the proof for one direction of the equivalence (the other is similar):

$$\frac{\frac{\frac{\frac{\llbracket 0 \rrbracket_\delta^c(\gamma) \vdash \text{wp } (\text{if true then } () \text{ else } 42(42)) \{ v. v = () \}}{\llbracket 0 \rrbracket_\delta^c(\gamma) \vdash \text{wp } (\text{if true then } () \text{ else } 42(42)) \{ \llbracket 1 \rrbracket_\delta \}} \text{unfold } \llbracket 1 \rrbracket}{\llbracket 0 \rrbracket_\delta^c(\gamma) \vdash \llbracket 1 \rrbracket_\delta^c(\text{if true then } () \text{ else } 42(42))} \text{unfold } \llbracket _ \rrbracket^e}{\llbracket 0 \rrbracket_\delta^c(\gamma) \vdash \llbracket 1 \rrbracket_\delta^c(\gamma(\text{if true then } () \text{ else } 42(42)))} \gamma = \emptyset \text{ by definition of } \llbracket 0 \rrbracket_\delta^c}{\frac{\frac{\frac{\vdash \forall \delta, \gamma. \llbracket 0 \rrbracket_\delta^c(\gamma) \multimap \llbracket 1 \rrbracket_\delta^c(\gamma(\text{if true then } () \text{ else } 42(42)))}{\vdash \Box (\forall \delta, \gamma. \llbracket 0 \rrbracket_\delta^c(\gamma) \multimap \llbracket 1 \rrbracket_\delta^c(\gamma(\text{if true then } () \text{ else } 42(42))))} \text{V-intro, } \multimap\text{-INTRO}}{\vdash \models (\text{if true then } () \text{ else } 42(42)) : 1} \Box\text{-INTRO} \quad \text{unfold } \models$$

This simple example already demonstrates the flexibility of semantic typing—while the expression $\text{if true then } () \text{ else } 42(42)$ cannot be typed syntactically due to the presence of the ill-typed

subexpression 42(42), it *can* be typed semantically because the subexpression 42(42) appears in the else branch, which never gets executed.

Composition of proofs of weakest preconditions. Iris provides two important rules to compose proofs of weakest preconditions:

$$\begin{aligned} & \text{wp } e \{w. \text{wp } K[w] \{\Phi\}\} \vdash \text{wp } K[e] \{\Phi\} & (\text{WP-BIND}) \\ & (\forall w. \Phi(w) \multimap \Psi(w)) * \text{wp } e \{\Phi\} \vdash \text{wp } e \{\Psi\} & (\text{WP-WAND}) \end{aligned}$$

The rule **WP-BIND** generalizes the sequencing rule of Hoare logic, and **WP-WAND** generalizes the rules of consequence and framing of separation logic. Concretely, **WP-BIND** expresses that proving a weakest precondition for $K[e]$ can be reduced to proving a weakest precondition for e , followed by a weakest precondition for the continuation $K[w]$, where w is the result value of e . The rule **WP-WAND** provides a form of “internal monotonicity”, which allows applying a wand in the postcondition of the weakest precondition. This rule is interderivable with the more conventional rules for “external monotonicity” and framing:

$$\begin{aligned} & \frac{\text{WP-MONO} \quad \forall w. (\Phi(w) \vdash \Psi(w))}{\text{wp } e \{\Phi\} \vdash \text{wp } e \{\Psi\}} & \text{WP-FRAME} \quad \frac{}{P * \text{wp } e \{\Phi\} \vdash \text{wp } e \{w. P * \Phi(w)\}} \end{aligned}$$

To see the rules **WP-MONO** and **WP-FRAME** in action, assume we have some expression e , for which we already have in hand a proof of the weakest precondition $\text{wp } e \{v. v = ()\}$. To show how this proof can be reused, suppose we want to establish the weakest precondition $\text{wp } ((\lambda x. x) e) \{v. v = ()\}$. This is done as follows:

$$\begin{aligned} & \frac{}{\vdash () = ()} \text{--refl} \\ & \frac{}{\vdash \text{wp } () \{v. v = ()\}} \text{WP-VAL} \\ & \frac{}{\vdash \multimap \text{wp } () \{v. v = ()\}} \multimap\text{-INTRO} \\ & \frac{}{\vdash \text{wp } ((\lambda x. x) ()) \{v. v = ()\}} \text{WP-PURE} \\ & \frac{}{\vdash \forall w. (w = ()) \multimap \text{wp } ((\lambda x. x) w) \{v. v = ()\}} \forall\text{-intro, } \multimap\text{-INTRO, --subst} \quad \frac{}{\vdash \text{wp } e \{v. v = ()\}} \dots \\ & \frac{}{\vdash (\forall w. (w = ()) \multimap \text{wp } ((\lambda x. x) w) \{v. v = ()\}) * \text{wp } e \{v. v = ()\}} \text{*-MONO} \\ & \frac{}{\vdash \text{wp } e \{w. \text{wp } ((\lambda x. x) w) \{v. v = ()\}\}} \text{WP-WAND} \\ & \frac{}{\vdash \text{wp } ((\lambda x. x) e) \{v. v = ()\}} \text{WP-BIND} \end{aligned}$$

Here, we use rule **WP-BIND** with the call-by-value evaluation context $K \triangleq (\lambda x. x) []$. This allows us to prove a weakest precondition for e , followed by a weakest precondition for $K[w] = (\lambda x. x) w$, where w is the result of e . After applying **WP-BIND**, we need to prove a weakest precondition for e , but the postcondition does not match up with the postcondition of the already proven weakest precondition $\text{wp } e \{v. v = ()\}$. We therefore apply the rule **WP-WAND**, after which we proceed using the rules **WP-PURE** and **WP-VAL**, as in the previous example.

Hoare-style specifications and stateful expressions. To explain Iris’s weakest precondition rules for stateful expressions, we first show how conventional Hoare style specifications are written in Iris. We then explain why we prefer the weakest-precondition style specifications in [Figure 6](#).

The standard Hoare triple $\{P\} e \{\Phi\}$ can be encoded as $P \vdash \text{wp } e \{\Phi\}$ (or $\Box(P \multimap \text{wp } e \{\Phi\})$ if one wants to allow nested Hoare-triples). The standard rules from separation logic for stateful expressions [[Reynolds 2002](#); [O’Hearn et al. 2001](#)] can be formulated as follows:

$$\begin{aligned} & \{\text{True}\} \text{ref } v \{u. u \in \text{Loc} * u \mapsto v\} & (\text{HOARE-ALLOC}) \\ & \{\ell \mapsto v\} !\ell \{u. u = v * \ell \mapsto v\} & (\text{HOARE-LOAD}) \\ & \{\ell \mapsto v\} \ell \leftarrow w \{u. u = () * \ell \mapsto w\} & (\text{HOARE-STORE}) \end{aligned}$$

The rule **HOARE-ALLOC** says that allocation can be performed in any context (precondition True) and produces a location that points to the right value ($u \mapsto v$ in the postcondition). The rule **HOARE-LOAD** says that, to read from a location ℓ , the location should exist on the heap (precondition $\ell \mapsto v$), and the value that is returned is equal to the stored value ($u = v$ in the postcondition). We emphasize that it is necessary to include $\ell \mapsto v$ in the postcondition as well because $\ell \mapsto v$ is an ephemeral proposition, which describes exclusive ownership of the location ℓ —if it were not included in the postcondition, we would lose the ownership, making it impossible to use ℓ afterwards. The rule **HOARE-STORE** is similar: it says that, to write to a location ℓ , the location needs to exist on the heap (precondition $\ell \mapsto v$), and the value is changed accordingly ($\ell \mapsto w$ in the postcondition).

While the above Hoare-style specifications are valid in Iris (and can be derived from the rules in Figure 6), they are inconvenient in proof trees and Coq proofs. As usual in a weakest-precondition style system [Dijkstra 1975], we prefer to write the rules so that the postcondition is an arbitrary predicate Φ and we can apply them in a bottom-up fashion. Inspired by the “backwards” rules for separation logic by Ishtiaq and O’Hearn [2001], our rules thus adopt the following template:

“precondition” $*$ (“postcondition” $\multimap$ wp “result” $\{\Phi\}$) $\vdash$ wp “expression” $\{\Phi\}$

The concrete instances for stateful expressions are:

$$\begin{aligned} \triangleright(\forall \ell. \ell \mapsto v \multimap \text{wp } \ell \{ \Phi \}) &\vdash \text{wp } \text{ref } v \{ \Phi \} && \text{(WP-ALLOC)} \\ \triangleright(\ell \mapsto v * (\ell \mapsto v \multimap \text{wp } v \{ \Phi \})) &\vdash \text{wp } !\ell \{ \Phi \} && \text{(WP-LOAD)} \\ \triangleright(\ell \mapsto v * (\ell \mapsto w \multimap \text{wp } () \{ \Phi \})) &\vdash \text{wp } (\ell \leftarrow w) \{ \Phi \} && \text{(WP-STORE)} \end{aligned}$$

These rules use the separating conjunction ($*$) to express that ownership of the precondition needs to be given up, and the magic wand ($\multimap$) to express that ownership of the postcondition is given back and one should continue proving the weakest precondition for the result. For example, the rule **WP-STORE** states that $\ell \mapsto v$ needs to be given up (where v is the old value stored at ℓ), and $\ell \mapsto w$ is given in return (where w is the new value stored at ℓ). Concerning the $\triangleright$ modality, it is used the same way here as in **WP-PURE**; for more details, see the paragraph on Löb induction below.

Let us see the rule **WP-STORE** in action:

$$\frac{\frac{P * \ell \mapsto v \vdash P * \ell \mapsto w}{P \vdash \ell \mapsto w \multimap (P * \ell \mapsto w)} \text{*-INTRO} \quad \frac{\ell \mapsto v \vdash \ell \mapsto v}{\ell \mapsto v \vdash \ell \mapsto v} \text{*-MONO}}{\frac{P * \ell \mapsto v \vdash (\ell \mapsto w \multimap (P * \ell \mapsto w)) * \ell \mapsto v}{P * \ell \mapsto v \vdash \ell \mapsto v * (\ell \mapsto w \multimap (P * \ell \mapsto w))} \text{*comm}} \text{*-INTRO} \quad \frac{P * \ell \mapsto v \vdash \triangleright(\ell \mapsto v * (\ell \mapsto w \multimap (P * \ell \mapsto w)))}{P * \ell \mapsto v \vdash \text{wp } (\ell \leftarrow w) \{ P * \ell \mapsto w \}} \text{WP-STORE with } \Phi(u) = P * \ell \mapsto w$$

After applying **WP-STORE**, we use ***-MONO** to give up the hypothesis $\ell \mapsto v$ —this is often called “framing out a hypothesis”—and then use ***-INTRO** to introduce $\ell \mapsto w$. To frame out $\ell \mapsto v$, we use commutativity of the separating conjunction ($*$) so as to ensure $\ell \mapsto v$ appears in the same position on both sides of the entailment relation ($\vdash$). In larger proofs we leave reasoning up to commutativity (and associativity) implicit because it quickly becomes tedious, and in practice, the Iris Proof Mode in Coq takes care of it automatically anyway.

We have previously seen how the rule **WP-WAND** makes it possible to compose proofs of pure expressions. Now let us see how that rule is used for stateful expressions. Suppose we have proved

$$\forall \ell, n. \ell \mapsto n \multimap \text{wp } (\text{inc } \ell) \{ w. (w = ()) * \ell \mapsto n+1 \}, \quad \text{(WP-INC)}$$

where $\text{inc} \triangleq \lambda x. x \leftarrow (!x + 1)$, and we wish to prove

$$\ell \mapsto n \multimap \text{wp } (\text{inc } \ell; \text{inc } \ell) \{ w. (w = ()) * \ell \mapsto n+2 \}.$$

A proof tree for this example is as follows:

$$\begin{array}{c}
\frac{}{\ell \mapsto n+1 \vdash \text{wp inc } \ell \{ \Phi_{n+2}^\ell \}} \text{WP-INC} \\
\frac{}{\ell \mapsto n+1 \vdash \triangleright \text{wp inc } \ell \{ \Phi_{n+2}^\ell \}} \triangleright\text{-INTRO} \\
\frac{}{\ell \mapsto n+1 \vdash \text{wp } ((); \text{inc } \ell) \{ \Phi_{n+2}^\ell \}} \text{WP-PURE} \\
\frac{}{w = () * \ell \mapsto n+1 \vdash \text{wp } (w; \text{inc } \ell) \{ \Phi_{n+2}^\ell \}} \text{subst } w \\
\frac{}{\Phi_{n+1}^\ell(w) \vdash \text{wp } (w; \text{inc } \ell) \{ \Phi_{n+2}^\ell \}} \text{unfold } \Phi_{n+1}^\ell \\
\frac{}{\vdash \forall w. \Phi_{n+1}^\ell(w) * \text{wp } (w; \text{inc } \ell) \{ \Phi_{n+2}^\ell \}} \forall\text{-intro, } *- \text{INTRO} \quad \frac{}{\ell \mapsto n \vdash \text{wp inc } \ell \{ \Phi_{n+1}^\ell \}} \text{WP-INC} \\
\frac{}{\ell \mapsto n \vdash (\forall w. \Phi_{n+1}^\ell(w) * \text{wp } (w; \text{inc } \ell) \{ \Phi_{n+2}^\ell \}) * \text{wp inc } \ell \{ \Phi_{n+1}^\ell \}} *- \text{MONO} \\
\frac{}{\ell \mapsto n \vdash \text{wp inc } \ell \{ w. \text{wp } (w; \text{inc } \ell) \{ \Phi_{n+2}^\ell \} \}} \text{WP-WAND} \\
\frac{}{\ell \mapsto n \vdash \text{wp } (\text{inc } \ell; \text{inc } \ell) \{ \Phi_{n+2}^\ell \}} \text{WP-BIND} \\
\frac{}{\vdash \ell \mapsto n * \text{wp } (\text{inc } \ell; \text{inc } \ell) \{ \Phi_{n+2}^\ell \}} *- \text{INTRO}
\end{array}$$

Here, we let $\Phi_m^\ell(w) \triangleq (w = ()) * \ell \mapsto m$. To use the specification of `inc`, we first apply **WP-BIND** and **WP-WAND**. We use ***-MONO** to split the separating conjunction ($*$), giving ownership of $\ell \mapsto n$ to the right branch and no ownership (*i.e.*, `True`) to the left. The right branch follows immediately from the specification of `inc`. In the left branch, we use ***-INTRO** to obtain $\ell \mapsto n+1$ in the context, so we can conclude the proof by using the specification of `inc` again.

Löb induction and recursive functions. An important feature of Iris is **LöB** induction:

$$\frac{\begin{array}{c} \text{LöB} \\ \triangleright P \vdash P \end{array}}{\vdash P}$$

This rule allows one to reason about recursive computations by a kind of implicit induction on the number of steps they take. Specifically, when proving a goal P , **LöB** induction allows one to assume $\triangleright P$, which denotes that P will hold one step of computation *later*. Correspondingly, weakest precondition rules for reasoning about expressions that take a step of computation—such as **WP-PURE** and **WP-STORE**—contain a later modality $\triangleright$ in the premise because the verification of the rest of the computation (after the first step) need only be valid “later”. As a consequence, after applying such a rule (backwards) in a program proof, the new goal (*i.e.*, the premise of the rule just applied) will have the form $\triangleright Q$. By the rule **\triangleright-MONO**, one can strip the $\triangleright$ off both the goal ($\triangleright Q$) and any hypothesis $\triangleright P$ that had been previously introduced by **LöB** induction. From that point on, the **LöB** induction hypothesis P can be used freely in the remainder of the proof.

To see **LöB** induction in action, let us prove a weakest precondition for a trivial program that loops forever, `loop` $\triangleq$ **rec** $f(x) = 1 + f x$. The specification is `wp loop () {False}`. We can prove the postcondition `False` because the program never returns a value. A proof tree for this example is:

$$\begin{array}{c}
\frac{}{\text{wp loop } () \{ \text{False} \} \vdash \text{wp } 1 + \text{loop } () \{ \text{False} \}} \dots \\
\frac{}{\triangleright \text{wp loop } () \{ \text{False} \} \vdash \triangleright \text{wp } 1 + \text{loop } () \{ \text{False} \}} \triangleright\text{-MONO} \\
\frac{}{\triangleright \text{wp loop } () \{ \text{False} \} \vdash \text{wp loop } () \{ \text{False} \}} \text{WP-PURE} \\
\frac{}{\vdash \text{wp loop } () \{ \text{False} \}} \text{LöB}
\end{array}$$

The **LöB** rule provides the goal under a later ($\triangleright$) as a hypothesis. By taking a step of computation (`loop ()` $\rightarrow_{\text{pure}}$ `1 + loop ()`) using **WP-PURE**, we obtain a new goal wrapped in a $\triangleright$ modality. Using **\triangleright-MONO** (instead of **\triangleright-INTRO**) we obtain the **LöB** induction hypothesis without $\triangleright$. The remainder of the proof is routine, using **WP-BIND** and **WP-WAND**.

The later modalities $\triangleright$ in the rules for weakest preconditions (Figure 6) make these rules strictly stronger. The later modalities signify that a step of computation has been taken, and thereby make

it possible to strip a later off all hypotheses, and in particular the **LÖB** induction hypothesis, as we have seen in the proof above. If it is not needed to strip off a later, versions of the weakest precondition rules without the $\triangleright$ can be derived using the rule **$\triangleright$ -INTRO**.

6.4 Monadic Rules for the Expression Interpretation

To prove the semantic typing rules in the coming sections (§6.5–§6.9), we typically proceed by unfolding the definition of the semantic typing judgment $\Gamma \vdash e : A$, the expression interpretation $\llbracket A \rrbracket_{\delta}^e(e)$, and the value interpretation $\llbracket A \rrbracket_{\delta}(v)$, to obtain an Iris proposition that we then prove using Iris’s proof rules. To streamline these proofs, we prove the following auxiliary rules, whose statements resemble the types of the monadic operators return and bind.

LEMMA 6.2 (THE MONADIC RULES FOR THE EXPRESSION INTERPRETATION).

$$\begin{aligned} \llbracket A \rrbracket_{\delta}(v) &\multimap \llbracket A \rrbracket_{\delta}^e(v) && \text{(LOGREL-VAL)} \\ \llbracket A \rrbracket_{\delta}^e(e) * (\forall v. \llbracket A \rrbracket_{\delta}(v) &\multimap \llbracket B \rrbracket_{\delta}^e(K[v])) &\multimap \llbracket B \rrbracket_{\delta}^e(K[e]) && \text{(LOGREL-BIND)} \end{aligned}$$

PROOF. The rule **LOGREL-VAL** follows by unfolding the expression interpretation and the rule **WP-VAL**. The rule **LOGREL-BIND** follows by unfolding the expression interpretation and a combination of **WP-BIND** and **WP-WAND**. The proof is visualized in the following proof tree:

$$\frac{\frac{\text{wp } e \{ \llbracket A \rrbracket_{\delta} \} * (\forall v. \llbracket A \rrbracket_{\delta}(v) \multimap \text{wp } K[v] \{ \llbracket B \rrbracket_{\delta} \}) \vdash \text{wp } e \{ v. \text{wp } K[v] \{ \llbracket B \rrbracket_{\delta} \} \}}{\text{wp } e \{ \llbracket A \rrbracket_{\delta} \} * (\forall v. \llbracket A \rrbracket_{\delta}(v) \multimap \text{wp } K[v] \{ \llbracket B \rrbracket_{\delta} \}) \vdash \text{wp } K[e] \{ \llbracket B \rrbracket_{\delta} \}} \text{WP-BIND} \quad \text{wp } e \{ \llbracket A \rrbracket_{\delta} \} * (\forall v. \llbracket A \rrbracket_{\delta}(v) \multimap \text{wp } K[v] \{ \llbracket B \rrbracket_{\delta} \}) \vdash \text{wp } e \{ v. \text{wp } K[v] \{ \llbracket B \rrbracket_{\delta} \} \} \text{WP-WAND}}{\llbracket A \rrbracket_{\delta}^e(e) * (\forall v. \llbracket A \rrbracket_{\delta}(v) \multimap \llbracket B \rrbracket_{\delta}^e(K[v])) \vdash \llbracket B \rrbracket_{\delta}^e(K[e])} \text{unfold } \llbracket _ \rrbracket_{\delta}^e \quad \square$$

The **LOGREL-BIND** rule is particularly useful in the proofs of semantic typing rules because it enables us to “zap” an unknown but semantically well-typed term to a semantically well-typed value and proceed with the proof. Specifically, suppose we are trying to establish a goal of the form $\llbracket A \rrbracket_{\delta}^e(e) * P \vdash \llbracket B \rrbracket_{\delta}^e(K[e])$. That is, we want to prove that $K[e]$ is semantically well-typed at type B , and we have by assumption that the first subexpression to be evaluated (e) is semantically well-typed at type A . Now, if we knew what e was, then we could proceed by evaluating it, but often when proving semantic typing rules, we *do not* know what e is (*i.e.*, it is universally quantified by the typing rule), so it may seem the proof is stuck. Fortunately, in these cases, we can instead apply the **LOGREL-BIND** rule to reduce the goal to one in which the occurrences of the unknown e are “zapped” to (*i.e.*, replaced by) an unknown value v :

$$\frac{\frac{\frac{\llbracket A \rrbracket_{\delta}(v) * P \vdash \llbracket B \rrbracket_{\delta}^e(K[v])}{P \vdash \forall v. \llbracket A \rrbracket_{\delta}(v) \multimap \llbracket B \rrbracket_{\delta}^e(K[v])} \forall\text{-intro, } \multimap\text{-INTRO}}{\llbracket A \rrbracket_{\delta}^e(e) * P \vdash \llbracket A \rrbracket_{\delta}^e(e) * (\forall v. \llbracket A \rrbracket_{\delta}(v) \multimap \llbracket B \rrbracket_{\delta}^e(K[v]))} \multimap\text{-MONO}}{\llbracket A \rrbracket_{\delta}^e(e) * P \vdash \llbracket B \rrbracket_{\delta}^e(K[e])} \text{LOGREL-BIND}$$

We can then proceed by unfolding the definition of $\llbracket A \rrbracket_{\delta}(v)$, which typically yields information about v that allows us to make progress in evaluating $K[v]$.

In the following sections, we will use the above proof pattern repeatedly and refer to it simply as “zap the goal using **LOGREL-BIND**”.

6.5 Variables and Ground Types

We are now ready to start proving semantic versions of the typing rules of **MyLang**. Let us begin with the rules for variables and ground types. The semantic typing rule for variables is as follows:

$$\frac{\text{S-VAR} \quad x : A \in \Gamma}{\Gamma \models x : A}$$

PROOF OF S-VAR. The proof follows almost immediately from the way we defined the semantic typing judgment. Unfolding $\Gamma \models x : A$, our goal becomes to show $\llbracket \Gamma \rrbracket_{\delta}^c(\gamma) \multimap \llbracket A \rrbracket_{\delta}^c(\gamma(x))$ for any semantic environment δ and closing substitution γ . From $x : A \in \Gamma$, we obtain that $\llbracket \Gamma \rrbracket_{\delta}^c(\gamma)$ entails $\llbracket A \rrbracket_{\delta}^c(\gamma(x))$. Since $\gamma(x)$ is a value, we conclude by **LOGREL-VAL** from **Lemma 6.2**. $\square$

Let us proceed with the ground types, whose value interpretations we recall from **Figure 5**:

$$\llbracket 1 \rrbracket_{\delta} \triangleq \lambda v. v = () \quad \llbracket Z \rrbracket_{\delta} \triangleq \lambda v. v \in \mathbb{Z} \quad \llbracket 2 \rrbracket_{\delta} \triangleq \lambda v. v \in \{\text{true}, \text{false}\}$$

As explained in §5.1, these interpretations are exactly what one would expect: the only value of the unit type **1** is the unit value $()$, the values of the Boolean type **2** are **true** and **false**, and the values of the integer type **Z** are the integer literals $\mathbb{Z}$. The semantic typing rules for introduction of these types are as follows:

$$\begin{array}{c} \text{S-UNIT} \\ \Gamma \models () : 1 \end{array} \quad \begin{array}{c} \text{S-INT} \\ n \in \mathbb{Z} \\ \hline \Gamma \models n : Z \end{array} \quad \begin{array}{c} \text{S-BOOL} \\ b \in \{\text{true}, \text{false}\} \\ \hline \Gamma \models b : 2 \end{array}$$

PROOF OF S-UNIT, S-INT AND S-BOOL. These semantic typing rules are proven by unfolding the definition of the semantic typing judgment and making use of the rule **LOGREL-VAL**. For example, $\Gamma \models () : 1$ unfolds to $\llbracket \Gamma \rrbracket_{\delta}^c(\gamma) \multimap \llbracket 1 \rrbracket_{\delta}^c(\gamma())$ for any semantic environment δ and closing substitution γ . Since the unit value $()$ is closed, we have $\gamma() = ()$. Hence the expression interpretation $\llbracket 1 \rrbracket_{\delta}^c(\gamma())$ can be reduced to $\llbracket 1 \rrbracket_{\delta}^c()$, which in turn can be reduced to the value interpretation $\llbracket 1 \rrbracket_{\delta}()$ by **LOGREL-VAL**. The value interpretation $\llbracket 1 \rrbracket_{\delta}()$ unfolds to $() = ()$, which is a tautology. $\square$

While the proofs of the preceding rules follow almost immediately from unfolding the definition of the typing judgment, the next rules are more interesting to prove:

$$\begin{array}{c} \text{S-IF} \\ \Gamma \models e : 2 \quad \Gamma \models e_1 : B \quad \Gamma \models e_2 : B \\ \hline \Gamma \models \text{if } e \text{ then } e_1 \text{ else } e_2 : B \end{array} \quad \begin{array}{c} \text{S-FORK} \\ \Gamma \models e : A \\ \hline \Gamma \models \text{fork } \{e\} : 1 \end{array}$$

PROOF OF S-IF. We first prove the following auxiliary result for closed expressions:

$$\llbracket 2 \rrbracket_{\delta}^c(e) * \llbracket B \rrbracket_{\delta}^c(e_1) * \llbracket B \rrbracket_{\delta}^c(e_2) \multimap \llbracket B \rrbracket_{\delta}^c(\text{if } e \text{ then } e_1 \text{ else } e_2)$$

Here is a proof tree for the auxiliary result:

$$\begin{array}{c} \frac{\text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \} * \text{wp } e_2 \{ \llbracket B \rrbracket_{\delta} \} \vdash \text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \}}{\text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \} * \text{wp } e_2 \{ \llbracket B \rrbracket_{\delta} \} \vdash \text{wp } \text{if true then } e_1 \text{ else } e_2 \{ \llbracket B \rrbracket_{\delta} \}} \text{WP-PURE, } \triangleright\text{-INTRO} \quad \dots \\ \hline \frac{\text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \} * \text{wp } e_2 \{ \llbracket B \rrbracket_{\delta} \} \vdash \text{wp } \text{if true then } e_1 \text{ else } e_2 \{ \llbracket B \rrbracket_{\delta} \}}{\text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \} * \text{wp } e_2 \{ \llbracket B \rrbracket_{\delta} \} \vdash \text{wp } \text{if } v \text{ then } e_1 \text{ else } e_2 \{ \llbracket B \rrbracket_{\delta} \}} \text{case for false} \\ \hline \frac{\text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \} * \text{wp } e_2 \{ \llbracket B \rrbracket_{\delta} \} \vdash \text{wp } \text{if } v \text{ then } e_1 \text{ else } e_2 \{ \llbracket B \rrbracket_{\delta} \}}{\llbracket 2 \rrbracket_{\delta}(\nu) * \text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \} * \text{wp } e_2 \{ \llbracket B \rrbracket_{\delta} \} \vdash \text{wp } \text{if } v \text{ then } e_1 \text{ else } e_2 \{ \llbracket B \rrbracket_{\delta} \}} \text{unfold } \llbracket 2 \rrbracket \\ \hline \frac{\llbracket 2 \rrbracket_{\delta}(\nu) * \text{wp } e_1 \{ \llbracket B \rrbracket_{\delta} \} * \text{wp } e_2 \{ \llbracket B \rrbracket_{\delta} \} \vdash \text{wp } \text{if } v \text{ then } e_1 \text{ else } e_2 \{ \llbracket B \rrbracket_{\delta} \}}{\llbracket 2 \rrbracket_{\delta}^c(e) * \llbracket B \rrbracket_{\delta}^c(e_1) * \llbracket B \rrbracket_{\delta}^c(e_2) \vdash \llbracket B \rrbracket_{\delta}^c(\text{if } e \text{ then } e_1 \text{ else } e_2)} \text{unfold } \llbracket _ \rrbracket^c \\ \hline \llbracket 2 \rrbracket_{\delta}^c(e) * \llbracket B \rrbracket_{\delta}^c(e_1) * \llbracket B \rrbracket_{\delta}^c(e_2) \vdash \llbracket B \rrbracket_{\delta}^c(\text{if } e \text{ then } e_1 \text{ else } e_2) \text{ LOGREL-BIND} \end{array}$$

Reading this proof tree bottom-up, we zap the goal using **LOGREL-BIND** (as discussed in §6.4) with evaluation context $K \triangleq \text{if } [] \text{ then } e_1 \text{ else } e_2$. This turns the premise $\llbracket 2 \rrbracket_{\delta}^c(e)$ into $\llbracket 2 \rrbracket_{\delta}(\nu)$, where

v is an unknown value, and leaves us with the subgoal $\llbracket B \rrbracket_\delta^e(\text{if } v \text{ then } e_1 \text{ else } e_2)$. After that, we unfold the definitions of $\llbracket 2 \rrbracket$ and $\llbracket B \rrbracket^e$, and perform a case analysis on $v \in \{\text{true}, \text{false}\}$. In both cases, we then use **WP-PURE** to take a pure step to either e_1 or e_2 , which satisfy $\llbracket B \rrbracket_\delta^e$ by assumption.

To prove the actual semantic typing rule **S-IF**, we unfold the definition of the semantic typing judgment $\Gamma \models \text{if } e \text{ then } e_1 \text{ else } e_2 : B$, which shows that we have to prove that

$$\llbracket \Gamma \rrbracket_\delta^c(\gamma) \multimap \llbracket B \rrbracket_\delta^e(\text{if } \gamma(e) \text{ then } \gamma(e_1) \text{ else } \gamma(e_2))$$

follows from the assumptions $\Gamma \models e : 2$ and $\Gamma \models e_1 : B$ and $\Gamma \models e_2 : B$, which unfold to

$$\llbracket \Gamma \rrbracket_\delta^c(\gamma) \multimap \llbracket 2 \rrbracket_\delta^c(\gamma(e)) \quad \text{and} \quad \llbracket \Gamma \rrbracket_\delta^c(\gamma) \multimap \llbracket B \rrbracket_\delta^c(\gamma(e_1)) \quad \text{and} \quad \llbracket \Gamma \rrbracket_\delta^c(\gamma) \multimap \llbracket B \rrbracket_\delta^c(\gamma(e_2)).$$

Since the interpretation $\llbracket \Gamma \rrbracket_\delta^c(\gamma)$ of typing contexts is persistent, we can duplicate it. Our goal then follows from the auxiliary result $\llbracket 2 \rrbracket_\delta^c(e) * \llbracket B \rrbracket_\delta^c(e_1) * \llbracket B \rrbracket_\delta^c(e_2) \multimap \llbracket B \rrbracket_\delta^c(\text{if } e \text{ then } e_1 \text{ else } e_2)$ for closed expressions that we proved above. $\square$

A note about proofs. For all the semantic typing rules that follow, we (1) prove an auxiliary result about closed expressions, and (2) prove the semantic typing rule as a corollary. Step (1) is the interesting part, whereas step (2) involves just threading through the context interpretation. From now on, we only show step (1) and omit step (2). On a related note, our Coq tactics perform step (2) mostly automatically, so for example, the mechanized proof of **S-IF** is only 4 lines long.

PROOF OF S-FORK. We prove the following auxiliary result for closed expressions, from which the semantic typing rule follows immediately:

$$\llbracket A \rrbracket_\delta^c(e) \multimap \llbracket 1 \rrbracket_\delta^c(\text{fork } \{e\})$$

Here is a proof tree for the auxiliary result:

$$\frac{\frac{\frac{\frac{}{\vdash () = ()} \text{--refl}}{\vdash \llbracket 1 \rrbracket_\delta()} \text{unfold } \llbracket 1 \rrbracket}}{\vdash \text{wp } () \{ \llbracket 1 \rrbracket_\delta \}} \text{WP-VAL} \quad \frac{\frac{}{\vdash \forall v. \llbracket A \rrbracket_\delta(v) \multimap \text{True}} \text{WP-WAND}}{\text{wp } e \{ \llbracket A \rrbracket_\delta \} \vdash \text{wp } e \{ \text{True} \}} \text{WP-WAND}}{\frac{\text{wp } e \{ \llbracket A \rrbracket_\delta \} \vdash \text{wp } () \{ \llbracket 1 \rrbracket_\delta \} * \text{wp } e \{ \text{True} \}}{\text{wp } e \{ \llbracket A \rrbracket_\delta \} \vdash \text{wp } \text{fork } \{e\} \{ \llbracket 1 \rrbracket_\delta \}} \text{*-MONO}} \text{WP-FORK, } \triangleright\text{-INTRO}}{\frac{\llbracket A \rrbracket_\delta^c(e) \multimap \llbracket 1 \rrbracket_\delta^c(\text{fork } \{e\})}{\llbracket A \rrbracket_\delta^c(e) \multimap \llbracket 1 \rrbracket_\delta^c(\text{fork } \{e\})} \text{unfold } \llbracket _ \rrbracket^c}$$

The key part of this proof relies on the fork rule of the Iris instance for **MyLang**:

$$\triangleright(\text{wp } () \{ \Phi \} * \text{wp } e \{ \text{True} \}) \vdash \text{wp } \text{fork } \{e\} \{ \Phi \} \quad (\text{WP-FORK})$$

This rule says that to prove a weakest precondition for **fork** $\{e\}$, we need to prove a weakest precondition $\text{wp } () \{ \Phi \}$ for the main thread separately from a weakest precondition $\text{wp } e \{ \text{True} \}$ for the forked-off thread. Note that a forked-off expression is allowed to return any value since its result is thrown away, hence the postcondition is simply **True**. $\square$

Neither the proof of **S-FORK**, nor that of other semantic typing rules, involves explicit reasoning about the thread-pool semantics. This kind of reasoning is hidden by working in the Iris logic.

6.6 Product, Sum, and Function Types

Recall from **Figure 5** the value interpretation for product, sum, and function types:

$$\begin{aligned} \llbracket A_1 \times A_2 \rrbracket_\delta &\triangleq \lambda v. \exists v_1, v_2. (v = (v_1, v_2)) * \llbracket A_1 \rrbracket_\delta(v_1) * \llbracket A_2 \rrbracket_\delta(v_2) \\ \llbracket A_1 + A_2 \rrbracket_\delta &\triangleq \lambda v. \bigvee_{i \in \{1,2\}} \exists w. (v = \text{inj}_i w) * \llbracket A_i \rrbracket_\delta(w) \\ \llbracket A \rightarrow B \rrbracket_\delta &\triangleq \lambda v. \square (\forall w. \llbracket A \rrbracket_\delta(w) \multimap \llbracket B \rrbracket_\delta^c(v w)) \end{aligned}$$

As explained in §5.1, values of $A_1 \times A_2$ are tuples (v_1, v_2) , where v_1 and v_2 are in the interpretations of A_1 and A_2 , respectively. Values of $A_1 + A_2$ are either $\text{inj}_1 w$ or $\text{inj}_2 w$, where w is in the interpretation of A_1 or A_2 , respectively. Values of $A \rightarrow B$ are functions v that map arguments in the interpretation of A to results $v w$ in the interpretation of B . Recall from §6.2 that the $\Box$ modality is used to make the interpretation of the function type $A \rightarrow B$ persistent.

For products and sums, we prove semantic typing rules corresponding to the syntactic typing rules **T-PAIR**, **T-PROJ**, **T-INJ**, **T-MATCH-SUM**. The proofs proceed in a similar way to the proofs we have seen in §6.5. More interesting are the rules for functions:

$$\begin{array}{c} \text{S-APP} \\ \frac{\Gamma \vdash e_1 : A \rightarrow B \quad \Gamma \vdash e_2 : A}{\Gamma \vdash e_1 e_2 : B} \end{array} \qquad \begin{array}{c} \text{S-REC} \\ \frac{\Gamma, x : A, f : A \rightarrow B \vdash e : B}{\Gamma \vdash \text{rec } f(x) = e : A \rightarrow B} \end{array}$$

PROOF OF S-APP. As in the proofs in §6.5, we show just the auxiliary result that we prove for closed expressions, from which **S-APP** follows immediately:

$$\llbracket A \rightarrow B \rrbracket_{\delta}^e(e_1) * \llbracket A \rrbracket_{\delta}^e(e_2) \multimap \llbracket B \rrbracket_{\delta}^e(e_1 e_2)$$

Here is a proof tree for this auxiliary result:

$$\begin{array}{c} \frac{\frac{\frac{\frac{\frac{\frac{\llbracket A \rrbracket_{\delta}(v_2) \multimap \llbracket B \rrbracket_{\delta}^e(v_1 v_2)}{\Box(\forall w. \llbracket A \rrbracket_{\delta}(w) \multimap \llbracket B \rrbracket_{\delta}^e(v_1 w)) * \llbracket A \rrbracket_{\delta}(v_2) \vdash \llbracket B \rrbracket_{\delta}^e(v_1 v_2)}{\llbracket A \rightarrow B \rrbracket_{\delta}(v_1) * \llbracket A \rrbracket_{\delta}(v_2) \vdash \llbracket B \rrbracket_{\delta}^e(v_1 v_2)}{\llbracket A \rightarrow B \rrbracket_{\delta}(v_1) * \llbracket A \rrbracket_{\delta}^e(e_2) \vdash \llbracket B \rrbracket_{\delta}^e(v_1 e_2)}{\llbracket A \rightarrow B \rrbracket_{\delta}^e(e_1) * \llbracket A \rrbracket_{\delta}^e(e_2) \vdash \llbracket B \rrbracket_{\delta}^e(e_1 e_2)} \text{LOGREL-BIND}}{\llbracket A \rightarrow B \rrbracket_{\delta}(v_1) * \llbracket A \rrbracket_{\delta}(v_2) \vdash \llbracket B \rrbracket_{\delta}^e(v_1 v_2)} \text{LOGREL-BIND}}{\Box(\forall w. \llbracket A \rrbracket_{\delta}(w) \multimap \llbracket B \rrbracket_{\delta}^e(v_1 w)) * \llbracket A \rrbracket_{\delta}(v_2) \vdash \llbracket B \rrbracket_{\delta}^e(v_1 v_2)} \text{unfold } \llbracket A \rightarrow B \rrbracket} \text{ } \end{array} \begin{array}{l} \text{*-ELIM on LHS} \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \end{array}$$

Reading this proof tree bottom-up, we zap the goal using **LOGREL-BIND** twice (following the scheme we described in §6.4), first for e_1 in context $K \triangleq [\] e_2$, and then for e_2 in context $K \triangleq v_1 [\]$. The last step truly demonstrates why “logical relations” are called “logical”—we use Iris’s modus ponens rule $(Q \multimap R) * Q \vdash R$ (***-ELIM**) to eliminate the magic wand that appears in the interpretation of the function type $A \rightarrow B$. $\square$

PROOF OF S-REC. The auxiliary result for closed expressions is as follows:

$$\Box(\forall w v. \llbracket A \rrbracket_{\delta}(w) * \llbracket A \rightarrow B \rrbracket_{\delta}(v) \multimap \llbracket B \rrbracket_{\delta}^e(e[w/x][v/f])) \multimap \llbracket A \rightarrow B \rrbracket_{\delta}^e(\text{rec } f(x) = e)$$

This says that $\text{rec } f(x) = e$ is in the interpretation of $A \rightarrow B$ if, for all values w in the interpretation of A and for all values v in the interpretation of the recursive call of $A \rightarrow B$, we have that $e[w/x][v/f]$ in the interpretation of $A \rightarrow B$.

Let us abbreviate $P \triangleq \Box(\forall w v. \llbracket A \rrbracket_{\delta}(w) * \llbracket A \rightarrow B \rrbracket_{\delta}(v) \multimap \llbracket B \rrbracket_{\delta}^e(e[w/x][v/f]))$. The proof of the auxiliary result is as follows:

$$\begin{array}{c} \frac{\frac{\frac{\frac{\frac{\frac{\llbracket B \rrbracket_{\delta}^e(e[w/x][\text{rec } f(x) = e/f]) \vdash \text{wp } e[w/x][\text{rec } f(x) = e/f] \{ \llbracket B \rrbracket_{\delta} \}}{P * \llbracket A \rrbracket_{\delta}(w) * \llbracket A \rightarrow B \rrbracket_{\delta}(\text{rec } f(x) = e) \vdash \text{wp } e[w/x][\text{rec } f(x) = e/f] \{ \llbracket B \rrbracket_{\delta} \}}{P * \llbracket A \rrbracket_{\delta}(w) * \Box(\text{rec } f(x) = e) \vdash \text{wp } e[w/x][\text{rec } f(x) = e/f] \{ \llbracket B \rrbracket_{\delta} \}}{P * \llbracket A \rrbracket_{\delta}(w) * \Box(\text{rec } f(x) = e) \vdash \text{wp } (\text{rec } f(x) = e) w \{ \llbracket B \rrbracket_{\delta} \}} \text{WP-PURE}}{\frac{P * \Box(\text{rec } f(x) = e) \vdash \Box(\forall w. \llbracket A \rrbracket_{\delta}(w) \multimap \text{wp } (\text{rec } f(x) = e) w \{ \llbracket B \rrbracket_{\delta} \})}{P * \Box(\text{rec } f(x) = e) \vdash \Box(\forall w. \llbracket A \rrbracket_{\delta}(w) \multimap \text{wp } (\text{rec } f(x) = e) w \{ \llbracket B \rrbracket_{\delta} \})} \text{ } \end{array} \begin{array}{l} \text{unfold } \llbracket _ \rrbracket^e \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \\ \text{ } \end{array}$$

The key step of this proof is the use of the rule **LöB** for Löb induction, using which we obtain the induction hypothesis (IH) $\Box(\text{rec } f(x) = e) \vdash \Box(\text{rec } f(x) = e)$. Subsequently, we proceed by unfolding

the value interpretation of the function type $A \rightarrow B$, after which we obtain the resulting goal $\Box (\forall w. \llbracket A \rrbracket_\delta(w) \multimap \text{wp}(\text{rec } f(x) = e) w \{ \llbracket B \rrbracket_\delta \})$. We then introduce the $\Box$ modality, universal quantifier, and magic wand, and use **WP-PURE** to reduce $(\text{rec } f(x) = e) w$ to $e[w/x][v/f]$ by performing a step of computation. As a result of that, we obtain a later modality $\triangleright$ in our goal, allowing us to use **$\triangleright$ -MONO** to obtain the IH *now* (i.e., without $\triangleright$). We then eliminate the magic wand connectives in the premise $P \triangleq \Box (\forall w v. \llbracket A \rrbracket_\delta(w) * \llbracket A \rightarrow B \rrbracket_\delta(v) \multimap \llbracket B \rrbracket_\delta^e(e[w/x][v/f]))$ to obtain $\llbracket B \rrbracket_\delta^e(e[w/x][\text{rec } f(x) = e/f])$, which matches our goal exactly. $\square$

A formal note. The primitive version of Iris's rule **LÖB** is restricted to the empty context (i.e., the LHS of the entailment $\vdash$ should be True). However, in the above proof, the context is non-empty (it contains P). We therefore in fact use the following derived rule:

$$\frac{P * \triangleright Q \vdash Q \quad \text{persistent}(P)}{P \vdash Q}$$

The primitive version of **$\triangleright$ -MONO** allows us to introduce a later if the entire context is below a later. In the above proof this is not the case, and we thus use the following derived rule:

$$\frac{P * Q \vdash R}{P * \triangleright Q \vdash \triangleright R}$$

This rule is derived by using **$\triangleright$ -INTRO** and **$\triangleright$ -SEP** to turn the context $P * \triangleright Q$ into $\triangleright(P * Q)$, which makes it possible to use the primitive **$\triangleright$ -MONO**.

6.7 Universal and Existential Types

Recall from **Figure 5** the value interpretation for universal and existential types:

$$\begin{aligned} \llbracket \alpha \rrbracket_\delta &\triangleq \delta(\alpha) \\ \llbracket \forall \alpha. A \rrbracket_\delta &\triangleq \lambda v. \Box (\forall (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(v\langle \rangle)) \\ \llbracket \exists \alpha. A \rrbracket_\delta &\triangleq \lambda v. \exists (\Psi : \text{Val} \rightarrow i\text{Prop}_\Box). \exists w. (v = \text{pack}\langle w \rangle) * \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(w) \end{aligned}$$

As explained in **§5.1**, the semantic environment δ maps the free type variables to their semantic value interpretations—hence, $\llbracket \alpha \rrbracket_\delta = \delta(\alpha)$. The value interpretations of $\forall \alpha. A$ and $\exists \alpha. A$ quantify over a semantic type $\Psi : \text{Val} \rightarrow i\text{Prop}_\Box$ using Iris's universal and existential quantifiers, respectively. Within the quantification, they extend the semantic environment δ of the value interpretation of A to map α to Ψ . Note that since the expression $\llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(v\langle \rangle)$ is not persistent (it is defined in terms of a weakest precondition, which is not persistent), we wrap the value interpretation of $\forall \alpha. A$ in a persistence modality $\Box$ to ensure it is persistent.

The proofs of the semantic typing rules corresponding to **T-TAPP**, **T-TLAM**, **T-PACK**, and **T-MATCH-EX** crucially rely on Iris's rules for quantifiers. We additionally need the standard substitution lemma for logical relations, which says that substitution in types corresponds to extending the semantic type environment.

LEMMA 6.3. $\llbracket A[B/\alpha] \rrbracket_\delta = \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket B \rrbracket_\delta}$ and $\llbracket A[B/\alpha] \rrbracket_\delta^e = \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket B \rrbracket_\delta}^e$.

PROOF. Both results are proven mutually by induction on the structure of A . $\square$

Now let us show the proofs for the elimination and introduction rules for universal types:

$$\begin{array}{c} \text{S-TAPP} \\ \Gamma \vdash e : \forall \alpha. A \\ \hline \Gamma \vdash e\langle \rangle : A[B/\alpha] \end{array} \qquad \begin{array}{c} \text{S-TLAM} \\ \Gamma \vdash e : A \\ \hline \Gamma \vdash \Lambda. e : \forall \alpha. A \end{array}$$

PROOF OF **S-TAPP**. Following the usual approach, we first prove an auxiliary result for closed expressions:

$$\llbracket \forall \alpha. A \rrbracket_{\delta}^e(e) \multimap \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket B \rrbracket_{\delta}}^e(e\langle \rangle)$$

The proof tree is as follows:

$$\frac{\frac{\frac{\square (\Psi : Val \rightarrow iProp_{\square}). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(v\langle \rangle) \vdash \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket B \rrbracket_{\delta}}^e(v\langle \rangle)}{\llbracket \forall \alpha. A \rrbracket_{\delta}^e(v) \vdash \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket B \rrbracket_{\delta}}^e(v\langle \rangle)} \text{ } \square\text{-ELIM, } \forall\text{-elim}}{\llbracket \forall \alpha. A \rrbracket_{\delta}^e(e) \vdash \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket B \rrbracket_{\delta}}^e(e\langle \rangle)} \text{unfold } \llbracket \forall \alpha. A \rrbracket \text{ LOGREL-BIND}$$

The key step of this proof is the elimination of the universally quantified semantic type Ψ , which employs the standard elimination rule of the universal quantifier in higher-order logic. This again demonstrates why “logical relations” are called “logical”.

To prove the actual semantic typing rule **S-TAPP**, we unfold the definition of the semantic typing judgment, and see that we must prove that

$$\llbracket \Gamma \rrbracket_{\delta}^c(\gamma) \multimap \llbracket A[B/\alpha] \rrbracket_{\delta}^c(\gamma(e)\langle \rangle)$$

follows from the assumption

$$\llbracket \Gamma \rrbracket_{\delta}^c(\gamma) \multimap \llbracket \forall \alpha. A \rrbracket_{\delta}^c(\gamma(e)).$$

This result follows by threading through $\llbracket \Gamma \rrbracket_{\delta}^c(\gamma)$, **Lemma 6.3**, and the auxiliary result for closed expressions that we proved above. $\square$

PROOF OF **S-TLAM**. The auxiliary result for closed expressions is as follows:

$$\square (\Psi : Val \rightarrow iProp_{\square}). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(e) \multimap \llbracket \forall \alpha. A \rrbracket_{\delta}^e(\Lambda. e)$$

The proof tree is as follows:

$$\frac{\frac{\frac{\llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(e) \vdash \text{wp } e \{ \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi} \}}{\llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(e) \vdash \text{wp } (\Lambda. e)\langle \rangle \{ \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi} \}} \text{unfold } \llbracket _ \rrbracket^e \text{ WP-PURE, } \triangleright\text{-INTRO}}{\frac{\square (\Psi : Val \rightarrow iProp_{\square}). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(e) \vdash \square (\Psi : Val \rightarrow iProp_{\square}). \text{wp } (\Lambda. e)\langle \rangle \{ \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi} \}}{\square (\Psi : Val \rightarrow iProp_{\square}). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(e) \vdash \llbracket \forall \alpha. A \rrbracket_{\delta}^e(\Lambda. e)} \text{ } \square\text{-MONO, } \forall\text{-mono} \text{ unfold } \llbracket \forall \alpha. A \rrbracket \text{ and } \llbracket _ \rrbracket^e \text{ on RHS LOGREL-VAL}$$

6.8 Recursive Types

Recall from **Figure 5** the value interpretation for recursive types:

$$\llbracket \mu \alpha. A \rrbracket_{\delta} \triangleq \mu (\Psi : Val \rightarrow iProp_{\square}). \lambda v. \exists w. (v = \text{fold } w) * \triangleright \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w)$$

As explained in **§5.1**, the interpretation of recursive types $(\mu \alpha. A)$ uses Iris’s *guarded fixed-point* operator $(\mu x. t)$, which can be used to define recursive predicates without a restriction on the variance of the recursive occurrences of x in t . Instead, all recursive occurrences of x must be *guarded*, i.e., they have to appear below a later modality $(\triangleright)$. In the above definition this means that $\llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w)$ must appear below a later. The later makes a proposition weaker—we have $P \vdash \triangleright P$ (see **▷-INTRO**), but not the inverse (indeed, that would make the logic inconsistent). However, having $\llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w)$ below a later is strong enough for proving the semantic typing rules:

$$\frac{\text{S-FOLD} \quad \Gamma \vdash e : A[\mu \alpha. A/\alpha]}{\Gamma \vdash \text{fold } e : \mu \alpha. A} \quad \frac{\text{S-UNFOLD} \quad \Gamma \vdash e : \mu \alpha. A}{\Gamma \vdash \text{unfold } e : A[\mu \alpha. A/\alpha]}$$

As we will see, the proof of **S-FOLD** uses rule **▷-INTRO**, *i.e.*, $P \vdash \triangleright P$, while the proof of **S-UNFOLD** crucially relies on the fact that a computation step is performed to strip off the later.

The proofs of the semantic typing rules use the following unfolding lemma.

LEMMA 6.4. $\llbracket \mu\alpha. A \rrbracket_\delta(v) = (\exists w. (v = \text{fold } w) * \triangleright \llbracket A[\mu\alpha. A/\alpha] \rrbracket_\delta(w)).$

PROOF. By definition of $\llbracket \mu\alpha. A \rrbracket$ and μ -UNFOLD we obtain $\llbracket \mu\alpha. A \rrbracket_\delta(v) = (\exists w. (v = \text{fold } w) * \triangleright \llbracket A \rrbracket_{\delta, \alpha \mapsto \llbracket \mu\alpha. A \rrbracket_\delta(w)})$, which in turn by Lemma 6.3 concludes the proof. $\square$

PROOF OF S-FOLD. The auxiliary result for closed expressions is as follows:

$$\llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}^e(e) \multimap \llbracket \mu\alpha. A \rrbracket_{\delta}^e(\text{fold } e)$$

Below there follows a proof tree for the auxiliary result:

$$\begin{array}{c}
\frac{}{[A[\mu\alpha. A/\alpha]]_{\delta}(v) \vdash [A[\mu\alpha. A/\alpha]]_{\delta}(v)} \text{ } \vdash\text{-INTRO} \\
\frac{[A[\mu\alpha. A/\alpha]]_{\delta}(v) \vdash \exists v. (\text{fold } v = \text{fold } w) * \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(w)}{[A[\mu\alpha. A/\alpha]]_{\delta}(v) \vdash \llbracket \mu\alpha. A \rrbracket_{\delta}(\text{fold } v)} \exists\text{-intro } (w \triangleq v), =\text{-refl} \\
\frac{[A[\mu\alpha. A/\alpha]]_{\delta}(v) \vdash \llbracket \mu\alpha. A \rrbracket_{\delta}(\text{fold } v)}{[A[\mu\alpha. A/\alpha]]_{\delta}(v) \vdash \llbracket \mu\alpha. A \rrbracket_{\delta}^e(\text{fold } v)} \text{Lemma 6.4 on RHS} \\
\frac{[A[\mu\alpha. A/\alpha]]_{\delta}(v) \vdash \llbracket \mu\alpha. A \rrbracket_{\delta}^e(\text{fold } v)}{[A[\mu\alpha. A/\alpha]]_{\delta}(v) \vdash \llbracket \mu\alpha. A \rrbracket_{\delta}^e(\text{fold } v)} \text{LOGREL-VAL} \\
\frac{[A[\mu\alpha. A/\alpha]]_{\delta}^e(e) \vdash \llbracket \mu\alpha. A \rrbracket_{\delta}^e(\text{fold } e)}{[A[\mu\alpha. A/\alpha]]_{\delta}^e(e) \vdash \llbracket \mu\alpha. A \rrbracket_{\delta}^e(\text{fold } e)} \text{LOGREL-BIND}
\end{array}$$

PROOF OF S-UNFOLD. The auxiliary result for closed expressions is as follows:

$$\llbracket \mu\alpha. A \rrbracket_{\delta}^e(e) \multimap \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}^e(\text{unfold } e)$$

Below there follows a proof tree for the auxiliary result:

$$\begin{array}{c}
\frac{}{\llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(w) \vdash \text{wp } w \{ \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta} \}} \text{WP-VAL} \\
\frac{}{\triangleright \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(w) \vdash \text{wp } w \{ \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta} \}} \triangleright\text{-MONO} \\
\frac{}{\triangleright \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(w) \vdash \text{wp } \text{unfold}(\text{fold } w) \{ \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta} \}} \text{WP-PURE} \\
\frac{}{\triangleright \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(w) \vdash \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}^{\epsilon}(\text{unfold}(\text{fold } w))} \text{unfold } \llbracket _ \rrbracket^{\epsilon} \\
\frac{\exists w. (v = \text{fold } w) * \triangleright \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(w) \vdash \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}^{\epsilon}(\text{unfold } v)}{\triangleright \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(v) \vdash \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}^{\epsilon}(\text{unfold } v)} \exists\text{-elim, } =\text{-subst} \\
\frac{}{\llbracket \mu\alpha. A \rrbracket_{\delta}^{\epsilon}(v) \vdash \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}^{\epsilon}(\text{unfold } v)} \text{Lemma 6.4 on LHS} \\
\frac{}{\llbracket \mu\alpha. A \rrbracket_{\delta}^{\epsilon}(e) \vdash \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}^{\epsilon}(\text{unfold } e)} \text{LOGREL-BIND}
\end{array}$$

The key step of this proof is the use of rule **WP-PURE**, whose premise contains a later, and thus allows stripping the later off of the hypothesis $\triangleright \llbracket A[\mu\alpha. A/\alpha] \rrbracket_{\delta}(\mathbf{w})$ using **$\triangleright$ -MONO**. $\square$

It is worth noting that neither the proofs in this section, nor the proofs of any other semantic typing rule, involve explicit reasoning about step-indices. The only place where step-indexed reasoning pops up is in a few judicious applications of the later modality.

6.9 Reference Types

Recall from [Figure 5](#) the value interpretation for reference types:

$$\llbracket \text{ref } A \rrbracket_\delta \triangleq \lambda v. \exists (\ell : \text{Loc}). (v = \ell) * \boxed{\exists w. \ell \mapsto w * \llbracket A \rrbracket_\delta(w)}^{\mathcal{N}_\ell}$$

As explained in §5.1, values of the reference type $\text{ref } A$ should be memory locations ℓ at which the value w stored may change over time but is always of type A . This definition uses the *points-to connective* $\ell \mapsto v$ (from vanilla separation logic), which asserts exclusive ownership of the location ℓ storing value v , and Iris’s *invariant assertion* $\boxed{P}^N$, which expresses that a proposition P holds *invariantly*—i.e., at all times. As explained in §6.2, $\ell \mapsto v$ asserts exclusive ownership and is thus an ephemeral (non-persistent) proposition. By wrapping it in an invariant, we obtain a persistent proposition, which is thus freely duplicable.

The formal rules for invariants in Iris (**INV-ALLOC**, **INV-PERSIST**, and **INV-OPEN-WP**) can be found in Figure 7. Before we go into detail about these rules, let us informally explain the high-level roadmap for how one reasons about invariants in Iris:

- (1) **Invariant allocation:** At any moment in an Iris proof, if one can assert ownership of a proposition P , one can give this up in exchange for creating an invariant $\boxed{P}^{\mathcal{N}}$ (the invariant namespace $\mathcal{N}$ can be ignored for now). This can be understood as a form of *ownership transfer*: P is being transferred from one's *private state* (i.e., the private state of the thread whose code one is verifying) to the *shared state* (i.e., state shared by all threads). This ownership transfer to obtain an invariant is called *allocating* an invariant.
- (2) **Invariant duplication:** The upside of creating an invariant is that it enables one to take an ephemeral proposition (describing exclusive ownership of some state) and make it accessible to multiple threads at the same time. As explained above, this is achieved by the fact that the invariant assertion $\boxed{P}^{\mathcal{N}}$ is persistent: after an invariant has been allocated, it can be freely duplicated and thus shared among multiple threads.
- (3) **Invariant access:** The downside of turning P into an invariant is that no thread has unfettered access to P anymore because it has become a shared resource. Rather, each thread may only access the resource governed by the invariant in a carefully restricted way: during any *atomic* step of computation, a thread may acquire exclusive ownership of P so long as it gives P back by the end of that step. Atomicity of invariant access is essential for soundness of invariants because, in between acquiring and releasing ownership of P , the thread *does* have exclusive ownership, so it may in fact temporarily break the invariant (by falsifying P). But since this temporary breaking of the invariant only occurs *within* the reasoning about an atomic step of computation, no other threads can observe it, so it does not cause any problems. We refer to the acquisition and release of the ownership of the contents of an invariant as the *opening* and *closing* of the invariant.

Opening and closing invariants. Iris's rule for opening invariants is **INV-OPEN-WP** (the other rule for opening invariants, **INV-OPEN-UPD**, will be discussed in §7):

$$\frac{\text{INV-OPEN-WP} \quad \text{atomic}(e) \quad \mathcal{N}^\uparrow \subseteq \mathcal{E}}{\boxed{P}^{\mathcal{N}} * \left(\triangleright P \multimap \text{wp}_{\mathcal{E} \setminus \mathcal{N}^\uparrow} e \{v. \triangleright P * \Phi(v)\} \right) \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}$$

This rule is quite a mouthful, so let us go over it piece by piece. When proving a weakest precondition of an atomic expression e , this rule allows one to temporarily acquire exclusive ownership of P for the duration of the atomic step. Using the magic wand, one acquires $\triangleright P$ as an additional resource that can be used for proving the weakest precondition. In turn, in the postcondition of the weakest precondition, one has to restore $\triangleright P$. The side-condition $\text{atomic}(e)$ makes sure the rule is only used for *physically atomic expressions*, i.e., expressions e that take at most one step of computation:²²

$$\text{atomic}(e) \triangleq \forall \sigma, \sigma', e'. (\sigma, e) \rightarrow_t (\sigma', e') \Rightarrow e' \in \text{Val}$$

Examples of physically atomic expressions are **ref** v , **!** ℓ , $\ell \leftarrow v$, **FAA**(ℓ, v), and **CAS**(ℓ, v_1, v_2).

Later modalities and impredicativity. The reader may rightly wonder about the appearance of the $\triangleright$ modality in the rule **INV-OPEN-WP**. It turns out this modality is crucial for ensuring soundness

²²Aside from physical atomicity, Iris also supports a notion of *logical atomicity*, inspired by the TaDA logic [da Rocha Pinto et al. 2014]. Logically atomic triples express a form of linearizability—i.e., that even though an expression might take multiple steps in the operational semantics, it still appears to behave atomically [Jung et al. 2015, 2020].

in the presence of impredicative invariants.²³ By *impredicativity* of invariants, we mean that the proposition P in $\boxed{P}^N$ can be *any* Iris proposition, including one that contains nested invariant assertions. Impredicativity in turn is crucial for modeling the combination of polymorphism and higher-order references: the interpretation of a type like $\forall \alpha \dots \text{ref } A \dots$ will quantify (universally) over an arbitrary predicate Ψ representing α , and then Ψ will appear inside the invariant modeling the reference type $\text{ref } A$. The later modality is largely not a problem in practice: after opening an invariant, one can use the step-taking weakest precondition rules (like **WP-PURE**, **WP-ALLOC**, **WP-LOAD**) in order to strip the $\triangleright$ modality off the assumed proposition $\triangleright P$, thus obtaining P for use “now” in proving the postcondition $\Phi(v)$. We will see an example of this in the proof of **S-LOAD** below.

Invariant namespaces and masks. Two other important Iris mechanisms—albeit largely administrative ones that serve to ensure soundness of invariant reasoning—are *invariant namespaces* $N \in \text{InvName}$ and *invariant masks* $\mathcal{E} \subseteq \text{InvName}$. Namespaces and masks are used to ensure that invariants cannot be opened twice in a nested fashion, *i.e.*, that a thread cannot acquire exclusive ownership of the contents of the same invariant twice during the same atomic step of computation—an issue often referred to as *reenetrancy*. To avoid reenetrancy, Iris annotates each invariant $\boxed{P}^N$ with a *namespace* N that identifies the invariant, and annotates weakest preconditions $\text{wp}_{\mathcal{E}} e \{\Phi\}$ with a *mask* $\mathcal{E}$ that keeps track of the invariants that may be opened. At the top level, we always consider weakest preconditions with the mask $\top$ (*i.e.*, the whole set InvName), meaning that all invariants are available to be opened. Opening an invariant $\boxed{P}^N$ removes the namespace N from the mask $\mathcal{E}$, ensuring that it cannot be opened in a nested fashion.

There is a minor but potentially confusing technical point here that is worth clarifying. Namespaces have a hierarchical structure and are like fully qualified module names (using “dot notation”) in conventional programming languages. This hierarchical structure is convenient in developing modular proofs. When we write $\boxed{P}^N$, what we therefore really mean is that P is enforced by some invariant whose name ι belongs to the namespace N (*i.e.*, N is a *prefix* of ι). For example, ι might be N , but it also might be $N.\text{foo}$. Consequently, when the rule **INV-OPEN-WP** is used to open an invariant $\boxed{P}^N$ in a proof of $\text{wp}_{\mathcal{E}} e \{\Phi\}$, we must remove from $\mathcal{E}$ all invariant names ι which have N as a prefix. The set of all such invariant names is denoted $N^\uparrow$, which explains why the mask $\mathcal{E} \setminus N^\uparrow$ appears on the left-hand side of the turnstile in **INV-OPEN-WP**.

In this paper, we suppress details of how namespaces are constructed since they are really a minor implementation detail. For example, to identify the invariant for each reference ℓ , we simply assume the existence of a namespace N_ℓ , defined so that distinct locations map to disjoint sets of invariant names—*i.e.*, if $\ell \neq \ell'$, then $N_\ell^\uparrow \cap N_{\ell'}^\uparrow = \emptyset$. A detailed description of namespaces can be found in Jung et al. [2018b, §7.1.2].

Allocation of invariants. Using the following rules, one can transfer exclusive ownership of $\triangleright P$ into an invariant $\boxed{P}^N$:

$$\triangleright P * \left(\boxed{P}^N \text{ } \text{--} * \text{wp}_{\mathcal{E}} e \{\Phi\} \right) \vdash \text{wp}_{\mathcal{E}} e \{\Phi\} \quad (\text{INV-ALLOC-WP})$$

Iris in fact provides a more flexible rule for invariant allocation, called **INV-ALLOC**, from which the above rule is derived. We will discuss this more flexible rule in §7.

²³Krebbers et al. [2017a] and Jung et al. [2018b] present a paradox showing that a removal of the $\triangleright$ from **INV-OPEN-UPD** makes the logic inconsistent (*i.e.*, allows proving False). By contrast, the two occurrences of the $\triangleright$ modality in **INV-OPEN-WP** are not strictly needed for ensuring consistency. In particular, recent work by Spies et al. [2022] on “later credits” shows that it is possible to remove the first occurrence at the expense of complicating the invariant allocation rule **INV-ALLOC**. The second occurrence simply makes the rule stronger by weakening the postcondition that must be proved for e .

$$\begin{array}{c}
\text{S-ALLOC} \\
\frac{\Gamma \models e : A}{\Gamma \models \text{ref } e : \text{ref } A}
\end{array}
\qquad
\begin{array}{c}
\text{S-LOAD} \\
\frac{\Gamma \models e : \text{ref } A}{\Gamma \models !e : A}
\end{array}
\qquad
\begin{array}{c}
\text{S-STORE} \\
\frac{\Gamma \models e_1 : \text{ref } A \quad \Gamma \vdash e_2 : A}{\Gamma \models e_1 \leftarrow e_2 : \mathbf{1}}
\end{array}$$
$$\llbracket A \rrbracket_{\delta}^e(e) \multimap \llbracket \text{ref } A \rrbracket_{\delta}(\text{ref } e)$$
[illegible]
$$\llbracket \text{ref } A \rrbracket_{\delta}^e(e) \multimap \llbracket A \rrbracket_{\delta}^e(!e)$$
$$\begin{array}{c}
\frac{\ell \mapsto w * \llbracket A \rrbracket_{\delta}(w) \vdash I_{\ell}}{} \text{unfold } I_{\ell}, \exists\text{-intro} \\
\frac{\ell \mapsto w * \llbracket A \rrbracket_{\delta}(w) \vdash I_{\ell}}{} \text{▷-INTRO} \\
\frac{}{\llbracket A \rrbracket_{\delta}(w) \vdash \llbracket A \rrbracket_{\delta}(w)} \\
\frac{\ell \mapsto w * \llbracket A \rrbracket_{\delta}(w) \vdash I_{\ell} * \llbracket A \rrbracket_{\delta}(w)}{} \text{*-MONO and □-DUP} \\
\frac{\ell \mapsto w * \llbracket A \rrbracket_{\delta}(w) \vdash \triangleright I_{\ell} * \llbracket A \rrbracket_{\delta}(w)}{} \text{WP-VAL} \\
\frac{\ell \mapsto w * \llbracket A \rrbracket_{\delta}(w) \vdash \text{wp}_{\mathcal{E}} w \{w. \triangleright I_{\ell} * \llbracket A \rrbracket_{\delta}(w)\}}{} \text{*-INTRO} \\
\frac{\llbracket A \rrbracket_{\delta}(w) \vdash \ell \mapsto w * \text{wp}_{\mathcal{E}} w \{w. \triangleright I_{\ell} * \llbracket A \rrbracket_{\delta}(w)\}}{} \text{▷-MONO, *-MONO} \\
\frac{\triangleright (\ell \mapsto w * \llbracket A \rrbracket_{\delta}(w)) \vdash \triangleright \left(\ell \mapsto w * \left(\ell \mapsto w * \text{wp}_{\mathcal{E}} w \{w. \triangleright I_{\ell} * \llbracket A \rrbracket_{\delta}(w)\} \right) \right)}{} \text{WP-LOAD} \\
\frac{\triangleright (\ell \mapsto w * \llbracket A \rrbracket_{\delta}(w)) \vdash \text{wp}_{\mathcal{E}} ! \ell \{w. \triangleright I_{\ell} * \llbracket A \rrbracket_{\delta}(w)\}}{} \text{▷-EXISTS, } \exists\text{-elim on LHS} \\
\frac{\triangleright (\exists w. \ell \mapsto w * \llbracket A \rrbracket_{\delta}(w)) \vdash \text{wp}_{\mathcal{E}} ! \ell \{w. \triangleright I_{\ell} * \llbracket A \rrbracket_{\delta}(w)\}}{} \text{unfold } I_{\ell} \\
\frac{\triangleright I_{\ell} \vdash \text{wp}_{\mathcal{E}} ! \ell \{w. \triangleright I_{\ell} * \llbracket A \rrbracket_{\delta}(w)\}}{} \text{INV-OPEN-WP} \\
\frac{\overline{I_{\ell}}^{N_{\ell}} \vdash \text{wp} ! \ell \{ \llbracket A \rrbracket_{\delta} \}}{} \exists\text{-elim, =-subst} \\
\frac{(\exists \ell. (v = \ell) * \overline{I_{\ell}}^{N_{\ell}}) \vdash \text{wp} ! v \{ \llbracket A \rrbracket_{\delta} \}}{} \text{unfold } [\text{ref } A] \\
\frac{\llbracket \text{ref } A \rrbracket_{\delta}(v) \vdash \text{wp} ! v \{ \llbracket A \rrbracket_{\delta} \}}{} \text{unfold } [_]^e \\
\frac{\llbracket \text{ref } A \rrbracket_{\delta}(v) \vdash \llbracket A \rrbracket_{\delta}^e(!v)}{} \text{LOGREL-BIND} \\
\frac{\llbracket \text{ref } A \rrbracket_{\delta}^e(e) \vdash \llbracket A \rrbracket_{\delta}^e(!e)}{}
\end{array}$$

Here, we let $\mathcal{E} \triangleq \top \setminus \mathcal{N}_\ell^\uparrow$, and again, we let $I_\ell \triangleq \exists w. \ell \mapsto w * \llbracket A \rrbracket_\delta(w)$. The most important part of this proof is the use of the invariant opening rule **INV-OPEN-WP** to obtain temporary ownership of I_ℓ , which is needed to prove the weakest precondition for the load operation. Rule **INV-OPEN-WP** only gives $\triangleright I_\ell$, but since the load instruction takes a step of computation, we can use **$\triangleright$ -MONO** to strip off the later. (We first distribute the later over the existential quantifier in I_ℓ to obtain the witness w needed to apply **WP-LOAD**.) Finally, note that, at the top of the proof we are free to duplicate $\llbracket A \rrbracket_\delta(w)$ because the value interpretation of types is persistent by construction. $\square$

The proof of **S-STORE** is similar to the proof of **S-LOAD**. The key part of the proof lies in the fact that the value w is existentially quantified in $\llbracket \exists w. \ell \mapsto w * \llbracket A \rrbracket_\delta(w) \rrbracket^{\mathcal{N}_\ell}$. This means that it is fine to update ℓ to a new w so long as it satisfies $\llbracket A \rrbracket_\delta$ (as the second premise of **S-STORE** guarantees). Similarly, we can prove semantic typing rules corresponding to **T-FAA** and **T-FORK**.

6.10 The Fundamental Theorem and Adequacy

With the semantic typing rules corresponding to all syntactic typing rules in hand, we obtain that syntactic typing implies semantic typing:

THEOREM 6.5 (FUNDAMENTAL THEOREM OF UNARY LOGICAL RELATIONS). *Every syntactically well-typed term is semantically well-typed. Formally, if $(\Gamma \vdash e : A)$, then $(\Gamma \models e : A)$.*

PROOF. This theorem is proven by induction on the type derivation $(\Gamma \vdash e : A)$. For each case we use the corresponding semantic typing rule. $\square$

THEOREM 6.6 (ADEQUACY OF UNARY LOGICAL RELATIONS). *Every closed semantically well-typed expression e is safe: If $(\emptyset \models e : A)$, then $\text{safe}(e)$.*

PROOF. From $(\emptyset \models e : A)$, we obtain $\llbracket A \rrbracket_\emptyset^c(e)$ by definition of the semantic typing relation, which in turn, by definition, is equivalent to a closed proof of $\text{wp } e \ \{ \llbracket A \rrbracket_\emptyset \}$. We now obtain $\text{safe}(e)$ by adequacy of weakest preconditions (**Theorem 6.1**). $\square$

COROLLARY 6.7 (SEMANTIC TYPE SOUNDNESS). *Every closed syntactically well-typed expression e is safe. Formally, if $(\emptyset \vdash e : A)$, then $\text{safe}(e)$.*

PROOF. Let us assume that we have $(\emptyset \vdash e : A)$. By the fundamental theorem (**Theorem 6.5**), we obtain $(\emptyset \models e : A)$. By adequacy (**Theorem 6.6**), we obtain $\text{safe}(e)$, which concludes the proof. $\square$

7 SAFE ENCAPSULATION OF UNSAFE FEATURES

In the previous section, we showed how the logical approach to type soundness in Iris can be used to establish the well-known type soundness theorem (**Corollary 6.7**): well-typed programs are safe. Of course, if all we wanted was to prove **Corollary 6.7**, logical/semantic type soundness would not be needed—syntactic type soundness would suffice. What the stronger logical/semantic type soundness affords us is the additional ability to ensure that our language provides proper support for *data abstraction*, and to exploit that data abstraction for modular reasoning.

Concretely, recall the “evil”, data abstraction-breaking **gremlin** operator from §3.1, which non-deterministically proceeds as a simple no-op or selects some memory location ℓ currently storing an integer value n , and it updates ℓ to store 0. Intuitively, it is easy to see that, although **gremlin** does not disturb syntactic type soundness, it *would* violate semantic type soundness because it is *not* semantically well-typed.²⁴ Suppose we tried to prove $\vdash \text{gremlin} : 1$. To do so, we would have to show a weakest precondition for **gremlin**, and the difficult case would be the one where

²⁴The argument given here for the semantic ill-typedness of **gremlin** is intuitive but informal. For a more formal argument, see the discussion of **gremlin** at the end of this section.

`gremlin` nondeterministically updates some arbitrary memory cell ℓ to 0. Of course, in a separation logic like Iris, one cannot simply modify a location ℓ that one does not own: it could be owned by another part of the program or governed by a shared invariant, and either way, updating it to 0 could break whatever invariant or ownership assertion is currently imposed on it. So with our Iris-based semantic typing judgment in hand, we can happily declare `gremlin` *persona non grata* in our programming language.

This is great news, but even better is the *positive* thing we obtain from the guaranteed absence of features like `gremlin`: namely, the ability to verify safety of abstractions that are implemented internally using unsafe (syntactically ill-typed) features. We will now demonstrate this additional power by verifying safety of the `symbol` ADT from §3.2.

Recall the implementation of the `symbol` ADT:

$$\begin{aligned} \text{symbol_type} &\triangleq \exists \alpha. (1 \rightarrow \alpha) \times (\alpha \rightarrow 2) \\ \text{symbol} &\triangleq \text{let } c = \text{ref } 0 \text{ in} \\ &\quad \text{pack} \left\langle \text{Z}, \left(\begin{array}{l} \lambda (). \text{FAA}(c, 1), \\ \lambda s. \text{assert } (s < !c) \end{array} \right) \right\rangle \text{ as symbol_type} \end{aligned}$$

As we already explained in §3, the implementation employs a private integer counter c , which is allocated when the expression defining `symbol` is evaluated. The counter c is used as a perpetual source of fresh symbols. When the `gensym` function (the first closure returned by `symbol`) is called, it uses the fetch-and-add (FAA) instruction to atomically increment the value of c and return the previous value. Thus, when called repeatedly, `gensym` will return 0, 1, 2, and so on.

The check function (the second closure returned by `symbol`) checks validity of its `symbol` argument by checking that it is less than the current value of the counter. For this, it uses **MyLang**'s unsafe `assert` operation; hence, check is only safe to execute (*i.e.*, will not get stuck) if $s < !c$ indeed evaluates to `true`. Due to **MyLang**'s support for data abstraction, $s < !c$ *does* always evaluate to `true` in all well-typed contexts. We will now formalize this informal argument by proving the following theorem:

THEOREM 7.1 (THE SYMBOL ADT IS SEMANTICALLY WELL-TYPED).

$$\models \text{symbol} : \text{symbol_type}$$

When proving that the `symbol` ADT is semantically well-typed at an existential type—here, $\text{symbol_type} = \exists \alpha. (1 \rightarrow \alpha) \times (\alpha \rightarrow 2)$ —the key step is to choose the right *semantic type* for modeling α , *i.e.*, an Iris predicate $\Psi : \text{Val} \rightarrow i\text{Prop}_{\square}$ that describes the valid values of the abstract type α . When the functions of the ADT are given a value v of type α , they can *assume* that v satisfies Ψ , and when they return a value v of type α , they must *establish* that v satisfies Ψ .

The difficulty in the case of the `symbol` ADT is that the valid values of type α *change* over time. Intuitively, at any given point during the execution of a program containing `symbol`, the valid values of type α will be the symbols that have been generated *so far*—these are represented by the integers that are smaller than the *current* value stored in the private integer counter c used in the implementation of `symbol`. But how do we describe this intuitive property as a (persistent) Iris predicate $\Psi : \text{Val} \rightarrow i\text{Prop}_{\square}$? It must be a *state-dependent* predicate, meaning that it grows dynamically to be satisfied by more and more values as the state of c increases over time. How can we even define such a thing in Iris?

Naively, one might think that the following definition of Ψ should do the trick:

$$\Psi \triangleq \lambda v. \exists n : \mathbb{N}. (v < n) * c \mapsto n$$

This definition appears to say that v is a valid symbol if it is less than the current value n pointed to by c . The problem is that semantic types must be *persistent*, but this definition is *not* persistent.

It asserts exclusive ownership of c , which persistent predicates may not do. Moreover, if a value v satisfies Ψ now, there is no guarantee that it will continue to do so even after c gets updated. Intuitively, to make Ψ persistent, we will need some way of ensuring that valid symbols *stay* valid, which means we will need some way of enforcing the invariant that the counter value pointed to by c only grows larger over time. Toward that end, we now introduce one more feature of Iris—in fact, one of its most powerful and defining features: *user-defined ghost state*.

User-defined ghost state in Iris. Modern separation logics provide a variety of mechanisms for ownership of auxiliary state, often called *ghost state*. Some well-known examples include ghost variables [O’Hearn 2007], permissions [Bornat et al. 2005], protocols [Dinsdale-Young et al. 2010; Svendsen and Birkedal 2014], and history/prophesy assertions about the past/future trace of execution [Fu et al. 2010; Jung et al. 2020]. These mechanisms do not denote ownership of *physical* state (e.g., a location in the heap); rather, they describe *logical* state—i.e., state that is useful to track in proofs, but which is not directly manifested in the physical state of the program being verified.

In this section, we will show how to use Iris’s support for ghost state to encode the logical state of the counter in the symbol ADT, along with the property that it only grows larger over time, so that we can formulate an appropriate *persistent* predicate Ψ with which to model the ADT’s abstract type α . To be as flexible as possible, Iris does not bake in a particular ghost state mechanism, but rather allows the user of the framework to “roll their own” form of *user-defined* ghost state. Rolling your own ghost state essentially involves choosing an appropriate “resource algebra” to represent the ghost state mechanism you want; once the resource algebra is chosen, the base proof rules of Iris allow you to derive a corresponding *ghost theory*—i.e., a set of abstract predicates describing ownership of ghost state, together with axioms for manipulating them—on top of it.

The details of resource algebras and how they can be used to derive ghost theories, are beyond the scope of this paper. We focus our attention on a handful of concrete instances of user-defined ghost theories and refer the reader to Jung et al. [2018b] for more details about how such theories can be derived from suitably chosen resource algebras within Iris.

We begin by presenting a ghost theory that is directly relevant to the proof of the symbol ADT—namely, a theory of *ghost counters*. The connectives for ghost counters are as follows:

$$\hookrightarrow_{=} : GName \rightarrow \mathbb{N} \rightarrow iProp \quad \hookrightarrow_{>} : GName \rightarrow \mathbb{N} \rightarrow iProp_{\Box}$$

Similar to locations $\ell \mapsto v$ in physical state, ghost counters $\gamma \hookrightarrow_{=} m$ and $\gamma \hookrightarrow_{>} n$ can be referred to by a name $\gamma \in GName$. The set of ghost names $GName$ is similar to the set of locations Loc ; it needs to be infinite so Iris can pick a fresh name for each ghost allocation. Ghost counters can be allocated at any time during a proof and come in pairs: $\gamma \hookrightarrow_{=} n$ is an ephemeral proposition that provides exclusive ownership of the ghost location and says its value is *exactly* n , while $\gamma \hookrightarrow_{>} m$ is a persistent proposition that says its value is *strictly greater than* m . Since $\gamma \hookrightarrow_{>} m$ provides *persistent* knowledge, the value of the ghost location γ can only ever be increased—decreasing it could result in an already-established persistent assertion $\gamma \hookrightarrow_{>} m$ becoming falsified, which is not something that Iris lets happen.

Ghost counters are used in the proof of semantic typing of the symbol ADT as follows:

$$\begin{aligned} I_{\gamma} &\triangleq \lambda \ell. \boxed{\exists n : \mathbb{N}. \ell \mapsto n * \gamma \hookrightarrow_{=} n}^{\mathcal{N}_{\text{sym}}} \\ \Psi_{\gamma} &\triangleq \lambda v. v \in \mathbb{N} * \gamma \hookrightarrow_{>} v \end{aligned}$$

The invariant $I_{\gamma}(\ell)$, which will be shared by the closures of the ADT, describes that the value stored in the physical location ℓ (the location to which c in symbol gets bound) matches up with the value stored in the ghost counter at all times. The predicate Ψ_{γ} , which is used for the interpretation of

$$\begin{array}{ll}
\text{True} \vdash \models_{\mathcal{E}} \exists \gamma. \gamma \hookrightarrow_{=} 0 & (\text{CNT-INIT}) \\
\gamma \hookrightarrow_{=} n \vdash \models_{\mathcal{E}} \gamma \hookrightarrow_{=} (n+1) * \gamma \hookrightarrow_{>} n & (\text{CNT-INC}) \\
\gamma \hookrightarrow_{>} m \vdash \Box(\gamma \hookrightarrow_{>} m) & (\text{CNT-PERSIST}) \\
\gamma \hookrightarrow_{=} n * \gamma \hookrightarrow_{>} m \vdash m < n & (\text{CNT-LT}) \\
\text{timeless}(\gamma \hookrightarrow_{=} n) \text{ and } \text{timeless}(\gamma \hookrightarrow_{>} n) & (\text{CNT-TIMELESS})
\end{array}$$

Fig. 8. Iris's rules for ghost counters.

the abstract type α , employs the persistent part of the ghost counter $\gamma \hookrightarrow_{>} m$ to ensure that the values of type α are integers m that are smaller than the value stored in the counter c .

The rules for ghost counters are given in Figure 8. These rules make use of a new connective called the *update modality* $\models_{\mathcal{E}}$, which (as the name suggests) is used to account for updates to ghost state. Before explaining it, let us provide the intuition for the rules for ghost counters:

- The rule **CNT-INIT** is used to allocate a new ghost counter. It provides exclusive ownership of $\gamma \hookrightarrow_{=} 0$, where γ is a fresh (*i.e.*, existentially quantified) name.
- The rule **CNT-INC** is used to increment the ghost counter. In addition to transforming $\gamma \hookrightarrow_{=} n$ into $\gamma \hookrightarrow_{=} (n+1)$, the rule also yields $\gamma \hookrightarrow_{>} n$, which provides the persistent knowledge that the ghost counter is strictly greater than n . (Note: $*$ binds more tightly than $\models_{\mathcal{E}}$, so “ $\models_{\mathcal{E}} P * Q$ ” means “ $\models_{\mathcal{E}} (P * Q)$ ”.)
- The rule **CNT-PERSIST** states that the connective $\gamma \hookrightarrow_{>} m$ is indeed persistent.
- The rule **CNT-LT** states that if we have exclusive ownership of γ (with current value n), along with the knowledge that γ 's value is greater than m , then we must know that $m < n$.

The update modality. The update modality $\models_{\mathcal{E}} Q$ has many similarities with the weakest precondition connective $\text{wp}_{\mathcal{E}} e \{w. Q\}$, but is used for reasoning about ghost state rather than physical state. Since ghost state is merely logical, there is no physical program e , and the postcondition is merely a proposition, not a value predicate (*i.e.*, it does not have a binder w for the return value). The update modality can be used for the following purposes:

- In order for clients of a ghost theory to make use of rules for allocating or updating ghost state (like **CNT-INIT** and **CNT-INC** in Figure 8), Iris provides the rules $\models_{\mathcal{E}}\text{-WP}$ and $\text{WP}\text{-}\models_{\mathcal{E}}$. These rules allow one to eliminate update modalities around weakest preconditions:

$$\begin{array}{ll}
\text{\textcolor{violet}{\models}\text{-WP}} & \text{WP}\text{-}\text{\textcolor{violet}{\models}} \\
\text{\textcolor{violet}{\models}}_{\mathcal{E}} \text{wp}_{\mathcal{E}} e \{ \Phi \} \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \} & \text{wp}_{\mathcal{E}} e \{ w. \text{\textcolor{violet}{\models}}_{\mathcal{E}} \Phi(w) \} \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}
\end{array}$$

(In a more traditional presentation with Hoare triples, these rules would correspond to a strengthened rule of consequence, in which the implications for adjusting the pre- and postcondition of the Hoare triple are additionally permitted to perform ghost updates.)

Apart from these rules, there are a number of administrative rules ($\models_{\mathcal{E}}\text{-MONO}$, $\models_{\mathcal{E}}\text{-INTRO}$, $\models_{\mathcal{E}}\text{-IDEMP}$, and $\models_{\mathcal{E}}\text{-FRAME}$), shown in Figure 7, which collectively establish that the update modality is a strong monad with respect to separating conjunction [Kock 1970, 1972]. Using these administrative rules, we can turn the rule **CNT-INC** into the following, more usable rule:

$$\begin{array}{c}
\text{CNT-INC}' \\
\hline
\gamma \hookrightarrow_{=} (n+1) * \gamma \hookrightarrow_{>} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \} \\
\hline
\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}
\end{array}$$

This rule says that if our context provides exclusive ownership of a ghost counter $\gamma \hookrightarrow_{=} n$ with name γ whose current value is exactly n , we can increase its value to $n + 1$. Additionally we obtain the persistent knowledge $\gamma \hookrightarrow_{>} n$ that the new value is strictly greater than n .

A proof tree for this rule is as follows:

$$\begin{array}{c}
 \frac{\gamma \hookrightarrow_{=} (n+1) * \gamma \hookrightarrow_{>} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}{\frac{\frac{\frac{\frac{\gamma \hookrightarrow_{=} (n+1) * \gamma \hookrightarrow_{>} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}{\frac{\frac{\gamma \hookrightarrow_{=} (n+1) * \gamma \hookrightarrow_{>} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}{\frac{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}}}{\frac{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}}}}}{\frac{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}}}}}{\frac{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}{\gamma \hookrightarrow_{=} n * Q \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}}}}}
 \end{array}
 \begin{array}{l}
 \text{⇒-MONO} \\
 \text{⇒-FRAME} \\
 \text{CNT-INC} \\
 \text{⇒-WP}
 \end{array}$$

This proof shows a typical pattern in Iris. We use a rule like **CNT-INC** in one of the hypotheses, and use **⇒-WP** to obtain a matching update modality in the goal. Then, using **⇒-FRAME** on the LHS, we move all other hypotheses below the update modality, and finally use **⇒-MONO** to strip the update modality off both the hypotheses and goal. When mechanizing Iris proofs in Coq, the Iris Proof Mode takes care of these administrative steps automatically.

- In order to build modular logical abstractions (as we will show in §8.3), Iris's update modality can also be used for allocating and opening invariants. Similar to weakest preconditions, the update modality $\text{⇒}_{\mathcal{E}}$ is thus equipped with a mask $\mathcal{E}$ that denotes which invariants may be opened. The rules **INV-ALLOC** and **INV-OPEN-UPD** for opening invariants around the update modality are as follows:

$$\begin{array}{c}
 \text{INV-ALLOC} \\
 \frac{}{\triangleright P \vdash \text{⇒}_{\mathcal{E}} \boxed{P}^N}
 \end{array}
 \quad
 \begin{array}{c}
 \text{INV-OPEN-UPD} \\
 \frac{\mathcal{N}^{\uparrow} \subseteq \mathcal{E}}{\boxed{P}^N * (\triangleright P * \text{⇒}_{\mathcal{E} \setminus \mathcal{N}^{\uparrow}} (\triangleright P * Q)) \vdash \text{⇒}_{\mathcal{E}} Q}
 \end{array}$$

The rule **INV-ALLOC** transfers ownership of a proposition P into an invariant $\boxed{P}^N$. The rule **INV-ALLOC-WP** for allocating invariants around weakest preconditions can be derived from **INV-ALLOC** and **⇒-WP**.

The rule **INV-OPEN-UPD** is very similar to the rule **INV-OPEN-WP** that we have seen in §6.9, except with update modalities instead of weakest precondition assertions. When proving an update to Q , this rule allows one to temporarily acquire exclusive ownership of P , assuming P is the content of an invariant named $\mathcal{N}$. Using the magic wand, one acquires $\triangleright P$ as an additional resource that can be used for proving the update to Q . In turn, one has to restore $\triangleright P$, and the mask keeps track of the fact that the invariant $\mathcal{N}$ cannot be opened in a nested fashion.

- When opening an invariant $\boxed{P}^N$ via the rules **INV-OPEN-UPD** and **INV-OPEN-WP**, one temporarily gets ownership of $\triangleright P$, where P is guarded by a later modality ($\triangleright$). As discussed in §6.9, the later modality is crucial for soundness in the presence of impredicative invariants; however, for the class of so-called *timeless* Iris propositions, one can acquire access to the contents of the invariant *without* a later:

$$\begin{array}{c}
 \text{INV-OPEN-UPD-TL} \\
 \frac{\mathcal{N}^{\uparrow} \subseteq \mathcal{E} \quad \text{timeless}(P)}{\boxed{P}^N * (P * \text{⇒}_{\mathcal{E} \setminus \mathcal{N}^{\uparrow}} (P * Q)) \vdash \text{⇒}_{\mathcal{E}} Q}
 \end{array}$$

$$\begin{array}{c}
 \text{INV-OPEN-WP-TL} \\
 \frac{\text{atomic}(e) \quad \mathcal{N}^{\uparrow} \subseteq \mathcal{E} \quad \text{timeless}(P)}{\boxed{P}^N * \left(P * \text{wp}_{\mathcal{E} \setminus \mathcal{N}^{\uparrow}} e \{ v. P * \Phi(v) \} \right) \vdash \text{wp}_{\mathcal{E}} e \{ \Phi \}}
 \end{array}$$

The formal definition of timelessness can be found in [Krebbers et al. 2017a; Jung et al. 2018b]; for present purposes, it is sufficient to think of timeless propositions as those whose meaning is independent of step-indexing (*i.e.*, is the same at every step-index). The class of timeless propositions is closed under the connectives of first-order logic (truth, falsehood, conjunction, disjunction, implication, and universal and existential quantification), separation logic (separating conjunction, magic wand, and the points-to connective), the persistence modality, and Iris's connectives for ghost ownership (*e.g.*, ghost counters). Propositions that are *not* timeless include invariant and weakest precondition assertions, as well as the later and update modalities. In practice, it is common in Iris to establish invariants $\Box^N P$ where P is timeless and hence the above $\triangleright$ -free rules apply.

Instead of baking in the rules **INV-OPEN-UPD-TL** and **INV-OPEN-WP-TL** as primitives, Iris provides the primitive rule **$\Rightarrow$ -TIMELESS** for removing a later from a timeless proposition:

$$\frac{\begin{array}{c} \Rightarrow\text{-TIMELESS} \\ \text{timeless}(P) \end{array}}{\triangleright P \vdash_{\varepsilon} P}$$

The rules **INV-OPEN-UPD-TL** and **INV-OPEN-WP-TL** follow from the ordinary rules for opening invariants (**INV-OPEN-UPD** and **INV-OPEN-WP**, respectively) and **$\Rightarrow$ -TIMELESS**. Apart from brevity, the rule **$\Rightarrow$ -TIMELESS** also provides more flexibility. When dealing with invariants that contain both a timeless and a non-timeless part, one can remove the later from just the timeless part.

We are now ready to proceed with the proof of semantic typing of the `symbol` ADT.

PROOF OF THEOREM 7.1. To prove $\vdash \text{symbol} : \text{symbol_type}$, we unfold the definition and see that we must first prove a result about the expression interpretation:

$$\llbracket \text{symbol_type} \rrbracket_{\delta}^e(\text{symbol})$$

The proof of this property is as follows:

$$\begin{array}{c} \frac{\text{proof of gensym} \quad \text{proof of check}}{I_Y(\ell) \vdash \llbracket 1 \rightarrow \alpha \rrbracket_{\delta, \alpha \mapsto \Psi_Y}(\lambda(). \text{FAA}(\ell, 1)) \quad I_Y(\ell) \vdash \llbracket \alpha \rightarrow 2 \rrbracket_{\delta, \alpha \mapsto \Psi_Y}(\lambda s. \text{assert}(s < !\ell))} \text{unfold } \llbracket _ \times _ \rrbracket \\ \frac{I_Y(\ell) \vdash \llbracket (1 \rightarrow \alpha) \times (\alpha \rightarrow 2) \rrbracket_{\delta, \alpha \mapsto \Psi_Y}(\lambda(). \text{FAA}(\ell, 1), \lambda s. \text{assert}(s < !\ell))}{I_Y(\ell) \vdash \exists \Phi. \llbracket (1 \rightarrow \alpha) \times (\alpha \rightarrow 2) \rrbracket_{\delta, \alpha \mapsto \Phi}(\lambda(). \text{FAA}(\ell, 1), \lambda s. \text{assert}(s < !\ell))} \exists\text{-intro} \\ \frac{I_Y(\ell) \vdash \exists \alpha. \llbracket (1 \rightarrow \alpha) \times (\alpha \rightarrow 2) \rrbracket_{\delta}(\text{pack}\langle \text{Z}, \dots \rangle)}{I_Y(\ell) \vdash \llbracket \text{symbol_type} \rrbracket_{\delta}(\text{pack}\langle \text{Z}, \dots \rangle)} \text{unfold } \llbracket \exists \alpha. _ \rrbracket \\ \frac{I_Y(\ell) \vdash \llbracket \text{symbol_type} \rrbracket_{\delta}(\text{pack}\langle \text{Z}, \dots \rangle)}{I_Y(\ell) \vdash \text{wp } \text{pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}} \text{WP-VAL} \\ \frac{I_Y(\ell) \vdash \text{wp } \text{pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}}{\ell \mapsto 0 * \gamma \hookrightarrow 0 \vdash \text{wp } \text{pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}} \text{INV-ALLOC} \\ \frac{\ell \mapsto 0 \vdash \text{wp } \text{pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}}{\ell \mapsto 0 \vdash \text{wp } \text{pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}} \text{CNT-INIT} \\ \frac{\ell \mapsto 0 \vdash \text{wp } \text{pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}}{\ell \mapsto 0 \vdash \text{wp } \text{let } c = \ell \text{ in pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}} \text{WP-PURE, } \triangleright\text{-INTRO} \\ \frac{\ell \mapsto 0 \vdash \text{wp } \text{let } c = \ell \text{ in pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}}{\ell \mapsto 0 \vdash \text{wp } \ell \{ \text{v. wp } \text{let } c = v \text{ in pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \} \}} \text{WP-VAL} \\ \frac{\ell \mapsto 0 \vdash \text{wp } \ell \{ \text{v. wp } \text{let } c = v \text{ in pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \} \}}{\vdash \text{wp } \text{ref } 0 \{ \text{v. wp } \text{let } c = v \text{ in pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \} \}} \text{WP-ALLOC, } \triangleright\text{-INTRO} \\ \frac{\vdash \text{wp } \text{let } c = \text{ref } 0 \text{ in pack}\langle \text{Z}, \dots \rangle \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}}{\vdash \text{wp } \text{symbol} \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}} \text{WP-BIND} \\ \frac{\vdash \text{wp } \text{symbol} \{ \llbracket \text{symbol_type} \rrbracket_{\delta} \}}{\vdash \llbracket \text{symbol_type} \rrbracket_{\delta}^e(\text{symbol})} \text{unfold } \llbracket _ \rrbracket^e\end{array}$$

Reading this proof tree bottom-up, as usual, we see that it comprises the following steps:

- (1) Symbolically execute the expression `symbol`, thereby obtaining exclusive ownership of the private counter location ℓ with initial value 0.

- (2) Use rule **CNT-INIT** to allocate ghost counter γ with initial value 0 (implicitly here, we use the same pattern as in the proof of **CNT-INC'** to eliminate the update modality).
- (3) Transfer ownership of both into a new counter invariant $I_\gamma(\ell)$, defined as follows:

$$I_\gamma \triangleq \lambda \ell. \boxed{\exists n : \mathbb{N}. \ell \mapsto n * \gamma \hookrightarrow_{=} n}^{\mathcal{N}_{\text{sym}}}$$

- (4) The remaining goal is to prove that the value produced by executing symbol inhabits the value interpretation of `symbol_type` $\triangleq \exists \alpha. (1 \rightarrow \alpha) \times (\alpha \rightarrow 2)$. Correspondingly, we choose as our interpretation of α the semantic type $\Psi_\gamma(v)$, defined as follows:

$$\Psi_\gamma \triangleq \lambda v. \exists m : \mathbb{N}. (v = m) * \gamma \hookrightarrow_{>} m$$

- (5) The proof then splits into the following two subgoals, corresponding to the gensym and check operations of the ADT.

$$\begin{aligned} I_\gamma(\ell) &\vdash \llbracket 1 \rightarrow \alpha \rrbracket_{\delta, \alpha \mapsto \Psi_\gamma} (\lambda (). \text{FAA}(\ell, 1)) \\ I_\gamma(\ell) &\vdash \llbracket \alpha \rightarrow 2 \rrbracket_{\delta, \alpha \mapsto \Psi_\gamma} (\lambda s. \text{assert}(s < !\ell)) \end{aligned}$$

We now proceed to prove these subgoals.

The proof of gensym.

$$\begin{array}{c} \frac{\ell \mapsto (n+1) * \gamma \hookrightarrow_{=} (n+1) \vdash (\exists m. \ell \mapsto m * \gamma \hookrightarrow_{=} m)}{\ell \mapsto (n+1) * \gamma \hookrightarrow_{=} (n+1) * \gamma \hookrightarrow_{>} n \vdash (\exists m. \ell \mapsto m * \gamma \hookrightarrow_{=} m) * \Psi_\gamma(n)} \text{WP-VAL} \\ \frac{\ell \mapsto (n+1) * \gamma \hookrightarrow_{=} (n+1) * \gamma \hookrightarrow_{>} n \vdash \text{wp}_E n \{v. (\exists m. \ell \mapsto m * \gamma \hookrightarrow_{=} m) * \Psi_\gamma(v)\}}{\ell \mapsto (n+1) * \gamma \hookrightarrow_{=} n \vdash \text{wp}_E n \{v. (\exists m. \ell \mapsto m * \gamma \hookrightarrow_{=} m) * \Psi_\gamma(v)\}} \text{CNT-INC'} \\ \frac{\ell \mapsto (n+1) * \gamma \hookrightarrow_{=} n \vdash \text{wp}_E n \{v. (\exists m. \ell \mapsto m * \gamma \hookrightarrow_{=} m) * \Psi_\gamma(v)\}}{\ell \mapsto n * \gamma \hookrightarrow_{=} n \vdash \text{wp}_E \text{FAA}(\ell, 1) \{v. (\exists m. \ell \mapsto m * \gamma \hookrightarrow_{=} m) * \Psi_\gamma(v)\}} \text{WP-FAA, } \triangleright\text{-INTRO} \\ \frac{\ell \mapsto n * \gamma \hookrightarrow_{=} n \vdash \text{wp}_E \text{FAA}(\ell, 1) \{v. (\exists m. \ell \mapsto m * \gamma \hookrightarrow_{=} m) * \Psi_\gamma(v)\}}{I_\gamma(\ell) \vdash \text{wp } \text{FAA}(\ell, 1) \{\Psi_\gamma\}} \text{WP-PURE, } \triangleright\text{-INTRO} \\ \frac{I_\gamma(\ell) \vdash \text{wp } ((\lambda (). \text{FAA}(\ell, 1)) ()) \{\Psi_\gamma\}}{I_\gamma(\ell) \vdash \square \forall v. (v = ()) \multimap \text{wp } ((\lambda (). \text{FAA}(\ell, 1)) v) \{\Psi_\gamma\}} \square\text{-MONO, } \multimap\text{-INTRO, subst } v \\ \frac{I_\gamma(\ell) \vdash \square \forall v. (v = ()) \multimap \text{wp } ((\lambda (). \text{FAA}(\ell, 1)) v) \{\Psi_\gamma\}}{I_\gamma(\ell) \vdash \llbracket 1 \rightarrow \alpha \rrbracket_{\delta, \alpha \mapsto \Psi_\gamma} (\lambda (). \text{FAA}(\ell, 1))} \text{unfold } \llbracket 1 \rightarrow \alpha \rrbracket \end{array}$$

Here, we let $\mathcal{E} \triangleq \top \setminus \mathcal{N}_{\text{sym}}^\uparrow$. The beginning of this proof is like the proofs we have seen before: we unfold the expression interpretation, after which we have to prove a weakest precondition. To prove the weakest precondition for `FAA`($\ell, 1$), we need to get temporary ownership of the points-to connective $\ell \mapsto n$, which we do by opening the invariant $I_\gamma(\ell)$. Since the invariant $I_\gamma(\ell)$ is timeless, we can use the rule **INV-OPEN-WP-TL** to acquire ownership of $I_\gamma(\ell)$ without the later modality. After we have used the rule **WP-FAA**, we use the rule **CNT-INC'** to update the ghost counter $\gamma \hookrightarrow_{=} n$ to $\gamma \hookrightarrow_{=} (n+1)$, as needed to restore the invariant $I_\gamma(\ell)$. By using **CNT-INC'**, we also get $\gamma \hookrightarrow_{>} n$, which we need to establish $\Psi_\gamma(n)$.

The proof of check.

$$\begin{array}{c}
\frac{}{\gamma \hookrightarrow k * k < n \vdash \text{true} \in \{\text{true}, \text{false}\}} \text{WP-VAL, unfold } \llbracket 2 \rrbracket \\
\frac{}{\gamma \hookrightarrow k * k < n \vdash \text{wp } \text{true} \{ \llbracket 2 \rrbracket \}} \text{WP-PURE, } \triangleright\text{-INTRO} \\
\frac{\ell \vdash n * \gamma \hookrightarrow n \vdash \exists m. \ell \vdash m * \gamma \hookrightarrow m}{\gamma \hookrightarrow k * k < n \vdash \text{wp } \text{assert}(k < n) \{ \llbracket 2 \rrbracket \}} * \text{-MONO} \\
\frac{\ell \vdash n * \gamma \hookrightarrow n * \gamma \hookrightarrow k * k < n \vdash (\exists m. \ell \vdash m * \gamma \hookrightarrow m) * \text{wp} \dots}{\ell \vdash n * \gamma \hookrightarrow n * \gamma \hookrightarrow k \vdash (\exists m. \ell \vdash m * \gamma \hookrightarrow m) * \text{wp} \dots} \text{CNT-LT} \\
\frac{\ell \vdash n * \gamma \hookrightarrow n * \gamma \hookrightarrow k \vdash \text{wp}_{\mathcal{E}} n \{ w. (\exists m. \ell \vdash m * \gamma \hookrightarrow m) * \text{wp} \dots \}}{\ell \vdash n * \gamma \hookrightarrow n * \gamma \hookrightarrow k \vdash \text{wp}_{\mathcal{E}} n \{ w. (\exists m. \ell \vdash m * \gamma \hookrightarrow m) * \text{wp} \dots \}} \text{WP-VAL} \\
\frac{\ell \vdash n * \gamma \hookrightarrow n * \gamma \hookrightarrow k \vdash \text{wp}_{\mathcal{E}} n \{ w. (\exists m. \ell \vdash m * \gamma \hookrightarrow m) * \text{wp} \dots \}}{\ell \vdash n * \gamma \hookrightarrow n * \gamma \hookrightarrow k \vdash \text{wp}_{\mathcal{E}} !\ell \{ w. (\exists m. \ell \vdash m * \gamma \hookrightarrow m) * \text{wp} \dots \}} \text{WP-LOAD, } \triangleright\text{-INTRO} \\
\frac{}{I_Y(\ell) * \gamma \hookrightarrow k \vdash \text{wp } !\ell \{ w. \text{wp } \text{assert}(k < w) \{ \llbracket 2 \rrbracket \} \}} \text{INV-OPEN-WP-TL} \\
\frac{}{I_Y(\ell) * \gamma \hookrightarrow k \vdash \text{wp } \text{assert}(k < !\ell) \{ \llbracket 2 \rrbracket \}} \text{WP-BIND} \\
\frac{}{I_Y(\ell) * \gamma \hookrightarrow k \vdash \text{wp } (\lambda s. \text{assert}(s < !\ell)) k \{ \llbracket 2 \rrbracket \}} \text{WP-PURE, } \triangleright\text{-INTRO} \\
\frac{}{I_Y(\ell) \vdash \forall v. \Psi_Y(v) * \text{wp } (\lambda s. \text{assert}(s < !\ell)) v \{ \llbracket 2 \rrbracket \}} *, \forall \text{ intro, unpack } \Psi_Y \\
\frac{}{I_Y(\ell) \vdash \Box \forall v. \Psi_Y(v) * \text{wp } (\lambda s. \text{assert}(s < !\ell)) v \{ \llbracket 2 \rrbracket \}} \Box \text{-MONO} \\
\frac{}{I_Y(\ell) \vdash \llbracket \alpha \rightarrow 2 \rrbracket_{\delta, \alpha \mapsto \Psi_Y} (\lambda s. \text{assert}(s < !\ell))} \text{unfold } \llbracket \alpha \rightarrow 2 \rrbracket
\end{array}$$

Here we let $\mathcal{E} = \top \setminus \mathcal{N}_{\text{sym}}^{\uparrow}$. This proof is similar structurally to the proof of gensym—to prove the weakest precondition for $!\ell$, we need temporary access to the points-to connective $\ell \mapsto n$, which we do by opening the invariant $I_Y(\ell)$ using **INV-OPEN-WP-TL**. Apart from the points-to connective, the invariant $I_Y(\ell)$ also provides temporary access to $\gamma \hookrightarrow n$, which, using **CNT-LT**, allows us to deduce that $k < n$ and hence that **assert**($k < n$) is safe.

This concludes the proof of **Theorem 7.1**: symbol is semantically well-typed and thus safe to use in all well-typed contexts, despite its use of an unsafe feature. $\square$

Semantic ill-typedness of gremlin. With the proof of **Theorem 7.1** in hand, it is worth circling back around to the evil **gremlin** operator. At the beginning of this section, we argued informally that any attempt to prove that **gremlin** is semantically well-typed is doomed to fail, but now we can state and prove this result formally:

THEOREM 7.2 (**gremlin** IS SEMANTICALLY ILL-TYPED).

$$\not\models \text{gremlin} : 1$$

PROOF. Suppose the opposite, i.e., that $\models \text{gremlin} : 1$. Then, by **Theorem 7.1**, along with the compatibility rules of **MyLang**, it is straightforward to show that **evil_client** from §3.1 is semantically well-typed as well. By Adequacy (**Theorem 6.6**), that means **evil_client** is safe to execute, and yet we know there is an execution of **evil_client** that gets stuck due to a failed assertion, thus yielding a contradiction. $\square$

Concluding remarks. In this section, we have demonstrated how the logical relation for semantic soundness, encoded in Iris, can be used to (1) enforce that language features respect data abstraction and (2) reason about the safe encapsulation of unsafe features. As shown in the RustBelt project [Jung et al. 2018a, 2021; Jung 2020; Dang et al. 2020], this approach scales up to much more complicated uses of unsafe features. Of course, when dealing with more complicated uses of unsafe features, one needs more complicated invariants and “ghost theories”, but the basic structure of the proofs nevertheless follows the template we have shown here.

8 REPRESENTATION INDEPENDENCE

In the previous sections, we have shown how the semantic approach can be used to prove type safety. However, the semantic approach is by no means limited to type safety—it can be used for

the verification of many program properties, including but not limited to compiler correctness [Benton and Hur 2009], capability safety (both for object capabilities [Devriese et al. 2016] and for capability machines [Georges et al. 2021]), and non-interference [Frumin et al. 2021a; Gregersen et al. 2021], as well as contextual refinement and representation independence, the topic of this section. Many of these properties are not about the execution of a single program—*i.e.*, they are not *unary* properties—but are rather about the relation between two runs of possibly different programs—*i.e.*, they are *binary* properties. In this section, we discuss how Iris can be used to apply the semantic approach to prove a particularly important binary program property—*contextual refinement*—and in particular, the instance of that property known as *representation independence*.

A program e is said to *contextually refine* a program e' , written $\Gamma \models e \leq_{ctx} e' : A$, if for all program contexts C with hole of type A , if $C[e]$ has some observable behavior, then so does $C[e']$. Contextual refinement is a strong notion: if e contextually refines e' , then whenever e' appears as part of a well-typed program (*i.e.*, plugged into a well-typed context C), e' can be replaced by e without changing the observable behavior of the program.

A particularly interesting application of contextual refinement is in relational reasoning about ADTs. Specifically, suppose we have two different ADTs, M_1 and M_2 , which have the same interface, *i.e.*, the same type, but with different implementations, *e.g.*, different representations of the abstract type or of the internal state managed by it. If we can prove M_1 refines M_2 , we know that any program that is written against the common interface of these ADTs can be linked with M_1 instead of M_2 , and this should not have any observable effect, *despite* the fact that the ADTs are implemented differently. This property is known as *representation independence*. In practice, representation independence can be used to show that it is sound to replace a less efficient reference implementation M_2 by an optimized implementation M_1 . Correspondingly, in a contextual refinement $\Gamma \models e \leq_{ctx} e' : A$, we often refer to e as the *implementation* and e' as the *specification*.

Contextual refinement (see Definition 8.1) is defined by quantifying over *all* possible program contexts C . This makes direct proofs of contextual refinement difficult in practice—carrying out a proof by induction on the context C is known to be tedious and complicated, and infeasible for even relatively small programs. A common approach to easing proofs of contextual refinement is to define a judgment $\Gamma \models e \leq_{log} e' : A$ for *logical refinement*, using *binary logical relations*. The high-level structure of the binary logical relations method is similar to the high-level structure of the unary method for semantic typing we have already seen.

- **Soundness.** We prove a *soundness* theorem, which states that the logical relation is sound with respect to contextual refinement:

$$\Gamma \models e \leq_{log} e' : A \text{ implies } \Gamma \models e \leq_{ctx} e' : A.$$

This is similar to the adequacy theorem for semantic typing, which says that logically typed programs are safe.

- **Compatibility lemmas.** We show that the logical relation is compatible with syntactic typing. For instance, for function applications (the typing rule **T-APP**) we show:

$$\frac{\Gamma \models e_1 \leq_{log} e'_1 : A \rightarrow B \quad \Gamma \models e_2 \leq_{log} e'_2 : A}{\Gamma \models e_1 e_2 \leq_{log} e'_1 e'_2 : B}$$

These compatibility lemmas are similar to the semantic typing rules.

These results can then be combined with manual proofs of logical refinements of ADTs to modularly prove contextual refinements of larger programs. For example, suppose we have manually proven $\emptyset \models e_1 \leq_{log} e_2 : A$, and suppose C is a (closed) context (of type B) with a hole of type A . It is

$$\begin{aligned}
C ::= & [] \mid \text{rec } f(x) = C \mid C \ e \mid e \ C \mid \Lambda. C \mid C \langle \rangle \mid C \odot e \mid e \odot C \mid \\
& \text{if } C \text{ then } e \text{ else } e \mid \text{if } e \text{ then } C \text{ else } e \mid \text{if } e \text{ then } e \text{ else } C \mid \\
& \text{fold } C \mid \text{unfold } C \mid \text{ref } C \mid !C \mid C \leftarrow e \mid e \leftarrow C \mid \\
& \text{CAS}(C, e, e) \mid \text{CAS}(e, C, e) \mid \text{CAS}(e, e, C) \mid \text{FAA}(C, e) \mid \text{FAA}(e, C) \mid \text{fork } \{C\} \mid \dots
\end{aligned}$$

$$\begin{array}{c}
\text{C-REC} \\
\frac{C : (\Gamma'; B') \rightsquigarrow (\Gamma, x : A, f : A \rightarrow B; B)}{\text{rec } f(x) = C : (\Gamma'; B') \rightsquigarrow (\Gamma; B)}
\end{array}
\qquad
\begin{array}{c}
\text{C-TLAM} \\
\frac{C : (\Gamma'; B') \rightsquigarrow (\Gamma; A)}{\Lambda. C : (\Gamma'; B') \rightsquigarrow (\Gamma; \forall \alpha. A)}
\end{array}$$

$$\begin{array}{c}
\text{C-APP}_1 \\
\frac{C : (\Gamma'; B') \rightsquigarrow (\Gamma; A \rightarrow B) \quad \Gamma \vdash e_2 : A}{C \ e_2 : (\Gamma'; B') \rightsquigarrow (\Gamma; B)}
\end{array}
\qquad
\begin{array}{c}
\text{C-APP}_2 \\
\frac{\Gamma \vdash e_1 : A \rightarrow B \quad C : (\Gamma'; B') \rightsquigarrow (\Gamma; A)}{e_1 \ C : (\Gamma'; B') \rightsquigarrow (\Gamma; B)}
\end{array}
\qquad
\begin{array}{c}
\text{C-TAPP} \\
\frac{C : (\Gamma'; B') \rightsquigarrow (\Gamma; \forall \alpha. A)}{C \langle \rangle : (\Gamma'; B') \rightsquigarrow (\Gamma; A[B/\alpha])}
\end{array}$$

Fig. 9. An excerpt of the grammar of program contexts and their syntactic typing rules.

an easy corollary of the above properties that we can obtain $\emptyset \models C[e_1] \leq_{\text{ctx}} C[e_2] : B$. We will see a more general version of this corollary in §8.5.

To define the binary logical relation for logical refinement, we follow the same pattern as we have used for the unary logical relation for semantic typing—with the main difference that we generalize all notions to pairs of values. That is, we define binary interpretations on pairs of closed values $\llbracket _ \rrbracket$, and pairs of closed expressions $\llbracket _ \rrbracket^e$, and then use these binary interpretations to define our logical relation for open programs, $\Gamma \models e \leq_{\log} e' : A$. The binary value interpretations are straightforward generalizations of their unary counterparts. For example, values of base type (unit, Boolean, and integer) are related if they are equal, and values of function type are related if, given related inputs, they have related results. The crux of the difference between the unary and binary logical relations is in the expression relation $\llbracket \tau \rrbracket^e(e, e')$, which expresses that e refines e' . To formalize this refinement in Iris, we use both weakest preconditions and Iris's ghost theory.

We proceed in this section with a formal definition of contextual refinement (§8.1). We then show how to generalize the value and expression interpretations to the binary case (§8.2 and 8.3). Subsequently, we prove the compatibility lemmas for the binary logical relation (§8.4), and prove that the logical relation is sound with respect to contextual refinement (§8.5). Finally, we demonstrate reasoning about representation independence of ADTs by proving that a fine-grained implementation of a concurrent stack refines a coarse-grained version (§8.6).

8.1 Contextual Refinement

To formally define the contextual refinement judgment $\Gamma \models e \leq_{\text{ctx}} e' : A$, we first need to define the notion of program contexts. Figure 9 shows an excerpt of the grammar and the syntactic typing rules for program contexts. We write $C : (\Gamma; A) \rightsquigarrow (\Gamma'; A')$ to say that the context C is a program of type A' (closed under Γ') with a hole that can be filled with any program of type A (closed under Γ). The typing rules for well-typed contexts in Figure 9 imply that whenever $\Gamma \vdash e : A$ and $C : (\Gamma; A) \rightsquigarrow (\Gamma'; A')$ hold, so does $\Gamma' \vdash C[e] : A'$, capturing the intuitive idea that well-typed contexts are just well-typed programs with a hole.

Definition 8.1 (Contextual refinement). We say e contextually refines e' , written $\Gamma \models e \leq_{ctx} e' : A$, if both $\Gamma \vdash e : A$ and $\Gamma \vdash e' : A$, and furthermore we have:

$$\forall C : (\Gamma; A) \leadsto (\emptyset; 1). C[e] \downarrow \implies C[e'] \downarrow$$

Here, an expression e is said to terminate, written $e \downarrow$, if $(\emptyset, e) \rightarrow_{tp}^* (\sigma, v; \vec{e})$ for some final state σ , value v , and additional threads $\vec{e}$ —i.e., if a program has e in its main (and initially only) thread, then there is an execution of that program in which its main thread terminates with a value.

The above definition of contextual refinement extends the standard definition for sequential languages. We follow [Turon et al. \[2013a\]](#) by only taking the termination behavior of the main thread into account, i.e., once the main thread of the implementation has terminated, the main thread of the specification should terminate, too.

At first glance, this definition of contextual refinement might appear weaker than it actually is since it only talks about termination and not about the resulting values of the programs being related. However, given two programs e and e' such that $\emptyset \models e \leq_{ctx} e' : \mathbb{Z}$, we can additionally conclude the following: whenever the computation of e results in some number $n \in \mathbb{Z}$, then so does the computation of e' . To see this, simply take the well-typed context `if [ ] =  $n$  then () else  $\Omega$` , where Ω is a program that does not terminate. Similar arguments can be employed to show that contextual refinement implies stronger properties—e.g., that related memory locations in the heaps of the two programs always store indistinguishable values.

8.2 The Binary Value Interpretation

Similar to the unary logical relation for semantic typing, we define the binary logical relation for logical refinements in two stages:

- (1) We mutually define the value interpretation $\llbracket A \rrbracket_\delta : Val \times Val \rightarrow iProp_\square$ and the expression interpretation $\llbracket A \rrbracket_\delta^e : Expr \times Expr \rightarrow iProp$, both over *closed* values/expressions, where $\delta : Tvar \rightarrow_{fin} (Val \times Val \rightarrow iProp_\square)$ is the interpretation for free type variables.
- (2) We define the logical refinement relation on *open* terms $\Gamma \models e \leq_{log} e' : A$ by lifting the value and expression relations to open terms using a closing substitution.

The value and expression interpretations are shown in [Figure 10](#). We begin by presenting the former; the latter will be presented in [§8.3](#).

The binary value interpretation is a generalization of its unary counterpart. Values of base types (1 , 2 , and $\mathbb{Z}$) are related if they are equal values of the respective type. Values of the product type are related if both are pairs of values, related component-wise by the value interpretations of the corresponding types. Values of the sum type are related if they are both constructed using the same injection with underlying values related at the corresponding type. Values of the function type are related if applying them to values related at the domain type produces expressions related at the codomain type. Values of the universal type are related if their specializations are related, regardless of which (persistent) predicate we take as the value interpretation of the quantified type. Values of the existential type are related if they are both ADTs such that there is a (persistent) predicate for the value interpretation of the quantified type. Values of the recursive type are related if both are a `fold` and their arguments are related one step later. Finally, values of the reference type are related if they are locations that always store related values.

As we did in the unary logical relation, we define the binary logical relation on open expressions using a closing substitution. For that, we first define the interpretation of typing contexts:

$$\begin{aligned} \llbracket \emptyset \rrbracket_\delta^c(\emptyset, \emptyset) &\triangleq \text{True} \\ \llbracket \Gamma, x : A \rrbracket_\delta^c((\gamma, x \mapsto w), (\gamma', x \mapsto w')) &\triangleq \llbracket \Gamma \rrbracket_\delta^c(\gamma, \gamma') * \llbracket A \rrbracket_\delta(w, w') \end{aligned}$$

$$\begin{aligned}
\llbracket A \rrbracket_{\delta}^e &\triangleq \lambda (e, e'). \forall j, K. \text{SpecCtx} * j \models K[e'] \multimap \text{wp } e \{v. \exists v'. j \models K[v'] * \llbracket A \rrbracket_{\delta}(v, v')\} \\
\llbracket \alpha \rrbracket_{\delta} &\triangleq \delta(\alpha) \\
\llbracket 1 \rrbracket_{\delta} &\triangleq \lambda (v, v'). v = v' = () \\
\llbracket 2 \rrbracket_{\delta} &\triangleq \lambda (v, v'). v = v' \in \{\text{true}, \text{false}\} \\
\llbracket Z \rrbracket_{\delta} &\triangleq \lambda (v, v'). v = v' \in \mathbb{Z} \\
\llbracket A_1 \times A_2 \rrbracket_{\delta} &\triangleq \lambda (v, v'). \exists v_1, v_2, v'_1, v'_2. (v = (v_1, v_2)) * (v' = (v'_1, v'_2)) * \llbracket A_1 \rrbracket_{\delta}(v_1, v'_1) * \llbracket A_2 \rrbracket_{\delta}(v_2, v'_2) \\
\llbracket A_1 + A_2 \rrbracket_{\delta} &\triangleq \lambda (v, v'). \bigvee_{i \in \{1, 2\}} \exists w, w'. (v = \text{inj}_i w) * (v' = \text{inj}_i w') * \llbracket A_i \rrbracket_{\delta}(w, w') \\
\llbracket A \rightarrow B \rrbracket_{\delta} &\triangleq \lambda (v, v'). \Box (\forall w, w'. \llbracket A \rrbracket_{\delta}(w, w') \multimap \llbracket B \rrbracket_{\delta}^e(v w, v' w')) \\
\llbracket \forall \alpha. A \rrbracket_{\delta} &\triangleq \lambda (v, v'). \Box (\forall (\Psi : \text{Val} \times \text{Val} \rightarrow i\text{Prop}_{\Box}). \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}^e(v \langle \rangle, v' \langle \rangle)) \\
\llbracket \exists \alpha. A \rrbracket_{\delta} &\triangleq \lambda (v, v'). \exists (\Psi : \text{Val} \times \text{Val} \rightarrow i\text{Prop}_{\Box}). \\
&\quad \exists w, w'. (v = \text{pack} \langle w \rangle) * (v' = \text{pack} \langle w' \rangle) * \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w, w') \\
\llbracket \mu \alpha. A \rrbracket_{\delta} &\triangleq \mu (\Psi : \text{Val} \times \text{Val} \rightarrow i\text{Prop}_{\Box}). \\
&\quad \lambda (v, v'). \exists w, w'. (v = \text{fold } w) * (v' = \text{fold } w') * \llbracket A \rrbracket_{\delta, \alpha \mapsto \Psi}(w, w') \\
\llbracket \text{ref } A \rrbracket_{\delta} &\triangleq \lambda (v, v'). \exists \ell, \ell'. (v = \ell) * (v' = \ell') * \boxed{\exists w, w'. \ell \mapsto w * \ell' \mapsto_s w' * \llbracket A \rrbracket_{\delta}(w, w')}^{\mathcal{N}_{\ell, \ell'}}
\end{aligned}$$

Fig. 10. The expression interpretation $\llbracket _ \rrbracket^e$ and value interpretation $\llbracket _ \rrbracket$ for logical refinement in **MyLang**.

We then define the binary logical relation for logical refinement, $\Gamma \models e \leq_{\log} e' : A$, as follows:

$$\Gamma \models e \leq_{\log} e' : A \triangleq \Box (\forall \delta, \gamma, \gamma'. \llbracket \Gamma \rrbracket_{\delta}^c(\gamma, \gamma') \multimap \llbracket A \rrbracket_{\delta}^e(\gamma(e), \gamma'(e')))$$

8.3 The Binary Expression Interpretation

While the binary value interpretation $\llbracket A \rrbracket_{\delta}(v, v')$ is an immediate generalization of the unary version, the binary expression interpretation $\llbracket A \rrbracket_{\delta}^e(e, e')$ requires more work since Iris has no built-in support for relational reasoning.²⁵ No matter: instead of extending Iris with primitive support for relational reasoning (e.g., a relational version of weakest preconditions), we will show how to use the idea of *specification resources* (due to [Turon et al. \[2013a\]](#)) to *encode* relational reasoning as a derived concept on top of ordinary Iris weakest preconditions.

To explain the idea of specification resources, recall that the expression interpretation $\llbracket A \rrbracket_{\delta}^e(e, e')$ describes a refinement between the *implementation* e and *specification* e' . Intuitively, $\llbracket A \rrbracket_{\delta}^e(e, e')$ says that for each terminating execution of the implementation e , there is a related terminating execution for the specification e' such that $\llbracket A \rrbracket_{\delta}(v, v')$ where v and v' are the values of e and e' , respectively. Following the approach by [Turon et al. \[2013a\]](#), this intuitive idea can be expressed using a weakest precondition on the implementation e with a pre- and postcondition that express the existence of a related execution for the specification e' . To describe that the execution of the specification is related to the execution of the implementation, we use *specification resources*:

- The *specification thread connective* $j \models e$ describes unique ownership of a thread (with thread identifier j) in the specification program, currently executing expression e .

²⁵Actually, the (unary) weakest-precondition connective in Iris is not built-in either—it is encoded from more primitive constructs. An exploration of how that works is outside the scope of this paper. See [Jung et al. \[2018b\]](#) for details.

$$\begin{aligned}
& \text{SpecCtx} * (e_1 \rightarrow_{\text{pure}} e_2) * j \Rightarrow K[e_1] \vdash \Rightarrow_{\mathcal{E}} j \Rightarrow K[e_2] & (\text{SPEC-PURE}) \\
& \text{SpecCtx} * j \Rightarrow K[\text{ref } v] \vdash \Rightarrow_{\mathcal{E}} \exists \ell. j \Rightarrow K[\ell] * \ell \mapsto_s v & (\text{SPEC-ALLOC}) \\
& \text{SpecCtx} * j \Rightarrow K[! \ell] * \ell \mapsto_s v \vdash \Rightarrow_{\mathcal{E}} j \Rightarrow K[v] * \ell \mapsto_s v & (\text{SPEC-LOAD}) \\
& \text{SpecCtx} * j \Rightarrow K[\ell \leftarrow w] * \ell \mapsto_s v \vdash \Rightarrow_{\mathcal{E}} j \Rightarrow K[()] * \ell \mapsto_s w & (\text{SPEC-STORE}) \\
& \text{SpecCtx} * j \Rightarrow K[\text{CAS}(\ell, v, w)] * \ell \mapsto_s v \vdash \Rightarrow_{\mathcal{E}} j \Rightarrow K[\text{true}] * \ell \mapsto_s w & (\text{SPEC-CAS-SUC}) \\
& \text{SpecCtx} * (v \neq w) * j \Rightarrow K[\text{CAS}(\ell, w, u)] * \ell \mapsto_s v \vdash \Rightarrow_{\mathcal{E}} j \Rightarrow K[\text{false}] * \ell \mapsto_s v & (\text{SPEC-CAS-FAIL}) \\
& \text{SpecCtx} * j \Rightarrow K[\text{FAA}(\ell, m)] * \ell \mapsto_s n \vdash \Rightarrow_{\mathcal{E}} j \Rightarrow K[n] * \ell \mapsto_s (n + m) & (\text{SPEC-FAA}) \\
& \text{SpecCtx} * j \Rightarrow K[\text{fork } \{e\}] \vdash \Rightarrow_{\mathcal{E}} \exists j'. j \Rightarrow K[()] * j' \Rightarrow e & (\text{SPEC-FORK})
\end{aligned}$$

Fig. 11. Rules for specification resources (we implicitly assume $\mathcal{N}_{\text{spec}}^{\uparrow} \subseteq \mathcal{E}$).

- The *specification points-to connective* $\ell \mapsto_s v$ describes unique ownership of a memory location ℓ in the specification program, currently storing value v .

Like the ghost counter in §7, specification resources are an instance of ghost state—they do not represent ownership of physical resources subject to weakest preconditions, but are rather there strictly for logical purposes. This means that the specification points-to connective $\ell \mapsto_s v$ should not be confused with the ordinary points-to connective $\ell \mapsto v$. The ordinary points-to connective $\ell \mapsto v$ describes ownership of a physical location that appears in the implementation program, whereas the specification points-to connective $\ell \mapsto_s v$ describes ownership of a logical location that appears in the specification program. Like all forms of ghost state in Iris, specification resources can be manipulated using the update modality $\Rightarrow$. The rules, given in Figure 11, basically express that one can update $j \Rightarrow e$ into $j \Rightarrow e'$ provided that e steps to e' in the operational semantics. For heap-manipulating operations (allocation, load, store, CAS, and FAA), one additionally has to update ownership of the required specification points-to connectives $\ell \mapsto_s v$. The assertion SpecCtx is there for administrative reasons (which we will discuss below).

Putting all this together, the expression interpretation can be formalized as follows:

$$[[A]]_{\delta}^e \triangleq \lambda (e, e'). \forall j, K. \text{SpecCtx} * j \Rightarrow K[e'] \multimap \text{wp } e \{v. \exists v'. j \Rightarrow K[v'] * [[A]]_{\delta}(v, v')\}$$

This definition reads as follows: assuming a specification thread (with identifier j) contains the expression e' in evaluation position (at context K), then for any execution of the implementation e that results in a value v , there is a related execution from e' to a related value v' . That the related execution obeys the operational semantics is guaranteed by the fact that the only way to manipulate specification resources is through the rules in Figure 11, which exactly correspond to what steps are allowed in the operational semantics.

The definition of specification resources. We now explain how specification resources are defined in Iris. This is done in two steps:

- (1) Using Iris's flexible ghost state mechanism, we obtain the connectives $j \Rightarrow e$ and $\ell \mapsto_s v$.
- (2) Using Iris's invariant mechanism, we ensure that these connectives are only manipulated in ways that obey the operational semantics of **MyLang**.

In the first step, we instantiate Iris with a suitable ghost theory (the details of which are beyond the scope of this paper) in order to establish the soundness of a number of primitive proof rules concerning the new connectives $j \Rightarrow e$ and $\ell \mapsto_s v$, together with an ephemeral proposition $\text{SpecCnf}(\sigma, \vec{e})$ that keeps track of the entire heap σ and the entire thread-pool $\vec{e}$ of the specification. These primitive rules are shown in Figure 12. Using these rules, one can basically manipulate $j \Rightarrow e$

$$\begin{array}{ll}
\text{SpecCnf}(\sigma, \vec{e}) * j \Rightarrow e \vdash e_j = e & (\text{STHREAD-AGREE}) \\
\text{SpecCnf}(\sigma, \vec{e}) * (j = \text{length}(\vec{e})) \vdash \text{SpecCnf}(\sigma, \vec{e} \, e) * j \Rightarrow e & (\text{STHREAD-ALLOC}) \\
\text{SpecCnf}(\sigma, \vec{e}) * j \Rightarrow e \vdash \models_{\mathcal{E}} \text{SpecCnf}(\sigma, \vec{e}[j \mapsto e']) * j \Rightarrow e' & (\text{STHREAD-UPD}) \\
\text{SpecCnf}(\sigma, \vec{e}) * \ell \mapsto_s v \vdash \sigma(\ell) = v & (\text{SHEAP-AGREE}) \\
\text{SpecCnf}(\sigma, \vec{e}) * (\ell \notin \text{dom}(\sigma)) \vdash \models_{\mathcal{E}} \text{SpecCnf}(\sigma \uplus \{(\ell, v)\}, \vec{e}) * \ell \mapsto_s v & (\text{SHEAP-ALLOC}) \\
\text{SpecCnf}(\sigma \uplus \{(\ell, v)\}, \vec{e}) * \ell \mapsto_s v \vdash \models_{\mathcal{E}} \text{SpecCnf}(\sigma \uplus \{(\ell, v')\}, \vec{e}) * \ell \mapsto_s v' & (\text{SHEAP-UPD}) \\
\text{timeless}(\text{SpecCnf}(\sigma, \vec{e})) \text{ and } \text{timeless}(j \Rightarrow e) \text{ and } \text{timeless}(\ell \mapsto_s v) & (\text{SPEC-TIMELESS})
\end{array}$$

Fig. 12. Primitive rules for specification resources.

and $\ell \mapsto_s v$ as long as that is done in sync with $\text{SpecCnf}(\sigma, \vec{e})$. The fact that $\text{SpecCnf}(\sigma, \vec{e})$ is in sync is witnessed by the rules [STHREAD-AGREE](#) and [SHEAP-AGREE](#), which say that if we own the thread connective $j \Rightarrow e$ (respectively, the points-to connective $\ell \mapsto_s v$), then the thread j is in fact in the thread-pool $\vec{e}$ (respectively, the location ℓ is in the heap σ), where it is mapped to e (respectively, v). When allocating or updating a thread connective $j \Rightarrow e$ (using [STHREAD-ALLOC](#) and [STHREAD-UPD](#)) or a points-to connective $\ell \mapsto_s v$ (using [SHEAP-ALLOC](#) and [SHEAP-UPD](#)), one has to change $\text{SpecCnf}(\sigma, \vec{e})$ in a corresponding fashion.

In the second step, we use Iris’s invariant mechanism to ensure that the thread and points-to connectives are manipulated in a way that obeys the operational semantics. For that, we use the following invariant:

$$\begin{aligned}
\text{SpecInv}(\sigma_{\text{init}}, \vec{e}_{\text{init}}) &\triangleq \boxed{\exists \sigma, \vec{e}. \text{SpecCnf}(\sigma, \vec{e}) * \left((\sigma_{\text{init}}, \vec{e}_{\text{init}}) \rightarrow_{\text{tp}}^* (\sigma, \vec{e}) \right)}^{\mathcal{N}_{\text{spec}}} \\
\text{SpecCtx} &\triangleq \exists \sigma_{\text{init}}, \vec{e}_{\text{init}}. \text{SpecInv}(\sigma_{\text{init}}, \vec{e}_{\text{init}})
\end{aligned}$$

Given some initial heap σ_{init} and initial thread-pool $\vec{e}_{\text{init}}$, the invariant $\text{SpecInv}(\sigma_{\text{init}}, \vec{e}_{\text{init}})$ expresses that the heap and thread-pool in $\text{SpecCnf}(\sigma, \vec{e})$ can always be reached by taking a sequence of steps in the operational semantics from $(\sigma_{\text{init}}, \vec{e}_{\text{init}})$. Most of the time, with the exception of the soundness proof in §8.5, we do not need to know the initial state. Hence, we define SpecCtx , which existentially quantifies the initial state.

With the above definitions in hand, we can now prove all the rules in [Figure 11](#). These follow from Iris’s rules for invariants, together with the primitive rules for specification resources in [Figure 12](#). Since specification resources are timeless, we can use the rule [INV-OPEN-UPD-TL](#) to open the invariant SpecInv without a later modality. Note that, due to the use of an invariant to define specification resources in Iris, we need the premise SpecCtx and side-condition $\mathcal{N}_{\text{spec}}^{\uparrow} \subseteq \mathcal{E}$ in the rules in [Figure 11](#).

8.4 Compatibility Lemmas

Just as we proved semantic typing rules for the unary logical relation in §6, we now prove *relational* semantic typing rules for **MyLang**. In logical relations jargon, the relational semantic typing rules are often referred to as *compatibility lemmas* (see e.g., [Pitts \[2005\]](#)) since they show how the binary logical relation is “compatible” with the various constructs of **MyLang**. A selection of the compatibility lemmas for **MyLang** are presented in [Figure 13](#). Below we discuss the proofs of a few of them. The proofs of other compatibility lemmas follow in a similar fashion, just as many of the semantic typing rules in §6 followed a common essential structure.

$$\begin{array}{c}
\text{RS-VAR} \\
\frac{x : A \in \Gamma}{\Gamma \models x \leq_{\log} x : A} \\
\\
\text{RS-UNIT} \\
\Gamma \models () \leq_{\log} () : 1 \\
\\
\text{RS-BOOL} \\
\frac{b \in \{\text{true}, \text{false}\}}{\Gamma \models b \leq_{\log} b : 2} \\
\\
\text{RS-INT} \\
\frac{n \in \mathbb{Z}}{\Gamma \models n \leq_{\log} n : \mathbb{Z}} \\
\\
\text{RS-REC} \\
\frac{\Gamma, x : A, f : A \rightarrow B \models e \leq_{\log} e' : A}{\Gamma \models (\text{rec } f(x) = e) \leq_{\log} (\text{rec } f(x) = e') : A \rightarrow B} \\
\\
\text{RS-APP} \\
\frac{\Gamma \models e_1 \leq_{\log} e'_1 : A \rightarrow B \quad \Gamma \models e_2 \leq_{\log} e'_2 : A}{\Gamma \models e_1 e_2 \leq_{\log} e'_1 e'_2 : B}
\end{array}$$

Fig. 13. An excerpt of relational semantic typing rules (compatibility lemmas).

Before we go on to discuss some of the compatibility lemmas, we prove the monadic rules for the binary expression relation, which are a generalization of the unary versions in Lemma 6.2

LEMMA 8.2 (THE MONADIC RULES FOR THE EXPRESSION INTERPRETATION).

$$\llbracket A \rrbracket_{\delta}(v, v') \multimap \llbracket A \rrbracket_{\delta}^e(v, v') \quad (\text{BIN-VAL})$$

$$\llbracket A \rrbracket_{\delta}^e(e, e') * (\forall v, v'. \llbracket A \rrbracket_{\delta}(v, v') \multimap \llbracket B \rrbracket_{\delta}^e(K[v], K'[v'])) \multimap \llbracket B \rrbracket_{\delta}^e(K[e], K'[e']) \quad (\text{BIN-BIND})$$

PROOF. The rule **BIN-VAL** follows immediately from **WP-VAL**. The proof for **BIN-BIND** is:

$$\begin{array}{c}
\frac{S * j \models K''[K'[v']] * \left(\frac{\forall j', K'''. S * j' \models K'''[K'[v']] \multimap \text{wp } K[v] \{ \Phi_{j', K''', B} \}}{\text{wp } K[v] \{ \Phi_{j', K''', B} \}} \right) \vdash \text{wp } K[v] \{ \Phi_{j, K'', B} \}}{\frac{S * j \models K''[K'[v']] * \llbracket B \rrbracket_{\delta}^e(K[v], K'[v']) \vdash \text{wp } K[v] \{ \Phi_{j, K'', B} \}}{\frac{S * j \models K''[K'[v']] * \llbracket A \rrbracket_{\delta}(v, v') * F \vdash \text{wp } K[v] \{ \Phi_{j, K'', B} \}}{\frac{S * \Phi_{j, K'' \circ K', A}(v) * F \vdash \text{wp } K[v] \{ \Phi_{j, K'', B} \}}{S * F \vdash \forall v. \Phi_{j, K'' \circ K', A}(v) \multimap \text{wp } K[v] \{ \Phi_{j, K'', B} \}}}} \quad \begin{array}{l} \text{wp-elim, } \multimap\text{-ELIM} \\ \text{unfold } \llbracket B \rrbracket^e \\ \text{unfold } F, \forall\text{-elim, } \multimap\text{-ELIM} \\ \text{unfold } \Phi, \exists\text{-elim} \end{array} \\
\frac{S * \text{wp } e \{ \Phi_{j, K'' \circ K', A} \} * F \vdash \text{wp } e \{ v. \text{wp } K[v] \{ \Phi_{j, K'', B} \} \}}{S * \text{wp } e \{ \Phi_{j, K'' \circ K', A} \} * F \vdash \text{wp } e \{ v. \text{wp } K[v] \{ \Phi_{j, K'', B} \} \}} \quad \begin{array}{l} \forall\text{-intro, } \multimap\text{-INTRO} \\ \text{WP-WAND} \end{array} \\
\frac{S * S * j \models K''[K'[e']] * \left(\frac{\forall j', K'''. S * j' \models K'''[e'] \multimap \text{wp } e \{ \Phi_{j', K''', A} \}}{\text{wp } e \{ \Phi_{j', K''', A} \}} \right) * F \vdash \text{wp } e \{ v. \text{wp } K[v] \{ \Phi_{j, K'', B} \} \}}{S * S * j \models K''[K'[e']] * \left(\frac{\forall j', K'''. S * j' \models K'''[e'] \multimap \text{wp } e \{ \Phi_{j', K''', A} \}}{\text{wp } e \{ \Phi_{j', K''', A} \}} \right) * F \vdash \text{wp } e \{ v. \text{wp } K[v] \{ \Phi_{j, K'', B} \} \}} \quad \begin{array}{l} \text{duplicate } S \\ \text{unfold } \llbracket A \rrbracket^e \\ \text{WP-BIND} \end{array} \\
\frac{S * j \models K''[K'[e']] * \llbracket A \rrbracket_{\delta}^e(e, e') * F \vdash \text{wp } e \{ v. \text{wp } K[v] \{ \Phi_{j, K'', B} \} \}}{S * j \models K''[K'[e']] * \llbracket A \rrbracket_{\delta}^e(e, e') * F \vdash \text{wp } K[e] \{ \Phi_{j, K'', B} \}} \quad \begin{array}{l} \forall\text{-intro, } \multimap\text{-INTRO} \\ \text{unfold } \llbracket B \rrbracket^e \end{array} \\
\frac{\llbracket A \rrbracket_{\delta}^e(e, e') * F \vdash \forall j, K''. S * j \models K''[K'[e']] \multimap \text{wp } K[e] \{ \Phi_{j, K'', B} \}}{\llbracket A \rrbracket_{\delta}^e(e, e') * F \vdash \llbracket B \rrbracket_{\delta}^e(K[e], K'[e'])} \quad \text{unfold } \llbracket B \rrbracket^e
\end{array}$$

We let $F \triangleq \forall v, v'. \llbracket A \rrbracket_{\delta}(v, v') \multimap \llbracket B \rrbracket_{\delta}^e(K[v], K'[v'])$ and $\Phi_{j, K, B}(w) \triangleq \exists w'. j \models K[w'] * \llbracket B \rrbracket_{\delta}(w, w')$ and $S \triangleq \text{SpecCtx}$. The proof starts as expected: we unfold the definition of $\llbracket B \rrbracket^e$ and introduce everything into our context. After using **WP-BIND**, we need to obtain a weakest precondition for e from our context, requiring us to unfold the definition of $\llbracket A \rrbracket^e$ and instantiate it accordingly. We let $K''' \triangleq K'' \circ K'$, where $\circ$ is the composition of two evaluation contexts. This step crucially relies on $K''[K'[e]] = (K'' \circ K')[e]$. Moreover, since **SpecCtx** is persistent, we duplicate it, which is needed so we can use it in future steps. We now use **WP-WAND**, requiring us to prove that the postcondition

of the weakest precondition in our context $(\Phi_{j,K'' \circ K', A}(v))$ implies the postcondition of the weakest precondition in the goal $(\text{wp } K[v] \{ \Phi_{j,K'', B} \})$. We now eliminate the existential in Φ , allowing us to instantiate F . We conclude the proof by unfolding $\llbracket B \rrbracket^e$ and instantiating K''' this time with K'' . $\square$

PROOF OF RS-VAR. By unfolding the definition of the logical refinement relation, we have to prove $\llbracket \Gamma \rrbracket_\delta^c(\gamma, \gamma') \multimap \llbracket A \rrbracket_\delta^e(\gamma(x), \gamma'(x))$. From $\llbracket \Gamma \rrbracket_\delta^c(\gamma, \gamma')$ and $x : A \in \Gamma$, we obtain $\llbracket A \rrbracket_\delta(\gamma(x), \gamma'(x))$. The result thus follows from **BIN-VAL**. $\square$

PROOF OF RS-APP. In order to prove the compatibility lemma **RS-APP**, we prove the following auxiliary result for closed expressions, from which the semantic typing rule on open expressions easily follows:

$$\llbracket A \rightarrow B \rrbracket_\delta^e(e_1, e'_1) * \llbracket A \rrbracket_\delta^e(e_2, e'_2) \multimap \llbracket B \rrbracket_\delta^e((e_1 \ e_2), (e'_1 \ e'_2))$$

Below there follows a proof tree for the auxiliary result:

$$\frac{\frac{\frac{\frac{\frac{\frac{\llbracket A \rrbracket_\delta(v_2, v'_2) \multimap \llbracket B \rrbracket_\delta^e((v_1 \ v_2), (v'_1 \ v'_2)) * \llbracket A \rrbracket_\delta(v_2, v'_2) \vdash \llbracket B \rrbracket_\delta^e((v_1 \ v_2), (v'_1 \ v'_2))}{\text{*-ELIM}}}{\llbracket \forall w, w'. \llbracket A \rrbracket_\delta(w, w') \multimap \llbracket B \rrbracket_\delta^e((v_1 \ w), (v'_1 \ w')) * \llbracket A \rrbracket_\delta(v_2, v'_2) \vdash \llbracket B \rrbracket_\delta^e((v_1 \ v_2), (v'_1 \ v'_2))}{\text{*-ELIM, } \forall\text{-elim}}}{\llbracket A \rightarrow B \rrbracket_\delta(v_1, v'_1) * \llbracket A \rrbracket_\delta(v_2, v'_2) \vdash \llbracket B \rrbracket_\delta^e((v_1 \ v_2), (v'_1 \ v'_2))}{\text{unfold } \llbracket A \rightarrow B \rrbracket^e}}}{\llbracket A \rightarrow B \rrbracket_\delta(v_1, v'_1) * \llbracket A \rrbracket_\delta^e(e_2, e'_2) \vdash \llbracket B \rrbracket_\delta^e((v_1 \ e_2), (v'_1 \ e'_2))}{\text{BIN-BIND}}}{\llbracket A \rightarrow B \rrbracket_\delta^e(e_1, e'_1) * \llbracket A \rrbracket_\delta^e(e_2, e'_2) \vdash \llbracket B \rrbracket_\delta^e((e_1 \ e_2), (e'_1 \ e'_2))}{\text{BIN-BIND}}$$

This proof tree is similar to the one for the semantic typing rule **S-APP** in the unary case. Reading this proof tree bottom-up, we start by using **BIN-BIND** twice (following the scheme we described in §6.4), first for expressions e_1 and e'_1 in contexts $K \triangleq [] \ e_2$ and $K' \triangleq [] \ e'_2$, and then for expressions e_2 and e'_2 in contexts $K \triangleq v_1 \ []$ and $K' \triangleq v'_1 \ []$. The last step again demonstrates why “logical relations” are called “logical”—we use Iris’s modus ponens rule $(Q \multimap R) * Q \vdash R$ (***-ELIM**) to eliminate the magic wand that appears in the interpretation of the function type $A \rightarrow B$.

The actual compatibility lemma **RS-APP** follows from this auxiliary lemma in the same way that **S-APP** followed from its corresponding auxiliary result in §6.5. $\square$

REMARK 8.3. *The remaining compatibility lemmas are proved in a similar way. One minor point of note is that, for the compatibility lemma for the CAS operation, the proof relies on the EqType(A) side-condition (in the typing rule T-CAS) in order to conclude that the CAS operation on the specification side returns the same Boolean value as the implementation side.*

8.5 The Fundamental Theorem and Soundness

THEOREM 8.4 (FUNDAMENTAL THEOREM OF BINARY LOGICAL RELATIONS). *Well-typed terms are related to themselves, i.e., if $\Gamma \vdash e : A$ then $\Gamma \models e \leq_{\log} e : A$.*

PROOF. By straightforward induction on the typing derivation $\Gamma \vdash e : A$. For each case in the induction proof, we use the corresponding compatibility lemma. $\square$

LEMMA 8.5 (CONGRUENCY OF BINARY LOGICAL RELATIONS). *The binary logical relation is closed under well-typed program contexts, i.e., if $\Gamma \models e \leq_{\log} e' : A$ and $C : (\Gamma; A) \rightsquigarrow (\Gamma'; A')$ then $\Gamma' \models C[e] \leq_{\log} C[e'] : A'$.*

PROOF. By straightforward induction on the derivation of $C : (\Gamma; A) \rightsquigarrow (\Gamma'; A')$. In each case, we apply the appropriate compatibility lemma and, when necessary, use the fundamental theorem (Theorem 8.4) to show that well-typed expressions are related to themselves. $\square$

LEMMA 8.6 (ADEQUACY OF BINARY LOGICAL RELATIONS). *The binary logical relation preserves termination, i.e., if $\emptyset \models e \leq_{\log} e' : A$, then $e \downarrow$ implies $e' \downarrow$.*

PROOF. To prove this lemma, we make use of a strengthened version of adequacy of weakest preconditions (Theorem 6.1). For brevity's sake, we do not consider the most general adequacy statement [Jung et al. 2018b, Theorem 7], but rather consider a version that is instantiated with the ghost theory for specification resources. That is, given a first-order proposition ϕ and a proof of

$$\text{SpecCnf}(\emptyset, \emptyset) \vdash \text{wp } e \{ \phi \},$$

if $e \downarrow$, then ϕ holds at the meta-level.

To prove our lemma, we pick $\phi \triangleq e' \downarrow$, which means we are done once we have proved $\text{SpecCnf}(\emptyset, \emptyset) \vdash \text{wp } e \{ e' \downarrow \}$. We prove this result in the following steps:

$$\text{SpecCnf}(\emptyset, \emptyset) \vdash \models \text{SpecInv}(\emptyset, e') * 1 \models e' \quad (\text{STEP1})$$

$$\text{SpecInv}(\emptyset, e') * 1 \models e' \vdash \text{wp } e \{ v. \exists v'. 1 \models v' * \llbracket A \rrbracket_{\delta}(v, v') \} \quad (\text{STEP2})$$

$$\text{SpecInv}(\emptyset, e') * 1 \models v' \vdash \models e' \downarrow \quad (\text{STEP3})$$

In STEP1, we allocate the invariant $\text{SpecInv}(\emptyset, e')$. We do this by first creating a specification thread resource $1 \models e'$ for the main thread (using `STHREAD-ALLOC`), and then transferring ownership of $\text{SpecCnf}(\emptyset, e')$ into the invariant $\text{SpecInv}(\emptyset, e')$ (using `INV-ALLOC`).

In STEP2, we make use of our premise $\emptyset \models e \leq_{\log} e' : A$. Since we are considering closed programs, by definition of the binary logical relation this premise is equivalent to $\llbracket A \rrbracket_{\delta}^e(e, e')$. By unfolding the expression interpretation we then get:

$$\forall j, K. \text{SpecInv}(\emptyset, e') * j \models K[e'] \multimap \text{wp } e \{ v. \exists v'. j \models K[v'] * \llbracket A \rrbracket_{\delta}(v, v') \}$$

Our result is obtained by specializing this statement by picking $K = []$ and $j = 1$:

$$\text{SpecInv}(\emptyset, e') * 1 \models e' \multimap \text{wp } e \{ v. \exists v'. 1 \models v' * \llbracket A \rrbracket_{\delta}(v, v') \}$$

In STEP3, we open the invariant $\text{SpecInv}(\emptyset, e')$ (using `INV-OPEN-UPD-TL`) to obtain that we have $(\emptyset, e') \rightarrow_{\text{tp}}^* (\sigma, \vec{e})$ for some heap σ and threadpool $\vec{e}$ with $\text{SpecCnf}(\sigma, \vec{e})$. Since we have $1 \models v'$, we obtain that the main thread of $\vec{e}$ is the value v' (by `STHREAD-AGREE`), which gives $e' \downarrow$ as desired.

The proof tree below shows how these steps lead to the final result:

$$\frac{\frac{\frac{\text{SpecInv}(\emptyset, e') * 1 \models v' * \llbracket A \rrbracket_{\delta}(v, v') \vdash \models e' \downarrow}{\text{SpecInv}(\emptyset, e') * \text{wp } e \{ v. \exists v'. 1 \models v' * \llbracket A \rrbracket_{\delta}(v, v') \} \vdash \text{wp } e \{ e' \downarrow \}}{\text{SpecInv}(\emptyset, e') * 1 \models e' \vdash \text{wp } e \{ e' \downarrow \}} \text{STEP2}}{\text{SpecCnf}(\emptyset, \emptyset) \vdash \text{wp } e \{ e' \downarrow \}} \text{STEP1, } \models\text{-WP, } \models\text{-MONO}$$

STEP3
WP-WAND, $\exists$ -elim, WP- $\models$

Note that, once established, we can keep the invariant $\text{SpecInv}(\emptyset, e')$ around throughout the proof since it is persistent (and thus duplicable). $\square$

THEOREM 8.7 (SOUNDNESS OF BINARY LOGICAL RELATIONS). *The binary logical relation is sound with respect to contextual refinement, i.e., if $\Gamma \vdash e : A$ and $\Gamma \vdash e' : A$ and $\Gamma \models e \leq_{\log} e' : A$, then $\Gamma \models e \leq_{\text{ctx}} e' : A$.*

PROOF. By definition of contextual refinement, in order to prove $\Gamma \models e \leq_{\text{ctx}} e' : A$, we are given a well-typed program context $C : (\Gamma; A) \rightsquigarrow (\emptyset; 1)$ and have to show that $C[e] \downarrow$ implies $C[e'] \downarrow$. By the assumption and congruency (Lemma 8.5), we have $\emptyset \models C[e] \leq_{\log} C[e'] : 1$, which gives that $C[e] \downarrow$ implies $C[e'] \downarrow$ by adequacy (Lemma 8.6). $\square$

<pre> stack_{fg} $\triangleq$ Λ. let s = ref (ref None) in let push_{fg} = rec f(x) = let z = !s in if CAS(s, z, ref (Some(x, fold z))) then () else f x in let pop_{fg} = rec f() = let z = !s in match !z with None $\Rightarrow$ None Some(hd, tl) $\Rightarrow$ if CAS(s, z, unfold tl) then Some(hd) else f () end in (push_{fg}, pop_{fg}) </pre>	<pre> stack_{cg} $\triangleq$ Λ. let s = ref None in let l = newlock () in let push_{cg} = λ x. acquire l; s $\leftarrow$ Some(x, fold !s); release l in let pop_{cg} = λ (). acquire l; let mx = match !z with None $\Rightarrow$ None Some(hd, tl) $\Rightarrow$ s $\leftarrow$ unfold tl; Some(hd) end in release l; mx in (push_{cg}, pop_{cg}) </pre>
---	--

Fig. 14. The source code of a fine-grained (left) and coarse-grained (right) concurrent stack.

$$\begin{aligned}
j \models K[\text{newlock } ()] &\vdash \models_{\mathcal{E}} \exists l. \text{isLock}_s(l, \text{false}) * j \models K[l] && (\text{SPEC-NEWLOCK}) \\
\text{isLock}_s(l, \text{false}) * j &\models K[\text{acquire } l] \vdash \models_{\mathcal{E}} \text{isLock}_s(l, \text{true}) * j \models K[()] && (\text{SPEC-ACQUIRE}) \\
\text{isLock}_s(l, \text{true}) * j &\models K[\text{release } l] \vdash \models_{\mathcal{E}} \text{isLock}_s(l, \text{false}) * j \models K[()] && (\text{SPEC-RELEASE}) \\
&\text{timeless}(\text{isLock}_s(l, b)) && (\text{SPEC-LOCK-TIMELESS})
\end{aligned}$$

Fig. 15. Rules for lock specification resources.

8.6 Representation Independence Proofs

The soundness theorem of our binary logical relation (Theorem 8.7) allows us to prove contextual refinement by means of logical refinement. As our logical relation is formalized on top of Iris, we have the entire power and support of Iris at our disposal when proving contextual refinement by means of logical refinement. In this subsection, we demonstrate this power by proving representation independence of two implementations of a concurrent stack displayed in Figure 14. Specifically, we prove the following refinement:

$$\models \text{stack}_{\text{fg}} \leq_{\text{ctx}} \text{stack}_{\text{cg}} : \forall \alpha. (\alpha \rightarrow \mathbf{1}) \times (\mathbf{1} \rightarrow \text{option } \alpha)$$

The stack ADT provides functions push and pop for pushing and popping elements on and off a stack. Since the stack is dynamically sized, the function push will always succeed. The function pop may fail by returning None if the stack is empty. Here, the option type is defined in the usual way using sums, *i.e.*, $\text{option } A \triangleq \mathbf{1} + A$, and the constructors are defined as $\text{None} \triangleq \text{inj}_1 ()$ and $\text{Some } v \triangleq \text{inj}_2 v$.

The two implementations of the stack ADT differ in the *granularity* of their concurrency: the first is *fine-grained*—it enforces atomicity at the level of individual instructions—whereas the second is *coarse-grained*—it enforces atomicity via a critical section, protected by a lock.

Concretely, the fine-grained implementation stack_{fg} employs a private reference s that points to the head of a linked list defined using the following recursive type:

$$\text{linkedList } A \triangleq \mu\alpha. \text{ref}(\text{option}(A \times \alpha))$$

The fine-grained implementation uses a technique known as optimistic concurrency to implement push and pop. It first reads the head reference s to the linked list. It then tries to update the head reference using the (atomic) compare-and-set instruction (**CAS**) to make sure it has not been modified in the meantime. If the **CAS** fails, another thread must have augmented the reference to the list; in that case, the operation starts over, trying to perform the push or pop again.

The coarse-grained implementation stack_{cg} stores the entire stack as a private reference s to a list defined using the following recursive type:

$$\text{list } A \triangleq \mu\alpha. \text{option}(A \times \alpha)$$

The coarse-grained implementation uses a lock l to make sure the push and pop instructions are carried out atomically. The operation `newlock` creates a new lock, which is initially in the unlocked state. The lock can be moved into the locked state using the `acquire` operation, which will block if another thread holds the lock. The lock can be put back into the unlocked state using the `release` operation. Release does not block, because only if one acquired the lock, it should release the lock.

Although **MyLang** does not have locks as primitives, they can easily be implemented using, say, a spin lock or a ticket lock. For the purpose of this paper, it does not matter what lock implementation is used—all that matters is that the implementation enjoys the logical rules in [Figure 15](#). (See [\[Frumin et al. 2021b, §5\]](#) for a proof that a spin lock implementation satisfies these logical rules.) Note that since locks are used for the specification side of the refinement, we have only included the rules in terms of specification resources, and not those in terms of weakest preconditions. Furthermore, note that the lock rules are similar to the rules for the heap operations we have seen in [Figure 11](#), but they involve a new predicate $\text{isLock}_s(l, b)$, where $b = \text{false}$ means that the lock l is in the unlocked state, and $b = \text{true}$ means it is in the locked state.

The proof of the stack refinement. In order to prove contextual refinement of the lock implementations, it suffices, by the soundness of the binary logical relations ([Theorem 8.7](#)), to prove the following, corresponding logical refinement:

$$\models \text{stack}_{\text{fg}} \leq_{\text{log}} \text{stack}_{\text{cg}} : \forall\alpha. (\alpha \rightarrow 1) \times (1 \rightarrow \text{option } \alpha)$$

The proof follows the same structure as the proof of safe encapsulation of the symbol ADT in [§7](#). We unfold the definition of the logical refinement judgment, and prove Iris weakest preconditions for the functions push and pop. The crux of the proof involves defining an invariant that relates the internal data structures used in both implementations. Since the stack ADT is polymorphic, this invariant should make sure that the values of both stacks are related by the binary value interpretation corresponding to the type α , which we call $\Phi : \text{Val} \times \text{Val} \rightarrow \text{iProp}$. To relate the internal data structures of both implementations we define the following Iris propositions:

$$\bar{\Phi} \triangleq \mu\bar{\Phi} : (\text{Val} \times \text{Val} \rightarrow \text{iProp}). \lambda (\ell, v').$$

$$(\ell \mapsto \text{None} * v' = \text{None}) \vee$$

$$(\exists w, w', \ell_{\text{tl}}, v'_{\text{tl}}. \ell \mapsto \text{Some}(w, \text{fold } \ell_{\text{tl}}) * v' = \text{Some}(w', \text{fold } v'_{\text{tl}}) * \Phi(w, w') * \triangleright \bar{\Phi}(\ell'_{\text{tl}}, v'_{\text{tl}}))$$

$$I \triangleq \boxed{\exists \ell, v'. s \mapsto \ell * s' \mapsto_s v' * \bar{\Phi}(\ell, v') * \text{isLock}_s(l', \text{false})}^{N_{\text{stk}}}$$

In prose, the invariant I states that:

- the private reference s of stack_{fg} always points to the head of a linked list ℓ ;
- the private reference s' of stack_{cg} always points to a functional list v' ;
- the values stored in the linked list at ℓ and function list v' are related by Φ ; and
- the lock of stack_{cg} is in unlocked state.

The relation $\bar{\Phi}(\ell, v')$ ensures that the linked list pointed by ℓ and the functional list v' have the same length and that the values in these lists are related by the value interpretation Φ . Note that, as far as the behavior of the ADTs is considered, the lock is never *observed* in the locked state. This is because all the `acquire` statements in the operations of the coarse-grained stack are followed by a `release`, and these operations are all executed atomically.

With the invariant in hand, the proof of the logical refinement is straightforward but lengthy. After we have allocated the resources of the stacks (the lists and locks), we create the Iris invariant I . Subsequently, we prove the refinements of `push` and `pop` using `LÖB` induction, where in each step we open and close the invariant I . A detailed and formal proof can be found in the accompanying Coq formalization (see §10 for the URL to the online repository with the Coq formalization).

Summary. We have shown how our logical relation can be used to show representation independence of two implementations of a concurrent stack. We note that since the coarse-grained stack uses locks to sequentialize accesses to the stack, one can understand our logical relations proof of contextual refinement as an alternative to the *linearizability* proof method for concurrent objects [Herlihy and Wing 1990]. A possible advantage of the logical relations approach shown here is that it also applies, *mutatis mutandis*, when the data structure in question involves higher-order functions; for example, one can easily extend the logical relations proof above to the case where the fine-grained and coarse-grained concurrent stacks include a higher-order iterator method. In contrast, linearizability has so far mostly been developed for first-order languages, although recent work by Murawski and Tzevelekos [2019] has extended linearizability to higher-order programming languages.

9 ADDITIONAL RELATED WORK

The “logical approach to type soundness” that we have advanced in this paper descends from multiple lines of prior work on the semantics of higher-order, imperative, and concurrent programming languages. In §4.2, we discussed earlier work on semantic type soundness and step-indexed models. In this section, we briefly survey some other key influences on our work, as well as closely related approaches.

Relational logics for richly typed languages. The most direct ancestor of our approach is the line of work on *relational logics*—logics for reasoning abstractly about relational program properties such as parametricity and representation independence in richly-typed languages. The primogenitor of this line is the seminal paper of Plotkin and Abadi [1993], who showed how to define logical relations for a polymorphic programming language in a second-order relational logic. Their approach was extended by Dreyer et al. [2011], who integrated the “later” ($\multimap$) modality into a Plotkin-Abadi style logic in order to define step-indexed logical relations for a language with polymorphic and recursive types. Dreyer et al.’s motivation was precisely to avoid the tedious step-indexed arithmetic that they had previously experienced when working directly with step-indexed models. Their method was extended further by Dreyer et al. [2010] to handle general (higher-typed) mutable references, using a second-order *relational separation logic*, inspired by [Yang 2007], with a notion of invariants (called “islands”, based on prior work of Ahmed et al. [2009]). Turon et al.

[2013a] later extended the logical approach to a language with concurrency (using a pre-Iris second-order *concurrent* separation logic called CaReSL), and Krogh-Jespersen et al. [2017] extended it (using Iris) to account for a region-based type-and-effect system.

Though directly continuing this line of work, our “logical approach to type soundness” goes beyond the aforementioned pre-Iris work on relational logics in several ways. First of all, we have shown that the approach of building logical relations in separation logic is useful not only in proving relational properties like representation independence but also in formalizing *semantic type soundness* results (a unary property) for richly-typed languages. We have also demonstrated the utility of the resulting semantic soundness theorems for verifying safe encapsulation of unsafe features. Second, our approach leverages a more modern concurrent separation logic, namely Iris, which offers a richer, more evolved, and still evolving logical language in which to encode logical relations models of types. Iris is also a language-agnostic framework, which can be instantiated for a wide variety of different languages so long as they can be formalized with a relatively standard style of operational semantics [Jung et al. 2018b, §7.3]. Last but not least, thanks to the Iris Proof Mode [Krebbers et al. 2017b], our approach has the key benefit of making it feasible to fairly rapidly develop both semantic soundness and representation independence proofs that are fully machine-checked in Coq. In contrast, none of the previous work on relational logics was mechanized in a proof assistant.

These benefits of the logical approach have already been demonstrated in a significant and growing set of papers that employ it for both unary and relational reasoning (see §10 for citations)—but, as we noted in the Introduction, these papers are not always the easiest on-ramps for newcomers wanting to learn the essential methodology. We hope that the present paper helps to fill the pedagogical gap by presenting the logical approach from first principles and in the setting of a simpler programming language.

“Semantic soundness” for compilers. A second key ancestor of our work is that of Benton and collaborators [Benton 2006; Benton and Zarfaty 2007; Benton and Tabareau 2009; Benton and Hur 2009]. Over a series of papers, they propounded the idea that “compiler correctness” ought to account for preservation of source-level relational reasoning down to the assembly level—and that, to realize this idea, one should build semantic models of high-level types as relational specifications on low-level code. Though superficially distinct from the kinds of results we have established in this paper, Benton et al.’s compiler correctness theorems were referred to as “semantic soundness” theorems, and indeed there is a strong kinship between theirs and ours. In particular, like our logical-relations models, theirs (1) were formulated using logical abstractions such as the later modality and separating conjunction to support higher-level reasoning, (2) were specifically used to verify that low-level, potentially unsafe code is well-behaved according to the semantic contracts of high-level types, and (3) were formalized in Coq.

Aside from the specific intended application, a key difference between our work and Benton et al.’s is that, although Benton et al.’s models make use of higher-level logical abstractions, the proofs about them are still conducted directly in the model of propositions (rather than in a *bona fide* logic like Iris) and without the rich tactical support for separation logic that the Iris Proof Mode provides, thus rendering them considerably lower-level than ours. This is quite understandable, given that Benton et al.’s work was conducted before a number of major advances in (higher-order concurrent) separation logic, which ultimately culminated in the development of Iris. In a sense, the work was ahead of its time. Nevertheless, it was a source of inspiration for us in how it used logical relations, along with techniques from step-indexing and separation logic, to carve out “well-behavedness” conditions on potentially unsafe code. Also inspiring to us were Benton et al.’s

observations concerning the limitations of syntactic type soundness, which were rather iconoclastic given the predominance of the syntactic approach at the time.

Simulation-based approaches. Around the same time as step-indexed models were being developed in the mid-2000s, there emerged an impressive series of papers on *(bi-)simulation* techniques for relational reasoning in higher-order, imperative, and concurrent languages—e.g., [Koutavas and Wand 2006; Sumii and Pierce 2007; Støvring and Lassen 2007; Lassen and Levy 2007; Sumii 2009]. We still lack a precise understanding of the relationship between logical relations and simulation-based methods—there are tradeoffs in terms of convenience of proof effort—but suffice it to say that, in terms of expressive power, both classes of techniques have proven capable (in principle) of supporting sophisticated relational reasoning in a range of different programming languages. For more details, we refer the reader to Hur et al. [2012].

There are, however, a number of differences between the simulation-based methods and the methods we have presented in this paper. First, since the simulation-based methods rely on coinduction (rather than step-indexing) to achieve reasoning about “circular” features (e.g., recursive types, higher-order state), they do not require anything comparable to the tedious reasoning about step-index arithmetic that we remarked upon in §4.3. However, as with direct reasoning in step-indexed models, the simulation-based methods *do* involve explicit, and sometimes low-level, reasoning about the global machine state and invariants on it (see points 2 and 3 in §4.3). Second, nearly all of the work on simulations has been focused exclusively on proving relational properties, not semantic type soundness. One exception is Sumii [2010], who explores the applicability of simulation-based methods to proving safe encapsulation of potentially unsafe deallocation operations in a sequential, *untyped* λ -calculus with higher-order state, but his approach has not seemingly been applied to a wider variety of languages. Lastly, with the exception of the line of work on “parametric simulations” [Hur et al. 2012; Neis et al. 2015], none of the simulation-based approaches have, to our knowledge, been formally mechanized in a proof assistant.

Syntactic type abstraction. In §3, we discussed the strengths and limitations of the syntactic approach to type soundness, noting in particular that syntactic type soundness has nothing to say about whether a language properly supports data abstraction. There is, however, at least one paper proposing a syntactic approach to reasoning about data abstraction properties of ADTs, namely that of Grossman et al. [2000]. Their approach involves introducing a syntactic notion of “principals” into their operational semantics in order to track which values arise from the implementation of an ADT vs. from its client. Although they develop their method in the presence of a wide range of features (including higher-order state), they only use it to prove a limited class of results: one stating that a value of some abstract type must have arisen from calling a specific operation provided by an ADT, and another stating that changing the integer representing a value of some abstract type will not affect client code. The proofs of those results eschew the complexities of semantic models but supplant them with arguments concerning highly intricate syntactic invariants (see for instance the proof of Theorem 3.13 in their paper). Moreover, it is not at all clear how one could apply their method to verify either the symbol ADT example from §7 or the representation independence example from §8.6, since those examples involve more complex invariants on state.

Hybrid syntactic/semantic approaches to type soundness. There have been a few approaches to type soundness that incorporate hybrids of syntactic and semantic/logical elements.

Tofte [1990] proposed an early approach to type inference (and type soundness) for an ML-like language combining polymorphism and mutable references. Tofte’s approach, which pre-dates the “progress and preservation” approach [Wright and Felleisen 1994; Harper 2016], defines a semantic typing relation, albeit using coinduction to handle circularities in the construction (rather than

step-indexing as we do), and using a syntactic heap typing to track the types of memory locations (rather than a semantic/logical model of heap typing as we formalize with Iris invariants). Tofte’s approach was adopted by several others in the early 1990s [Leroy and Weis 1991; Talpin and Jouvelot 1994], when a number of researchers were investigating how best to make ML-style type inference play well with mutable references. However, it fell out of favor after much simpler methods were proposed: the “value restriction” [Wright 1995] for safely combining ML-style polymorphism with references (which was ultimately integrated into both Standard ML and OCaml), and progress and preservation for proving type soundness. Moreover, Tofte’s approach was limited to predicative polymorphism (see the discussion in [Tofte 1990, p. 21]), which neither syntactic nor logical type soundness are; and due to its reliance on syntactic techniques, it suffers from the same limitations of syntactic type soundness that we laid out in §3.

Mezzo [Balabonski et al. 2016] is a recently proposed programming language with (broadly speaking) similar goals to Rust: supporting low-level, fine-grained control over the representation and access of data in memory, while preserving type and memory safety. Also like Rust, Mezzo employs a substructural type system in order to track aliasing and ownership of memory. The soundness proof for Mezzo is clearly syntactic, following the tradition of progress-and-preservation proofs. However, in order to support a more modular presentation of the soundness proof, Balabonski et al. formalize a notion of “resource” using something called a “monotonic resource algebra” (which is closely related to, but not the same as, the “cameras” and “resource algebras” used in Iris—see the discussion in [Jung et al. 2018b, §9.3]). These resources definitely give their soundness proof a separation-logic flavor; yet it remains syntactic, and thus does not offer a way to reason about data abstraction or safe encapsulation of unsafe features.

10 CONCLUSION

In this paper, we have demonstrated that semantic type soundness is a more useful result than syntactic type soundness, and we have shown how to prove it at a higher level of abstraction than in prior work by exploiting the features of a modern separation logic, Iris. We conclude in this section by illustrating that our logical approach to type soundness is eminently scalable and practical. We do so by describing a general recipe for extending the logical approach to different languages, type systems, and program properties. Additionally, we offer a brief discussion of recent work that has employed the logical approach in practice, and provide references to papers and online tutorials that show how to mechanize logical type soundness proofs in Coq.

Applying the logical approach to other languages. We have studied the logical approach here in the context of the simple programming language **MyLang**, which exhibits a fairly pedestrian set of features. To apply the logical approach to a different language or a different type system, one roughly has to follow the following three steps:

1. *Instantiate Iris with the language.* The most common way to instantiate Iris with a programming language of choice is to start by defining its syntax and operational semantics. As explained in §6.1, Iris’s program logic (whose primary component is the connective for weakest preconditions) is parametric in the types of expressions, values, and states, and in a reduction relation.

Instead of defining the syntax and operational semantics of a language from scratch, Iris’s default language *HeapLang* could be reused. *HeapLang* is similar to **MyLang**, but comes with a number of additional features, such as arrays. Some programming languages are well-suited to be defined as a shallow embedding on top of *HeapLang* (see e.g., Hinrichsen et al. [2021] for a language with message-passing primitives *à la* session types).

2. *Define reasoning principles for the language.* After having defined the syntax and operational semantics of the programming language, one needs to define logical connectives for ownership of physical resources (like the points-to connective $\ell \mapsto v$) and program reasoning (like the weakest precondition connective $\text{wp } e \{ \Phi \}$).

If the programming language fits into Iris’s common format for small-step operational semantics, then Iris’s generic program logic and definition of weakest preconditions can be used. Additionally, if the language has a simple memory model like **MyLang**’s, Iris’s generic library for the points-to connective $\ell \mapsto v$ can be used. If the language has a more sophisticated memory model (e.g., a block/offset-based memory model like CompCert’s [Leroy and Blazy 2008], as used in RustBelt [Jung et al. 2018a, 2021; Jung 2020] and RefinedC [Sammler et al. 2021]), or if it has additional physical resources (e.g., a program counter and registers as in a low-level capability machine [Georges et al. 2021]), then custom connectives for ownership of physical resources need to be defined, either by combining existing libraries or rolling one’s own library using ghost state. Such a library can then be plugged into Iris’s generic weakest precondition connective.

Alternatively, instead of reusing Iris’s generic program logic and weakest preconditions, one can define a *custom* program logic with the help of Iris’s base logic (which includes $*$, $\multimap$, $\Box$, $\Rightarrow$, and $\triangleright$). This is useful, for example, to obtain a weakest precondition in big-step rather than small-step style (see e.g., Timany et al. [2018]; Gregersen et al. [2021]), to establish program properties that are out of scope of Iris’s generic program logic (e.g., security, see Frumin et al. [2021a]; Gregersen et al. [2021]), or to consider programming languages with non-local control (e.g., effect handlers, see de Vilhena and Pottier [2023]). One can go even further and only rely on Iris’s basic constructs for step-indexing (and the Iris Proof Mode in Coq) to develop a custom model of separation logic. For instance, Jacobs et al. [2024] develop a linear (instead of affine) variant of Iris for deadlock-freedom of message-passing programs, which they use to give a semantic model of linear session types.

There is a middle ground as well: rather than building a custom program logic from scratch, one can instead define new notions of weakest precondition *on top of* Iris’s generic weakest preconditions. iGPS [Kaiser et al. 2017] and RustBelt Relaxed [Dang et al. 2020] employ this methodology to build a weakest precondition for relaxed memory concurrency. Timany and Birkedal [2019] use this methodology to develop a notion of *context-local* weakest precondition to reason about concurrent programs with first-class continuations, and Timany et al. [2024] develop a notion of *well-bracketed* weakest precondition for exploiting the absence of continuations. Interestingly, Timany et al. [2024] then use their logic to build unary and binary logical relations which are almost verbatim the same as the ones we have developed in this article, the only difference being that they model their expression relation using a well-bracketed weakest precondition instead of the generic Iris one.

3. *Define ghost theories for modeling the type system.* Once reasoning principles for the programming language have been set up, one needs to define suitable ghost theories for modeling the features of the type system. In this paper, we have seen three instances of ghost theories: impredicative invariants for modeling higher-order references (§6.9), ghost counters for monotonically increasing counters as used in the symbol ADT example (§7), and specification resources for proving representation independence (§8.3). Other examples of ghost theories are RustBelt’s lifetime logic for modeling Rust’s lifetime and borrowing mechanism [Jung et al. 2018a, 2021; Jung 2020; Dang et al. 2020], and Actris’s dependent separation protocol mechanism for modeling session types [Hinrichsen et al. 2020, 2022, 2021].

Ghost theories are defined using Iris’s mechanism of *higher-order ghost state* [Jung et al. 2016]. This mechanism is based on PCMs (partial commutative monoids)—as found in many separation logics—but generalizes them with a step-indexed notion of equality. Iris’s generalized PCMs are

called *step-indexed resource algebras*, or *cameras* for short. While we have presented impredicative invariants as a primitive of Iris, they are in fact defined in terms of Iris’s higher-order ghost state mechanism. Some work (e.g., [Giarrusso et al. \[2020\]](#)) uses higher-order ghost state directly instead of invariants.

Recent work that employs the logical approach. In recent years, the logical approach to type soundness has been deployed in a variety of applications. One of the earliest examples we are aware of is the work of [Gordon et al. \[2012\]](#), who studied a type system for safe parallelism based on reference immutability and uniqueness. They proved (pen-and-paper) semantic soundness of their type system by modeling its typing judgment in an early version of the Views separation-logic framework [[Dinsdale-Young et al. 2013](#)], a key precursor of Iris.

Since the development of the Iris Proof Mode for Coq [[Krebbers et al. 2017b](#)], Iris has become the *lingua franca* for machine-checked proofs of logical type soundness. For instance, Iris has been used for a machine-checked proof of type soundness of a significant subset of the Rust programming language [[Jung et al. 2018a, 2021](#); [Jung 2020](#); [Dang et al. 2020](#)], an extension of Scala’s core type system DOT [[Giarrusso et al. 2020](#)], session types [[Hinrichsen et al. 2021](#); [Jacobs et al. 2024](#)], and refinement types for the C programming language [[Sammler et al. 2021](#)]. Aside from type soundness, it has also been used to prove robust safety [[Swasey et al. 2017](#); [Sammler et al. 2020](#); [Georges et al. 2021](#); [Rao et al. 2023](#)], various forms of representation independence and program refinement [[Krogh-Jespersen et al. 2017](#); [Tassarotti et al. 2017](#); [Timany et al. 2018](#); [Timany and Birkedal 2019](#); [Frumin et al. 2018, 2021b](#); [Jacobs et al. 2021](#); [Timany et al. 2024](#)], and various security properties [[Frumin et al. 2021a](#); [Gregersen et al. 2021](#); [Georges et al. 2021](#)]. It has even recently been used to build logical relations on top of a *denotational*, rather than operational, semantics [[Frumin et al. 2024](#)]. Instead of discussing these applications in detail, we highlight some interesting differences between our presentation of the logical approach and theirs.

In this paper we have considered a “standard” (unrestricted) type system, in which variables can be used any number of times and types are not used to enforce a discipline of resource ownership. However, since Iris is a separation logic, it is in fact designed to reason about ownership and is thus ideally suited to applying the logical approach to substructural type systems. For example, [Jung et al. \[2018a, 2021\]](#), [Jung \[2020\]](#), and [Dang et al. \[2020\]](#) have used Iris to model Rust’s type system, which employs a strict ownership discipline to guarantee memory safety and data-race freedom in the context of low-level programming paradigms. [Tassarotti et al. \[2017\]](#), [Hinrichsen et al. \[2021\]](#), and [Jacobs et al. \[2024\]](#) have used Iris to model session types, which employ ownership to enforce protocol compliance in message-passing communication.

The crucial difference between unrestricted and substructural type systems is whether to make the value interpretation $\llbracket A \rrbracket$ persistent or not. In an unrestricted type system (such as the type system for **MyLang** in this paper), the value interpretation $\llbracket A \rrbracket$ is persistent for *any* type A . In a substructural type system (such as Rust or session types), the value interpretation $\llbracket A \rrbracket$ is *not* persistent for types that denote ownership (such as mutable references and channels), while it is persistent for “copyable” types (such as integers, Booleans and shared references).

With regard to refinement proofs, let us point out two differences from some of the above-cited papers. First, in this paper we had to unfold the logical refinement judgment and carry out a proof in terms of its definition in Iris (see §8.6 for how this is done for the example of concurrent stacks). In contrast, [Frumin et al. \[2018, 2021b\]](#) present a *logic* for doing such refinement proofs at a higher level of abstraction and show how it simplifies reasoning about refinements.

Second, in this paper, we have used Iris to prove termination-*insensitive* program refinement, in which any non-terminating program is a contextual refinement of any other program. In contrast, [Tassarotti et al. \[2017\]](#) develop a version of Iris to prove termination-*preserving* refinements. While

their approach establishes a stronger version of refinement, it also has some limitations—it can only be used in the context of languages with countable non-determinism (instead of concurrency), and for refinement proofs that involve finite stuttering. Recent work by Spies et al. [2021] overcomes these limitations in a non-concurrent setting by employing a *transfinite* version of step-indexing, where steps are modeled using ordinals instead of natural numbers. An interesting direction for future work is to scale this approach to the concurrent setting.

Coq material. To develop Iris proofs in practice, nearly all Iris users make use of the Iris Proof Mode [Krebbers et al. 2017b, 2018], which provides tactics and other infrastructure for carrying out separation logic proofs in Coq. While a presentation of the Iris Proof Mode is beyond the scope of this paper, we provide some references to relevant online materials.

First of all, there is the Coq development accompanying the present paper:

- <https://gitlab.mpi-sws.org/iris/examples> (directory `logrel/F_mu_ref_conc`)
This development contains a mechanization of the semantic type soundness proof (§5–§6) and representation independence proof (§8) for **MyLang**, as well as the proofs that `symbol` is semantically well-typed and `gremlin` is not (§7).

In addition, here are links to several other tutorial materials with different emphases:

- <https://gitlab.mpi-sws.org/iris/tutorial-popl20/>
This tutorial (which was presented at the POPL’20 conference) demonstrates how to prove logical type soundness in Iris/Coq. The structure of this tutorial largely follows §5–§7, but uses Iris’s default language **HeapLang** (and the infrastructure that Iris provides for HeapLang) instead of the language **MyLang**. This tutorial comes with exercises.
- <https://gitlab.mpi-sws.org/iris/tutorial-popl21/>
This tutorial (which was presented at the POPL’21 conference, and is based on an earlier version at POPL’18) does not specifically target logical type soundness, but provides an introduction to reasoning about concurrent programs in Iris. This tutorial comes with exercises.
- <https://github.com/tchajed/iris-simp-lang/>
This tutorial demonstrates how to instantiate Iris with a custom language, based on a stripped-down version of HeapLang.
- <https://gitlab.mpi-sws.org/FP/semantics-course/>
This development accompanies lecture notes from a course on Semantics taught periodically at Saarland University (<https://plv.mpi-sws.org/semantics-course/lecturenotes.pdf>). The first half of the notes cover semantic type soundness and logical relations for a language similar to **MyLang**, formalized directly in the traditional, explicitly step-indexed style; the second half of the notes offer a tutorial on Iris, with the ulterior motive of showing how to re-implement the semantic models of the first half in the logical style. Both the lecture notes and the accompanying Coq development come with many exercises.

ACKNOWLEDGMENTS

We wish to thank our many collaborators on the Iris project for helpful discussions, the anonymous reviewers for their extremely thorough and constructive reviews, and Ron Garcia for his feedback on an earlier draft of this paper. This research was supported in part by the Dutch Research Council (NWO), project 016.Veni.192.259, in part by a Villum Investigator grant (no. 25804), Center for Basic Research in Program Verification (CPV), from the VILLUM Foundation, in part by a European Research Council (ERC) Consolidator Grant for the project “RustBelt”, funded under the European Union’s Horizon 2020 Framework Programme (grant agreement no. 683289), and in part by generous gifts from Google. Amin Timany was a postdoctoral fellow of the Flemish research fund (FWO) during parts of this project.

REFERENCES

- Martín Abadi and Gordon D. Plotkin. 1990. A PER model of polymorphism and recursive types. In *LICS*. 355–365. <https://doi.org/10.1109/LICS.1990.113761>
- Amal Ahmed. 2004. *Semantics of types for mutable state*. Ph.D. Dissertation. Princeton University.
- Amal Ahmed, Andrew W. Appel, Christopher D. Richards, Kedar N. Swadi, Gang Tan, and Daniel C. Wang. 2010. Semantic foundations for typed assembly languages. *TOPLAS* 32, 3 (2010), 7:1–7:67. <https://doi.org/10.1145/1709093.1709094>
- Amal Ahmed, Andrew W. Appel, and Roberto Virga. 2002. A stratified semantics of general references. In *LICS*. 75–86. <https://doi.org/10.1109/LICS.2002.1029818>
- Amal Ahmed, Derek Dreyer, and Andreas Rossberg. 2009. State-dependent representation independence. In *POPL*. 340–353. <https://doi.org/10.1145/1480881.1480925> Technical appendix: <http://www.ccs.neu.edu/home/amal/papers/sdri/main-long.pdf>.
- Amal J. Ahmed. 2006. Step-indexed syntactic logical relations for recursive and quantified types. In *ESOP (LNCS, Vol. 3924)*. 69–83. https://doi.org/10.1007/11693024_6 Technical appendix: <http://www.ccs.neu.edu/home/amal/papers/lr-recquant-techrpt.pdf>.
- Stuart Allen. 1987. *A non-type-theoretic semantics for type-theoretic language*. Ph.D. Dissertation. Cornell University.
- Andrew W. Appel. 2001. Foundational proof-carrying code. In *LICS*. 247–256. <https://doi.org/10.1109/LICS.2001.932501>
- Andrew W. Appel and Amy P. Felty. 2000. A Semantic Model of Types and Machine Instructions for Proof-Carrying Code. In *POPL*. 243–253. <https://doi.org/10.1145/325694.325727>
- Andrew W. Appel and David A. McAllester. 2001. An indexed model of recursive types for foundational proof-carrying code. *TOPLAS* 23, 5 (2001), 657–683. <https://doi.org/10.1145/504709.504712>
- Andrew W. Appel, Paul-André Melliès, Christopher D. Richards, and Jérôme Vouillon. 2007. A very modal model of a modern, major, general type system. In *POPL*. 109–122. <https://doi.org/10.1145/1190216.1190235>
- E. S. Bainbridge, Peter J. Freyd, Andre Scedrov, and Philip J. Scott. 1990. Functorial polymorphism. *TCS* 70, 1 (1990), 35–64. [https://doi.org/10.1016/0304-3975\(90\)90151-7](https://doi.org/10.1016/0304-3975(90)90151-7)
- Thibaut Balabonski, François Pottier, and Jonathan Protzenko. 2016. The design and formalization of Mezzo, a permission-based programming language. *TOPLAS* 38, 4 (2016), 14:1–14:94. <https://doi.org/10.1145/2837022>
- Hendrik Pieter Barendregt. 1985. *The lambda calculus – its syntax and semantics*. Studies in logic and the foundations of mathematics, Vol. 103. North-Holland.
- Nick Benton. 2006. Abstracting allocation: The new new thing. In *CSL*. 182–196. https://doi.org/10.1007/11874683_12
- Nick Benton and Chung-Kil Hur. 2009. Biorthogonality, step-indexing and compiler correctness. In *ICFP*. 97–108. <https://doi.org/10.1145/1596550.1596567>
- Nick Benton and Nicolas Tabareau. 2009. Compiling functional types to relational specifications for low level imperative code. In *TLDI*. 3–14. <https://doi.org/10.1145/1481861.1481864>
- Nick Benton and Uri Zarfaty. 2007. Formalizing and verifying semantic type soundness of a simple compiler. In *PPDP*. 1–12. <https://doi.org/10.1145/1273920.1273922>
- Lars Birkedal, Ales Bizjak, and Jan Schwinghammer. 2013. Step-indexed relational reasoning for countable nondeterminism. *LMCS* 9, 4 (2013). [https://doi.org/10.2168/LMCS-9\(4:4\)2013](https://doi.org/10.2168/LMCS-9(4:4)2013)
- Lars Birkedal and Robert Harper. 1999. Relational interpretations of recursive types in an operational setting. *Information and Computation* 155, 1-2 (1999), 3–63. <https://doi.org/10.1006/inco.1999.2828>
- Lars Birkedal, Bernhard Reus, Jan Schwinghammer, Kristian Støvring, Jacob Thamsborg, and Hongseok Yang. 2011. Step-indexed Kripke models over recursive worlds. In *POPL*. 119–132. <https://doi.org/10.1145/1926385.1926401>
- Lars Birkedal, Filip Sieczkowski, and Jacob Thamsborg. 2012. A concurrent logical relation. In *CSL (LIPIcs, Vol. 16)*. 107–121. <https://doi.org/10.4230/LIPIcs.CSL.2012.107>
- Lars Birkedal, Kristian Støvring, and Jacob Thamsborg. 2010a. The category-theoretic solution of recursive metric-space equations. *TCS* 411, 47 (2010), 4102–4122. <https://doi.org/10.1016/j.tcs.2010.07.010>
- Lars Birkedal, Kristian Støvring, and Jacob Thamsborg. 2010b. Realisability semantics of parametric polymorphism, general references and recursive types. *MSCS* 20, 4 (2010), 655–703. <https://doi.org/10.1017/S0960129510000162>
- Richard Bornat, Cristiano Calcagno, Peter W. O’Hearn, and Matthew J. Parkinson. 2005. Permission accounting in separation logic. In *POPL*. 259–270. <https://doi.org/10.1145/1040305.1040327>
- Stephen Brookes. 2007. A semantics for concurrent separation logic. *TCS* 375, 1-3 (2007), 227–270. <https://doi.org/10.1016/j.tcs.2006.12.034>
- Kim B. Bruce and John C. Mitchell. 1992. PER models of subtyping, recursive types and higher-order polymorphism. In *POPL*. 316–327. <https://doi.org/10.1145/143165.143230>
- Robert L. Constable, Stuart F. Allen, Mark Bromley, Rance Cleaveland, J. F. Cremer, Robert Harper, Douglas J. Howe, Todd B. Knoblock, N. P. Mendler, Prakash Panangaden, James T. Sasaki, and Scott F. Smith. 1986. *Implementing mathematics with the Nuprl proof development system*. Prentice Hall. <http://dl.acm.org/citation.cfm?id=10510>

- Karl Crary and Robert Harper. 2007. Syntactic logical relations for polymorphic and recursive types. *ENTCS* 172 (2007), 259–299. <https://doi.org/10.1016/j.entcs.2007.02.010>
- Pedro da Rocha Pinto, Thomas Dinsdale-Young, and Philippa Gardner. 2014. TaDA: A Logic for Time and Data Abstraction. In *ECOOP (LNCS, Vol. 8586)*, 207–231. https://doi.org/10.1007/978-3-662-44202-9_9
- Hoang-Hai Dang, Jacques-Henri Jourdan, Jan-Oliver Kaiser, and Derek Dreyer. 2020. RustBelt meets relaxed memory. *PACMPL* 4, POPL (2020), 34:1–34:29. <https://doi.org/10.1145/3371102>
- N.G de Bruijn. 1972. Lambda calculus notation with nameless dummies, a tool for automatic formula manipulation, with application to the Church-Rosser theorem. *Indagationes Mathematicae (Proceedings)* 75, 5 (1972), 381–392. [https://doi.org/10.1016/1385-7258\(72\)90034-0](https://doi.org/10.1016/1385-7258(72)90034-0)
- Paulo Emílio de Vilhena and François Pottier. 2023. A Type System for Effect Handlers and Dynamic Labels. In *ESOP (LNCS, Vol. 13990)*, 225–252. https://doi.org/10.1007/978-3-031-30044-8_9
- Dominique Devriese, Lars Birkedal, and Frank Piessens. 2016. Reasoning about object capabilities with logical relations and effect parametricity. In *EuroS&P*, 147–162. <https://doi.org/10.1109/EuroSP.2016.22>
- Edsger W. Dijkstra. 1975. Guarded Commands, Nondeterminacy and Formal Derivation of Programs. *CACM* 18, 8 (1975), 453–457. <https://doi.org/10.1145/360933.360975>
- Thomas Dinsdale-Young, Lars Birkedal, Philippa Gardner, Matthew J. Parkinson, and Hongseok Yang. 2013. Views: compositional reasoning for concurrent programs. In *POPL*, 287–300. <https://doi.org/10.1145/2429069.2429104>
- Thomas Dinsdale-Young, Mike Dodds, Philippa Gardner, Matthew J. Parkinson, and Viktor Vafeiadis. 2010. Concurrent abstract predicates. In *ECOOP (LNCS, Vol. 6183)*, 504–528. https://doi.org/10.1007/978-3-642-14107-2_24
- Derek Dreyer. 2018. Milner Award Lecture: The type soundness theorem that you really want to prove (and now you can). Keynote talk at POPL 2018, https://www.youtube.com/watch?v=8Xyk_dGcAwk&ab_channel=POPL2018.
- Derek Dreyer, Amal Ahmed, and Lars Birkedal. 2011. Logical step-indexed logical relations. *LMCS* 7, 2 (2011). [https://doi.org/10.2168/LMCS-7\(2:16\)2011](https://doi.org/10.2168/LMCS-7(2:16)2011)
- Derek Dreyer, Georg Neis, and Lars Birkedal. 2012. The impact of higher-order state and control effects on local relational reasoning. *JFP* 22, 4-5 (2012), 477–528. <https://doi.org/10.1017/S095679681200024X> Technical appendix: <http://www.mpi-sws.org/tr/2012-001.pdf>.
- Derek Dreyer, Georg Neis, Andreas Rossberg, and Lars Birkedal. 2010. A relational modal logic for higher-order stateful ADTs. In *POPL*, 185–198. <https://doi.org/10.1145/1706299.1706323>
- Derek Dreyer, Simon Spies, Lennard Gäher, Ralf Jung, Jan-Oliver Kaiser, Hoang-Hai Dang, David Swasey, Jan Menz, Niklas Mück, and Benjamin Peters. 2022. Semantics of type systems (Lecture notes). Available at <https://plv.mpi-sws.org/semantics-course/>.
- Matthias Felleisen and Robert Hieb. 1992. The revised report on the syntactic theories of sequential control and state. *TCS* 103, 2 (1992), 235–271. [https://doi.org/10.1016/0304-3975\(92\)90014-7](https://doi.org/10.1016/0304-3975(92)90014-7)
- Dan Frumin, Robbert Krebbers, and Lars Birkedal. 2018. ReLoC: A mechanised relational logic for fine-grained concurrency. In *LICS*, 442–451. <https://doi.org/10.1145/3209108.3209174>
- Dan Frumin, Robbert Krebbers, and Lars Birkedal. 2021a. Compositional non-interference for fine-grained concurrent programs. In *S&P*, 1416–1433. <https://doi.org/10.1109/SP40001.2021.00003>
- Dan Frumin, Robbert Krebbers, and Lars Birkedal. 2021b. ReLoC Reloaded: A mechanized relational logic for fine-grained concurrency and logical atomicity. *LMCS* 17, 3 (2021). [https://doi.org/10.46298/lmcs-17\(3:9\)2021](https://doi.org/10.46298/lmcs-17(3:9)2021)
- Dan Frumin, Amin Timany, and Lars Birkedal. 2024. Modular Denotational Semantics for Effects with Guarded Interaction Trees. *Proc. ACM Program. Lang.* 8, POPL (2024), 332–361. <https://doi.org/10.1145/3632854>
- Ming Fu, Yong Li, Xinyu Feng, Zhong Shao, and Yu Zhang. 2010. Reasoning about optimistic concurrency using a program logic for history. In *CONCUR (LNCS, Vol. 6269)*, 388–402. https://doi.org/10.1007/978-3-642-15375-4_27
- Aïna Linn Georges, Armaël Guéneau, Thomas Van Strydonck, Amin Timany, Alix Trieu, Sander Huyghebaert, Dominique Devriese, and Lars Birkedal. 2021. Efficient and provable local capability revocation using uninitialized capabilities. *PACMPL* 5, POPL (2021), 1–30. <https://doi.org/10.1145/3434287>
- Paolo G. Giarrusso, Léo Stefanescu, Amin Timany, Lars Birkedal, and Robbert Krebbers. 2020. Scala step-by-step: Soundness for DOT with step-indexed logical relations in Iris. *PACMPL* 4, ICFP (2020), 114:1–114:29. <https://doi.org/10.1145/3408996>
- Jean-Yves Girard. 1972. *Interpretation fonctionnelle et élimination des coupures de l'arithmétique d'ordre supérieur*. Ph.D. Dissertation. Université Paris VII.
- Colin S. Gordon, Matthew J. Parkinson, Jared Parsons, Aleks Bromfield, and Joe Duffy. 2012. Uniqueness and reference immutability for safe parallelism. In *OOPSLA*, 21–40. <https://doi.org/10.1145/2384616.2384619>
- Simon Oddershede Gregersen, Johan Bay, Amin Timany, and Lars Birkedal. 2021. Mechanized logical relations for termination-insensitive noninterference. *PACMPL* 5, POPL (2021), 1–29. <https://doi.org/10.1145/3434291>
- Dan Grossman, Greg Morrisett, and Steve Zdancewic. 2000. Syntactic type abstraction. *TOPLAS* 22, 6 (2000), 1037–1080. <https://doi.org/10.1145/371880.371887>

- Robert Harper. 2016. *Practical Foundations for Programming Languages (2nd. Ed.)*. Cambridge University Press. <https://www.cs.cmu.edu/~rwh/pfpl/index.html>
- Maurice Herlihy and Jeannette M. Wing. 1990. Linearizability: A correctness condition for concurrent objects. *TOPLAS* 12, 3 (1990), 463–492. <https://doi.org/10.1145/78969.78972>
- Jonas Kastberg Hinrichsen, Jesper Bengtson, and Robbert Krebbers. 2020. Actris: Session-type based reasoning in separation logic. *PACMPL* 4, POPL (2020), 6:1–6:30. <https://doi.org/10.1145/3371074>
- Jonas Kastberg Hinrichsen, Jesper Bengtson, and Robbert Krebbers. 2022. Actris 2.0: Asynchronous session-type based reasoning in separation logic. *LMCS* 18, 2 (2022). [https://doi.org/10.46298/lmcs-18\(2:16\)2022](https://doi.org/10.46298/lmcs-18(2:16)2022)
- Jonas Kastberg Hinrichsen, Daniël Louwink, Robbert Krebbers, and Jesper Bengtson. 2021. Machine-checked semantic session typing. (2021), 178–198. <https://doi.org/10.1145/3437992.3439914>
- Chung-Kil Hur and Derek Dreyer. 2011. A Kripke logical relation between ML and assembly. In *POPL*. 133–146. <https://doi.org/10.1145/1926385.1926402>
- Chung-Kil Hur, Derek Dreyer, Georg Neis, and Viktor Vafeiadis. 2012. The marriage of bisimulations and Kripke logical relations. In *POPL*. 59–72. <https://doi.org/10.1145/2103656.2103666>
- Samin S. Ishtiaq and Peter W. O’Hearn. 2001. BI as an Assertion Language for Mutable Data Structures. In *POPL*. 14–26. <https://doi.org/10.1145/360204.375719>
- Jules Jacobs, Jonas Kastberg Hinrichsen, and Robbert Krebbers. 2024. Deadlock-Free Separation Logic: Linearity Yields Progress for Dependent Higher-Order Message Passing. *PACMPL* 8, POPL (2024), 1385–1417. <https://doi.org/10.1145/3632889>
- Koen Jacobs, Amin Timany, and Dominique Devriese. 2021. Fully abstract from static to gradual. *PACMPL* 5, POPL (2021), 1–30. <https://doi.org/10.1145/3434288>
- Achim Jung and Jerzy Tiuryn. 1993. A new characterization of lambda definability. In *TLCA (LNCS, Vol. 664)*. 245–257. <https://doi.org/10.1007/BFb0037110>
- Ralf Jung. 2020. *Understanding and evolving the Rust programming language*. Ph.D. Dissertation. Saarland University. <https://publikationen.sulb.uni-saarland.de/handle/20.500.11880/29647>
- Ralf Jung, Jacques-Henri Jourdan, Robbert Krebbers, and Derek Dreyer. 2018a. RustBelt: Securing the foundations of the Rust programming language. *PACMPL* 2, POPL (2018), 66:1–66:34. <https://doi.org/10.1145/3158154>
- Ralf Jung, Jacques-Henri Jourdan, Robbert Krebbers, and Derek Dreyer. 2021. Safe systems programming in Rust. *CACM* 64, 4 (2021), 144–152. <https://doi.org/10.1145/3418295>
- Ralf Jung, Robbert Krebbers, Lars Birkedal, and Derek Dreyer. 2016. Higher-order ghost state. In *ICFP*. 256–269. <https://doi.org/10.1145/2951913.2951943>
- Ralf Jung, Robbert Krebbers, Jacques-Henri Jourdan, Ales Bizjak, Lars Birkedal, and Derek Dreyer. 2018b. Iris from the ground up: A modular foundation for higher-order concurrent separation logic. *JFP* 28 (2018), e20. <https://doi.org/10.1017/S0956796818000151>
- Ralf Jung, Rodolphe Lepigre, Gaurav Parthasarathy, Marianna Rapoport, Amin Timany, Derek Dreyer, and Bart Jacobs. 2020. The future is ours: Prophecy variables in separation logic. *PACMPL* 4, POPL (2020), 45:1–45:32. <https://doi.org/10.1145/3371113>
- Ralf Jung, David Swasey, Filip Sieczkowski, Kasper Svendsen, Aaron Turon, Lars Birkedal, and Derek Dreyer. 2015. Iris: Monoids and invariants as an orthogonal basis for concurrent reasoning. In *POPL*. 637–650. <https://doi.org/10.1145/2676726.2676980>
- Jan-Oliver Kaiser, Hoang-Hai Dang, Derek Dreyer, Ori Lahav, and Viktor Vafeiadis. 2017. Strong logic for weak memory: reasoning about release-acquire consistency in Iris. In *ECOOP (LIPIcs, Vol. 74)*. 17:1–17:29. <https://doi.org/10.4230/LIPIcs.ECOOP.2017.17>
- Anders Kock. 1970. Monads on symmetric monoidal closed categories. *Archiv der Mathematik* 21, 1 (1970), 1–10.
- Anders Kock. 1972. Strong functors and monoidal monads. *Archiv der Mathematik* 23, 1 (1972), 113–120.
- Vasileios Koutavas and Mitchell Wand. 2006. Small bisimulations for reasoning about higher-order imperative programs. In *POPL*. 141–152. <https://doi.org/10.1145/1111037.1111050>
- Robbert Krebbers, Jacques-Henri Jourdan, Ralf Jung, Joseph Tassarotti, Jan-Oliver Kaiser, Amin Timany, Arthur Charguéraud, and Derek Dreyer. 2018. MoSeL: A general, extensible modal framework for interactive proofs in separation logic. *PACMPL* 2, ICFP (2018), 77:1–77:30. <https://doi.org/10.1145/3236772>
- Robbert Krebbers, Ralf Jung, Ales Bizjak, Jacques-Henri Jourdan, Derek Dreyer, and Lars Birkedal. 2017a. The essence of higher-order concurrent separation logic. In *ESOP*. 696–723. https://doi.org/10.1007/978-3-662-54434-1_26
- Robbert Krebbers, Amin Timany, and Lars Birkedal. 2017b. Interactive proofs in higher-order concurrent separation logic. In *POPL*. 205–217. <https://doi.org/10.1145/3093333.3009855>
- Neelakantan R. Krishnaswami and Nick Benton. 2011. Ultrametric semantics of reactive programs. In *LICS*. 257–266. <https://doi.org/10.1109/LICS.2011.38>

- Neelakantan R. Krishnaswami, Aaron Turon, Derek Dreyer, and Deepak Garg. 2012. Superficially substructural types. In *ICFP*. 41–54. <https://doi.org/10.1145/2364527.2364536>
- Jean-Louis Krivine. 1994. Classical logic, storage operators and second-order lambda-calculus. *APAL* 68, 1 (1994), 53–78. [https://doi.org/10.1016/0168-0072\(94\)90047-7](https://doi.org/10.1016/0168-0072(94)90047-7)
- Morten Krogh-Jespersen, Kasper Svendsen, and Lars Birkedal. 2017. A relational model of types-and-effects in higher-order concurrent separation logic. In *POPL*. 218–231. <https://doi.org/10.1145/3093333.3009877>
- Søren B. Lassen and Paul Blain Levy. 2007. Typed normal form bisimulation. In *CSL (LNCS, Vol. 4646)*. 283–297. https://doi.org/10.1007/978-3-540-74915-8_23
- Xavier Leroy and Sandrine Blazy. 2008. Formal verification of a C-like memory model and its uses for verifying program transformations. *JAR* 41, 1 (2008), 1–31. <https://doi.org/10.1007/s10817-008-9099-0>
- Xavier Leroy and Pierre Weis. 1991. Polymorphic type inference and assignment. In *POPL*. 291–302. <https://doi.org/10.1145/99583.99622>
- David B. MacQueen. 1984. Modules for Standard ML. In *LFP*. 198–207. <https://doi.org/10.1145/800055.802036>
- David B. MacQueen, Gordon D. Plotkin, and Ravi Sethi. 1986. An ideal model for recursive polymorphic types. *Information and Control* 71, 1/2 (1986), 95–130. [https://doi.org/10.1016/S0019-9958\(86\)80019-5](https://doi.org/10.1016/S0019-9958(86)80019-5)
- Robin Milner. 1978. A theory of type polymorphism in programming. *JCSS* 17, 3 (1978), 348–375. [https://doi.org/10.1016/0022-0000\(78\)90014-4](https://doi.org/10.1016/0022-0000(78)90014-4)
- John C. Mitchell. 1986. Representation independence and data abstraction. In *POPL*. 263–276. <https://doi.org/10.1145/512644.512669>
- John C. Mitchell and Gordon D. Plotkin. 1988. Abstract types have existential type. *TOPLAS* 10, 3 (1988), 470–502. <https://doi.org/10.1145/44501.45065>
- Greg Morrisett, Matthias Felleisen, and Robert Harper. 1995. Abstract models of memory management. In *FPCA*. 66–77. <https://doi.org/10.1145/224164.224182>
- Andrzej S. Murawski and Nikos Tzevelekos. 2019. Higher-order linearisability. *J. Log. Algebraic Methods Program.* 104 (2019), 86–116. <https://doi.org/10.1016/j.jlamp.2019.01.002>
- Hiroshi Nakano. 2000. A modality for recursion. In *LICS*. 255–266. <https://doi.org/10.1109/LICS.2000.855774>
- Emeric Nasi. 2011. Modify any Java class field using reflection. Website: <https://blog.sevagas.com/Modify-any-Java-class-field-using-reflection>.
- Georg Neis, Derek Dreyer, and Andreas Rossberg. 2009. Non-parametric parametricity. In *ICFP*. 135–148. <https://doi.org/10.1145/1596550.1596572>
- Georg Neis, Chung-Kil Hur, Jan-Oliver Kaiser, Craig McLaughlin, Derek Dreyer, and Viktor Vafeiadis. 2015. Pilsner: A compositionally verified compiler for a higher-order imperative language. In *ICFP*. 166–178. <https://doi.org/10.1145/2784731.2784764>
- Peter W. O’Hearn. 2007. Resources, concurrency, and local reasoning. *TCS* 375, 1-3 (2007), 271–307. <https://doi.org/10.1016/j.tcs.2006.12.035>
- Peter W. O’Hearn, John C. Reynolds, and Hongseok Yang. 2001. Local reasoning about programs that alter data structures. In *CSL (LNCS, Vol. 2142)*. 1–19. https://doi.org/10.1007/3-540-44802-0_1
- Peter W. O’Hearn and Robert D. Tennent. 1992. Semantics of local variables. In *Applications of Categories in Computer Science*. London Mathematical Society Lecture Note Series, Vol. 177. 217–238.
- Benjamin C. Pierce. 2002. *Types And Programming Languages*. MIT Press.
- Andrew M. Pitts. 1996. Relational properties of domains. *Information and Computation* 127, 2 (1996), 66–90. <https://doi.org/10.1006/inco.1996.0052>
- Andrew M. Pitts. 2005. Typed operational reasoning. In *Advanced Topics in Types and Programming Languages*, B. C. Pierce (Ed.). The MIT Press, Chapter 7, 245–289.
- Andrew M. Pitts and Ian Stark. 1998. Operational reasoning for functions with local state. In *HOOTS*.
- Gordon D. Plotkin and Martín Abadi. 1993. A logic for parametric polymorphism. In *TLCA (LNCS, Vol. 664)*. 361–375. <https://doi.org/10.1007/BFb0037118>
- Xiaojia Rao, Aina Linn Georges, Maxime Legoupil, Conrad Watt, Jean Pichon-Pharabod, Philippa Gardner, and Lars Birkedal. 2023. Iris-Wasm: Robust and Modular Verification of WebAssembly Programs. *PACMPL* 7, PLDI (2023), 1096–1120. <https://doi.org/10.1145/3591265>
- John C. Reynolds. 1974. Towards a theory of type structure. In *Programming Symposium, Proceedings Colloque sur la Programmation, Paris (LNCS, Vol. 19)*. 408–423. https://doi.org/10.1007/3-540-06859-7_148
- John C. Reynolds. 1983. Types, abstraction and parametric polymorphism. In *Information Processing 83*. 513–523.
- John C. Reynolds. 2002. Separation logic: A logic for shared mutable data structures. In *LICS*. 55–74. <https://doi.org/10.1109/LICS.2002.1029817>
- Andreas Rossberg, Claudio V. Russo, and Derek Dreyer. 2014. F-ing modules. *JFP* 24, 5 (2014), 529–607. <https://doi.org/10.1017/S0956796814000264>

- Michael Sammler, Deepak Garg, Derek Dreyer, and Tadeusz Litak. 2020. The high-level benefits of low-level sandboxing. *PACMPL* 4, POPL (2020), 32:1–32:32. <https://doi.org/10.1145/3371100>
- Michael Sammler, Rodolphe Lepigre, Robbert Krebbers, Kayvan Memarian, Derek Dreyer, and Deepak Garg. 2021. RefinedC: Automating the foundational verification of C code with refined ownership types. In *PLDI*. 158–174. <https://doi.org/10.1145/3453483.3454036>
- Jan Schwinghammer, Lars Birkedal, François Pottier, Bernhard Reus, Kristian Støvring, and Hongseok Yang. 2013. A step-indexed Kripke model of hidden state. *MSCS* 23, 1 (2013), 1–54. <https://doi.org/10.1017/S0960129512000035>
- Simon Spies, Lennard Gäher, Daniel Gratzer, Joseph Tassarotti, Robbert Krebbers, Derek Dreyer, and Lars Birkedal. 2021. Transfinite Iris: Resolving an existential dilemma of step-indexed separation logic. In *PLDI*. 80–95. <https://doi.org/10.1145/3453483.3454031>
- Simon Spies, Lennard Gäher, Joseph Tassarotti, Ralf Jung, Robbert Krebbers, Lars Birkedal, and Derek Dreyer. 2022. Later credits: Resourceful reasoning for the later modality. *PACMPL* 6, ICFP (2022), 283–311. <https://doi.org/10.1145/3547631>
- Kristian Støvring and Søren B. Lassen. 2007. A complete, co-inductive syntactic theory of sequential control and state. In *POPL*. 161–172. <https://doi.org/10.1145/1190216.1190244>
- Eijiro Sumii. 2009. A complete characterization of observational equivalence in polymorphic λ -calculus with general references. In *CSL (LNCS, Vol. 5771)*. 455–469. https://doi.org/10.1007/978-3-642-04027-6_33
- Eijiro Sumii. 2010. A bisimulation-like proof method for contextual properties in untyped λ -calculus with references and deallocation. *TCS* 411, 51–52 (2010), 4358–4378. <https://doi.org/10.1016/j.tcs.2010.09.009>
- Eijiro Sumii and Benjamin C. Pierce. 2007. A bisimulation for type abstraction and recursion. *JACM* 54, 5 (2007), 26. <https://doi.org/10.1145/1284320.1284325>
- Kasper Svendsen and Lars Birkedal. 2014. Impredicative concurrent abstract predicates. In *ESOP (LNCS, Vol. 8410)*. 149–168. https://doi.org/10.1007/978-3-642-54833-8_9
- Kasper Svendsen, Lars Birkedal, and Matthew J. Parkinson. 2013. Modular reasoning about separation of concurrent data structures. In *ESOP (LNCS, Vol. 7792)*. 169–188. https://doi.org/10.1007/978-3-642-37036-6_11
- David Swasey, Deepak Garg, and Derek Dreyer. 2017. Robust and compositional verification of object capability patterns. *PACMPL* 1, OOPSLA (2017), 89:1–89:26. <https://doi.org/10.1145/3133913>
- Jean-Pierre Talpin and Pierre Jouvelot. 1994. The type and effect discipline. *Information and Computation* 111, 2 (1994), 245–296. <https://doi.org/10.1006/inco.1994.1046>
- Joseph Tassarotti, Ralf Jung, and Robert Harper. 2017. A higher-order logic for concurrent termination-preserving refinement. In *ESOP (LNCS, Vol. 10201)*. 909–936. https://doi.org/10.1007/978-3-662-54434-1_34
- Jacob Thamsborg and Lars Birkedal. 2011. A Kripke logical relation for effect-based program transformations. In *ICFP*. 445–456. <https://doi.org/10.1145/2034773.2034831>
- Amin Timany. 2018. *Contributions in programming languages theory*. Ph.D. Dissertation. KU Leuven. <https://lirias.kuleuven.be/retrieve/510052>
- Amin Timany and Lars Birkedal. 2019. Mechanized relational verification of concurrent programs with continuations. *PACMPL* 3, ICFP (2019), 105:1–105:28. <https://doi.org/10.1145/3341709>
- Amin Timany, Armaël Guéneau, and Lars Birkedal. 2024. The Logical Essence of Well-Bracketed Control Flow. *PACMPL* 8, POPL (2024), 575–603. <https://doi.org/10.1145/3632862>
- Amin Timany, Léo Stefanescu, Morten Krogh-Jespersen, and Lars Birkedal. 2018. A logical relation for monadic encapsulation of state: Proving contextual equivalences in the presence of runST. *PACMPL* 2, POPL (2018), 64:1–64:28. <https://doi.org/10.1145/3158152>
- Mads Tofte. 1990. Type inference for polymorphic references. *Inf. Comput.* 89, 1 (1990), 1–34. [https://doi.org/10.1016/0890-5401\(90\)90018-D](https://doi.org/10.1016/0890-5401(90)90018-D)
- Aaron Turon, Derek Dreyer, and Lars Birkedal. 2013a. Unifying refinement and Hoare-style reasoning in a logic for higher-order concurrency. In *ICFP*. 377–390. <https://doi.org/10.1145/2500365.2500600>
- Aaron Turon, Jacob Thamsborg, Amal Ahmed, Lars Birkedal, and Derek Dreyer. 2013b. Logical relations for fine-grained concurrency. In *POPL*. 343–356. <https://doi.org/10.1145/2429069.2429111> Technical appendix: <https://people.mpi-sw.s.org/~dreyer/papers/relcon/appendix.pdf>.
- Andrew K. Wright. 1995. Simple imperative polymorphism. *Lisp Symb. Comput.* 8, 4 (1995), 343–355. <https://doi.org/10.1007/BF01018828>
- Andrew K. Wright and Matthias Felleisen. 1994. A syntactic approach to type soundness. *Information and Computation* 115, 1 (1994), 38–94. <https://doi.org/10.1006/inco.1994.1093>
- Hongseok Yang. 2007. Relational separation logic. *TCS* 375, 1–3 (2007), 308–334. <https://doi.org/10.1016/j.tcs.2006.12.036>