

Interaction Trees and Verified Compilation

Extended Abstract

Paolo Torrini, INRIA Sophia-Antipolis, France
Benjamin Gregoire, INRIA Sophia-Antipolis, France
Paolo.Torrini@inria.fr
Benjamin.Gregoire@inria.fr

Abstract

We have used Interaction Trees (ITrees) to formalize and verify the front-end of the Jasmin compiler, relying on an integration of coinductive denotational semantics with Relational Hoare Logic. A crucial element in our approach is the definition of an appropriate relation to compare programs, which makes it possible to simplify the reasoning about undefined behaviour.

ITrees [11] can be used to reason denotationally about possibly diverging programs in an executable, modular way, thus providing an attractive alternative to the operational semantics used in major verified compilation projects such as CompCert [7] and CakeML [6]. The coinductive datatype $\text{itree } E \ V$, parameterized by a type family of events E and a final result type V , represents computations as trees with possibly infinite branches. A leaf $\text{Ret } v$ is termination with value $v : V$. A unary node $\text{Tau } t$ is a silent transition step which continues with $t : \text{itree } E \ V$. A multi-branching node $\text{Vis } A \ e \ k$ is a computation that executes the event $e : E \ A$ and feeds the answer to the continuation $k : A \rightarrow \text{itree } E \ V$. Being monad themselves, ITrees can be executed by layering monadic interpreters, each associated with a handler for a subfamily of events. Generalized recursion can be represented as corecursive interpretation of call events relying on a monadic iterator. Failures can be represented as events that return in the null type, to be interpreted using an error monad transformer. ITrees facilitate coinductive reasoning through a rich equational theory [1]. Modularity is supported by disjoint unions of event families (e.g. $E_1 + E_2$) and interpretation layers, possibly associated with monad transformers.

Jasmin [2, 3] is a comparatively low-level language intended for security applications, designed to give the programmer explicit control on the generation of assembly code. It includes operators and instructions on memory and registers, conditional statements, while loops, internal and external function calls. It is formalized in Rocq with a fully verified compiler based on a big-step semantics for terminating programs and a concrete memory model. In order to lift the termination restriction, we defined a new semantics with ITrees and used it to verify so far the compiler front-end, based on a language that remains syntactically the same while being semantically interpreted differently after each pass (more than twenty ones, including constant propagation, dead code elimination, inlining and stack allocation).

Reasoning about compilation relies on the definition of behavioural equivalence between source and target, which can be used either for forward or backward reasoning. Some subtlety is involved in dealing with abnormal termination. Compilation is meant to preserve normal behaviour, but in general the compiler may not worry about what happens when the source throws a run-time error. Indeed, compiler correctness proofs are often restricted to *safe* executions, i.e. those that do not terminate abnormally. Verified compilation is meant to preserve safety as part of behaviour preservation, and therefore it suffices to assume that the initial source is safe. However, the safety restriction makes proofs less general and can complicate compositionality. In order to lift this restriction, in our verification we use a notion of program equivalence which makes it possible to express that the compiled code behaves like the source only up to events called *cutoffs*. Cutoff events can be used to represent source errors which are meant to be matched with any behaviour in the target, as causal instances of undefined behaviour (UB). Notice that unlike C and LLVM, Jasmin has a deterministic semantics, but interpreting failure as UB helps to make the semantics simpler.¹

The ITree library [1] provides generalizations of bisimilarity, both strong ($\cong_R$) and weak ($\approx_R$, called *equivalence up-to-tau*), parameterized by a relation R between result types, allowing for trees which are heterogeneous in the returned value types to be related, provided they use the same events. Instantiating R with equality gives us the usual strong and weak bisimilarity. The fact that, regardless of R , related events must be equal, makes it hard to directly relate function calls across a translation while keeping into account differences in the types of call arguments and results. A way around this problem, showcased in [4, 11–13] consists in layering interpretations, separating a common ground of ITrees based on the same events which can be compared directly, from higher-level layers which need to be compared in terms of more complex notions of refinement, also using predicates on ITrees [4, 13] to account for non-determinism (a problem we do not need to deal with). A more intensional way to relate programs is given by the event-heterogeneous relation introduced in [10], used in [4, 9] and here denoted $\approx_R^e$, as a generalization of $\approx_R$ which is parameterized by two additional relations: a precondition specifying which events

¹Unlike [4] we do not consider time-travelling UB.

can be related, inclusive of their inputs, and a postcondition which is expected to hold between the outputs of event execution. $\approx_R^e$ makes it possible to capture structural similarity more directly, yet it does not let us interpret source errors as UB, as it does not allow relating a terminating program with a silently non-terminating one. While needed with normal termination, this restriction is unnecessary when dealing with an abnormal one.

We can answer this problem by a generalization of $\approx_R^e$, denoted $\approx_R^u$, which we call *equivalence up-to-cutoff*. As much as $\approx_R$ and $\approx_R^e$, $\approx_R^u$ is coinductively defined as the greatest fixed point of an inductive relation. Its definition is illustrated in Figure 1, where a single bar means that the hypothesis is used inductively, while a double bar means that it is used coinductively (or simply non-inductively). $\approx_R^u$ can relate wholly heterogeneous trees $t_1 : \text{itree } E_1 \ V_1$ and $t_2 : \text{itree } E_2 \ V_2$. The rules RET, TAU, TAU_l, TAU_r, and VIS are those of $\approx_R^e$. RET allows relating return values by R . In the VIS rule, Φ acts as relational precondition and Ψ as postcondition for the event handlers used to interpret the events. As in [10, 11] the fact that rules TAU_l and TAU_r rely on their hypothesis inductively ensures that the related trees can only differ up to a finite number of silent steps. Specifically to $\approx_R^u$, there are as additional parameters two Boolean predicates C^l on E_1 events and C^r on E_2 ones, deciding which are the cutoffs on each side, and two additional rules CUT_l and CUT_r, which allow for a cutoff to be related with any tree.² $\approx_R^u$ supports setoid rewriting, enjoys structural properties, monotonicity, transitivity, and it is preserved by interpretation (Intr) for related handlers (h_1, h_2):

$$t_1 \approx_R^u t_2 \rightarrow (\forall e_1 e_2. h_1 e_1 \approx_R^u h_2 e_2) \rightarrow \text{Intr}_{h_1} t_1 \approx_R^u \text{Intr}_{h_2} t_2$$

Our intended semantics for Jasmin is to relate a source failure to any behaviour of the target program beyond that point. Using $\approx_R^u$, it suffices to declare source failures as cutoff events. We write $\approx_R^u \Phi \Psi C^l C^r$ for the fully parameterized relation. We also write $[p_i]_{s_i}$, depending on state $s_i : S_i$, for the semantic interpretation of program $p_i : \mathcal{L}$ (after the i -th pass) as an ITree of type $\text{itree } (F +' G) \ S_i$ where F is the type family of failure events. Further writing R_i for the relation between states across the pass, Φ_i for the precondition that holds between events and Ψ_i for the corresponding postcondition, C_F for the predicate that sets F events as the only cutoffs and C_0 as the predicate that sets no cutoffs, the verification of a two-pass compilation (where related events may differ even if they belong to the same family) can be built transitively out of two components:

$$\frac{\vdash_{\Gamma} [p_0]_{s_0} \approx_{R_1 \Phi_1 \Psi_1 C_F C_0}^u [p_1]_{s_1} \quad \vdash_{\Gamma} [p_1]_{s_1} \approx_{R_2 \Phi_2 \Psi_2 C_F C_0}^u [p_2]_{s_2}}{\vdash_{\Gamma} [p_0]_{s_0} \approx_{(R_1 \circ R_2) (\Phi_1 \circ \Phi_2) (\Psi_1 \circ \Psi_2) C_F C_0}^u [p_2]_{s_2}}$$

This principle holds regardless of the events in G (which are still uninterpreted). Nonetheless, our current verification

²A more specific form of this relation was already used in [4].

$$\begin{array}{c} \frac{r_1 \ R \ r_2}{\text{Ret}(r_1) \approx^u \text{Ret}(r_2)} \text{RET} \qquad \frac{t_1 \approx^u t_2}{\text{Tau}(t_1) \approx^u \text{Tau}(t_2)} \text{TAU} \\[10pt] \frac{C^l(e_1)}{\text{Vis}(e_1, k_1) \approx^u t_2} \text{CUT}_l \qquad \frac{C^r(e_2)}{t_1 \approx^u \text{Vis}(e_2, k_2)} \text{CUT}_r \\[10pt] \frac{t_1 \approx^u t_2}{\text{Tau}(t_1) \approx^u t_2} \text{TAU}_l \qquad \frac{t_1 \approx^u t_2}{t_1 \approx^u \text{Tau}(t_2)} \text{TAU}_r \\[10pt] \frac{e_1 \ \Phi \ e_2 \quad \forall v_1 v_2. (e_1, v_1) \ \Psi \ (e_2, v_2) \implies k_1(v_1) \approx^u k_2(v_2)}{\text{Vis}(e_1, k_1) \approx^u \text{Vis}(e_2, k_2)} \text{VIS} \end{array}$$

Figure 1. Equivalence up-to-cutoff ($\approx^u$), parameterized by R, Φ, Ψ, C^l and C^r .

makes limited use of event layering, as we leaned significantly on reuse of existing low-level proofs. Indeed, on top of $\approx_R^u$, we relied on a form of Relational Hoare Logic, originally introduced in [5] which we specialized to ITrees, allowing us to obtain proofs that are comparatively similar to those originally made with the inductive big-step semantics. With the notable exception of inlining, it was possible to make our proofs relying on coinductive lemmas that were either provided by [1] or did not pose particular challenges.

Our compiler correctness statements, which we basically proved by induction on source p for the single passes, then gluing the pieces together by transitivity, have general form:

$$\forall p \ s \ s', \ s \ R \ s' \rightarrow [p]_s \approx_{R \Phi \Psi C_F C_0}^u [p']_{s'}$$

While our proofs follow in the footsteps of forward simulation given the determinism of the target, the double-sided character of $\approx_R^u$ can equally support backward reasoning, allowing us to backward-match cutoff-set target errors with any further source behaviour. We proved a generalized form of transitivity:

$$\frac{\begin{array}{l} \vdash_{\Gamma} \forall e, \neg(C_1^r e \wedge C_2^l e) \quad \vdash_{\Gamma} \forall ee', e \ \Phi_1 \ e' \wedge C_2^l e' \rightarrow C_1^l e \\ \vdash_{\Gamma} \forall ee', e \ \Phi_2 \ e' \wedge C_1^r e \rightarrow C_2^r e' \end{array} \quad \frac{\vdash_{\Gamma} t \approx_{R_1 \Phi_1 \Psi_1 C_1^l C_1^r}^u t' \quad \vdash_{\Gamma} t' \approx_{R_2 \Phi_2 \Psi_2 C_2^l C_2^r}^u t''}{\vdash_{\Gamma} t \approx_{(R_1 \circ R_2) (\Phi_1 \circ \Phi_2) (\Psi_1 \circ \Psi_2) C_1^l C_2^r}^u t''}}$$

As far as UB-related non-determinism goes, this property could help us in combining together backward and forward reasoning, possibly relaxing safety preservation to integrate with safety analysis on targets.

As future work, we want to cover the Jasmin back-end (which includes linearization). $\approx_R^u$ can also be used to relate trees up to cutoff-determined prefixes, possibly making step-indexing easier, and indeed we are interested in the comparison with step-indexing and fuel-based inductive semantics [8], as well as in safety analysis.

References

- [1] 2025. *Interaction Trees, GitHub Repository*. <https://github.com/DeepSpec/InteractionTrees>
- [2] 2025. *Jasmin, GitHub Repository*. <https://github.com/jasmin-lang/jasmin>
- [3] José Bacelar Almeida, Manuel Barbosa, Gilles Barthe, Arthur Blot, Benjamin Grégoire, Vincent Laporte, Tiago Oliveira, Hugo Pacheco, Benedikt Schmidt, and Pierre-Yves Strub. 2017. Jasmin: High-Assurance and High-Speed Cryptography. In *Proc. ACM Computer and Communication Security (CCS 2017)*, Bhavani Thuraisingham, David Evans, Tal Malkin, and Dongyan Xu (Eds.). ACM, 1807–1823. <https://doi.org/10.1145/3133956.3134078>
- [4] Calvin Beck, Irene Yoon, Hanxi Chen, Yannick Zakowski, and Steve Zdancewic. 2024. A Two-Phase Infinite/Finite Low-Level Memory Model: Reconciling Integer-Pointer Casts, Finite Space, and undef at the LLVM IR Level of Abstraction. *Proc. ACM Program. Lang.* 8, ICFP (2024), 789–817. <https://doi.org/10.1145/3674652>
- [5] Nick Benton. 2004. Simple relational correctness proofs for static analyses and program transformations. In *Proceedings of the 31st ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages, POPL 2004, Venice, Italy, January 14-16, 2004*, Neil D. Jones and Xavier Leroy (Eds.). ACM, 14–25. <https://doi.org/10.1145/964001.964003>
- [6] Ramana Kumar, Magnus O. Myreen, Michael Norrish, and Scott Owens. 2014. CakeML: a verified implementation of ML. In *The 41st Annual ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages, POPL '14, San Diego, CA, USA, January 20-21, 2014*, Suresh Jagannathan and Peter Sewell (Eds.). ACM, 179–192. <https://doi.org/10.1145/2535838.2535841>
- [7] Xavier Leroy. 2009. Formal verification of a realistic compiler. *Commun. ACM* 52, 7 (2009), 107–115.
- [8] Scott Owens, Magnus O. Myreen, Ramana Kumar, and Yong Kiam Tan. 2016. Functional Big-Step Semantics. In *Proceedings of the 25th European Symposium on Programming Languages and Systems - Volume 9632*. Springer-Verlag, Berlin, Heidelberg, 589–615. https://doi.org/10.1007/978-3-662-49498-1_23
- [9] Lucas Silver, Eddy Westbrook, Matthew Yacavone, and Ryan Scott. 2023. Interaction Tree Specifications: A Framework for Specifying Recursive, Effectful Computations That Supports Auto-Active Verification. In *Proc. ECOOP 2023 (Leibniz International Proceedings in Informatics (LIPIcs), Vol. 263)*, Karim Ali and Guido Salvaneschi (Eds.). Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 30:1–30:26. <https://doi.org/10.4230/LIPIcs.ECOOP.2023.30>
- [10] Lucas Silver and Steve Zdancewic. 2021. Dijkstra monads forever: termination-sensitive specifications for interaction trees. *Proc. ACM Program. Lang.* 5, POPL (2021), 28 pages. <https://doi.org/10.1145/3434307>
- [11] Li-yao Xia, Yannick Zakowski, Paul He, Chung-Kil Hur, Gregory Malecha, Benjamin C. Pierce, and Steve Zdancewic. 2020. Interaction trees: representing recursive and impure programs in Coq. *Proc. ACM Program. Lang.* 4, POPL (2020), 51:1–51:32. <https://doi.org/10.1145/3371119>
- [12] Irene Yoon, Yannick Zakowski, and Steve Zdancewic. 2022. Formal reasoning about layered monadic interpreters. *Proc. ACM Program. Lang.* 6, ICFP (2022), 254–282. <https://doi.org/10.1145/3547630>
- [13] Yannick Zakowski, Calvin Beck, Irene Yoon, Ilia Zaichuk, Vadim Zaliva, and Steve Zdancewic. 2021. Modular, compositional, and executable formal semantics for LLVM IR. *Proc. ACM Program. Lang.* 5, ICFP (2021), 1–30. <https://doi.org/10.1145/3473572>