

Cazamariposas: Automated Instability Debugging in SMT-based Program Verification

Yi Zhou[✉], Amar Shah[✉], Zhengyao Lin[✉], Marijn Heule[✉], and
Bryan Parno[✉]

Carnegie Mellon University, Pittsburgh, PA, USA
{yizhou5, amarshah, zhengyal, mheule, bparno}@cs.cmu.edu

Abstract. Program verification languages such as Dafny and F^{*} often rely heavily on Satisfiability Modulo Theories (SMT) solvers for proof automation. However, SMT-based verification suffers from instability, where semantically irrelevant changes in the source program can cause spurious proof failures. While existing mitigation techniques emphasize preemptive measures, we propose a complementary approach that focuses on diagnosing and repairing specific instances of instability-induced failures. Our key technique is a novel differential analysis to pinpoint problematic quantified formulas in an unstable query. We implement this technique in Cazamariposas, a tool that automatically identifies such quantified formulas and suggests fixes. We evaluate Cazamariposas on multiple large-scale systems verification projects written in three different program verification languages. Our results demonstrate Cazamariposas’ effectiveness as an instability debugger. In the majority of cases, Cazamariposas successfully isolates the issue to a single problematic quantifier, while providing a stabilizing fix.

Keywords: SMT · Program Verification · Proof Instability.

1 Introduction

Satisfiability Modulo Theories (SMT) solvers play a prominent role in automated program verification. In verifiers such as Dafny [22], F^{*} [30], or Verus [20], SMT solvers can often discharge complex verification conditions automatically, despite the generally undecidable program properties. In this way, SMT solvers significantly reduce the need for manual proof steps, facilitating the verification of large-scale software systems [15, 19, 26].

However, the solvers must resort to incomplete heuristics to reason about undecidable program properties. Consequently, SMT-based program verifiers suffer from a persistent problem of *proof instability* [16, 34], where trivial, non-semantic changes to the source program result in spurious proof failures. For instance, an SMT solver may fail to verify (i.e., return an inconclusive result on) a previously-proven verification condition after a simple variable renaming, even though the program’s semantics remain unchanged.

Instability is a major impediment to the industrial deployment of automated verification [12]. In particular, instability disrupts the usual development work

flow, incurring time and resource costs to debug spurious failures. Worse yet, the solver provides little insight as to why it rejects the modified query. As a result, the developer has limited recourse beyond blindly modifying their source code in the hopes of nudging the solver into an accepting state [31].

While prior work [1, 21, 33] has tried to preemptively mitigate instability (with partial success—see Sec. 2), we propose a complementary approach to debug and repair unstable SMT queries. In more detail, we automatically provide the developer with query-specific *repair strategies*, targeting the quantified formulas causing instability. In fact, for 61% the queries in our study, there is only a *single* quantified formula to blame! In other words, we can repair most of the unstable proofs by limiting the impact of just one ill-behaved quantified formula.

We introduce a *differential analysis* to identify such problematic formulas. We leverage the fact that by definition [34], semantically equivalent variants of an unstable query create a mix of verification successes and failures. We thus compare the quantifier-instantiation profiles between succeeding and failing variants, highlighting the formulas that are over-/under-instantiated in the failing variants. We further refine the analysis using novel proof and trace mining techniques, exploiting the causal relation between the instantiations.

We implement our analysis in Cazamariposas, an SMT-level instability debugging tool. Cazamariposas takes as input an unstable SMT query Φ and provides a repair strategy as output. For example, suppose our analysis points to an over-instantiated formula ϕ within Φ . Cazamariposas confirms that removing ϕ from Φ stabilizes the query, and then it encourages the developer to remove the source-level construct that introduced ϕ to Φ . Crucially, we ensure that the repair strategy is sound with respect to program-verification semantics (Sec. 4).

We evaluate Cazamariposas on a diverse set of benchmarks, with 615 unstable SMT queries collected from 12 system verification projects written in Dafny [22], F* [30], or Verus [20]. We find that Cazamariposas successfully repairs 70% of the unstable queries, of which 87% involve a single quantified formula (Sec. 5).

2 Notation and Related Work

2.1 Notation

We use a conjunctive formula $\Phi = \bigwedge_{i=0}^n \psi_i$ to represent an SMT query, where each ψ_i is an assertion. We slightly abuse the notation here by treating Φ as a set of assertions. For example, we use $\Phi \setminus \{\psi\}$ and $\Phi \cup \{\psi\}$ to denote a new query with an assertion ψ removed or added, respectively.

We assume the *goal-axioms* structure in program verification queries [33]. In particular, $\psi_0 = \neg\theta$ is the negation of the properties of the procedure under verification. Meanwhile, $\Lambda = \bigwedge_{i=1}^n \psi_i$ is a conjunction of *axioms* encoding the semantics of: (1) the verification language’s constructs, and (2) other developer-written procedures that have already been verified. By checking that $\Phi = \Lambda \wedge \neg\theta$ is unsatisfiable, the SMT solver confirms that θ is a logical consequence of Λ .

The use of quantifiers is common in program verification. For ease of exposition, we use single-variable quantified formulas (e.g., $\phi = \forall x.\varphi$) as examples as

long as it is clear how the method under discussion generalizes to quantification over multiple variables. We use Ω to represent the set of quantified formulas (including all the nested ones) in Φ .

For a universally quantified formula $\phi = \forall x.\varphi$, we use $\varphi[x \mapsto t]$ to denote the result of capture-free substitution of some ground term t for all free occurrences of x in φ . We refer to $\varphi[x \mapsto t]$ as an *instantiation* of ϕ , and t as the *instantiating term*. We use the calligraphic $\mathcal{I}^\phi$ to denote a set of instantiations of ϕ . For convenience, we define $\mathcal{I}^\phi = \{\}$ if ϕ is existentially quantified.

We analyze the quantifier reasoning process based on solver-generated logs. Concretely, for a given query Φ , we leverage the Z3 [11] solver to produce a pair of proof and trace $(\mathbf{p}, \mathbf{t})$ logs where $\mathbf{p}$ is a proof tree of unsatisfiability [6], and $\mathbf{t}$ records all the solver-discovered instantiations. We use $\mathcal{I}_{\mathbf{p}}^\phi$ and $\mathcal{I}_{\mathbf{t}}^\phi$ to denote the set of instantiations of ϕ in $\mathbf{p}$ and $\mathbf{t}$, respectively.

SMT-based verification languages rely heavily on pattern-based quantifier instantiation [23, 24]. Each universally quantified formula $\forall x.\varphi$ is associated with (at least) one syntactic *pattern* π , where π would be a ground term expect for the free variable x in it. The body φ remains hidden until a ground term $\pi[x \mapsto t]$ enters the solver’s context, at which point the solver creates the instantiation $\varphi[x \mapsto t]$. We refer to $\pi[x \mapsto t]$ as the *triggering term*.

2.2 Related Work

Proof instability is an obstacle to wide-scale industrial adoption of SMT-based verification. Galois highlights the “fragility of proofs” which can be “highly sensitive to minor changes in logical terms” [12]. Similarly, Amazon complains about the serious challenge of “lack of monotonicity and stability in runtimes” [28]. Numerous other large-scale verification projects cite SMT instability as a key pain point [2, 9, 10, 13, 16, 19, 27].

In prior work, researchers applied source-program-level analysis to choose better syntactic patterns for user-introduced quantified formulas, hoping to produce more stable SMT queries [21]. More recently, Bordis and Leino propose using *free facts* [7], where quantified axioms are replaced by specific instantiations before the query is dispatched to the solver. Both techniques have shown some improvement, although they rely on ad hoc measures of instability.

In our Mariposa work [34], we presented a statistically rigorous approach to characterizing proof instability. At a high level, the Mariposa tool takes in an SMT query-solver pair (Φ, s) , and outputs whether the query Φ is stable, unstable, or unsolvable when run with solver s . Intuitively, Mariposa generates semantically-equivalent *mutants* of Φ and if the solver’s performance varies significantly across the mutants, Φ is unstable on s .

Mariposa creates the mutants by (1) reordering assertions, (2) α -renaming variables, or (3) changing the random seed. Therefore, for two mutants Φ_1 and Φ_2 of Φ , there is a one-to-one correspondence in terms of the function symbols and the formulas between Φ_1 and Φ_2 . It is thus possible to compare how the “same” quantified formula ϕ is instantiated in Φ_1 and Φ_2 , even though ϕ may appear in different forms.

As part of the Mariposa project, we also curated a collection of program verification queries. In particular, we included several large-scale systems-verification projects written in Dafny [22] and F* [30] as a part of the Mariposa benchmark suite, where we found non-trivial amounts of instability.

In our follow-up work, we demonstrate that proof stability is strongly connected to the relevance of axioms [33]. In general, given a verification goal θ , if Λ is populated with unnecessary or irrelevant axioms, then the query is more likely to experience instability. We then introduce SHAKE, a context-pruning technique that reduces the number of irrelevant axioms in program verification queries, mitigating the instability on the Mariposa benchmark suite by 29% on Z3 and 41% on cvc5.

Amrollahi et al. preprocess an SMT query to put it into a canonical form [1] so as to normalize away semantically irrelevant source-level changes. They demonstrate mixed results, reducing instability on one Mariposa benchmark by 20%, while increasing it on another by 76%. Cumulatively, their approach increased instability by 4%.

In contrast to these approaches based on preemptive preprocessing, in this work, we propose to diagnose and repair specific instances of unstable proofs. Our approach draws inspiration from the field of automated theorem proving [18,29]. In particular, our query-specific diagnosis draws a parallel to the axiom selection problem [17], where the goal is to select a small subset of axioms from a large set of axioms to prove a theorem. Oftentimes, axiom selection is done using techniques from machine learning [5].

3 Motivating Examples

When a programmer encounters instability, fixing it often requires considerable effort. When a proof in an automated verification language fails, conventional wisdom suggests various manual debugging techniques [25,31,32], including:

1. Adding source-level assertions to help guide the solver towards deriving important intermediate facts, and to trigger important quantifiers.
2. Adding source-level annotations to hide function definitions that are unnecessary for the proof. These annotations cause the verifier to encode the function such that the SMT solver treats it as uninterpreted.

We observe that these techniques target different problems at the SMT level. When the solver quickly returns **unknown** because of insufficient information, this may be addressed at the source level by adding source-level assertions. If the solver “times out” because it has spent too much time exploring extraneous parts of the proof space, this may be addressed at the source level by hiding unnecessary functions definitions.

In Sec. 4.3, we describe how Cazamariposas is able to automatically differentiate between the two cases above and suggest the appropriate fixes. In contrast, developers often struggle to find such fixes. To illustrate this, below we present

two examples ¹ from the Verus benchmarks (Sec. 5.2) where Cazamariposas identifies a fix that the original programmers missed.

3.1 Diagnosing and Repairing Unnecessary Instantiations

Our first example, a lemma `lemma_from_after` written in Verus, has an unstable proof. In the original source code, the developer added a Verus annotation asking the solver to spin off a separate SMT query just for this lemma (normally Verus proves many goals incrementally in the same SMT context). Verus developers often use this annotation to try to improve stability, but Cazamariposas’ measurements indicate it is still unstable.

To fix this proof’s instability, the developer can try to hide various functions definitions using Verus’s `hide` keyword. However, from the developer’s perspective, it is not obvious which function might be to blame for the instability. The file containing `lemma_from_after` has 6 function definitions, and it imports 27 other Rust crates, each of which contributes many more functions that might be causing the instability. A priori, it is not even clear that hiding a function definition is the correct fix. The developer might instead guess that adding assertions to the body of `lemma_from_after` will improve stability.

In contrast, Cazamariposas explains that `make_stateful_set` is the cause, an unrelated function imported from a completely different crate. Cazamariposas suggests that the developer hide `make_stateful_set`, which produces a stable SMT query.

3.2 Diagnosing and Repairing Missing Instantiations

In the Verus code below, the function `entry_alive_wraps` on line 7 proves that every index `i` from `low` to `high` is alive if and only if `BUFFER_SIZE + i` is alive. Liveness is defined via `entry_alive`, which among other operations, performs a division. At the SMT level, Verus represents division with the function `Euc_Div`, which is guarded to prevent division by 0. The original SMT query for `entry_alive_wraps` is unstable because the solver gets stuck going from `Euc_Div` to SMT-LIB’s built-in division operator. As a developer, diagnosing such is quite difficult, let alone finding a fix.

```

1  const BUFFER_SIZE = ...
2
3  fn entry_alive(i: int) -> bool {
4    (i / BUFFER_SIZE) % 2 == 0
5  }
6
7  fn entry_alive_wraps(low: nat, high: nat)
8    ensures forall|i: nat|low <= i < high ==>
9      entry_alive(i, BUFFER_SIZE) == entry_alive(i + BUFFER_SIZE, BUFFER_SIZE)
10 {}

```

In contrast, Cazamariposas automatically diagnoses the problem at the SMT level and suggests a fix. At the SMT level, the `ensures` clause on

¹ We have simplified the examples here and placed the full versions in our technical report [35].

`entry_alive_wraps` is negated, turning the universal quantifier on line 8 into an existential. Cazamariposas experiments with Skolemizing the quantified variable (i.e., turning it into a constant) and then searches for universal quantifiers in the query that might benefit from instantiation with the new Skolem constant. In this case, it identifies an instantiation of a quantifier about `Euc_Div` that stabilizes the proof. It suggests this fix to the developer, who can then use appropriate Verus-level syntax to apply the fix.

4 The Cazamariposas Methodology

Cazamariposas produces SMT-level *query edits* as repair strategies for proof instability. For example, given an unstable query Φ , Cazamariposas may pinpoint a quantified axiom ϕ_i (for $i \neq 0$) such that $\Phi^* = \Phi \setminus \{\phi_i\}$ is stable. (We note that removing ϕ_i is sound since Φ^* maintains the original verification goal.) The developer can then apply an analogue of this SMT-level edit as a source-level change (e.g., hiding the procedure axiomatized by ϕ_i in this case), stabilizing the proof.

Cazamariposas follows an “edit-and-test” scheme to identify stabilizing repair strategies. Conceptually, for each quantified axiom $\phi_i \in (\Lambda \cap \Omega)$, Cazamariposas goes through the following:

- Hypothesize that quantifier reasoning over ϕ_i is the cause of instability.
- Select a *query edit* on Φ to reduce the reasoning obligations over ϕ_i .
- Apply the query edit to create a *candidate query* Φ^* .
- Test the stability of Φ^* using Mariposa.
- If Φ^* is not stable, dismiss the hypothesis for now.
- If Φ^* is stable, report ϕ_i as a cause of instability.

In [Sec. 4.1](#), we discuss how we ensure that a query edit **(1)** weakens a particular target axiom, and **(2)** preserves the rest of the query context. With this design, if Φ^* is stable, the edit is also a sound repair strategy.

While the “edit-and-test” scheme is conceptually simple, the vast number of quantified axioms makes it impractical to exhaustively test the stability impact of each ϕ_i individually. Therefore, as illustrated in [Figure 1](#) and described in [Sec. 4.2-4.4](#), Cazamariposas must efficiently prioritize the likely suspects. In particular, we leverage the fact that an unstable query Φ has at least one passing mutant Φ_s and one failing mutant Φ_f . This allows us to compare the instantiations of ϕ_i between Φ_f and Φ_s , as we explained in [Sec. 2](#).

More concretely, we obtain from the solver a trace $\log \mathbf{t}$ for Φ_f , and a proof $\log \mathbf{p}$ for Φ_s . In theory, our method generalizes to multiple traces and proofs. In practice, collecting even one pair of $(\mathbf{t}, \mathbf{p})$ can entail difficulties ([Sec. 5](#)). Hence, we focus our discussion on one such pair. We use $\mathcal{I}_{\mathbf{t}}^{\phi_i}$ and $\mathcal{I}_{\mathbf{p}}^{\phi_i}$ to denote the instantiation set of ϕ_i in $\mathbf{t}$ and $\mathbf{p}$, respectively.

We first triage ([Sec. 4.2](#)) which kind of instability-induced failure we face: either a *quick unknown* (QU) or a *slow timeout* (TO). We then compute key metrics to distinguish the potentially problematic quantified axioms ([Sec. 4.3](#)). Next, using the metrics and the failure mode, we select the query edits most

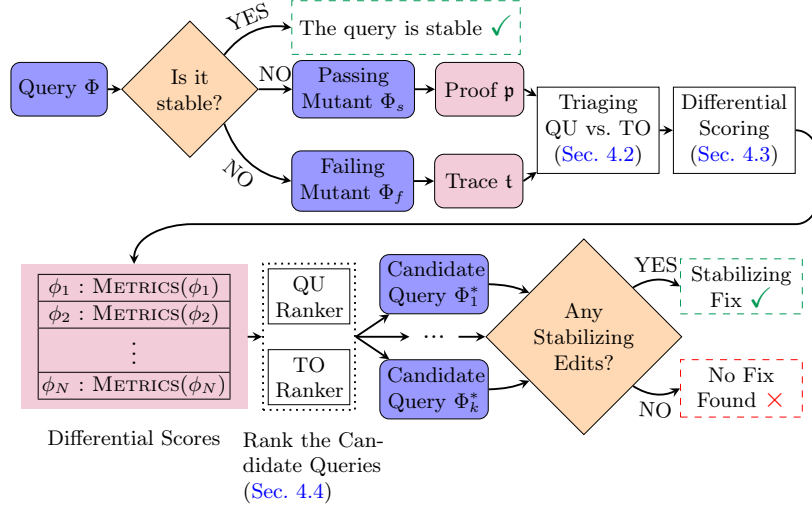


Fig. 1: The Design of Cazamariposas. We have colored the **SMT queries**, **data**, and **Mariposa calls**. We use $\text{METRICS}(\phi_i)$ to represent the metrics in Sec. 4.3. For simplicity, we have omitted doubleton edits.

likely to stabilize the query (Sec. 4.4). Finally, we use Mariposa to evaluate the edits, and if any succeed in stabilizing the query, we report success and suggest the corresponding edit as a repair.

4.1 Testing the Axioms for Stability

We start with our methodology to evaluate the stability impact of an individual axiom using query edits leveraging the proof log $\mathbf{p}$. More formally, we define a *singleton edit* as a pair (ϕ_i, a_i) , where $\phi_i \in (\Lambda \cap \Omega)$, and a_i is an *action* among the following:

- **del**: delete ϕ_i from the query.
- **inst**: augment ϕ_i with its proof instantiations (i.e., $\mathcal{I}_{\mathbf{p}}^{\phi_i}$) as new axioms.
- **inst-del**: replace ϕ_i with $\mathcal{I}_{\mathbf{p}}^{\phi_i}$ in the query.
- **sk**: Skolemize an existentially quantified assertion, ϕ_i .

We formally define the edits in Table 1. We use $f_{\phi_{ix}}$ to denote a Skolem constant from some existential assertion $\phi_i = \exists x.\varphi$.

Action a_i	Applicability ϕ_i	Edited Candidate Φ^*
del	$\phi_i \in (\Lambda \cap \Omega)$	$\Phi \setminus \{\phi_i\}$
inst	$\phi_i \in \Lambda, \phi_i = \forall x.\varphi$	$\Phi \cup \mathcal{I}_{\mathbf{p}}^{\phi_i}$
inst-del		$(\Phi \cup \mathcal{I}_{\mathbf{p}}^{\phi_i}) \setminus \{\phi_i\}$
sk		$(\Phi \cup \{\varphi[x \mapsto f_{\phi_{ix}}]\}) \setminus \{\phi_i\}$

Table 1: Cazamariposas Query Edits

Intuitively, the edits are meant to reduce or eliminate a solver’s reasoning over ϕ_i , so a stabilized candidate also points to ϕ_i as a cause of instability. We now discuss some basic properties of the edit actions, including soundness, which ensures that a stabilizing edit is also a valid repair strategy.

Soundness. We define the soundness of the candidate Φ^* as:

$$\Phi^* \vdash \perp \implies \Phi \vdash \perp$$

We demonstrate soundness with a case analysis. When $a_i = \mathbf{sk}$, the queries Φ^* and Φ are equivalent, so soundness trivially holds. Since we have restricted $\mathbf{sk}$ to be the only potential action on the goal, θ remains unchanged for the rest of the actions. Therefore, we could instead show that:

$$\Lambda^* \vdash \theta \implies \Lambda \vdash \theta$$

which holds as long as Λ^* is no stronger than Λ . When $a_i = \mathbf{inst}$, because the elements of $\mathcal{I}_p^{\phi_i}$ are tautological consequences of ϕ_i , Λ^* is as strong as Λ . When $a_i \in \{\mathbf{del}, \mathbf{inst-del}\}$, Λ^* might be weaker than Λ .

Completeness. We define completeness of the axiom set as follow:

$$\Lambda \vdash \theta \implies \Lambda^* \vdash \theta$$

The proof instantiation set $\mathcal{I}_p^{\phi_i}$ is sufficient to establish θ by definition, so we maintain completeness. However, since the edits may weaken the axioms, we do sacrifice a broader sense of completeness. Specifically, $\mathbf{del}$ and $\mathbf{inst-del}$ may remove a quantified axiom ϕ_i , while the $\mathcal{I}_p^{\phi_i}$ is only a finite subset of all possible instantiations of ϕ_i . Therefore do not guarantee that $\Lambda^* \vdash \phi_i$.

Composability. We note that if we perform a series of singleton edits $\Delta = \langle \dots, (\phi_i, a_i), \dots \rangle$, we also maintain soundness and completeness. Intuitively, when instability arises from the interaction of multiple quantified axioms, singleton edits (i.e., $\|\Delta\| = 1$) might fail to capture the cause, and thus we need to consider $\|\Delta\| \geq 2$. In the case where $\|\Delta\| = 2$, we call Δ a *doubleton edit*.

Practicality. In Cazamariposas, we focus on the SMT level to ensure applicability to multiple languages. Eventually, we would like to apply the edit actions to the source code, which we leave as future work. The query edits do generally correspond to source-level features in Dafny, F*, and Verus:

- **del** corresponds to source-level visibility control mechanisms. For example, Dafny’s `opaque` keyword hides the definition of a function.
- **inst** corresponds to quantifier instantiations as source-level annotations.
- **sk** corresponds to Hilbert’s choice as a language construct. For example, Dafny’s `var x :| P(x)` assigns x an arbitrary value such that $P(x)$ holds.

However, the translation might not always be straightforward. As we discussed in [Sec. 2](#), the axioms may also encode the semantics of language constructs. For example, **del** on an axiom for higher-order functions has no direct source-level equivalent. More specific to **inst**, if the repair adds a large number of instantiations to the source code, it is arguably impractical due to maintenance and readability concerns. Nevertheless, we have some empirical evidence

that the repairs are often practical, which would make it interesting to explore automatic translation in the future.

Complexity. Another more pressing concern is the complexity of the search space. Consider a query Φ with n applicable singleton edits. The total number of potential candidates is roughly:

$$\sum_{i=1}^{\|\Delta\|} \binom{n}{i} = \binom{n}{1} + \binom{n}{2} + \dots + \binom{n}{\|\Delta\|}$$

The combinatorial explosion makes it infeasible to test all candidates. Meanwhile, if a stabilizing edit involves too many axioms, it also become less realistic to reenact the repair at the source level.

Given these considerations, we limit our experiments to two classes: *singleton* and *doubleton* edits, i.e., $\|\Delta\| \leq 2$. Nevertheless, a massive search space remains for each class, with typically thousands of quantified axioms in a query. We thus further introduce a parameter k to limit the number of candidates we test for full stability. Specifically, we first test k singleton edits, and then k doubleton edits if the former fails to stabilize the query.

4.2 Triaging the Failure Mode

Given our methodology to test the stability of individual axioms, we now discuss how to select the most promising candidates. We begin by triaging the high-level cause of the underlying instability. Specifically, we observe two distinct modes of an instability-induced failure, which we name *quick unknown* (QU) and *slow timeout* (TO). In a QU, the solver quickly terminates (e.g., in < 1 second) with an **unknown** result, despite being given a generous timeout and resource limit. Meanwhile, for a TO, the solver runs on the query until it runs out of its time budget.

As we discuss in depth in [Sec. 5.3](#), the two failure modes correspond to different instantiation profiles, signifying different underlying causes of instability. In a QU, the solver only explores a small number of instantiations before giving up. We therefore hypothesize that the solver is failing because it is missing key instantiations to prove the goal θ . In contrast, a TO is associated with orders of magnitude more instantiations, suggesting that the solver is spending significant time and resources on quantifier reasoning irrelevant to θ .

4.3 Calculating the Differential Metrics

We now introduce measures to quantify the degree of insufficient or excessive instantiation. We define the metrics for quantified formula $\phi_i \in \Omega$, leveraging the proof instantiation set $\mathcal{I}_p^{\phi_i}$ and the trace instantiation set $\mathcal{I}_t^{\phi_i}$. (In [Sec. 4.4](#), we discuss how to aggregate the metrics to rank potential edits over the quantified axioms.) Specifically, we compute the following metrics here:

- $\text{DEFICIT}(\phi_i, \mathbf{p}, \mathbf{t}) = \|\mathcal{I}_{\mathbf{p}-\mathbf{t}}^{\phi_i}\|$, where $\mathcal{I}_{\mathbf{p}-\mathbf{t}}^{\phi_i} = \mathcal{I}_{\mathbf{p}}^{\phi_i} \setminus \mathcal{I}_{\mathbf{t}}^{\phi_i}$. Intuitively, $\mathcal{I}_{\mathbf{p}-\mathbf{t}}^{\phi_i}$ is the set of instantiations in the proof but not in the trace. When $\|\mathcal{I}_{\mathbf{p}-\mathbf{t}}^{\phi_i}\|$ is large, the solver may be missing instantiations of ϕ_i that are crucial to reaching **unsat**.
- $\text{EXCESS}(\phi_i, \mathbf{p}, \mathbf{t}) = \|\mathcal{I}_{\mathbf{t}-\mathbf{p}}^{\phi_i}\|$, where $\mathcal{I}_{\mathbf{t}-\mathbf{p}}^{\phi_i} = \mathcal{I}_{\mathbf{t}}^{\phi_i} \setminus \mathcal{I}_{\mathbf{p}}^{\phi_i}$. Intuitively, $\mathcal{I}_{\mathbf{t}-\mathbf{p}}^{\phi_i}$ is the set of instantiations in the trace but not in the proof. When $\|\mathcal{I}_{\mathbf{t}-\mathbf{p}}^{\phi_i}\|$ is large, the solver may be wasting time and resources instantiating ϕ_i without making progress towards proving the goal.

When ϕ_i is existentially quantified, $\mathcal{I}_{\mathbf{t}}^{\phi_i}$ and $\mathcal{I}_{\mathbf{p}}^{\phi_i}$ are both empty, so **DEFICIT** and **EXCESS** are trivially 0. In that case, we introduce the following metric based on the instantiations that depend on ϕ_i :

- $\text{CONTINGENCY}(\phi_i, \mathbf{p}) = \sum_{\phi_j \in \Omega} \|\{I \mid I \in \mathcal{I}_{\mathbf{p}}^{\phi_j}, f_{\phi_{ix}} \sqsubseteq I\}\|$, where ϕ_i creates the the Skolem constant $f_{\phi_{ix}}$, $\phi_j \in \Omega$ is some (universally) quantified formula, $I \in \mathcal{I}_{\mathbf{p}}^{\phi_j}$ is some instantiation of ϕ_j containing $f_{\phi_{ix}}$, and $f_{\phi_{ix}} \sqsubseteq I$ denotes that $f_{\phi_{ix}}$ is a sub-term of I . Intuitively, the metric reflects the proof instantiations that depend on $f_{\phi_{ix}}$. When ϕ_i has a high **CONTINGENCY** score, other quantified formulas cannot be sufficiently instantiated until ϕ_i is Skolemized.

Naively, we could already start prioritizing the quantified axioms based on these scores alone, but this would provide only crude guidance without considering the dependence between the formulas. Thus, we next discuss how we aggregate the scores and choose the most promising edit actions for each axiom.

4.4 Ranking the Candidate Queries

As mentioned in [Sec. 4.1](#), we use a parameter k to limit the number of candidates we test for full stability, where we first consider top- k singleton edits, and then doubleton edits if none of the singleton edits is stabilizing. We describe how to compute the scores for the singleton and doubleton edits in this section. The output of this stage are two partial maps, **SSCORE** and **DScore**.

1. We score each axiom ϕ_i , and then select an appropriate edit for it. When there are multiple possible actions on ϕ_i , we commit to one that is likely stabilizing. More formally, we create a partial map:

$$\text{SScore} = \{(\phi_i, a_i) \mapsto s_i \mid \phi_i \in \Phi\}$$

where $a_i \in \{\mathbf{del}, \mathbf{inst}, \mathbf{inst-del}, \mathbf{sk}\}$ is the chosen edit action.

2. We then score ordered pairs of quantified assertions, along with the most promising actions for each assertion. More formally, we create another partial map:

$$\text{DScore} = \{\langle(\phi_i, a_i), (\phi_d, a_j)\rangle \mapsto s_{ij} \mid \phi_i, \phi_j \in \Phi\}$$

where $a_i, a_j \in \{\mathbf{del}, \mathbf{inst}, \mathbf{inst-del}, \mathbf{sk}\}$ are the chosen edit actions. $\langle(\phi_i, a_i), (\phi_j, a_j)\rangle$ is the doubleton edit we apply (in order). We note that both maps are partial because we may not find an applicable action for certain assertions.

We split the discussion on the ranking of edits based on the hypothesized failure mode (QU or TO), as the two failure modes require different strategies.

Ranking Edits for QU We start with how we handle QU failures, which is more straightforward. At the general triage (Sec. 4.2) stage, we hypothesize that the QU failures are due to the absence of certain instantiations. Intuitively, we are looking for under-instantiated axioms, where **inst** is applicable, i.e.,

$$\text{SScore} = \{(\phi_i, \mathbf{inst}) \mapsto s_i \mid \phi_i \in \Phi\}$$

There are various ways to use the differential scores to set s_i . Plausible contenders include:

1. (DEFICIT, $-\text{EXCESS}$)
2. ($-\text{EXCESS}$, DEFICIT)
3. $\kappa \cdot \text{DEFICIT} - \text{EXCESS}$ (for some constant κ)
4. DEFICIT/EXCESS

We experimented with multiple examples of each of these heuristics. Eventually we settled on the first one, using DEFICIT as the primary metric.

However, the picture becomes complicated when instantiations contain Skolem constants. In that case, we cannot fully materialize all of ϕ_i 's proof instantiations $\mathcal{I}_{\mathbf{p}}^{\phi_i}$, unless all its Skolem dependencies are met. If the actual materializable instantiation count is 0 (indicating that no instantiations can be created without Skolemization), then we drop ϕ_i in the singleton phase.

We address this issue in the doubleton stage. Specifically, we use the CONTINGENCY score to select the first axiom ϕ_i for **sk**; i.e., the quantified assertion with the most “contingent” instantiations depending on its Skolem constant. When choosing the second axiom ϕ_j , we only consider ϕ_j candidates that depend on the Skolem constant $f_{\phi_{ix}}$ in their instantiations, and we can apply **inst** to ϕ_j . More formally,

$$\text{DScore} = \{((\phi_i, \mathbf{sk}), (\phi_j, \mathbf{inst})) \mapsto s_{ij} \mid \phi_i, \phi_j \in \Phi\}$$

where $s_{ij} = (\text{CONTINGENCY}(\phi_i, \mathbf{p}), \text{DEFICIT}(\phi_j, \mathbf{p}, \mathbf{t}))$, and $\exists I \mid I \in \mathcal{I}_{\mathbf{p}}^{\phi_j}, f_{\phi_{ix}} \sqsubseteq I$.

Ranking Edits for TO In the general triage stage (Sec. 4.2), we hypothesize that in a TO failure, the solver is spending significant time and resources on irrelevant quantified formulas. For this failure mode, we focus on the *quantified axioms* in Λ as targets. Intuitively, we want to suppress the excessive instantiations in order to stabilize the query, under the constraint that we cannot edit the goal itself (since we want our edits to preserve soundness).

A natural choice would be to use the EXCESS for SScore, and then apply **del** to the axiom ϕ_i with the highest EXCESS. However, the situation is more complex than QU in two ways. **(1)** We cannot simply delete an arbitrary ϕ_i with a high EXCESS. The axiom may be necessary for the proof, and deleting it will render the goal un-provable (i.e., creating incompleteness). **(2)** Even if ϕ_i is indeed unnecessary, other excessively instantiated axioms may also be contributing to the instability.

Problem (1) is easier to address. We use the **inst-del** edit action, replacing the axiom ϕ_i with its instantiations from the successful proof trace $\mathcal{I}_{\mathbf{p}}^{\phi_i}$. Intuitively, this eliminates the need (and the ability) for the solver to instantiate ϕ_i :

since $\mathcal{I}_p^{\phi_i}$ is sufficient for the proof, this action works around the incompleteness issue.

Problem (2) is more challenging. Anecdotally, if we focus solely on the EXCESS score, the debugging process turns into a “whack-a-mole” situation, where we delete one axiom, only to find another axiom with high EXCESS taking its place, and we fail to stabilize the query. Hence, to successfully repair the query, we need a mechanism to identify the underlying cause of the excessive instantiations.

Dependency Analysis In order to locate the root cause of TO instability, we further analyze the causal relations between the instantiations. Our notion of causality extends the *instantiation graph* from the SMTSCOPE (formerly the Axiom Profiler) [4], a tool to analyze instantiation loops and other sources of poor performance in pattern-based SMT solvers. Below, we describe Axiom Profiler’s approach and then our extension.

The instantiation graph is a directed acyclic graph over the terms (instantiations) in a trace log $\mathbf{t}$. More formally, we model this as a graph G_0 with the node set:

$$\{(I, \phi_i) \mid I \in \mathcal{I}_t^{\phi_i}, \phi_i \in \Omega\}$$

where each instantiation is labeled with its quantified formula ϕ_i . Edges in the graph indicate the causal relations, which include the following:

- **Instantiating Dependency:** an instantiation causes another one to materialize due to a matched pattern. Let (I_s, ϕ_s) and (I_d, ϕ_d) be two nodes in G_0 , where $\phi_d = \forall x. \varphi_j$ is guarded by the pattern π_j . Suppose a sub-term of I_s matches π_j , i.e., $\pi_j[x \mapsto t] \sqsubseteq I_s$ for some ground term t . This match triggers the creation of $I_d = \varphi_j[x \mapsto t]$, corresponding to an edge $(I_s, \phi_s) \rightarrow (I_d, \phi_d)$ in G_0 .
- **Equational Dependency:** an equational rewrite (from one instantiation) contributes to another instantiation. Continuing the example above, I_s may only trigger π_j after additional equality rewrites. Consider a quantified formula $\phi_{eq} = \forall x. p(x) \cong q(x)$ and one of its instantiations $I_{eq} = p(a) \cong q(a)$. The solver might have to rewrite I_s with I_{eq} first, and then the rewrite result, $I_s[p(a) \mapsto q(a)]$, triggers the creation of I_d . In that case, there is also an edge $(I_{eq}, \phi_{eq}) \rightarrow (I_d, \phi_d)$ in G_0 .

We further extend this graph G_0 from prior work into a graph G_1 to capture two additional types of dependencies.

- **Skolemizing Dependency:** a Skolem constant is a sub-term of an instantiation. Consider the existentially quantified $\phi_i = \exists x. \varphi_i$ with Skolem constant $f_{\phi_{ix}}$. There might be some node (I_d, ϕ_d) in G_1 such that $f_{\phi_{ix}} \sqsubseteq I_d$. In that case, we add the node $(f_{\phi_{ix}}, \phi_s)$, and the edge $(f_{\phi_{ix}}, \phi_s) \rightarrow (I_d, \phi_d)$ to G_1 . This form of dependency follows the same intuition as in our definition of CONTINGENCY, except we apply it to the trace log here.
- **Nesting Dependency:** an instantiation is a (previously-nested) quantified formula, which creates further instantiations. For example, consider (I_s, ϕ_s) , where $\phi_s = \forall x. (f(x) \wedge \forall y. g(x, y))$, and $I_s = f(t) \wedge \forall y. g(t, y)$ for some ground term t . Let $\phi_d = \forall y. g(t, y)$ be the nested quantified formula. Intuitively, I_s is

the reason why ϕ_d exists at all. We thus add an edge from (I_s, ϕ_s) to every (I_d, ϕ_d) , where $I_d \in \mathcal{I}_t^{\phi_d}$.

The graph G_1 captures the four types of dependencies we discussed above, which offers a rather low-level view of the instantiation reasoning in the trace. We further process G_1 so that it reflects the relation between the quantified formulas.

1. We collapse G_1 into a multi-edge graph G_2 . We initialize G_2 with Ω as its nodes. For each edge $(I_s, \phi_s) \rightarrow (I_d, \phi_d)$ in G_1 , we create an edge $\phi_s \rightarrow \phi_d$.
2. We reduce G_2 into a weighted simple graph G_3 . For each pair of neighboring nodes ϕ_s and ϕ_d with $m_{s,d}$ parallel edges in G_2 , we keep one edge $\phi_s \rightarrow \phi_d$ in G_3 with the weight $m_{s,d}$.
3. We normalize the edge weights in G_3 , where we set the weight for $\phi_s \rightarrow \phi_d$ in G_3 to:

$$w_{s,d} = \frac{m_{s,d}}{\sum_{\phi_i \rightarrow \phi_d} m_{i,d}}$$

Intuitively, $w_{s,d}$ reflects the normalized “impact” of ϕ_s on ϕ_d over all the in-coming edges (via other ϕ_i) to ϕ_d .

Hence the output of our dependency analysis is a directed simple graph G_3 over Ω , where each edge weight $w_{s,d}$ captures (or rather, approximates) the normalized impact of ϕ_s over ϕ_d . For example, $w_{s,d} = 0.5$ signifies that ϕ_s has an *immediate impact* on 50% of the instantiations of ϕ_d .

We then compute the transitive impact through fixed-point iterations. Concretely, for $\phi_i \in \Lambda$, we consider the reachable subgraph G_{ϕ_i} in G_3 . We initialize a ratio $r_d = 0$ for each ϕ_d in G_{ϕ_i} , except for ϕ_i , where we set $r_i = 1$. We then update each ratio $r_d = \sum_s r_s \cdot w_{s,d}$. After the fixed-point computation terminates, we use the weighted sum of EXCESS as the final score for ϕ_i :

$$\text{SScore} = \{(\phi_i, a_i) \mapsto \sum_{\phi_j \in \Omega} \text{EXCESS}(\phi_j, \mathbf{t}, \mathbf{p}) \cdot r_j \mid \phi_i \in \Lambda\}$$

The fixed-point computation is non-decreasing by transitivity. However, there is no theoretical guarantee that it will converge. In particular, when G_{ϕ_i} contains a cycle, a node’s ratio may approach a limit at an exponentially decaying rate. Nevertheless, this is not a threat to practical usage. In particular, since floating point numbers represent the ratios, the convergence criteria must be threshold-based. For our implementation, we consider a ratio to have converged if its increment from the previous iteration is $\leq 10^{-4}$.

Now that we have the scores for each axiom, we proceed to choose the singleton edit action. We do so with a simple heuristic: if we can delete an axiom without causing incompleteness, we choose **del**. Otherwise, we instantiate the axiom with its proof instantiations, using **inst-del**. However, if there is Skolemization dependency preventing us from fully materializing the proof instantiations, we choose **inst** instead. Finally, if we have no other choice beyond Skolemization (**sk**), we select it.

Given this setup, ranking the doubleton edits is simple. We use the fixed-point computation to estimate the impact of each pair of axioms; i.e., we initialize

$r_i = 1, r_j = 1$ for the pair (ϕ_i, ϕ_j) , and then iterate over the nodes in G_3 to update the ratios. We then use the same weighted sum of EXCESS to calculate the final score for each pair. We also use the same heuristic to choose the edit actions for each pair.

5 Evaluation

In this section we perform an evaluation of Cazamariposas. We start with a brief overview of the implementation in [Sec. 5.1](#), followed by a discussion of the benchmarks used in our evaluation [Sec. 5.2](#). We then present the results of our evaluation, structured around two research questions: (1) Does the experimental data support our hypothesis about the different failure modes? and (2) How effective is Cazamariposas at identifying stabilizing edits?

5.1 Implementation

Our implementation of Cazamariposas is in 4,730 lines of Python, publicly available on GitHub [\[8\]](#), as a part of the Mariposa tool-chain. Here we also discuss our use of other tools and the configuration settings.

Mariposa We use Mariposa twice in our workflow, first to test whether the initial query Φ is stable and finally to test whether our various fixes are stable. We run Mariposa with its standard configuration, creating 180 total mutants for each query for full stability test. We use the default timeout of 60 seconds for the Mariposa benchmarks, but for the Verus benchmarks, we align the timeout limit with the project artifacts, which use 10 seconds.

Z3 We use a recent version of the Z3 solver (4.13.0). We evaluate Cazamariposas with Z3, as Dafny, F*, and Verus are designed with Z3 in mind. Past work [\[34\]](#) has shown that other solvers such as cvc5 [\[3\]](#) and VAMPIRE [\[18\]](#) are not optimized for these verification languages, resulting in large numbers of unsolvable queries.

Z3 provides proof-production functionality, which can be complicated to use in practice. Enabling proof production often causes Z3 to take a different path, completely failing on an otherwise solvable query. To work around this, we have employed the following strategies: (1) use Z3 to find an unsatisfiable core first and then produce a proof from the core, (2) use 4 different versions of Z3 solver, (3) use an extended timeout of 1 hour, and (4) use up to 256 mutants. We note that these configuration are for finding proofs, not for testing stability. Despite all of these strategies, we were unable to get proofs for 4 Mariposa queries and 3 Verus queries. In our evaluation, we count these cases as if Cazamariposas *fails to find a fix*. Ideally the proof-production failures can be addressed within the SMT solver.

SMTSCOPE We created a fork of SMTSCOPE [\[14\]](#) to parse the Z3 trace logs, and then we enhance their instantiation graph as described in [Sec. 4.4](#).

5.2 Verus Benchmark Set

In addition to the Mariposa benchmark with 545 unstable queries (from five systems projects written in Dafny and F*), we curate a new benchmark set comprised of SMT queries from ten Verus verification projects. Between these Verus projects, there are a total of 7,584 queries of which 7,514 (~99%) are stable and 70 (~1%) are unstable. We refer the readers to a detailed breakdown of projects in our technical report [35].

We conduct all of our stability tests, for benchmark creation and evaluation, all on the same set of machines with an Intel Core i9-9900K (max 5.00 GHz) CPU, 128 GB of RAM, and Ubuntu 20.04.3.

In summary, we evaluate Cazamariposas on 70 unstable Verus queries and 545 unstable Mariposa queries, for a total of unstable 615 queries.

5.3 Failure Mode Distinction

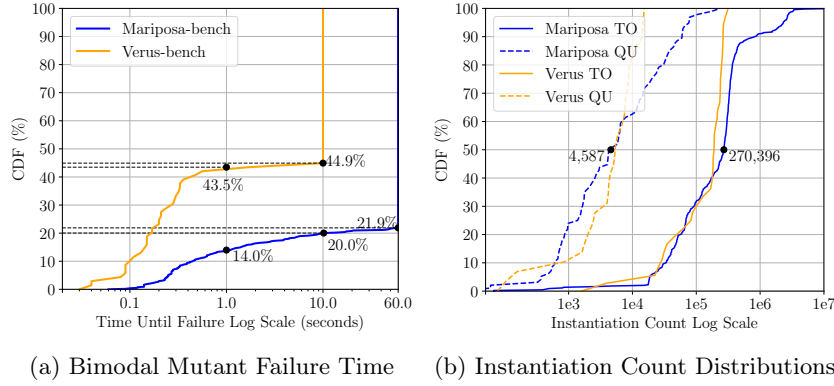


Fig. 2: **Failure Mode Distinction.** In (a) we plot the runtime of the failed mutants Φ_f , **before** the triage. In (b), we plot the instantiation count for the failed mutants Φ_f based on the failure modes triage.

Cazamariposas’s first step is to triage a query’s failure mode into either quick unknown (QU) or slow timeout (TO). Here we present empirical evidence for the distinction. In Figure 2a, for each original query Φ in the benchmarks, we report the runtime of its failed mutant Φ_f . The plot is in log scale, and we observe that the distribution for each benchmark is bimodal. For Verus-bench, ~43% of the failures occur within 1 second, barely any occur between 1 and 10 seconds, and the rest fail at 10 seconds. For Mariposa-bench, the distribution is more spread out, but the separation is still clear, where ~19% queries fail within 10 seconds, and ~78% time out after 60 seconds. There is almost no middle ground between the two modes. The x-axis of Figure 3 shows how many queries from each benchmark Cazamariposas ultimately classifies as QU versus TO.

In Figure 2b, we examine our hypothesis that QU failures are due to insufficient instantiation, and TO failures are due to excessive instantiation. We perform our triage, and then plot the instantiation counts based on the failure modes. Note the log-scale on x-axis, which highlights that the TO failures have orders of magnitude more instantiations than the QU failures. For example, in Mariposa TO, the median instantiation count is 270,396 while in Mariposa QU, the median is 4,587. The separation is also clear within Verus benchmark.

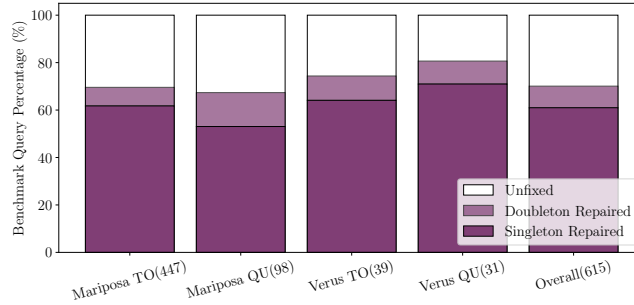


Fig. 3: Percentage of Benchmark Queries Repaired.

5.4 Stabilizing Edits Found

Next, we evaluate Cazamariposas’ ability to automatically identify stabilizing edits. Recall that Cazamariposas first tries $k = 10$ singleton edits, and if none works, it tries 10 doubleton edits. Overall, Cazamariposas repairs 431/615 ($\approx 70\%$) of the benchmark queries. Figure 3 provides more details, reporting Cazamariposas’ performance on the two benchmarks, subdivided by the underlying failure type (TO vs. QU). We note that Mariposa TO accounts for the largest absolute number of queries. Cazamariposas appears to be more effective on Verus queries in either failure type. Nevertheless, Cazamariposas repairs approximately 69% of the Mariposa queries and 77% of the Verus queries. This compares favorably to the best results from prior work (Sec. 2), which stabilized 29% of the Mariposa benchmark. We also observe that 375/615 ($\approx 61\%$) queries can be stabilized with a single edit. Doubleton edits subsequently provide a small but noticeable boost.

We also evaluate how well Cazamariposas ranks the stabilizing edits. First, in Figure 4a, we show the distribution of the number of quantified formulas in the original queries. For example, in Mariposa TO, the median count is 5,965, which is a large search space for possible edits. The median count is lower in Verus QU, making it potentially more tractable to fully explore.

Now that we have sense of the search space, we evaluate how well Cazamariposas identifies the useful edits. In Figure 4b, we report the minimal rank of

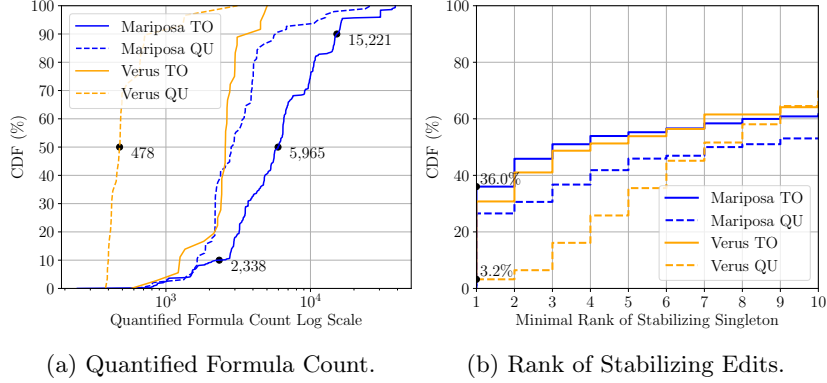


Fig. 4: Finding Repairs Among Large Number of Quantified Formulas.

the stabilizing singleton edits. Specifically in singletons, given a query, Cazamariposas produces a ranked list of 10 edits, and we report the rank of the first stabilizing edit within this list. We note the endpoints of the CDFs on the y-axis. It is the probability that Cazamariposas finds a stabilizing edit within the first 10 singleton edits, which corresponds to Figure 3. We note the start points of the CDFs on the y-axis. This is the probability that the first edit Cazamariposas tries would directly work. For a Mariposa TO query, Cazamariposas has a 36% chance of finding a stabilizing edit with one shot.

6 Conclusions

Proof instability is a major impediment to industrial use of automated program verification tools. Hence, we propose a new approach that targets specific instances of instability. Using a novel differential analysis, we automatically classify the type of instability a query is experiencing, rank the problematic quantifiers in the query, and then efficiently identify targeted query edits that stabilize the query. We implement this approach in Cazamariposas and evaluate it on SMT queries from numerous verification projects written in three automated verification languages. Cazamariposas successfully repairs 70% of the unstable queries.

Acknowledgments

We thank Jay Bosamiya and the anonymous reviewers for their valuable feedback on the paper. We thank Jonáš Fiala for his help with using and modifying SMTSCOPE.

This work was supported in part by the National Science Foundation (NSF) under grant 2224279, funding from AFRL and DARPA under Agreement FA8750-24-9-1000, and the Future Enterprise Security initiative at Carnegie Mellon CyLab (FutureEnterprise@CyLab).

Disclosure of Interests: The authors have no competing interests to declare that are relevant to the content of this article.

References

1. Amrollahi, D., Preiner, M., Niemetz, A., Reynolds, A., Charikar, M., Tinelli, C., Barrett, C.: Using normalization to improve SMT solver stability (2024), <https://arxiv.org/abs/2410.22419>
2. Backes, J., Bolignano, P., Cook, B., Dodge, C., Gacek, A., Luckow, K., Rungta, N., Tkachuk, O., Varming, C.: Semantic-based automated reasoning for AWS access policies using SMT. In: Formal Methods in Computer Aided Design (FMCAD) (2018). <https://doi.org/10.23919/FMCAD.2018.8602994>
3. Barbosa, H., Barrett, C., Brain, M., Kremer, G., Lachnitt, H., Mann, M., Mohamed, A., Mohamed, M., Niemetz, A., Nötzli, A., et al.: cvc5: A Versatile and Industrial-Strength SMT Solver. In: Tools and Algorithms for the Construction and Analysis of Systems (TACAS) (2022). https://doi.org/10.1007/978-3-030-99524-9_24
4. Becker, N., Müller, P., Summers, A.J.: The axiom profiler: Understanding and debugging SMT quantifier instantiations. In: Tools and Algorithms for the Construction and Analysis of Systems (TACAS) (2019)
5. Blanchette, J.C., Kühlwein, D., Geschke, S., Loewe, B., Schlicht, P.: A survey of axiom selection as a machine learning problem. In: Infinity, Computability, and Metamathematics: Festschrift Celebrating the 60th Birthdays of Peter Koepke and Philip Welch. Tributes, College Publications (2014)
6. Böhme, S.: Proof reconstruction for Z3 in Isabelle/HOL. In: International Workshop on Satisfiability Modulo Theories (2009)
7. Bordis, T., Leino, K.R.M.: Free facts: An alternative to inefficient axioms in Dafny. In: Formal Methods: 26th International Symposium (FM) (2024). https://doi.org/10.1007/978-3-031-71162-6_8
8. Cazamariposas. <https://github.com/secure-foundations/mariposa>, accessed Feb. 2025
9. Chakarov, A., Geldenhuys, J., Heck, M., Hicks, M., Huang, S., Jaloyan, G.A., Joshi, A., Leino, R., Mayer, M., McLaughlin, S., Mritunjai, A., Claudel, C.P., Porncharoenwase, S., Rabe, F., Rapoport, M., Reger, G., Roux, C., Rungta, N., Salkeld, R., Schlaipfer, M., Schoepe, D., Schwartzentruber, J., Tasiran, S., Tomb, A., Torlak, E., Tristan, J., Wagner, L., Whalen, M., Willems, R., Xiang, J., Byun, T.J., Cohen, J., Wang, R., Jang, J., Rath, J., Syeda, H.T., Wagner, D., Yuan, Y.: Formally verified cloud-scale authorization. In: International Conference on Software Engineering (ICSE) (2025), <https://www.amazon.science/publications/formally-verified-cloud-scale-authorization>
10. Cutler, J.W., Disselkoen, C., Eline, A., He, S., Headley, K., Hicks, M., Hietala, K., Ioannidis, E., Kastner, J., Mamat, A., McAdams, D., McCutchen, M., Rungta, N., Torlak, E., Wells, A.M.: Cedar: A new language for expressive, fast, safe, and analyzable authorization. In: Proceedings of the ACM Conference on Object-Oriented Programming Systems, Languages, and Applications (OOPSLA) (Apr 2024). <https://doi.org/10.1145/3649835>
11. De Moura, L., Bjørner, N.: Z3: An efficient SMT solver. In: Tools and Algorithms for the Construction and Analysis of Systems (TACAS) (2008). https://doi.org/10.1007/978-3-540-78800-3_24

12. Dodds, M.: Formally Verifying Industry Cryptography. *IEEE Security and Privacy Magazine* (2022). <https://doi.org/10.1109/MSEC.2022.3153035>
13. Ferraiuolo, A., Baumann, A., Hawblitzel, C., Parno, B.: Komodo: Using Verification to Disentangle Secure-Enclave Hardware from Software. In: *Proceedings of the ACM Symposium on Operating Systems Principles (SOSP)* (2017). <https://doi.org/10.1145/3132747.3132782>
14. Fiala, J.: SMTSCOPE (2025), <https://github.com/viperproject/smt-scope>, accessed Feb. 2025
15. Hawblitzel, C., Howell, J., Kapritsos, M., Lorch, J.R., Parno, B., Roberts, M.L., Setty, S., Zill, B.: IronFleet: Proving practical distributed systems correct. In: *Proceedings of the ACM Symposium on Operating Systems Principles (SOSP)* (2015)
16. Hawblitzel, C., Howell, J., Lorch, J.R., Narayan, A., Parno, B., Zhang, D., Zill, B.: Ironclad Apps: End-to-end security via automated full-system verification. In: *Proceedings of the USENIX Symposium on Operating Systems Design and Implementation (OSDI)* (October 2014)
17. Hoder, K., Voronkov, A.: Sine qua non for large theory reasoning. In: *International Conference on Automated Deduction*. pp. 299–314. Springer (2011)
18. Kovács, L., Voronkov, A.: First-order theorem proving and Vampire. In: *Computer Aided Verification (CAV)* (2013)
19. Lattuada, A., Hance, T., Bosamiya, J., Brun, M., Cho, C., LeBlanc, H., Srinivasan, P., Achermann, R., Chajed, T., Hawblitzel, C., Howell, J., Lorch, J., Padon, O., Parno, B.: Verus: A practical foundation for systems verification. In: *Proceedings of the ACM Symposium on Operating Systems Principles (SOSP)* (November 2024)
20. Lattuada, A., Hance, T., Cho, C., Brun, M., Subasinghe, I., Zhou, Y., Howell, J., Parno, B., Hawblitzel, C.: Verus: Verifying Rust programs using linear ghost types. In: *Proceedings of the ACM Conference on Object-Oriented Programming Systems, Languages, and Applications (OOPSLA)* (December 2023)
21. Leino, K.R.M., Pit-Claudel, C.: Trigger selection strategies to stabilize program verifiers. In: Chaudhuri, S., Farzan, A. (eds.) *Proceedings of the International Conference on Computer Aided Verification (CAV)* (2016)
22. Leino, K.R.M.: Dafny: An automatic program verifier for functional correctness. In: *Logic for Programming, Artificial Intelligence, and Reasoning (LPAR)* (2010)
23. Moskal, M.: Programming with triggers. In: *Proceedings of the Workshop on Satisfiability Modulo Theories* (2009)
24. de Moura, L., Bjørner, N.: Efficient e-matching for SMT solvers. In: *Conference on Automated Deduction (CADE)* (2007)
25. Profiling Z3 and solving proof performance issues. https://fstar-lang.org/tutorial/book/under_the_hood/uth_smt.html#profiling-z3-and-solving-proof-performance-issues
26. Protzenko, J., Parno, B., Fromherz, A., Hawblitzel, C., Polubelova, M., Bhargavan, K., Beurdouche, B., Choi, J., Delignat-Lavaud, A., Fournet, C., Kulatova, N., Ramananandro, T., Rastogi, A., Swamy, N., Wintersteiger, C., Zanella-Beguelin, S.: EverCrypt: A Fast, Verified, Cross-Platform Cryptographic Provider. In: *Proceedings of the IEEE Symposium on Security and Privacy* (May 2020)
27. Reitz, A., Fromherz, A., Protzenko, J.: StarMalloc: Verifying a modern, hardened memory allocator. *Proc. ACM Program. Lang.* (Oct 2024). <https://doi.org/10.1145/3689773>
28. Rungta, N.: A billion SMT queries a day (invited paper). In: Shoham, S., Vizel, Y. (eds.) *Proceedings of the International Conference on Computer Aided Verification (CAV)* (2022)

29. Schulz, S., Cruanes, S., Vukmirović, P.: Faster, higher, stronger: E 2.3. In: Fontaine, P. (ed.) Proc. of the 27th CADE, Natal, Brasil. pp. 495–507. No. 11716 in LNAI, Springer (2019)
30. Swamy, N., Hrițcu, C., Keller, C., Rastogi, A., Delignat-Lavaud, A., Forest, S., Bhargavan, K., Fournet, C., Strub, P.Y., Kohlweiss, M., Zinzindohoue, J.K., Zanella-Béguelin, S.: Dependent Types and Multi-Monadic Effects in F*. In: Proceedings of the ACM Symposium on Principles of Programming Languages (POPL) (2016)
31. Tomb, A., Tristan, J.B.: Avoiding verification brittleness in Dafny. <https://dafny.org/blog/2023/12/01/avoiding-verification-brittleness/> (2023)
32. Verification debugging when verification fails. <https://dafny.org/dafny/DafnyRef/DafnyRef#sec-brittle-verification>
33. Zhou, Y., Bosamiya, J., Li, J., Heule, M., Parno, B.: Context pruning for more robust SMT-based program verification. In: Proceedings of the Formal Methods in Computer-Aided Design (FMCAD) Conference (October 2024)
34. Zhou, Y., Bosamiya, J., Takashima, Y., Li, J., Heule, M., Parno, B.: Mariposa: Measuring SMT instability in automated program verification. In: Proceedings of the Formal Methods in Computer-Aided Design (FMCAD) (October 2023)
35. Zhou, Y., Shah, A., , Lin, Z., Heule, M., Parno, B.: Cazamariposas: Automated Instability Debugging in SMT-based Program Verification (Technical Report). Tech. rep., Carnegie Mellon University (June 2025), doi.org/10.1184/R1/29207189