



Fair Operational Semantics

DONGJAE LEE*, Seoul National University, Republic of Korea

MINKI CHO*, Seoul National University, Republic of Korea

JINWOO KIM, Seoul National University, Republic of Korea

SOONWON MOON, Inha University, Republic of Korea

YOUNGJU SONG, MPI-SWS, Germany

CHUNG-KIL HUR, Seoul National University, Republic of Korea

Fairness properties, which state that a sequence of bad events cannot happen infinitely before a good event takes place, are often crucial in program verification. However, general methods for expressing and reasoning about various kinds of fairness properties are relatively underdeveloped compared to those for safety properties.

This paper proposes FOS (Fair Operational Semantics), a theory capable of expressing arbitrary notions of fairness as an operational semantics and reasoning about these notions of fairness. In addition, FOS enables thread-local reasoning about fairness by providing thread-local simulation relations equipped with separation-logic-style resource algebras. We verify a ticket lock implementation and a client of the ticket lock under weak memory concurrency as an example, which requires reasoning about different notions of fairness including fairness of a scheduler, fairness of the ticket lock implementation, and even fairness of weak memory. The theory of FOS, as well as the examples in the paper, are fully formalized in Coq.

CCS Concepts: • **Theory of computation** → **Operational semantics**; **Separation logic**; **Concurrency**.

Additional Key Words and Phrases: Fairness, Fair Operational Semantics, Fairness Logic

ACM Reference Format:

Dongjae Lee, Minki Cho, Jinwoo Kim, Soonwon Moon, Youngju Song, and Chung-Kil Hur. 2023. Fair Operational Semantics. *Proc. ACM Program. Lang.* 7, PLDI, Article 139 (June 2023), 24 pages. <https://doi.org/10.1145/3591253>

1 INTRODUCTION

Although safety properties (*i.e.*, something bad never happens under a certain condition) have been major goals of verification, fairness properties (*i.e.*, something bad cannot happen indefinitely without anything good occurring) are also often crucial in verification, in particular, for concurrent programs. For example, fairness assumptions about the underlying system—such as fairness about the scheduler (*i.e.*, a thread cannot be delayed by the scheduler indefinitely), or fairness about weak memory (*i.e.*, memory cannot delay committing a write indefinitely)—are often required to verify concurrent programs. Moreover, customized concepts of fairness are required when working with

*Dongjae Lee and Minki Cho contributed equally to this work.

Authors' addresses: Dongjae Lee, Seoul National University, Seoul, Republic of Korea, dongjae.lee@sf.snu.ac.kr; Minki Cho, Seoul National University, Seoul, Republic of Korea, minki.cho@sf.snu.ac.kr; Jinwoo Kim, Seoul National University, Seoul, Republic of Korea, jinwoo.kim@sf.snu.ac.kr; Soonwon Moon, Inha University, Incheon, Republic of Korea, damhiya@inha.edu; Youngju Song, MPI-SWS, Saarbrücken, Germany, youngju@mpi-sws.org; Chung-Kil Hur, Seoul National University, Seoul, Republic of Korea, gil.hur@sf.snu.ac.kr.



This work is licensed under a Creative Commons Attribution 4.0 International License.

© 2023 Copyright held by the owner/author(s).

2475-1421/2023/6-ART139

<https://doi.org/10.1145/3591253>

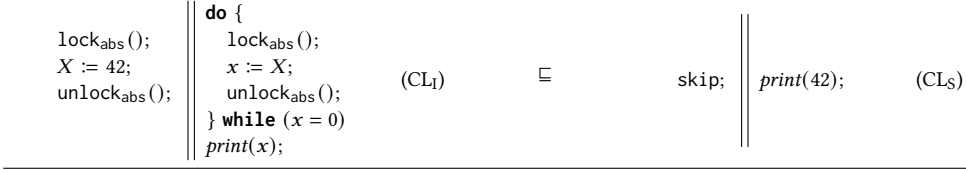


Fig. 1. An example of concurrent program verification.

custom libraries: for example, fairness about acquiring a lock provided by a custom ticket lock library (*i.e.*, a lock request cannot fail indefinitely given that the lock is available).

However, general methods for abstractly expressing various concepts of fairness and reasoning about them are relatively underdeveloped. For example, separation logics such as Iris [Jung et al. 2015] provide a flexible and powerful mechanism for specifying and reasoning about arbitrary safety properties. In contrast, existing work in the fairness domain such as TaDA-Live [D’Ousaldo et al. 2021] and LiLi [Liang and Feng 2016] are comparatively limited in their power, foremost in that they can only handle fixed notions of fairness instead of providing a general mechanism for expressing and reasoning about custom fairness properties.

Clearly, it is desirable to have a general mechanism capable of (i) capturing various kinds of fairness, and (ii) providing reasoning principles to both validate that a certain implementation is fair and verify client code assuming fairness of components. Consider Fig. 1 as an example, where a client program CL_I consisting of two threads utilizes a library-provided lock to protect a shared location X initialized to 0. The first thread sets X to 42, while the second thread repeatedly reads from X until reading a non-zero value, upon which it will print the read value.

Intuitively, CL_I *refines* the specification CL_S (where the first thread does nothing and the second prints 42) *assuming* that the scheduler and lock are ‘fair’. For example, CL_I fails to refine CL_S if the scheduler is unfair and schedules thread 2 only; even if the scheduler is fair, an unfair lock giving the lock only to thread 2, would result in an execution trace of CL_I that cannot be captured by CL_S .

Thus to prove that CL_I refines CL_S , one requires (i) a flexible mechanism for *expressing* the fairness of schedulers and locks, and (ii) *reasoning principles* for reasoning about such flexible fairness properties. Moreover, the ability to express and reason about general fairness properties enables a proof to be split in two ways: (i) one may *exploit* the fairness assumptions when verifying client code (*e.g.*, proving that CL_I refines CL_S), and (ii) separately *validate* that a specific implementation satisfies the fairness assumption (*e.g.*, proving that a specific scheduler or a lock is actually fair).

FOS, a Theory for Fairness. In this paper, we present **FOS** (Fair Operational Semantics), a theory for *expressing and reasoning about fairness* as an operational semantics. FOS provides:

- The standard notion of operational semantics extended with a special kind of events, called *fairness events*, that allows expressing arbitrary kinds of custom fairness properties.
- A simple global simulation relation that allows one to validate (for libraries) and exploit (for clients) fairness when proving refinements.
- A thread-local version of the aforementioned simulation and a logic capable of employing separation-logic style reasoning to perform thread-local reasoning and simplify proofs.

We demonstrate the generality and power of FOS by (i) specifying fairness of the scheduler, locks, and weak memory; (ii) proving fairness of the scheduler for a simple scheduler implementation and fairness of the lock for a ticket lock implementation; and (iii) exploiting these three kinds of fairness properties to verify client code (*e.g.*, $\text{CL}_I \sqsubseteq \text{CL}_S$). Note that existing work for reasoning about fairness is based on simple sequentially consistent concurrency, whereas we use FOS to

model a (fair) weak memory concurrency model and verify our examples under the model. The theory of FOS is mechanized in the Coq proof assistant [The Coq Development Team 2021].

Paper Structure. The remainder of the paper is structured as follows. §2 first illustrates how fairness properties can be expressed and reasoned about in a global fashion in FOS. §3 extends these ideas to show how FOS expresses fairness under concurrency, and shows how thread-local reasoning about fairness can be achieved in FOS. §4 and §5 formalize the high-level ideas presented in §2 and §3. §6 presents the module system of FOS, used to modularize proofs and define specifications capturing various notions of fairness, and §7 presents the details of the aforementioned thread-local reasoning. §8 presents a high-level overview of the verification of our motivating example. §9 concludes with related work.

2 MAIN IDEAS: FORMALLY EXPRESSING FAIRNESS

In this section, we aim to illustrate the key abilities of FOS: (i) How can we abstractly encode assumptions about fairness? (§2.1), (ii) How can one prove that a client program refines its specification assuming fairness about a library? (§2.2), and (iii) How can one prove that a library implementation refines its specification, validating the fairness assumption about the library? (§2.3).

We illustrate these abilities via the following example **LOT**, which draws a Boolean value via a lottery function `lottery()`, and prints "win!" whenever that value is *true*:

```
loop { if lottery() then print("win!") else skip } (LOT)
```

Note that **LOT** is *sequential* even though fairness commonly comes into play when reasoning about concurrent programs. Nevertheless, **LOT** is expressive enough to illustrate the key ideas behind FOS. The fairness assumption in **LOT** is that `lottery()` cannot return *false* indefinitely without returning *true*. Under this assumption, **LOT** prints infinitely many "win!"s, thus *refining* the following program:

```
loop { print("win!") } (WIN)
```

The goal of this section is to first show how the assumption that `lottery()` is fair can be encoded as a semantics of code, which doubles as a specification (§2.1), prove that **LOT** indeed refines **WIN** under this assumption (§2.2), and finally show how a concrete implementation of `lottery()` can be shown to validate this assumption (§2.3).

2.1 Fairness as a Semantics

As briefly mentioned in §1, fairness is the concept of ensuring that only a finite number of ‘bad events’ happen before a ‘good event’ takes place. FOS captures this concept by defining two *fairness events*: **good** and **bad**. Intuitively, **good** is triggered when a good event happens (e.g., `lottery()` triggers **good** when returning *true*), and **bad** is triggered when a bad event happens (e.g., `lottery()` triggers **bad** when returning *false*).

Observe that in a general setting, there are typically many *entities* for which fairness should be ensured (e.g., each thread in a scheduler). In such a setting, a specific event (e.g., thread 1 being scheduled) may be **good** for certain entities (e.g., thread 1) and **bad** for other entities (e.g., the other threads). FOS formalizes this intuition by triggering fairness events according to a **fairness map** `fmap`, which maps a *fairness id* to a fairness event **good** or **bad**. A fairness id is a unique identifier for each entity vying for fairness, *one for each kind of fairness* one wishes to consider: for example, each thread would have two fairness ids in a scenario considering both lock and scheduler fairness.

In this example, there is only one entity that needs fairness to be ensured (the program), and only one kind of fairness (the lottery), hence there is only one fairness id `lot`. Thus our `fmap` is of type `fmap : {lot}  $\mapsto$  {good, bad}`.¹

In order to actually trigger the fairness events within a fairness map, one requires a construct to trigger these events in code. This role is fulfilled by the fairness constructor `FAIR` in `FOS`, whose basic semantics is to take as argument an `fmap` and trigger all fairness events within the map for the appropriate fairness ids.

Modelling fairness through `good` and `bad` events that are triggered explicitly throughout the program allows us to give a formal definition of what it means for a program trace to be fair:

Definition 2.1 (Fair Trace). Consider a program P equipped with a set of fairness ids ID , and a trace t obtained by executing P . For a fairness id $id \in ID$, let t_{id} be a sub-sequence of events of t obtained by taking only the fairness events that are associated with id .

We say that t is a *fair trace* of P iff every t_{id} does not contain an infinite sequence of `bad` events that precedes a `good` event. Otherwise, we say that t is an *unfair trace* of P . If all possible traces of P are fair traces, then we say that P is *fair*.

Following Definition 2.1, consider the following specification of `lottery()` that encodes fairness of `lottery()` for `lot` (`PICK(B)` nondeterministically picks a Boolean value):

```
def lottery()spec = if PICK(B) then FAIR([lot  $\mapsto$  good]); ret true
                      else FAIR([lot  $\mapsto$  bad]); ret false (SPEC)
```

Here `FAIR` takes the `fmap` `[lot  $\mapsto$  good]` when `lottery()` returns `true`, which captures that a good event for `lot` happens. On the other hand, `FAIR` takes `[lot  $\mapsto$  bad]` on the false branch, capturing that a bad event for `lot` has occurred.

As stated in Definition 2.1, a fair execution (trace) of `SPEC` is a trace where each fairness id does not accumulate an infinite number of `bad` events before a `good` event. In tandem with triggering fairness events, the fairness constructor `FAIR` also *filters out* such unfair traces. A good way to understand `FAIR` is as a monitor that keeps track of the fairness events that each fairness id sees via the supplied fairness maps, and filters unfair executions out. One can then understand that the following unfair execution trace is *not* a valid trace of `SPEC`, as `FAIR` *prohibits* such unfair traces:

`[lot  $\mapsto$  bad] :: [lot  $\mapsto$  bad] :: [lot  $\mapsto$  bad] :: [lot  $\mapsto$  bad] :: ...`

On the other hand, traces such as the following are allowed, as each `bad` has a following `good`:

`[lot  $\mapsto$  bad] :: [lot  $\mapsto$  good] :: [lot  $\mapsto$  bad] :: [lot  $\mapsto$  good] :: ...`

In the end, the fairness events `good` and `bad` are silent events that have no effect on the actual behavior of the program; they are merely used to encode fairness via `FAIR`. Thus `SPEC` becomes a valid fair specification for `lottery()` that captures exactly the ‘fair’ behavior of a lottery.

As shown, `FOS` allows one to encode desired notions of fairness *directly as a piece of code but still in an abstract form* through the use of fairness maps and `FAIR`. This puts us in a very favorable position for our next task, to prove that `LOT` refines `WIN` assuming that `lottery()` is fair, as one may now utilize `SPEC` directly to build the refinement proof.

2.2 Proving that LOT Refines WIN: Exploiting Fairness

Based on our encoding of fairness as a semantics, we now wish to prove that `LOT` where the semantics of `lottery()` are given by `SPEC` (which we denote as `LOT[SPEC]`) refines `WIN`. In

¹A formal definition of `fmap` includes ϵ (an empty event) in the codomain, but we choose to ignore this for the time being.

proving refinement, we take the standard approach of constructing a simulation that matches arbitrary program steps from the target to program steps of the source.

What makes refinement in FOS special is the semantics of FAIR, which filters out any unfair execution. Because we are left only with fair executions, we call refinement in FOS as *fair refinement*, as it only maps fair target behaviors to fair source behaviors. Then FOS develops a simulation that ensures fair refinement; given that traditional simulations prove refinement with stepwise rules, FOS also aims for a simulation providing such stepwise rules, without having to reason about the entire execution trace. The main obstacle in developing such a simulation is the very presence of FAIR, which decides the fairness based on the entire trace. We overcome this by developing stepwise simulation rules that are capable of dealing with the FAIR constructs directly, by correctly reflecting their semantics (filtering out unfair traces) in the simulation.

Fairness Counter. To reason about FAIR in our simulation rules, we introduce the concept of a **fairness counter**, which intuitively counts the number of **bad** events that happen before a **good** event. Formally, a fairness counter is a per-program (source / target) map whose domain is identical to the fairness map of the program (that is, the entities for which fairness must be ensured) and whose range is a set equipped with a well-founded relation. In this example, we take our fairness counter c to be of type $\text{cmap} : \{\text{lot}\} \rightarrow \mathbb{N}$; i.e., a map that assigns a natural number to **lot**.

The key idea behind the fairness counter is that each entity is mapped to a value which *cannot decrease indefinitely*. As stated, this counter counts the maximum number of **bad** events that an entity may see *before* encountering a **good** event: a value cannot decrease indefinitely under a well-founded relation, thus *limiting* the number of **bad** events that may happen to a finite value. The simulation keeps track of the fairness counter along with the continuation of the program, and will update the fairness counter when dealing with FAIRs in the source or target program—a **bad** event decreases the counter, while a **good** event will *reset* this counter to some arbitrary value.

Formally, we capture this update mechanism as a relation $c \hookrightarrow_f c'$, and say that c' updates c given the fairness map f (the fairness map is required to determine the fairness ids for which the fairness counter should decrease and be reset). The exact definition of $c \hookrightarrow_f c'$ depends on the set of entities considered; in this example, $c \hookrightarrow_f c'$ can be defined as:

$$c \hookrightarrow_f c' \triangleq \text{match } f(\text{lot}) \text{ with } \mid \text{good} \Rightarrow \top \mid \text{bad} \Rightarrow c'(\text{lot}) < c(\text{lot})$$

This definition states that c' updates c if (i) **lot** triggers **good**, in which case *any* c' updates c , or (ii) **lot** triggers **bad**, in which case the counter of **lot** must *decrease*. Because the counter of **lot** cannot decrease below 0, updating the fairness counter ensures that unfair traces in which **lot** triggers **bad** infinitely are not considered when constructing our simulation; this corresponds to the semantics of FAIR which filters out unfair traces, allowing the fairness counter to capture the semantics of FAIR within a simulation.

Exploiting Fairness. Now, we wish to construct a simulation that relies on `lottery()` being fair to prove that **LOT[SPEC]** refines **WIN**. The following rule enables this by dealing with FAIR in the target program when constructing a simulation:

$$\frac{\text{SIMFT} \quad \forall c'. (c \hookrightarrow_f c') \rightarrow (c', k) \lesssim \mathbb{S}}{(c, \text{FAIR}(f); k) \lesssim \mathbb{S}}$$

where k denotes the continuation, $\mathbb{S}$ the source, and the relation $\lesssim$ denotes that the right-hand argument of $\lesssim$ simulates the left-hand argument. Note that **SIMFT** (and our simulation rules in general) relates a *pair* of a fairness counter and a continuation. The rule consumes a FAIR from the *target*, and intuitively *applies the effect* of the FAIR construct by updating the fairness counter from c to c' . In particular, **SIMFT** states that we *only need to consider* cases where c' updates c under f :

as we will illustrate below, this has the effect of *discarding* unfair executions of the target in the simulation, just as FAIR filters out unfair executions.

As an example, consider an application of **SIMFT** in order to prove the following simulation:

$$([\text{lot} \mapsto i], \text{FAIR}([\text{lot} \mapsto \text{bad}]); k) \lesssim \mathbb{S}$$

SIMFT applied to this proof goal requires that we prove a simulation for any c' that updates $[\text{lot} \mapsto i]$ when lot triggers **bad**: that is, any $[\text{lot} \mapsto i']$ for $i' < i$. If lot keeps on losing due to an unfair execution, the simulation proceeds until the following proof goal is reached:

$$([\text{lot} \mapsto 0], \text{FAIR}([\text{lot} \mapsto \text{bad}]); k) \lesssim \mathbb{S}$$

At this point, observe that there does *not exist* any c' such that c' updates $[\text{lot} \mapsto 0]$, as $\nexists i' \in \mathbb{N}. i' < 0$. Thus the premise of this proof goal becomes a vacuous truth, and it follows that the simulation trivially holds for such unfair executions: this is the effect of **SIMFT** discarding unfair target executions in the simulation, similar to how FAIR filters out unfair executions!

On the other hand, suppose that lot encounters a **good** event and wishes to prove the following:

$$([\text{lot} \mapsto i], \text{FAIR}([\text{lot} \mapsto \text{good}]); k) \lesssim \mathbb{S}$$

In this case, the updated counter c' may be an *arbitrary* natural number: this ensures that fair traces are not discarded, no matter how many **bad** events precede a **good** event. Since a **good** event returns *true*, the simulation will then be able to match a *print* from **LOT**[SPEC] with a *print* from **WIN**, and further apply the same reasoning inductively to prove that **LOT**[SPEC] indeed does refine **WIN**.

The power of **SIMFT** allowing us to discard unfair traces of the target is what makes constructing a simulation that considers only fair behavior possible: we call this property **fairness exploitation**.

2.3 Proving that lottery() is Fair: Validating Fairness

To wrap this section up, we consider constructing a simulation in the opposite direction to §2.2: how can we deal with FAIRs in the *source* program? Such scenarios emerge when one is trying to prove that a certain implementation is indeed fair: *e.g.*, when proving that a specific implementation of **lottery()** refines **SPEC**.

Consider the following two implementations of **lottery()**:

```
def lotfair() = if x then x := false; ret true else x := true; ret false
def lotunfair() = ret false
```

(IMPL)

Before discussing the fairness of lot_{fair} and $\text{lot}_{\text{unfair}}$ directly, observe that both lot_{fair} and $\text{lot}_{\text{unfair}}$ are trivially ‘fair’ in themselves as their executions are finite. However, when used with **LOT**, clearly $\text{lot}_{\text{unfair}}$ becomes unfair while lot_{fair} is still fair. As illustrated, the *context* in which an implementation is used has an effect on whether the implementation is fair or not; we will thus prove that **LOT**[lot_{fair}] refines **LOT**[SPEC] in order to prove that lot_{fair} is a fair implementation.

SIMFS is the rule for dealing with FAIRs from the source in the simulation ($\mathbb{T}$ denotes the target):

$$\frac{\text{SIMFS} \quad \exists c'. (c \hookrightarrow_f c') \wedge \mathbb{T} \lesssim (c', k)}{\mathbb{T} \lesssim (c, \text{FAIR}(f); k)}$$

Like **SIMFT** for consuming FAIRs in the target, **SIMFS** consumes FAIRs from the *source* and models the effect of FAIR in the simulation. In contrast to **SIMFT**, **SIMFS** states that there *must exist* a fair update of the fairness counter c : this ensures that there must exist a fair behavior of the *target* corresponding to the fair source behavior in order for the simulation to hold.

To see this effect, let us attempt to prove that **LOT**[$\text{lot}_{\text{unfair}}$] fairly refines **LOT**[SPEC] (which is clearly untrue). Assume this time that we are starting with the following proof goal, where FAIR

occurs in the source $\text{LOT}[\text{SPEC}]$:

$$\mathbb{T} \lesssim ([\text{lot} \mapsto i], \text{FAIR}([\text{lot} \mapsto \text{bad}])); k)$$

Observe that because the *target* $\text{LOT}[\text{lot}_{\text{unfair}}]$ has a trace where it cannot print any "win!"s due to $\text{lot}_{\text{unfair}}$ returning *false* indefinitely, the source must also be able to match this trace for a simulation to hold. Clearly, lot must trigger *bad* indefinitely within this source-trace as well: thus the fairness map is fixed to $[\text{lot} \mapsto \text{bad}]$, and applying SIMFS multiple times in this fashion will eventually again yield a proof goal where the fairness counter of lot is zero:

$$\mathbb{T} \lesssim ([\text{lot} \mapsto 0], \text{FAIR}([\text{lot} \mapsto \text{bad}])); k)$$

Here, observe that we *cannot* apply SIMFS to consume the source-FAIR anymore as there no longer exists a fairness counter that can update $[\text{lot} \mapsto 0]$. In essence, the fairness counter limits the behavior of the *source* to be fair, similar to how FAIR in the source filters out unfair traces of the source. The fact that the target $\text{LOT}[\text{lot}_{\text{unfair}}]$ has unfair behavior, which cannot be simulated by the only-fair behavior of the source, is reflected by SIMFS becoming no longer applicable when constructing the simulation. Because we can no longer consume the FAIR of the source, the simulation can no longer proceed and thus we cannot prove that $\text{LOT}[\text{lot}_{\text{unfair}}]$ refines $\text{LOT}[\text{SPEC}]$!

In contrast, suppose now that we are correctly trying to prove that $\text{LOT}[\text{lot}_{\text{fair}}]$ refines $\text{LOT}[\text{SPEC}]$. Because $\text{LOT}[\text{lot}_{\text{fair}}]$ is fair and emits "win!" every other iteration, the source program can also let lot trigger *good* when constructing the simulation:

$$\mathbb{T} \lesssim ([\text{lot} \mapsto i], \text{FAIR}([\text{lot} \mapsto \text{good}])); k)$$

At this point, applying SIMFS allows us to update c' to any value; the simulation can thus choose a value larger than the number of *bad* events it will see before seeing another *good* event (here, even 1 will suffice as the target $\text{LOT}[\text{lot}_{\text{fair}}]$ alternates between printing "win!" and not). One can successfully prove that $\text{LOT}[\text{lot}_{\text{fair}}]$ refines $\text{LOT}[\text{SPEC}]$ by applying this reasoning inductively.

The reasoning principle enabled by SIMFS stands in direct contrast with the reasoning principle enabled by SIMFT , which allowed us to discard unfair target traces by corresponding them to vacuous truths. In contrast, SIMFS requires that the target must *only exhibit* fair behavior for it to be able to refine a (fair) source: we call this principle **fairness validation**.

3 MAIN IDEAS: CONCURRENCY AND THREAD-LOCAL REASONING

Having established the basic reasoning principles behind FOS, we now illustrate how fairness under concurrency with schedulers and locks can be modeled in FOS (§3.1), and then how refinement for concurrent programs can be proved in FOS, particularly in a thread-local manner (§3.2).

The running example in this section is identical to the motivating example from §1 (memory locations and variables in the paper, stated otherwise, are initialized to 0):

$$\begin{array}{l|l} \text{Y}; \text{lock}_{\text{abs}}(); & \text{do } \{ \\ \text{Y}; X := 42; & \text{Y}; \text{lock}_{\text{abs}}(); \\ \text{Y}; \text{unlock}_{\text{abs}}(); & \text{Y}; x := X; \\ \text{Y}; & \text{Y}; \text{unlock}_{\text{abs}}(); \text{Y}; \\ & \} \text{ while } (x = 0) \\ & \text{Y}; \text{print}(x); \text{Y}; \end{array} \quad (\text{CL}_1)$$

We assume sequential consistency as the memory model for CL_1 ; later in §6.1, we will encode memory fairness to allow the use of weak memory models as well. As discussed, CL_1 is a worker (thread 1) and a waiter (thread 2) process, which relies on a lock to protect non-atomic accesses to the shared memory location X . In this paper, we define the semantics of yielding such that a thread

will yield if and only if it meets an explicit *yield instruction* Υ in the program: this models threads to be altruistic, and also has the effect of treating program fragments between Υ s as atomic.

As discussed in §1, we wish to prove that CL_I fairly refines the following specification CL_S :

$$\Upsilon; \text{skip}; \Upsilon; \parallel \Upsilon; \text{print}(42); \Upsilon; \quad (\text{CL}_S)$$

To prove this refinement, one must be capable of establishing that thread 1 eventually acquires the lock, which in turn depends on fairness of the scheduler and the lock. We thus formally express scheduler and lock fairness first in §3.1, then illustrate the key argument required for establishing that thread 1 eventually acquires the lock during a fair execution in §3.2.

3.1 Fair Specifications for Concurrent Components

To prove that CL_I fairly refines CL_S , one first requires a specification that encodes the fairness of the scheduler and the lock. In a concurrent setting, this means that fairness should be ensured for all threads, which are vying to get scheduled or acquire the lock. We can express these two distinct kinds of fairness by setting the fairness ids ID as the disjoint union of th_i , the ids for scheduler fairness, and lk_i , the ids for lock fairness, where i is a thread id. Thus ID becomes the domain of the fairness map fmap and the fairness counter cmap :

$$\text{f} \in \text{fmap}_{\text{ID}} \triangleq \text{ID} \rightarrow \{\text{good}, \text{bad}\} \quad \text{c} \in \text{cmap}_{\text{ID}} \triangleq \text{ID} \rightarrow \mathbb{N}$$

The events **good** and **bad** again model good and bad events: in this case, a **good** event is when a thread is scheduled or successfully acquires a lock.

Scheduler Fairness. In this paper, concurrency is modeled via thread interleaving, where a thread pool is a finite map from a thread id (in $\mathbb{N}$) to a code. Then the scheduler manages a thread pool by kicking out threads when they terminate; we call the remaining threads in the pool valid threads, and denote their ids as $\mathbb{N}_v$. Under this thread interleaving model, a scheduler may be viewed as a program that picks a valid thread id and executes the code of that thread; afterward, a yield Υ will return execution to the scheduler which will proceed to schedule the next thread.

A fair spec of a scheduler may be written as follows, where P represents the thread pool:

$$\lambda P. \text{while}(\mathbb{N}_v \neq \emptyset) \{ n \leftarrow \text{PICK}(\mathbb{N}_v); \text{FAIR}([\text{th}_n \mapsto \text{good}, \{\text{th}_m \mid m \in \mathbb{N}_v \wedge m \neq n\} \mapsto \text{bad}]); \text{exec}(P(n)) \} \quad (\text{SCH})$$

In **SCH**, **PICK** nondeterministically picks a valid thread id n to execute. The fairness of **SCH** is guaranteed by the following **FAIR** constructor, whose fairness map assigns **good** to the scheduled thread and **bad** to all other threads: following the semantics of the **FAIR** constructor, it is clear that any fair execution of **SCH** will result in valid threads becoming scheduled infinitely often. This captures exactly the concept of ‘scheduler fairness’ that programmers rely on when writing programs: that running threads will not be starved indefinitely by the scheduler.

Throughout the paper, we assume that schedulers satisfy **SCH** if not stated otherwise; this allows us to exploit scheduler fairness. We emphasize that **SCH** represents only the *fair* schedulers. Therefore, if a system’s correctness relies on **SCH**, one should also validate the fairness of the scheduler implementation employed by the system to guarantee the full correctness of the system.

Lock Fairness. In addition to scheduler fairness, proving that CL_I refines CL_S requires fairness of the *lock*: clearly CL_I will fail to terminate if, e.g., the lock implementation is unfair and never grants the lock to thread 1. Encoding this fairness requirement once again requires a specification for the lock, which in addition to guaranteeing fairness as desired, should ideally abstract away implementation details such as data structures or memory accesses. The second desideratum allows different *concrete implementations* to refine the specification, as done in §2.3.

def lock _{abs} () =		$W \subseteq \{lk_i \mid i \in \mathbb{N}\}$ $own \in \mathbb{B}$
$n = \text{GetTid}; W = W \cup \{lk_n\};$	//L1	def unlock _{abs} () =
loop { if own then $\bar{Y}$; else break }	//L2	assume ($own = \text{true}$); //L1
$own = \text{true}; W = W \setminus \{lk_n\}; \text{FAIR}([lk_n \mapsto \text{good}, W \mapsto \text{bad}])$	//L3	$own = \text{false};$ //L2
ret	//L4	ret //L3

Fig. 2. A fair spec of fair locks, ABSLock.

Fig. 2 presents such a specification for fair locks. Consider the lock function lock_{abs}() : here, each thread with thread id i is represented using a new entity lk_i , which is a new fairness id for each thread required to ensure that a **good** obtained via a scheduling event does not also result in a **good** of the lock, and vice versa. The specification is straightforward: a thread will add itself to the waiting set W and enter the loop to wait for a lock, then trigger a **good** for itself and **bad**s for all other threads in W upon acquiring a lock. This implies that a thread *waiting* for the lock must eventually acquire the lock if the lock is available, as it will accumulate an infinite number of **bad**s otherwise (assuming scheduler fairness). On the other hand, a thread outside of the waiting set W is unaffected by the fairness events triggered by the lock: this captures the common concept of ‘lock fairness’, in that a thread attempting to acquire an available lock will eventually succeed.

(T, S : globally fixed memory locations)

def lock_{tk}() = $t := \text{FAI}(T); \text{do } \{ s := S; \} \text{ while}(t \neq s)$ (TicketLock)

def unlock_{tk}() = $s := S; s = s + 1; S := s;$

To conclude the discussion on fair concurrent specifications, let us briefly consider how the ticket lock implementation in **TicketLock** can validate the fair specification given in Fig. 2. **TicketLock** issues a ticket counter t to every thread that attempts to acquire the lock; this counter is guaranteed to strictly increase due to the *fetch-and-increment* instruction **FAI**(T), which increments the value stored at the memory location T by 1 and returns the old value. When a thread calls unlock_{tk}(), the service counter stored in the memory location S is increased: this allows the ticket lock to service the next thread that holds the ticket corresponding to the current service counter.

TicketLock is fair under sequential consistency and scheduler fairness, refining the fair specification from Fig. 2. This is because the ticketing scheme ensures that only a finite number of threads may be queued before a thread to acquire a lock; this means that a thread can only trigger a finite number of **bad**s before successfully acquiring the lock, ensuring that the trace is fair.

3.2 Exploiting Fairness Under Concurrency via Thread-Local Reasoning

Given the fairness properties we have encoded in §3.1, we now illustrate how one may actually prove that **CL_I** fairly refines **CL_S**. The key intuition in this section is that one may perform induction on the fairness counters when constructing a simulation to discard unfair behavior. Furthermore, these fairness counters may be treated as *shared state* between threads, which enables *thread-local reasoning* when constructing a simulation.

Global Proof by Induction. Before illustrating how thread-local reasoning can be performed in FOS, we first establish a global argument to prove that **CL_I** fairly refines **CL_S**; we will later show that this argument extends naturally to thread-local reasoning.

One of the main reasoning patterns enabled by fairness is the ability to argue that something **good** must happen eventually in a fair execution—e.g., thread 1 eventually acquires the lock in **CL_I**. In **CL_I**, this is required to ensure that the loop in thread 2 of **CL_I** eventually exits in a fair execution, which allows one to match the final *print* with the *print* in **CL_S** when constructing a simulation.

The key idea in proving that a **good** event happens is to construct a value that *decreases* but cannot decrease infinitely, until the **good** event is triggered. The fairness counter c (from the global

simulation relation in §2) captures perfectly this required decreasing value. Recall that the fairness counter also allows one to discard unfair executions of the target: thus by performing induction on the fairness counter, one can focus only on fair executions when constructing the simulation.

To see this strategy in detail, reconsider the fact that we must show thread 1 of CL_1 eventually acquires the lock (a *good* event) when constructing a simulation. Intuitively, lk_1 , the fairness counter for thread 1 acquiring the lock, can serve as the decreasing value for this *good* event, as thread 2 acquiring the lock instead will trigger a *bad* for lk_1 . Then by performing induction on lk_1 , one can establish that either (i) thread 1 acquires the lock and writes 42 to X , or (ii) the execution being considered is unfair, as lk_1 cannot decrease indefinitely.

To be more formal about the argument, one must also consider the fact that thread 2 progresses and eventually releases the lock (otherwise thread 1 will not accumulate a *bad*). This can be done by introducing an additional counter n , which keeps track of the remaining lines of code (*i.e.*, the number of remaining yields) until the lock is released. Of course, progress of thread 2 relies on thread 2 being scheduled—which can in turn be captured by the fairness counter th_2 . Thus the *actual* decreasing value for thread 1 to acquire the lock is given as a tuple (lk_1, n, th_2) , where the order of the tuples is given by lexicographic order: lk_1 is the most significant as the subsequent two counters serve to ensure that lk_1 decreases. A formal proof would perform induction on this tuple to ensure progress as opposed to merely lk_1 .

Thread-Local Reasoning for Fair Refinement. Given the aforementioned proof strategy based on induction on a decreasing value, we now show how this reasoning may be performed in a *thread-local* manner to prove that CL_1 fairly refines CL_5 without considering all thread interleavings.

As mentioned, the key idea that enables local reasoning for FOS is that the fairness counter—which captures the concept of fair behavior in the simulation—can naturally be treated as *shared state* between threads. This allows each thread to perform local reasoning via a **shared protocol**, in which threads *rely* on other threads upholding the protocol when resuming execution from a yield Υ , and conversely *guaranteeing* that the protocol is satisfied before yielding. Through this rely-guarantee reasoning, one can perform induction thread-locally to prove that a refinement between threads holds using almost the same argument that was applied globally, *without* having to consider all thread interleavings in a global simulation.

In our example, consider the following protocol which states that threads satisfy one of the given three states on a yield. This protocol is used to ensure that thread 1 eventually acquires the lock:

- Either the lock is unowned ($own = false$),
- Thread 1 holds the lock, or
- Thread 2 holds the lock and the tuple (lk_1, n, th_2) decreases.

Here, it is true that n is local to thread 2 and inaccessible to thread 1; we will temporarily assume that n is a ‘ghost’ variable that is managed by thread 2 on every yield for the sake of presentation.

Consider this protocol in the context of thread 2: if thread 1 does not hold the lock, then (i) thread 2 yields in a state where $own = false$ (directly after releasing the lock), or (ii) it will have decreased the tuple (lk_1, n, th_2) . This is because thread 2 acquiring the lock will either trigger a *bad* event for thread 1, or decrease n by progressing within the loop and reducing the number of remaining yields. Thus thread 2 *guarantees* that it upholds the protocol at each point it encounters a yield, *without* considering the behavior of thread 1 at all.

On the other hand, thread 1 may *rely* on the protocol being upheld when resuming execution at a yield point. If $own = false$, then thread 1 acquires the lock and the simulation can make progress. If $own = true$ but thread 1 does not have the lock, then thread 1 decreases the given tuple as th_2 decreases due to thread 1 winning in the scheduler. Thus thread 1 can also *guarantee* that the protocol is upheld, without considering the behavior of thread 2.

$ID : Type$	$flag \triangleq \{\text{good}, \text{bad}, \epsilon\}$	$fmap_{ID} \in ID \rightarrow flag$	$R : Type$
$FL_{ID,R} \stackrel{\text{coind}}{=} $	$FAIR(f \in fmap_{ID}) \gg (k \in () \rightarrow FL) \mid PICK(X : Type) \gg (k \in X \rightarrow FL)$ $ Obs(fn \in string, args \in list Val) \gg (k \in Val \rightarrow FL) \mid \text{ret}(r \in R) \mid \text{stuck}$		
$SilEv \triangleq $	$\delta(f \in fmap)$		
$ObsEv \triangleq $	$obs(fn \in string, args \in list Val, v \in Val)$		
$Trace \stackrel{\text{coind}}{=} $	$(e \in SilEv \sqcup ObsEv) :: (tr \in Trace) \mid Term(r \in R) \mid Error$		
$Behavior \stackrel{\text{coind}}{=} $	$(o \in ObsEv) :: (tr \in Behavior) \mid Diverge \mid Term(r \in R) \mid Error$		
$FairTr \in \mathbb{P}(Trace) \triangleq \{$	$tr \mid \forall i \in ID. FairTr_i(tr) \}$		
$FairTr_{i \in ID} \in \mathbb{P}(Trace) \triangleq$	$vX. \mu Y. \nu Z. \lambda tr. //X, Y, Z \in \mathbb{P}(Trace)$		
$\text{match } tr \text{ with}$	$ Term\ r \Rightarrow \top \mid Error \Rightarrow \top$ $ obs(fn, args, v) :: tl \Rightarrow Z(tl)$ $ \delta(f) :: tl \Rightarrow \text{match } f(i) \text{ with}$		
	$ \epsilon$	$\Rightarrow Z(tl)$	$//Z: \text{inner coinductive}$
	$ \text{bad}$	$\Rightarrow Y(tl)$	$//Y: \text{middle inductive}$
	$ \text{good}$	$\Rightarrow X(tl)$	$//X: \text{outer coinductive}$

Fig. 3. Core definitions of the language FL, behavior, and fair trace.

As the shared protocol is guaranteed by both threads, one may perform induction locally in a thread to ensure progress. For example, in thread 1, applying induction using the protocol results in that either (i) thread 1 acquires the lock, or (ii) the execution is *unfair*, as the tuple failing to decrease indicates that the fairness counters lk_1 or th_2 cannot decrease, implying unfairness. Thus thread-local reasoning based on this protocol shows that executions where thread 1 cannot acquire the lock and does not terminate are unfair—and are thus discarded, allowing one to construct a thread-local simulation showing that thread 1 of CL_1 fairly refines thread 1 of CL_5 !

As shown above, protocols for thread-local reasoning about fairness often require the use of ghost values such as n , which may not be easy to expose to other threads. Fortunately, existing work on modern concurrent separation logic [Jung et al. 2015] allows us to express such protocols as invariants. FOS blends naturally with this idea, resulting in a simulation technique where true thread-local reasoning is enabled via separation logic; this technique is formalized in §7.

4 CORE DEFINITIONS OF FOS

We now formalize the ideas presented so far in Sections 4 to 6. In this section, we present (whole-program) fair semantics, refinement, and a simulation relation, which are straightforward formalizations of the ideas in §2; concurrency and the module system are formalized in §5 and §6.

4.1 Definitions of Fair Operational Semantics

To write fair programs (e.g., those in §2), we first define a language FL (fair language) in Fig. 3. FL is coinductively defined with three constructors FAIR, PICK, and Obs, in addition to **ret** for termination and an explicit **stuck** to indicate an error, i.e., *undefined behavior*. The fairness constructor FAIR invokes fairness events via the fairness map $fmap$. $fmap$'s range is $flag$, which now includes ϵ to express that the fairness event is undefined for that index (usually omitted for brevity). PICK non-deterministically picks a value from any given set X , passing it to the continuation. Finally, Obs invokes an observable effect, e.g., a system call, represented by a function name fn and arguments $args$, and passes the return value to the continuation. Terms in FL within the Coq development are defined using *interaction trees* [Xia et al. 2019], as opposed to directly embedded as Coq functions. This allows us to reuse programming constructs such as **match** ... **with** for branches, $x \leftarrow p$; k , $\gg$ for monadic bind, and **loop** for loops, thereby keeping our language minimal yet expressive.

The semantics of an FL program p is defined by the set of its possible *behaviors* $Beh(p)$. For this, we first derive a set of possible *traces* from a program where a *trace* is a stream of silent (δ) or observable (obs) events and can possibly terminate with a return value or an error. The

$$\begin{array}{c}
\text{cmap}_{\text{ID}, C_{<}} \in \text{ID} \rightarrow C_{<} \quad < \in \mathbb{P}(C_{<} \times C_{<}) \text{ is well-founded} \quad \lesssim \in \mathbb{P}((\text{cmap} \times \text{FL}) \times (\text{cmap} \times \text{FL})) \\
\text{c} \hookrightarrow_f \text{c}' \in \mathbb{P}(\text{cmap} \times \text{fmap} \times \text{cmap}) \triangleq \forall i \in \text{ID}. \text{match } f(i) \text{ with } \mid \text{good} \Rightarrow \top \mid \text{bad} \Rightarrow \text{c}'(i) < \text{c}(i) \mid \epsilon \Rightarrow \text{c}'(i) = \text{c}(i) \\
\begin{array}{ccc}
\text{SIMRET} & \text{SIMSTUCK} & \text{SIMPT} \quad \text{SIMPS} \\
\frac{r_t = r_s}{(\text{c}_t, \text{ret } r_t) \lesssim (\text{c}_s, \text{ret } r_s)} & \frac{}{\top \lesssim (\text{c}, \text{stuck})} & \frac{\forall x \in X. (\text{c}, k(x)) \lesssim \mathbb{S}}{(\text{c}, \text{PICK}(X) \ggg k) \lesssim \mathbb{S}} \quad \frac{\exists x \in X. \top \lesssim (\text{c}, k(x))}{\top \lesssim (\text{c}, \text{PICK}(X) \ggg k)} \\
\frac{\text{SIMFT} \quad \forall \text{c}' \in \text{cmap}. (\text{c} \hookrightarrow_f \text{c}') \rightarrow (\text{c}', k()) \lesssim \mathbb{S}}{(\text{c}, \text{FAIR}(f) \ggg k) \lesssim \mathbb{S}} & \frac{\text{SIMFS} \quad \exists \text{c}' \in \text{cmap}. (\text{c} \hookrightarrow_f \text{c}') \wedge \top \lesssim (\text{c}', k())}{\top \lesssim (\text{c}, \text{FAIR}(f) \ggg k)}
\end{array}
\end{array}$$

Fig. 4. Definitions of cmap and selected rules of our simulation relation $\lesssim$ (slightly simplified).

language construct **FAIR**(f) emits $\delta(f)$, **PICK** emits $\delta([_ \mapsto \epsilon])$, **Obs** emits corresponding obs, and **ret**(r)/**stuck** terminates the trace correspondingly. Then, a behavior is defined as an observable summary of a trace in the sense that it drops all silent events and leaves only observable events. An infinite trace with only silent events results in a silent divergence [Leroy 2006].

Now, for a set of traces, we derive a set of behaviors by (i) filtering out *unfair* traces with help from δ s, and (ii) erasing the now meaningless δ s. A *fair* trace, defined by FairTr_i , is a trace that satisfies FairTr_i for every index i in ID. FairTr_i allows only finite **bad**s until the next **good** for i , captured by the mixed coinductive-inductive-coinductive definition: the inner coinductive (Z) captures the possibly infinite fairness-irrelevant trace, the middle inductive (Y) ensures only finite **bad**s are encountered until a **good**, which can appear infinitely as expressed by the outer coinductive (X).

Then whole-program refinement between two FL programs p_t and p_s is defined as follows:

$$p_t \sqsubseteq p_s \triangleq \text{Beh}(p_t) \subseteq \text{Beh}(p_s)$$

$\sqsubseteq$ is *transitive*, *reflexive*, and preserves termination (i.e., if the target has (fair) divergence, so does the source). Note that the source and the target may have different IDs.

4.2 Simulation Relation

The simulation ($\lesssim$) presented in §2.2 is formalized in Fig. 4, which relates two FL programs. The rules are identical except that **cmap** is now parameterized over ID, the set of fairness indices, and $C_{<}$, a well-founded order.² Expectedly, the simulation satisfies the following adequacy theorem:

THEOREM 4.1 (ADEQUACY). *For a pair of programs p_t, p_s and a well-founded set $C_{<}$, we have:*

$$(\forall \text{c}_t \in \text{cmap}_{(\text{ID}_t, \mathbb{N})}. \exists \text{c}_s \in \text{cmap}_{(\text{ID}_s, C_{<})}. (\text{c}_t, p_t) \lesssim (\text{c}_s, p_s)) \implies p_t \sqsubseteq p_s$$

5 FORMALIZING FAIR CONCURRENCY

In this section, we present the formal definitions of our model of concurrency and scheduler fairness following the ideas in §3.1, and show that a simple round-robin scheduler satisfies scheduler fairness.

5.1 Semantics of Concurrency

In this paper, concurrency is modeled by *interleaving* semantics with *thread-yields*, meaning that the scheduler interleaves the execution of the threads and each thread yields to the scheduler. To write concurrent programs, we define the thread language TFL and the scheduler language SFL. Then threads and a scheduler are together interpreted as a FL program, so the semantics (behavior) of a concurrent system (threads and a scheduler) is naturally defined by the semantics of FL (Fig. 5).

²We omit ID and $C_{<}$ for brevity whenever they are clear from the context.

$\text{th} \triangleq \mathbb{N} \quad R, ST : \text{Type} \quad \text{TPool}_{\text{ID}, R, ST} \triangleq \text{th} \xrightarrow{\text{fin}} \text{TFL}_{\text{ID}, R, ST} \quad \text{Sch} \in \text{scheduler}_R \triangleq (\text{th} \times \text{Fin}(\text{th})) \rightarrow \text{SFL}_R$	
$\text{TFL}_{\text{ID}, R, ST} \stackrel{\text{coind}}{=} \mid \text{Y} \gg (k \in () \rightarrow \text{TFL}) \mid \text{GetTid} \gg (k \in \text{th} \rightarrow \text{TFL}) \mid \text{Put}(st \in ST) \gg (k \in () \rightarrow \text{TFL})$ $\mid \text{Get} \gg (k \in ST \rightarrow \text{TFL}) \mid \text{FAIR}(f) \gg k \mid \text{PICK}(X) \gg k \mid \text{Obs}(\dots) \gg k \mid \text{ret } r \mid \text{stuck}$	
$\text{SFL}_{\text{ID}=\text{th}, R} \stackrel{\text{coind}}{=} \mid \text{Exec}(n \in \text{th}) \gg (k \in R? \rightarrow \text{SFL}) \mid \text{FAIR}(f) \gg k \mid \dots$	
$R^\vee \triangleq \mid \blacktriangle(r \in R) \mid \vee(t \in \text{TFL})$ $\text{TI}(n \in \text{th}, t \in \text{TFL}, st \in ST) \in \text{FL}_{\text{ID}, R^\vee \times ST} \triangleq$ $\text{match } t \text{ with } \mid \dots \mid \text{ret } r \Rightarrow \text{ret } (\blacktriangle r, st)$ $\mid \text{Y} \gg k \Rightarrow \text{ret } (\vee k(), st)$ $\mid \text{GetTid} \gg k \Rightarrow \text{TI}(n, k(n), st)$ $\mid \text{Put}(st') \gg k \Rightarrow \text{TI}(n, k(), st')$ $\mid \text{Get} \gg k \Rightarrow \text{TI}(n, k(st), st)$	$\text{CI}(P, \text{Sch}, st) \triangleq \text{SI}(P, \text{Sch}(0, \text{dom}(P) \setminus \{0\}), st)$ $\text{SI}(P \in \text{TPool}, S \in \text{SFL}, st \in ST) \in \text{FL} \triangleq$ $\text{match } S \text{ with } \mid \dots \mid \text{ret } r \Rightarrow \text{ret } r$ $\mid \text{Exec}(n) \gg k \Rightarrow \text{match } P(n) \text{ with } \mid \text{None} \Rightarrow \text{stuck}$ $\mid \text{Some } t \Rightarrow x \leftarrow \text{TI}(n, t, st); \text{match } x \text{ with}$ $\mid \blacktriangle r, st' \Rightarrow \text{SI}(P[n \mapsto], k(\text{Some } r), st')$ $\mid \vee t_k, st' \Rightarrow \text{SI}(P[n \mapsto t_k], k(\text{None}), st')$

Fig. 5. Definitions of the thread/scheduler language TFL/SFL and the interpreters CI, SI, TI.

Threads. We define a thread as a procedure in TFL, and a thread pool as a finite map from thread ids to TFL. TFL extends FL with *shared states* among threads (e.g., memory) and concurrency features. The state ST is parameterized for flexibility, and TFL defines the Get/Put constructors to handle shared state. TFL also has constructors for concurrency, Y (Yield) and GetTid: Y enables a thread to yield to the scheduler, and GetTid returns the thread id of the current thread.

Scheduler. We define a scheduler as a *program* in SFL, parameterized by the initial thread id and a finite set of thread ids, $\text{Fin}(\text{th})$. SFL extends FL with the $\text{Exec}(n)$ constructor, which *executes* the thread with id n . When a thread executes, it proceeds until it yields with Y or terminates with **ret**, upon which control is returned to the scheduler; this process repeats until the thread pool is empty (or a thread gets **stuck**). One advantage of modeling the scheduler as a program is that we can implement various *scheduling policies* in SFL, and we can fix an *ideal fair scheduler* for an operational spec of fair schedulers.

Interpreting Concurrency. A concurrent system is *interpreted* into a FL program by the scheduler interpreter SI and the thread interpreter TI. The interpreters leave the common constructors (FAIR, PICK, Obs) as they are and only interpret the added constructors, such as Get, Put, and Y . We focus on the most important detail: how the thread interleavings are realized by the interpreters.

The thread interpreter TI enables thread interleaving using the *decorated return type* $R^\vee$. It is a disjoint union of two cases, (i) termination with a value r ($\blacktriangle r$) and (ii) yield with a continuation t ($\vee t$). By distinguishing the two cases, the scheduler can correctly update the thread pool, as described in the interpreting rule for $\text{Exec}(n)$ in SI: First, the scheduler executes thread n and waits for it to yield, which returns a decorated value in $R^\vee$ together with an updated shared state. Then the scheduler inspects the decorated value to (i) kick out a terminated thread ($\blacktriangle$) or (ii) update the thread pool with the continuation of the yielded thread ($\vee$). Therefore, the scheduler tracks the continuation of each thread, being able to interleave the execution of threads in the pool. All in all, the concurrency interpreter CI interprets a concurrent system into a FL program.

5.2 Scheduler Fairness, Operationally

Scheduler fairness guarantees that every thread *eventually* gets scheduled infinitely often. We define scheduler fairness operationally by defining a fair specification for a scheduler FAIRSCh (Fig. 6), and say that a scheduler guarantees scheduler fairness when it refines FAIRSCh.

A Spec for Fair Schedulers. FAIRSCh abstracts what one would expect from a scheduler through nondeterminism (PICK) and mathematical sets: it is a loop that terminates when all the threads have terminated (L6), executing a single thread each iteration. What makes it an *abstract and fair*

$\text{FAIRSch}(n \in \text{th}, \text{ths} \in \text{Fin}(\text{th})) \in \text{SFL} \triangleq$ $\text{loop } \{ x \leftarrow \text{Exec}(n); \text{match } x \text{ with} \quad //L1$ $\quad \text{None} \Rightarrow n' \leftarrow \text{PICK}(\{m \mid m \in \text{ths} \cup \{n\}\}); \quad //L2$ $\quad \quad \text{ths}' := \text{ths} \cup \{n\} \setminus \{n'\}; \quad //L3$ $\quad \quad \text{FAIR}([n' \mapsto \text{good}, \text{ths}' \mapsto \text{bad}]); \quad //L4$ $\quad \quad n := n'; \text{ths} := \text{ths}'; \quad //L5$ $\quad \text{Some } r \Rightarrow \text{if } \text{ths} = \emptyset \text{ then ret } r \quad //L6$ $\quad \quad \text{else } n' \leftarrow \text{PICK}(\{m \mid m \in \text{ths}\}); \quad //L7$ $\quad \quad \text{ths}' := \text{ths} \setminus \{n'\}; \quad //L8$ $\quad \quad \text{FAIR}([n' \mapsto \text{good}, \text{ths}' \mapsto \text{bad}]); \quad //L9$ $\quad \quad n := n'; \text{ths} := \text{ths}'; \quad //L10$	$\text{FIFOSch}(n \in \text{th}, \text{ths} \in \text{Fin}(\text{th})) \in \text{SFL} \triangleq$ $q := \text{set2queue}(\text{ths});$ $\text{loop } \{ x \leftarrow \text{Exec}(n); \text{match } x \text{ with}$ $\quad \text{None} \Rightarrow (n, q) := \text{pop}(\text{push}(n, q));$ $\quad \text{Some } r \Rightarrow \text{if } q = [] \text{ then ret } r$ $\quad \quad \text{else } (n, q) := \text{pop}(q); \}$ $\text{set2queue} \in \text{Fin}(\text{th}) \rightarrow \text{queue th}$ $\text{pop} \in \text{queue } A \rightarrow A \times \text{queue } A$ $\text{push} \in A \times \text{queue } A \rightarrow \text{queue } A$
--	---

Fig. 6. Definition of a spec of fair schedulers FAIRSch and a simple round-robin scheduler FIFOSch.

$\text{OTFL}_{\text{ID}, R, \text{ST}} \stackrel{\text{coind}}{=} \text{Call}(fn \in \text{string}, \text{args} \in \text{list Val}) \gg (k \in \text{Val} \rightarrow \text{OTFL}) \dots$ $M \in \text{Mod}_{\text{ID}, \text{ST}} \triangleq \{(\text{init}, \text{funs}) \in \text{ST} \times (\text{string} \xrightarrow{\text{fin}} (\text{list Val} \rightarrow \text{OTFL}_{\text{ID}, \text{Val}, \text{ST}}))\}$ $\text{Config} \triangleq \text{th} \xrightarrow{\text{fin}} (\text{string} \times \text{list Val}) \quad \text{Load} \in \text{Config} \rightarrow \text{Mod} \rightarrow \text{TPool}$
$\leq_m \in \mathbb{P}(\text{Mod} \times \text{Mod}) \quad M_1 \circ M_2 \in \text{Mod}_{\text{ID}_1 + \text{ID}_2, \text{ST}_1 \times \text{ST}_2} \quad M_1[M_2] \in \text{Mod}_{\text{ID}_1 + \text{ID}_2, \text{ST}_1 \times \text{ST}_2}$ $M_1 \circ M_2 \leq_m M_2 \circ M_1 \quad M_1 \leq_m M_2 \rightarrow M \circ M_1 \leq_m M \circ M_2 \quad M_t \leq_m M_s \rightarrow M_c[M_t] \leq_m M_c[M_s]$

Fig. 7. Definitions of our module system and properties of the module simulation.

spec is the fact that it schedules threads *non-deterministically*, hiding implementation details such as queues (L2, L7); and ensures fairness by invoking **bad** for the unscheduled threads (L4, L9).

We now present our formalization of *scheduler fairness*:

Definition 5.1 (Scheduler Fairness). We say a scheduler *Sch* guarantees *scheduler fairness* when $\text{IsFairSch}(Sch)$ holds. IsFairSch is defined as follows:

$$\text{IsFairSch}(Sch) \triangleq \forall P \text{ st. } \text{CI}(P, Sch, st) \sqsubseteq \text{CI}(P, \text{FAIRSch}, st)$$

Example. To illustrate a concrete scheduler that guarantees scheduler fairness, we present a simple round-robin scheduler FIFOSch (Fig. 6). This scheduler uses a *queue* to schedule the threads in a first-in-first-out manner. Since every thread always gets scheduled after other threads in the thread pool get scheduled, which happens only finitely many times, it is clear that this scheduler guarantees fairness. We formally establish such a guarantee using the above definition:

THEOREM 5.2 (FIFOSch IS FAIR). $\text{IsFairSch}(\text{FIFOSch})$ holds.

This theorem states that FIFOSch fairly refines FAIRSch, which corresponds to fairness validation: we wish to show that FIFOSch satisfies the fairness requirements set by FAIRSch using simulation. Thus the key rule application in the proof of Theorem 5.2 becomes **SIMFS** presented in §2.3, and the main challenge is to find a suitable value of c' to update the fairness counter with. In the case of FIFOSch, c' can be set to the maximum number of threads, denoted M : intuitively, the FIFO queue cannot contain more than M threads whenever a thread is enqueued. Thus a thread scheduled via FIFOSch can only trigger less than M **bad** events within the proof before it gets scheduled and triggers a **good** event: using M as the value of c' when applying **SIMFS** captures this intuition.

6 MODULE SYSTEM AND FAIR SPECS

For reusability and modularity, FOS provides a module system (Fig. 7) inspired by [Gu et al. 2015] that comprises (i) module type Mod , consisting of module-local state and (possibly open) module functions, (ii) linking operation $\circ$ for modules, and (iii) close operation $M_1[M_2]$ that closes open functions in a module M_1 upon getting a module M_2 . Open module functions are written in OTFL, which extends TFL with Call ; the close operation $M_1[M_2]$ interprets $\text{Call}(fn, \text{args})$ s in M_1 by

substituting them with the corresponding TFLs in M_2 . Also, we define a configuration *Config*, which maps each thread id to a function name and arguments. *Load* takes a configuration and a module, initiates TFL for each thread by the function name and arguments, and outputs a thread pool.

What underlies the power of modularization provided by our module system is the module simulation $\lesssim_m$. Module simulation requires the user to prove *thread-local simulation* (see §7) for each pair of functions within the source-target modules, and establishes the refinement. Since the module linking and close operations respect $\lesssim_m$ (Fig. 7), this result implies *contextual refinement*.

THEOREM 6.1 (ADEQUACY OF MODULE SIMULATION). *For a pair of modules M_t and M_s , if $M_t \lesssim_m M_s$ holds, then for any configuration $p \in \text{Config}$, refinement under scheduler fairness holds:*

$$\text{CI}(\text{Load } p M_t, \text{FAIRSCh}, M_t.\text{init}) \sqsubseteq \text{CI}(\text{Load } p M_s, \text{FAIRSCh}, M_s.\text{init})$$

6.1 Memory Modules and Memory Fairness

One key example of the module system is *memory modules* inspired by [Song et al. 2023]. Defining memory models as modules grants us the flexibility required to instantiate concurrent programs with different kinds of memory models, e.g., with a SC memory module or a weak memory module.

Then, as discussed in §3.1, one can prove that the ticket lock module under either SC / weak memory refines the spec. However, proving that the refinement holds under weak memory requires an extra fairness assumption: *memory fairness* [Lahav et al. 2021]. This is because, under weak memory, a value written to memory must be propagated for other threads to read it (e.g., by a cache coherence protocol)—which is not always guaranteed: a thread may never read the most recent update under a buggy protocol, which may render a thread unable to acquire a lock because it cannot read the most recent service value! Memory systems that rule out such cases are said to guarantee memory fairness, and threads eventually obtain newer values under memory fairness.

To model memory fairness in FOS, we develop a fair weak memory module FWMM under *view semantics*, a “promise-free” fragment of the promising semantics [Kang et al. 2017; Lee et al. 2020] without fences, which can be also seen as a fragment of an operational version of RC11 [Lahav et al. 2017]³. In view semantics, each memory location has a *history* of written values with *timestamps*. Each thread has its own *view of memory*; intuitively, the view points to the most recently propagated value to that thread and increases monotonically following execution. A thread can only read the same or newer values than its current view, and can write to memory with newer timestamps than the current view. The following shows an example memory view for some thread k and location X :

$$\begin{array}{cccccccccccc} X : & \boxed{v_0} & \boxed{v_1} & \boxed{v_2} & \boxed{v_3} & \boxed{v_4} & \boxed{v_5} & \boxed{v_6} & \dots & \boxed{v_n} \\ \text{timestamp :} & t_0 & t_1 & t_2 & t_3 & t_4 & t_5 & t_6 & \dots & t_n \end{array}$$

In the figure, k ’s current view for X is t_2 and there are newer values that have not yet propagated to k ($t_3 \sim t_n$). When thread k accesses X , it can update its timestamp to one of t_2 to t_n .

Based on the figure, we demonstrate how we encode memory fairness in our model. Suppose that thread k reads from X , which resulted in updating its timestamp at X from t_2 to t_4 . At this moment, **bad** is invoked for every unpropagated timestamp, i.e., $\text{FAIR}(\{t_i \mid 5 < i \leq n\} \mapsto \text{bad})$ is triggered. This guarantees that every timestamp is eventually propagated to the thread k upon continuous access to X since infinite continuous **bad** is unfair.

Returning to the ticket lock example from §3, we can prove that a ticket lock refines *ABSLock* thanks to the fairness guarantee of FWMM:

³In the Coq development, we adopt the model developed in [Cho et al. 2022].

THEOREM 6.2 (TICKETLOCK IS FAIR). $\text{TicketLock}_{\text{FWMM}}$, a ticket lock module under FWMM, is simulated by the spec ABSLock : $\text{TicketLock}_{\text{FWMM}} \lesssim_m \text{ABSLock}$. Thus it contextually refines ABSLock .⁴

6.2 Expressing Various Concepts of Fairness

So far, we have demonstrated that FOS can model fair semantics of various systems, such as concurrency, library specifications, and memory models. However, the expressiveness of FOS reaches much further, being capable of describing various concepts of fairness proposed in the literature. We provide three examples, all assuming scheduler fairness: starvation/deadlock freedom, strong/weak fairness, and readers-writers problem.

Starvation/Deadlock Freedom. Starvation/Deadlock free concurrent objects guarantee certain progress of threads accessing the object: *starvation* freedom guarantees that *every* thread eventually makes progress, and *deadlock* freedom guarantees that *some* thread eventually makes progress [Herlihy and Shavit 2011]. We can express these properties in FOS, as illustrated by the following ($X++$ is executed atomically):

$\text{def incr}_{\max}() = \text{ret } X++$	$\text{def incr}_{\min}() =$	$\text{while}(\text{PICK}(\mathbb{B})) \{ \text{FAIR}([\alpha \mapsto \text{bad}]); \text{Y}; \}$ $\text{FAIR}([\alpha \mapsto \text{good}]); \text{ret } X++$	(INCR)
--	------------------------------	---	--------

One can easily see that $\text{incr}_{\max}()$ guarantees *starvation freedom*, since any thread calling it immediately returns. More interesting is $\text{incr}_{\min}()$, which guarantees *deadlock freedom* by imposing fairness to single index α : if every callee thread gets stuck in the loop, $[\alpha \mapsto \text{bad}]$ accumulates indefinitely. Hence by fair semantics, *some* thread will eventually get out of the loop, returning with the wanted value. However, that thread invokes a $[\alpha \mapsto \text{good}]$ before returning, which *resets* the accumulated up **bad**s; so other threads stuck in the loop are *not* guaranteed to escape the loop.

Strong/Weak Fairness. Strong/Weak fairness is about progress of threads under constraints: *strong* fairness guarantees that "every thread that is *enabled infinitely often* gets its turn infinitely often", while *weak* fairness guarantees that "every thread that is *continuously enabled* gets its turn infinitely often" [Baier and Katoen 2008]. In general, this means that threads can make progress (*cf.* gets its turn) only when some condition is satisfied (*cf.* enabled). Then, each constraint can be expressed in FOS as the following ($X--$ is executed atomically):

$\text{def decr}_{\text{st}}() =$ $n = \text{GetTid}; W = W \cup \{n\};$ $\text{while}(X \leq 0) \{ \text{Y}; \} W = W \setminus \{n\};$ $\text{FAIR}([n \mapsto \text{good}, W \mapsto \text{bad}]); \text{ret } X--$	$\text{def decr}_{\text{wk}}() =$ $\text{while}(X \leq 0) \{ \text{Y}; \} \text{ret } X--$		(DECR)
---	---	--	--------

These functions guarantee that they (i) decrement X and return if some increment function (*e.g.*, **INCR**) increments X to a nonnegative number, or (ii) loop infinitely—they are *enabled* only when $X > 0$. Then it is easy to see that $\text{decr}_{\text{wk}}()$ guarantees *weak* fairness: if $X > 0$ remains true (*e.g.*, the increment function is always called in between) from some point, namely *continuously enabled*, every thread escapes the loop and returns. However, if X becomes 0 infinitely often, the thread may not be able to escape the loop. On the other hand, $\text{decr}_{\text{st}}()$ guarantees *strong* fairness: if $X > 0$ is true *infinitely often* (*e.g.*, the increment function is called infinitely often), some thread escapes the loop, decreasing X . The escaping thread also invokes $[W \mapsto \text{bad}]$ for the threads waiting to get enabled, thereby ensuring that any waiting thread eventually gets its turn. Note that ABSLock also guarantees strong fairness since the lock is eventually acquired if it is freed infinitely often; one can observe that the codes have a similar pattern.

⁴This ABSLock needs slight modifications from the one under an SC memory, since we need to pass around the view. However, this detail is solely due to the view semantics, orthogonal to our discussion regarding fairness.

The Readers-Writers Problem. The Readers-Writers problem is the problem of the mutual exclusion of several threads accessing a shared resource, where "readers" share the resource with other readers and "writers" require exclusive access [Courtois et al. 1971]. We simplify the problem to one in which there is a single shared location X , where readers read from and writers increment. There are two kinds of problems: Problem 1 states that *readers should not be blocked unless a writer is writing*. Using FOS, we can specify an abstract spec that captures problem 1:

<pre>def read₁() = n = GetTid; R = R ∪ {n}; while(PICK($\mathbb{B}$)) { FAIR([n ↦ bad]); Υ; } R = R \ {n}; FAIR([n, W ↦ good]); ret X</pre>	<pre>def write₁() = n = GetTid; W = W ∪ {n}; while(PICK($\mathbb{B}$)) { FAIR([n ↦ bad]); Υ; } W = W \ {n}; FAIR([n, R ↦ good]); X++; ret</pre>
--	--

(P1)

The spec employs two sets to hold reader(R)/writer(W) threads, and those threads wait in a loop, invoking a **bad** for itself for each iteration. What makes this spec interesting is the **good** events, each invoked by the *opposite* class; *readers* invoke $[W \mapsto \text{good}]$ and *writers* invoke $[R \mapsto \text{good}]$. Intuitively, this means that each class *interrupts* the other class from making progress, since a **good** *resets* the **bad**s. Therefore, if readers are holding the resource, every reader will eventually make progress while any writers are blocked, and vice versa when a writer is holding the resource.

On the other hand, problem 2 states that *writers should write as soon as possible*. In other words, writers have priority over readers in accessing the resource. As with the previous case, we can specify an abstract spec that captures problem 2:

<pre>def read₂() = n = GetTid; R = R ∪ {n}; while(PICK($\mathbb{B}$)) { FAIR([n ↦ bad]); Υ; } R = R \ {n}; FAIR([n ↦ good]); ret X</pre>	<pre>def write₂() = FAIR([R ↦ good]); X++; ret</pre>
---	--

(P2)

This time, any writer should be able to write as soon as it wants, so it is not blocked by anyone. However, readers can still be blocked by the writers; thus all writers invoke $[R \mapsto \text{good}]$, where R is the set of blocked readers, to *reset* the accumulated **bad**s of the readers. Note that [D'Ousualdo et al. 2021] specifies these kinds of properties (some operations can prevent termination of others while other operations do not) under the name "impedance".

7 A PROGRAM LOGIC FOR FAIRNESS

In this section, we present a more *modular* and *abstract* interface to the simulation technique presented in §3, which we call **Fairness Logic**.

Fairness logic tackles two issues with the rudimentary simulation technique presented in §3.2. First, while the rudimentary technique achieves mostly thread-local reasoning by condensing all the needed "global" information in the form of `cmap`, the `cmap` itself remains a global object and reasoning around it remains global. Second, the proof of fair refinement often depends on conditions on ghost variables as shown in §3.2.

Fairness logic addresses these issues with the help of *separation logic*. Modern separation logic⁵ solves the first issue by allowing modular reasoning on global state via the notion of ownership and ownership transfer. Separation logic also allows one to introduce user-defined ghost variables conforming to certain protocols defined via the theory of Partial Commutative Monoids (PCM), solving the second issue: our examples indeed reuse a large body of theory previously developed around PCMs [Jung et al. 2018].

⁵We use a non-step-indexed variant of Iris [Jung et al. 2015] resource algebra, developed in CCR [Song et al. 2023].

RULES FOR SOURCE FAIR		
$\succeq (i \in \text{ID}, o \in \text{Ordinal})$	WIN-SRC	LOSE-SRC
$\succeq\text{-SEP}$ $\succeq (i, o_0 \oplus o_1) \dashv\vdash \succeq (i, o_0) * \succeq (i, o_1)$	$\succeq (i, o) \multimap \text{sim}_I(Q, k_s, k_t)$	$\succeq (i, 1) * \text{sim}_I(Q, k_s, k_t)$
MONO $(o \geq o') * \succeq (i, o) \vdash \dot{\vdash} \succeq (i, o')$	$\text{sim}_I(Q, \text{FAIR}([i \mapsto \text{good}]); k_s, k_t)$	$\text{sim}_I(Q, \text{FAIR}([i \mapsto \text{bad}]); k_s, k_t)$
RULES FOR TARGET FAIR		
$\blacklozenge_{q \in (0,1]} (i \in \text{ID}, n \in \mathbb{N})$ $\blacklozenge (i \in \text{ID}) \quad \blacklozenge_{\text{th}} \triangleq \forall i. \blacklozenge (\text{th}_i)$ $\blacklozenge (i) \triangleq \exists n. \blacklozenge_1 (i, n)$	$\blacklozenge\text{-SEP}$ $\blacklozenge_{q_0+q_1} (i, \min(n_0, n_1)) \dashv\vdash \blacklozenge_{q_0} (i, n_0) * \blacklozenge_{q_1} (i, n_1)$	DEC $\blacklozenge_q (i, n) * \blacklozenge (i)$ $\dot{\vdash} \exists n'. \blacklozenge_q (i, n') * (n' < n)$
LOSE-TGT $\blacklozenge (i) \multimap \text{sim}_I(Q, k_s, k_t)$	WIN-TGT $\blacklozenge (i) * (\blacklozenge (i) \multimap \text{sim}_I(Q, k_s, k_t))$	
$\text{sim}_I(Q, k_s, \text{FAIR}([i \mapsto \text{bad}]); k_t)$	$\text{sim}_I(Q, k_s, \text{FAIR}([i \mapsto \text{good}]); k_t)$	
RULES FOR Υ		
YIELD-SRC $\text{sim}_I(Q, k_s, \Upsilon; k_t)$	YIELD-TGT $b \in \mathcal{B} \quad I * (I \multimap \blacklozenge(\text{th}_{tid}) * (\blacklozenge(\text{th}_{tid}) \multimap \blacklozenge_{\text{th}} \multimap \text{sim}_I(Q, (b ? \Upsilon : \text{skip}); k_s, k_t)))$	
$\text{sim}_I(Q, \Upsilon; k_s, \Upsilon; k_t)$	$\text{sim}_I(Q, \Upsilon; k_s, \Upsilon; k_t)$	

Fig. 8. Core rules of fairness logic.

7.1 Core Rules of Fairness Logic

Core rules of fairness logic are presented in Fig. 8, comprising: rules for (i) executing FAIR() in the source, (ii) executing FAIR() in the target, and (iii) executing Υ .

Simulation Weakest Precondition. A central notion in our rules is “simulation weakest precondition” [Gäher et al. 2022]: $\text{sim}_I(Q, k_s, k_t)$ denotes the weakest precondition to simulate k_t (target) against k_s (source) with postcondition Q , under a relational invariant I shared among threads. In the reasoning, one can *rely* on that I holds when it receives control (*i.e.*, at the beginning of the function and after Υ) and should *guarantee* that I holds when it transfers control (*i.e.*, at the end of the function and before Υ). Our simulation weakest precondition satisfies all the standard rules for simulation argument and weakest precondition, which we omit here.

Rules for executing source FAIR. We start by presenting the fairness assertions used by the source rules. For the source fairness counters, we provide $\succeq(i, o)$ denoting the *right* (or, ownership) to decrement the counter of id i by an ordinal o : this assertion is *local* as it only concerns the id i instead of the global cmap . Also, the $\succeq\text{-SEP}$ rule gives equivalence between $\succeq(i, o_0 \oplus o_1)$ and its decomposition $\succeq(i, o_0) * \succeq(i, o_1)$ where $\oplus$ is the natural sum [Hessenberg 1906] of ordinals. Such a split resource can then be distributed across threads, ensuring local reasoning even when multiple threads are accessing the same i . $\succeq$ also enjoys monotonicity on its second argument (MONO).

Now, the WIN-SRC rule says that one gets $\succeq(i, o)$ for o of its choice when winning. Such a rule is designed with stress on usability: it does not require *any* $\succeq$ as a precondition. This in turn means $\succeq$ in your frame remains valid after the rule, and this is still sound because under the hood it *increments* the counter by o , instead of setting it exactly into o . The LOSE-SRC rule says that it consumes a right to decrement i by one, and actually decrements the counter by one under the hood. Note how these rules together imposes **fairness validation** as expected.

Rules for executing target FAIR. Rules for executing target FAIR follow a similar spirit to those for the source. Nonetheless, the rules are not exactly symmetric because of a different nature between fairness validation and fairness exploitation.

Fairness assertions that the target rules will use are twofold: (i) $\blacklozenge_q(i, n)$ for an id i , a fraction q , and a natural number n denotes knowledge that the counter for i is at most n and a fractional

ownership q to execute **good**, and (ii) $\diamond(i)$ for an id i denotes a “receipt” for decreasing the counter for i by one. As before, we have a rule ($\blacklozenge$ -SEP) decomposing $\blacklozenge$ but with a twist: since n does not refer to the value of the counter but instead means the upper bound for the counter, n does not get split into two when splitting and we take the minimum of the two n s when merging two $\blacklozenge$ s. The most interesting rule in the target side is the DEC rule, concerning **fairness exploitation**: it consumes a $\blacklozenge$ and a $\diamond$ to produce a $\blacklozenge$ with a decreased number. Such a rule coincides with intuitive interpretation of $\blacklozenge$ and $\diamond$, and allows fairness exploitation since a number cannot decrease indefinitely.

Now, the LOSE-TGT rule says that one gets $\diamond(i)$ when id i triggers a bad event. Again, such a rule is designed with usability in mind: a direct, lower-level reasoning would dictate the exact counter value for i before/after execution, but would not be ideal for local reasoning (e.g., when multiple threads trigger bad events on the same id, some synchronization has to be made to track the latest value). Here the rule requires *nothing* in the precondition: this is because it does not need the exact value of i , but only the fact that a decrement has been made. The WIN-TGT rule consumes the full fraction of $\blacklozenge$ with any value n and returns the full fraction of $\blacklozenge$ with an unknown value n . Note the difference with WIN-SRC: the counter value for i is updated to an unknown value, and the rule contains the full fraction in the precondition, which is required for soundness.

Rules for executing Υ . Finally, we give rules about executing Υ that are capable of reasoning about *scheduler fairness* in a thread-local fashion.

The YIELD-SRC rule allows executing Υ in the source as if it is “skip”, in the presence of the Υ in the target. Now, consider scheduler fairness: the rule needs to somehow impose fairness validation regarding scheduler fairness events. For this, instead of baking in WIN-SRC and LOSE-SRC rules into this rule as-is, we give a much simpler yet expressive enough interface: following Simuliris [Gäher et al. 2022], we use an inductive-coinductive definition that allows using the YIELD-SRC rule only finitely unless coinductive progress is made in the YIELD-TGT rule (below). This allows, in the majority of verification scenarios, the user to completely ignore proof obligations regarding (scheduler) fairness validation.

On the other hand, the YIELD-TGT rule executes Υ in the target. When b is false, it simply executes Υ on both sides in lock-step with the usual rely-guarantee principle on I (it demands I to hold before and gives it back after). When b is true, it just executes Υ in the target (keeping Υ in the source): this flexibility is useful as it allows one to execute multiple target Υ s with a single source Υ . Now, consider scheduler fairness. Recall that after the current thread, tid , takes the control back, a fairness event for winning tid and losing everyone else is invoked. The rule has applications of WIN-TGT and LOSE-TGT to this event baked-in, which appears as $\blacklozenge(\text{th}_{tid})$ and $\diamond_{\text{th}} \cdot \diamond_{\text{th}}$ is simply a conjunction of $\diamond$ for *all* possible thread ids.

Adequacy. Fairness logic satisfies the following two adequacy theorems.

THEOREM 7.1 (CONTEXTUAL ADEQUACY). *For a pair of module M_t, M_s , and a relational invariant I , if the initial states of the modules satisfy I , we have:*

$$(\forall tid f v. I * \blacklozenge(\text{th}_{tid}) \vdash \text{sim}_I(I * \blacklozenge(\text{th}_{tid}), M_s.\text{funs } f v, M_t.\text{funs } f v)) \implies M_t \lesssim_m M_s$$

Here, $I * \blacklozenge(\text{th}_{tid})$ plays essentially the same rely-guarantee reasoning with the YIELD-TGT rule.

THEOREM 7.2 (WHOLE PROGRAM ADEQUACY). *For a pair of modules M_t, M_s , a relational invariant I , a whole-program configuration $p \in \text{Config}$, and a precondition P_{tid} for each thread id tid , if the initial states of the modules satisfy $*_{tid \in \text{dom}(p)} \blacklozenge(\text{th}_{tid}) \multimap (I * *_{tid \in \text{dom}(p)} P_{tid})$, we have:*

$$(\forall tid f v. (p \text{ tid} = (f, v)) * I * P_{tid} \vdash \text{sim}_I(I * \blacklozenge(\text{th}_{tid}), M_s.\text{funs } f v, M_t.\text{funs } f v)) \implies \text{CI}(\text{Load } p M_t, \text{FAIRSCh}, M_t.\text{init}) \sqsubseteq \text{CI}(\text{Load } p M_s, \text{FAIRSCh}, M_s.\text{init})$$

This theorem additionally allows each thread to have its precondition, P_{tid} , but can only be used for whole-program refinement, not for contextual refinement.

We conclude this section with the following remark. While the fairness logic contains only the minimal core rules, we are gathering confidence that it is powerful enough to handle various interesting examples: our flagship example, presented in the next section, involves non-trivial reasoning that spans multiple different notions of fairness. We believe further abstract constructs [D’Oswaldo et al. 2021] could be derived on top of fairness logic and leave it as an interesting future work.

7.2 Example using the Fairness Logic

We demonstrate how the fairness logic is applied with a simplified but still illustrative example. The example focuses on *exploiting* fairness; for the case of validating fairness, we believe that the explanation in the previous section should be adequate for actual applications. This section assumes some familiarity with modern separation logic, e.g., Iris [Jung et al. 2015], and uses the usual separation logic predicates with minimal explanations.

To show how to use the fairness logic in proofs relying on exploiting fairness, we inspect a simplified version of CL_I , which assumes that memory access is atomic and removes the locks:

$$\Psi; X := 42; \Psi; \parallel \text{do } \{ \Psi; x := X; \Psi; \} \text{ while } (x = 0) \Psi; \text{print}(x); \Psi; \quad (\text{SCL}_I)$$

This program refines CL_S , and the reasoning underlying the proof is similar to the one presented in §3.2: we construct a tuple that decreases throughout the program execution and perform induction on it. However, with the power of fairness logic, we can carry out the proof in a thread-local fashion—the most interesting proof obligation now becomes constructing a shared invariant that captures the rely-guarantee reasoning.

To begin with, we present an invariant one would set up, but without the fairness assertions:

$$(X \mapsto 0) \vee (X \mapsto 42 * \text{Ex})$$

where $\mapsto$ is the usual points-to and Ex is a token for exclusive ownership in separation logic. This invariant states that the value stored in the memory location X is either 0 or 42, and in the latter case it also holds the token Ex , which means that thread 1 has written to X : thread 1 starts with Ex and hands it over to the invariant right after the write to X . Unfortunately, this invariant is too weak for the proof since it does not capture scheduler fairness; for the **while** loop in thread 2 to always terminate, it should be guaranteed that thread 1 will eventually write 42 to X , which in turn relies on scheduler fairness.

This intuition is precisely captured with a fairness assertion $\blacklozenge_1(\text{th}_1, n)$, representing that the fairness counter for thread 1 (th_1) is at most n . Then a correct invariant would be:

$$(X \mapsto 0 * \exists n. (\blacklozenge_1(\text{th}_1, n) * =(n))) \vee (X \mapsto 42 * \text{Ex}) \quad (\text{INV})$$

where $=(n)$, together with $\leq(m)$, are predicates for monotonicity satisfying the following rules:

$$=(n) \vdash =(n) * \leq(n), \quad =(n) * \leq(m) \vdash (n \leq m), \quad =(m) * (n \leq m) \vdash \dot{=}(n)$$

Intuitively, $=(n)$ denotes the exact value of n , which can monotonically decrease, and $\leq(m)$ denotes that m is an upper bound of n [Timany and Birkedal 2021]. Then INV states that the value at X is either 0 or 42, and when it is 0, there is a value n that represents how many times thread 1 can be starved by the scheduler ($\blacklozenge_1(\text{th}_1, n)$) and monotonically decreases ($=(n)$).

With INV , we demonstrate the core of the proof that SCL_I refines CL_S —that the while-loop in thread 2 always terminates. Let us focus on **do** $\{ \Psi; x := X; \Psi; \}$ **while** $(x = 0)$. For this piece of code, we stutter the source with the first Ψ during simulation. Then at $x := X;$, we destruct INV and get two cases: (i) $X \mapsto 42$ case terminates the loop by making the loop condition false, leaving (ii) $X \mapsto 0$ case, with a variable n , $\blacklozenge_1(\text{th}_1, n)$, and $=(n)$. In this case, we perform **(strong) induction**

on n to conclude the proof, and we first obtain $\leq(n)$ from $=(n)$. Next, we proceed to execute Υ using YIELD-TGT and obtain $\diamond_{\text{th}}$ —note that thread 2 owns $\diamond(\text{th}_2)$ from the start, and we can easily show the $X \mapsto 0$ case to guarantee the invariant **INV**. Then the loop iterates, almost bringing us to the state satisfying the induction hypothesis (after a trivial application of YIELD-TGT).

This time, we have $\leq(n)$ and $\diamond_{\text{th}}$ along with **INV**, allowing us to finish the induction: we can obtain n' which satisfies $n' < n$ and the induction hypothesis. After destructing the invariant, we again inspect the $X \mapsto 0$ case, where it gives a fresh variable m , $\diamond_1(\text{th}_1, m)$, and $=(m)$. First, $=(m)$ and $\leq(n)$ gives $m \leq n$ and case analysis leaves us with $m = n$ case since $m < n$ case concludes the induction. Next, substituting m with n , we apply DEC to $\diamond_1(\text{th}_1, n)$ and $\diamond_{\text{th}}$, obtaining $\diamond_1(\text{th}_1, n')$ where $n' < n$. Finally, after updating $=(n)$ to $=(n')$, we can conclude the induction since $n' < n$.

This line of reasoning is seamlessly extended to a tuple instead of a single value. To illustrate this, consider the following modification:

$\Upsilon; \text{skip}; \Upsilon; X := 42; \Upsilon; \parallel \text{do } \{ \Upsilon; x := X; \Upsilon; \} \text{ while } (x = 0) \Upsilon; \text{print}(x); \Upsilon; \quad (\text{SCL}'_1)$

Because of the inserted skip, we now need to incorporate the number of remaining Υ s in thread 1, denoted l , into the induction (as we did in §3.2). To encode l in the invariant, we use the usual authoritative assertions $\bullet(l)$ and $\circ(l)$, satisfying the following rules [Jung et al. 2018]:

$$\bullet(a) * \circ(b) \vdash (a = b), \quad \bullet(a) * \circ(b) \vdash \Rrightarrow \bullet(c) * \circ(c)$$

Then an invariant for proving that **SCL'**₁ refines **CL**_S is:

$$(X \mapsto 0 * \exists n, l. (\diamond_1(\text{th}_1, n) * \bullet(l) * =(l, n))) \vee (X \mapsto 42 * \text{Ex}) \quad (\text{INV}')$$

The other half, $\circ(l)$, is owned by thread 1 and tracks the number of remaining Υ s in thread 1. Using **INV'**, we can prove the goal with induction; we omit the details, which is similar to the one above.

8 CASE STUDY

Composing all of our results, we can prove the motivating example: **CL**_I using TicketLock_{FWM} and under FIFOSch refines **CL**_S under FAIRSCh. To state this, we first wrap up each as a module, **CL**_{I, tk} and **CL**_S, which contains appropriate functions and an initial state. Then we load the modules with a configuration p that maps threads 1 and 2 to corresponding functions, and state the desired refinement, which we prove by assembling the refinement results by the transitivity of refinement:

$$\begin{aligned} \text{CI}(\text{Load } p \text{ CL}_{I, \text{tk}}, \text{FIFOSch}, \text{CL}_{I, \text{tk}}.\text{init}) &\sqsubseteq \text{CI}(\text{Load } p \text{ CL}_{I, \text{tk}}, \text{FAIRSCh}, \text{CL}_{I, \text{tk}}.\text{init}) \\ &\sqsubseteq \text{CI}(\text{Load } p \text{ CL}_{I, \text{abs}}, \text{FAIRSCh}, \text{CL}_{I, \text{abs}}.\text{init}) \\ &\sqsubseteq \text{CI}(\text{Load } p \text{ CL}_S, \text{FAIRSCh}, \text{CL}_S.\text{init}) \end{aligned}$$

where the first and second refinements are direct applications of Theorems 5.2 and 6.2. Theorem 6.2 and the final refinement are proved using fairness logic (Theorems 7.1 and 7.2). Also, we remark that client-library modular reasoning manifests itself in the application of Theorem 6.2 at the second refinement. This result, including the whole theory of FOS, is fully mechanized in Coq.

9 RELATED WORK AND DISCUSSION

Various concepts of fairness have been studied within the literature, including scheduler fairness [Lamport 1977; Lehmann et al. 1981], progress properties for concurrent objects [Courtois et al. 1971; Herlihy and Shavit 2011; Liang and Feng 2017], weak memory models [Lahav et al. 2021], and model checking [Baier and Katoen 2008].

As shown throughout the paper, the definition of fairness described in Definition 2.1 is general enough to capture all of these concepts of fairness, which in turn makes these concepts expressible in FOS. In this section, we describe some pieces of previous work in more detail and compare the advantages (and disadvantages) that FOS provides as a framework compared to existing work.

Progress Properties of Concurrent Programs. An interesting line of work in establishing progress properties of a concurrent program is to prove that a program refines some operational specification that encodes the desired progress property [Gotsman and Yang 2011]. A prime example is LiLi [Liang and Feng 2016, 2017], which provides a method to express such specs, then allows users to prove that a program contextually refines the spec (thereby ensuring that a progress property holds).

However, the specs that LiLi provides do not express fairness directly, and instead employ an explicit queue of definite actions within the semantics to capture the idea of fairness. LiLi also does not provide any machinery for verifying clients that rely on fairness assumptions outside of scheduler fairness—and thus cannot exploit fairness, as we have done in FOS.

Liang et al. [2013] studies relations between all common progress properties and operational definitions, based on contextual refinements. They prove that each progress property is equivalent to some specific type of contextual refinement for linearizable objects. We believe that this approach can be applied to our theory, such that proving contextual refinements can formally establish desired progress properties, which we leave for future work.

Liveness and Temporal Logics. Fairness shares many similarities with *liveness*, which states that a ‘good thing’ must eventually happen at some point during the execution of a program: for example, that a process must ‘get scheduled’ eventually. Such liveness properties are typically encoded via temporal logics [Owicki and Lamport 1982]—most prominently, linear temporal logic (LTL) [Kesten et al. 1998]—which state properties that an entire trace of a program execution must satisfy.

We note that while fairness and liveness are similar concepts with many overlapping applications, the definition of fairness presented in this paper does have differences with the standard notion of liveness: liveness requires that a *good* event *must* happen, while with fairness, it is fine that a *good* event does not occur as long as each fairness id only accumulates a finite number of *bad* events.

Concurrent Separation Logic. Concurrent separation logics [Brookes 2007; O’Hearn 2007] extend modern separation logics with the goal of proving safety properties of concurrent programs. One notable feature of concurrent separation logics is that they focus on thread-local and compositional reasoning, in order to reduce proof burdens.

One notable work in this area is TaDA-Live [D’Ousualdo et al. 2021], which extends [Rocha Pinto et al. 2016], a concurrent separation logic for verifying total correctness of client programs using concurrent objects. TaDA-Live provides fairness-aware specifications of concurrent objects in the style of Hoare triples, and exploits scheduler fairness during verification, with a focus on proving that a set of threads terminate under scheduler fairness. Such termination proofs often rely on induction, where one must identify a decreasing value to perform induction upon (such as the tuple of values from §3). While abstracted away in this paper, identifying and constructing these values presents the main challenge in FOS proofs; a major contribution of TaDA-Live is that it provides a high degree of abstraction to hide this complexity, providing users with a simple proof interface.

However, TaDA-Live cannot provide any guarantees for nonterminating programs: for example, TaDA-Live cannot be used to express progress properties for programs that loop indefinitely. TaDA-Live is also only capable of unary reasoning and cannot be used for proving refinement.

There has been work on proving termination-preserving refinement under a fair scheduler using relational separation logic [Gäher et al. 2022; Tassarotti et al. 2017]. However, such approaches do not support reasoning about fairness directly, or exploiting fairness.

ACKNOWLEDGMENTS

Supported by Samsung Research Funding Center of Samsung Electronics under Project Number SRFC-IT2102-03. Chung-Kil Hur is the corresponding author. We thank Sung-Hwan Lee and Ori Lahav for feedback and initial discussions on this paper.

DATA AVAILABILITY STATEMENT

The Coq formalization of this paper may be found at [Lee et al. 2023].

REFERENCES

- Christel Baier and Joost-Pieter Katoen. 2008. *Principles of Model Checking*. The MIT Press.
- Stephen Brookes. 2007. A semantics for concurrent separation logic. *Theoretical Computer Science* 375, 1-3 (2007), 227–270.
- Minki Cho, Sung-Hwan Lee, Dongjae Lee, Chung-Kil Hur, and Ori Lahav. 2022. Sequential Reasoning for Optimizing Compilers under Weak Memory Concurrency. In *Proceedings of the 43rd ACM SIGPLAN International Conference on Programming Language Design and Implementation* (San Diego, CA, USA) (PLDI 2022). Association for Computing Machinery, New York, NY, USA, 213–228. <https://doi.org/10.1145/3519939.3523718>
- P. J. Courtois, F. Heymans, and D. L. Parnas. 1971. Concurrent Control with “Readers” and “Writers”. *Commun. ACM* 14, 10 (oct 1971), 667–668. <https://doi.org/10.1145/362759.362813>
- Emanuele D’Osualdo, Julian Sutherland, Azadeh Farzan, and Philippa Gardner. 2021. TaDA Live: Compositional Reasoning for Termination of Fine-Grained Concurrent Programs. *ACM Trans. Program. Lang. Syst.* 43, 4, Article 16 (nov 2021), 134 pages. <https://doi.org/10.1145/3477082>
- Lennard Gäher, Michael Sammler, Simon Spies, Ralf Jung, Hoang-Hai Dang, Robbert Krebbers, Jeehoon Kang, and Derek Dreyer. 2022. Simuliris: A Separation Logic Framework for Verifying Concurrent Program Optimizations. *Proc. ACM Program. Lang.* 6, POPL, Article 28 (jan 2022), 31 pages. <https://doi.org/10.1145/3498689>
- Alexey Gotsman and Hongseok Yang. 2011. Liveness-Preserving Atomicity Abstraction. In *Proceedings of the 38th International Conference on Automata, Languages and Programming - Volume Part II* (Zurich, Switzerland) (ICALP’11). Springer-Verlag, Berlin, Heidelberg, 453–465.
- Ronghui Gu, Jérémie Koenig, Tahina Ramananandro, Zhong Shao, Xiongnan (Newman) Wu, Shu-Chun Weng, Haozhong Zhang, and Yu Guo. 2015. Deep Specifications and Certified Abstraction Layers. In *Proceedings of the 42nd ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages* (POPL 2015).
- Maurice Herlihy and Nir Shavit. 2011. On the Nature of Progress. In *Principles of Distributed Systems*, Antonio Fernández Anta, Giuseppe Lipari, and Matthieu Roy (Eds.). Springer Berlin Heidelberg, Berlin, Heidelberg, 313–328.
- Gerhard Hessenberg. 1906. *Grundbegriffe der mengenlehre*. Vol. 1. Vandenhoeck & Ruprecht.
- Ralf Jung, Robbert Krebbers, Jacques-Henri Jourdan, Aleš Bizjak, Lars Birkedal, and Derek Dreyer. 2018. Iris from the ground up: A modular foundation for higher-order concurrent separation logic. *Journal of Functional Programming* 28 (2018).
- 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. *ACM SIGPLAN Notices* 50, 1 (2015), 637–650.
- Jeehoon Kang, Chung-Kil Hur, Ori Lahav, Viktor Vafeiadis, and Derek Dreyer. 2017. A Promising Semantics for Relaxed-Memory Concurrency. In *Proceedings of the 44th ACM SIGPLAN Symposium on Principles of Programming Languages* (Paris, France) (POPL ’17). Association for Computing Machinery, New York, NY, USA, 175–189. <https://doi.org/10.1145/3009837.3009850>
- Yonit Kesten, Amir Pnueli, and Li-on Raviv. 1998. Algorithmic verification of linear temporal logic specifications. In *Automata, Languages and Programming: 25th International Colloquium, ICALP’98 Aalborg, Denmark, July 13–17, 1998 Proceedings* 25. Springer, 1–16.
- Ori Lahav, Egor Namakonov, Jonas Oberhauser, Anton Podkopaev, and Viktor Vafeiadis. 2021. Making Weak Memory Models Fair. *Proc. ACM Program. Lang.* 5, OOPSLA, Article 98 (oct 2021), 27 pages. <https://doi.org/10.1145/3485475>
- Ori Lahav, Viktor Vafeiadis, Jeehoon Kang, Chung-Kil Hur, and Derek Dreyer. 2017. Repairing Sequential Consistency in C/C++ 11. In *38th ACM SIGPLAN Conference on Programming Language Design and Implementation*. ACM, 618–632.
- L. Lamport. 1977. Proving the Correctness of Multiprocess Programs. *IEEE Transactions on Software Engineering* SE-3, 2 (1977), 125–143. <https://doi.org/10.1109/TSE.1977.229904>
- Dongjae Lee, Minki Cho, Jinwoo Kim, Soonwon Moon, Youngju Song, and Chung-Kil Hur. 2023. PLDI23 artifact for Fair Operational Semantics. <https://doi.org/10.5281/zenodo.7711063>
- Sung-Hwan Lee, Minki Cho, Anton Podkopaev, Soham Chakraborty, Chung-Kil Hur, Ori Lahav, and Viktor Vafeiadis. 2020. Promising 2.0: Global Optimizations in Relaxed Memory Concurrency. In *Proceedings of the 41st ACM SIGPLAN Conference on Programming Language Design and Implementation* (London, UK) (PLDI 2020). Association for Computing Machinery, New York, NY, USA, 362–376. <https://doi.org/10.1145/3385412.3386010>
- D. Lehmann, A. Pnueli, and J. Stavi. 1981. Impartiality, justice and fairness: The ethics of concurrent termination. In *Automata, Languages and Programming*, Shimon Even and Oded Kariv (Eds.). Springer Berlin Heidelberg, Berlin, Heidelberg, 264–277.
- Xavier Leroy. 2006. Formal Certification of a Compiler Back-end or: Programming a Compiler with a Proof Assistant. In *Proceedings of the 33rd ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages* (POPL 2006).

- Hongjin Liang and Xinyu Feng. 2016. A Program Logic for Concurrent Objects under Fair Scheduling. In *Proceedings of the 43rd Annual ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages* (St. Petersburg, FL, USA) (POPL '16). Association for Computing Machinery, New York, NY, USA, 385–399. <https://doi.org/10.1145/2837614.2837635>
- Hongjin Liang and Xinyu Feng. 2017. Progress of Concurrent Objects with Partial Methods. *Proc. ACM Program. Lang.* 2, POPL, Article 20 (dec 2017), 31 pages. <https://doi.org/10.1145/3158108>
- Hongjin Liang, Jan Hoffmann, Xinyu Feng, and Zhong Shao. 2013. Characterizing Progress Properties of Concurrent Objects via Contextual Refinements. In *Proceedings of the 24th International Conference on Concurrency Theory* (Buenos Aires, Argentina) (CONCUR'13). Springer-Verlag, Berlin, Heidelberg, 227–241. https://doi.org/10.1007/978-3-642-40184-8_17
- Peter W. O'Hearn. 2007. Resources, Concurrency, and Local Reasoning. *Theor. Comput. Sci.* 375, 1–3 (apr 2007), 271–307. <https://doi.org/10.1016/j.tcs.2006.12.035>
- Susan Owicki and Leslie Lamport. 1982. Proving Liveness Properties of Concurrent Programs. *ACM Trans. Program. Lang. Syst.* 4, 3 (jul 1982), 455–495. <https://doi.org/10.1145/357172.357178>
- Pedro Rocha Pinto, Thomas Dinsdale-Young, Philippa Gardner, and Julian Sutherland. 2016. Modular Termination Verification for Non-Blocking Concurrency. In *Proceedings of the 25th European Symposium on Programming Languages and Systems - Volume 9632*. Springer-Verlag, Berlin, Heidelberg, 176–201.
- Youngju Song, Minki Cho, Dongjae Lee, Chung-Kil Hur, Michael Sammler, and Derek Dreyer. 2023. Conditional Contextual Refinement. *Proc. ACM Program. Lang.* 7, POPL, Article 39 (jan 2023), 31 pages. <https://doi.org/10.1145/3571232>
- Joseph Tassarotti, Ralf Jung, and Robert Harper. 2017. A Higher-Order Logic for Concurrent Termination-Preserving Refinement. In *Programming Languages and Systems: 26th European Symposium on Programming, ESOP 2017, Held as Part of the European Joint Conferences on Theory and Practice of Software, ETAPS 2017, Uppsala, Sweden, April 22–29, 2017, Proceedings* (Uppsala, Sweden). Springer-Verlag, Berlin, Heidelberg, 909–936. https://doi.org/10.1007/978-3-662-54434-1_34
- The Coq Development Team. 2021. The Coq Proof Assistant 8.13.2 Reference Manual. <https://coq.github.io/doc/V8.13.2/refman/>.
- Amin Timany and Lars Birkedal. 2021. Reasoning about Monotonicity in Separation Logic. In *Proceedings of the 10th ACM SIGPLAN International Conference on Certified Programs and Proofs* (Virtual, Denmark) (CPP 2021). Association for Computing Machinery, New York, NY, USA, 91–104. <https://doi.org/10.1145/3437992.3439931>
- Li-yao Xia, Yannick Zakowski, Paul He, Chung-Kil Hur, Gregory Malecha, Benjamin C. Pierce, and Steve Zdancewic. 2019. Interaction Trees: Representing Recursive and Impure Programs in Coq. *Proc. ACM Program. Lang.* 4, POPL, Article 51 (Dec. 2019), 32 pages. <https://doi.org/10.1145/3371119>

Received 2022-11-10; accepted 2023-03-31