

Validation of CHC Satisfiability with ATHENA

RODRIGO OTONI, USI Lugano, Switzerland

MARTIN BLICHA, USI Lugano, Switzerland and Charles University, Czech Republic

PATRICK EUGSTER, USI Lugano, Switzerland

NATASHA SHARYGINA, USI Lugano, Switzerland

Formal verification tooling increasingly relies on logic solvers as automated reasoning engines. A commonality among these solvers is the high complexity of their codebases, which makes bug occurrence disturbingly frequent. Tool competitions have showcased many examples of state-of-the-art solvers disagreeing on the satisfiability of logic formulas, be it solvers for Boolean satisfiability (SAT), satisfiability modulo theories (SMT), or constrained Horn clauses (CHC). The validation of solvers' results is thus of paramount importance, in order to increase the confidence not only in the solvers themselves, but also in the tooling which they underpin. Among the formalisms commonly used by modern verification tools, CHC is one that has seen, at the same time, extensive practical usage and very little efforts in result validation. We propose a two-layered validation approach for witnesses of CHC satisfiability that validates CHC models via proof-backed SMT queries. We developed a modular evaluation framework, ATHENA, and assessed the approach's viability via large scale experimentation, comparing three CHC solvers, five SMT solvers, and five proof checkers. Our results indicate that the approach is feasible, with potential to be incorporated into CHC-based tooling, and also confirm the need for validation, with fourteen bugs being found in the tools used.

CCS Concepts: • **Theory of computation** → **Automated reasoning**; **Logic and verification**.

Additional Key Words and Phrases: validation, constrained Horn clauses, SMT proofs

Published in Formal Aspects of Computing (FAC)
DOI: <https://doi.org/10.1145/3716505>

1 Introduction

First-order logic (FOL) is a formalism capable of representing many interesting verification problems, ranging from simple integer overflow to elaborate safety and liveness properties. Different fragments of FOL are suitable to aid in specific verification tasks, with one fragment of particular practical interest being constrained Horn clauses (CHC) [32]. The CHC fragment has been shown to be a match for Hoare logic [38] with practical uses [11], aiding in reasoning about the behaviour of procedural [11] and functional [31] programs, as well as concurrent systems [40] and smart contracts [50], to name a few examples.

To enable the automated reasoning of FOL formulas different logic solvers can be used, depending on the fragment selected. The most common categories of such tools are arguably Boolean satisfiability (SAT) and satisfiability modulo theories (SMT) solvers [44], which respectively solve formulas

Authors' Contact Information: [Rodrigo Otoni](mailto:otonir@usi.ch), otonir@usi.ch, USI Lugano, Lugano, Switzerland; [Martin Blicha](mailto:blichm@usi.ch), blichm@usi.ch, USI Lugano, Lugano, Switzerland and Charles University, Prague, Czech Republic; [Patrick Eugster](mailto:eugstp@usi.ch), eugstp@usi.ch, USI Lugano, Lugano, Switzerland; [Natasha Sharygina](mailto:sharygin@usi.ch), sharygin@usi.ch, USI Lugano, Lugano, Switzerland.

in the propositional fragment of FOL and in extensions of it with theories such as arithmetics, arrays, and bit vectors. CHC solvers are also available, e.g., ELDARICA [41], GOLEM [13], and SPACER [43], serving, for instance, as back-end reasoning engines of verification tools targeting programs written in C/C++ [33], Java [42], Rust [47], and Solidity [1], as well as Android applications [19].

1.1 Need for Validation

Despite their extensive usage in verification, logic solvers are themselves not immune to bugs. To illustrate this point, in the 2023 edition of the annual SMT competition there were 300 benchmarks in which at least two state-of-the-art solvers disagreed on the results, with in total 10 competing solvers producing unsound results [16]. In light of this, having guarantees about solvers' results is of paramount importance. One approach to achieve this goal is to formally verify the solvers' code, as has been done for read-eval-print loop (REPL) [45] and garbage collector [54] implementations. Despite the strong guarantees provided, this approach incurs a high cost to verify the existing codebase and any future modifications to it, as well as potentially preventing many code optimizations to be made, which are essential for solver performance. Another, less invasive, approach, is to validate solvers' outputs, rather than verifying the solvers themselves. This requires a solver, in addition to producing its standard output, to also produce a witness that can be used by an independent tool to validate the given result. Currently, the community is moving towards the second approach, with many witness formats being proposed to validate the outputs of SAT [4, 21, 35, 37, 56, 60] and SMT [22, 39, 48, 55, 57] solvers, with codebase verification being applied instead to the validation tools [36, 46], which are much less complex and easier to maintain. Both the annual SAT and SMT competitions now incorporate witness usage in some capacity¹ and the organisers of the annual CHC competition, CHC-COMP, have the validation of witnesses as a goal for future editions [27].

1.2 CHC Witnesses

The input of a CHC solver, detailed in Section 3, is a conjunction of logical implications containing uninterpreted predicates, with the task of the solver being to decide if *false* can be derived or not. If it can the input is considered unsatisfiable, UNSAT for short, and if it cannot the input is considered satisfiable, SAT for short, not to be confused with Boolean satisfiability. A witness for an UNSAT result, called *UNSAT proof*, should contain an explanation of how *false* can be derived, while a witness for a SAT result, called *SAT model*, should contain interpretations for all the predicates in such a way that all clauses evaluate to *true*, entailing that *false* cannot be derived; for the remainder of the article UNSAT proofs and SAT models will be referred to simply as *proofs* and *models*.

The production of witnesses is a common feature of modern CHC solvers, but efforts in witnesses validation are limited at present. The validation of models is done via SMT queries, and is currently supported only by an ad hoc validator tied to the SMT solver Z3 [23]². Unlike the case for models, however, the proofs produced cannot be independently validated with available tooling, given that, to the best of our knowledge, no proof checking approach currently exists.

1.3 Two-Layered Model Validation

Since CHC model validation is underpinned by SMT solving, the same concern regarding the correctness of CHC solvers' results is put on the validation itself, i.e., on the correctness of SMT solvers' results. To address this, we propose a two-layered validation approach to provide additional guarantees about the results obtained, illustrated in Figure 1. The first layer, consisting of the

¹The SAT competition requires its competitors to produce witnesses since its 2013 edition, while the SMT competition started an exhibition track for this, still separate from the main tracks, in its 2022 edition.

²See <https://github.com/chc-comp/chc-tools/blob/master/chctools/chcmodel.py>.

SMT queries responsible for model validation, is enhanced by a second layer, consisting of the production and checking of SMT proofs, with the result obtained being forwarded to the user or tool interacting with the CHC solver. The approach is generic w.r.t. FOL theories and solvers, and is also very modular, enabling different SMT solvers to be used in the validation, further increasing assurances.

1.4 The ATHENA Framework

To assess the viability of practical model validation we developed a modular evaluation framework, called ATHENA, capable of catering to different combinations of state-of-the-art CHC and SMT solvers, and used it to conduct a large scale evaluation. Concretely, we used our framework to validate the models produced by three CHC solvers, ELDARICA [41], GOLEM [13], and SPACER [43], with each model produced being separately validated, in Layer 1, by five proof producing SMT solvers, CVC5 [5], OPENSMIT [17], SMT-INTERPOL [20], VERiT [15], and Z3 [23].

In addition, we checked, in Layer 2, all the proofs produced in the proof formats currently supported by automated proof checkers, by using the checkers ALFC [52], CARCARA [2], LFSC checker [57], SMTINTERPOL checker [39], and TSWC [48].

1.5 Contributions

In addition to the detailed description of our validation approach and the ATHENA framework, we also report the results of an extensive evaluation targeting two FOL theories, namely the linear integer arithmetic (LIA) theory and the theory of arrays combined with LIA (ALIA). We used all LIA and ALIA benchmarks from CHC-COMP 2024 in our evaluation, 450 containing only linear Horn clauses, i.e., implications with a single uninterpreted predicate in the implicant, and 600 containing nonlinear Horn clauses, i.e., implications with multiple uninterpreted predicates in the implicant. The benchmarks used led to 63994 SMT instances and 299658 SMT proofs being produced as part of the validation process.

Three observations can be made from the results obtained. First, the proposed proof-backed model validation approach is viable in practice, with the majority of the models being validated with available tooling. This means that any CHC-based tool, e.g., the SEAHORN [33] and SolCMC [1] model checkers, can in principle benefit from the guarantees provided by model validation. In particular, invariant generation in CHC-based tools, which commonly relies on models, can be made more robust. Second, model and proof sizes, which reached the order of gigabytes in our experiments, are a concern and a potential limitation to the practical use of validation. Producing compact models and proofs is thus an important goal, with compression, recently investigated in the context of unsatisfiability proofs for SAT solvers [51], being a potential complementary goal.

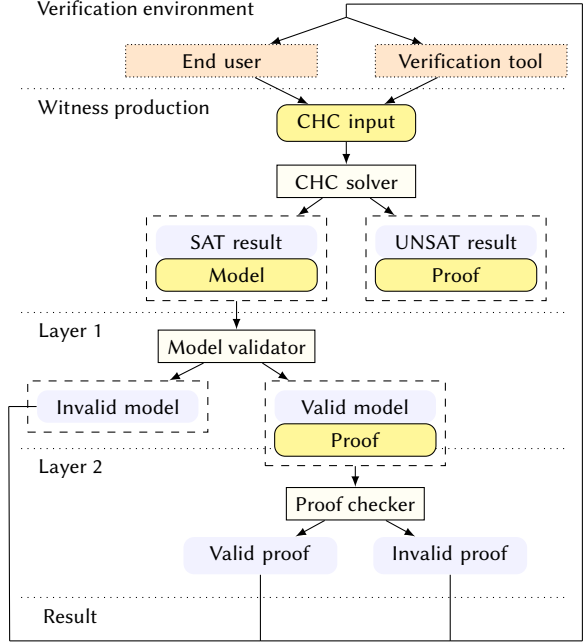


Fig. 1. Overview of our two-layered approach for validation of CHC satisfiability. CHC models are validated through proof-backed SMT queries. Although capable of being produced, CHC proofs cannot be checked currently.

Table 1. Brief descriptions of the bugs found during the evaluation.

		Bug effect
CHC solvers	ELDARICA	Invalid model produced ⁹
	ELDARICA	Crash due to quantifiers in predicates ⁶
	SPACER	Invalid model produced ¹¹
	GOLEM	Syntactically malformed model produced (x2) ⁸
	GOLEM	Crash during model production ⁷
SMT solvers	CVC5	Invalid proof produced ¹⁵
	CVC5	Crash during proof production (x2) ¹³
	OPENSMT	Crash during sort inference ¹²
	SMTINTERPOL	Invalid proof produced ¹⁴
Proof checkers	ALFC	Crash during proof checking ¹⁸
	CARCARA	Parsing error due to unknown attribute ¹⁶
	LFSC checker	Crash during type inference ¹⁷

Lastly, model validation provides a useful way to generate new and interesting SMT instances. Our evaluation uncovered bugs not only in the selected CHC solvers, which are our main focus, but also in three SMT solvers and three proof checkers for SMT proofs. The bugs found, listed in Table 1, range from parsing errors to invalid models being produced. They were all acknowledged by the developers and are detailed in Section 6. In addition to aiding in tool development, these bugs confirm the need for additional guarantees to be provided to modern verification tooling.

To summarise, our contributions are the following:

- (1) Proposal of a two-layered approach to validate CHC satisfiability results.
- (2) Development of a modular evaluation framework, ATHENA, to make the approach practical.
- (3) Staging of a large scale evaluation to determine the viability of our approach.

An earlier version of this work was published as a conference paper [49]. The present article substantially extends and improves that previous one, in the following fivefold manner:

- (I) Addition of support for the theory of arrays with linear integer arithmetic to ATHENA, which demonstrates our approach’s generality in regards to FOL theories.
- (II) Integration to our framework of (i) a postprocessing step for aggregation of layers’ results per input file and of (ii) a new SMT proof format, ALETHELF, and of its checker, ALFC, which together provide more information and validation options to the user.
- (III) Extension of our evaluation to include 1050 new benchmarks, timeouts of 30 minutes, and the latest stable versions of thirteen state-of-the-art tools, which combined provide an extremely detailed picture of CHC satisfiability validation in practice.
- (IV) Reporting of five new bugs uncovered during our extended evaluation, including the production of invalid proofs by SMTINTERPOL.
- (V) Improvement of explanations in many parts of the manuscript, including the addition of an algorithmic description of our validation approach (in Section 4).

1.6 Roadmap

The remainder of the paper is structured as follows. Related witness validation approaches are covered in Section 2. The necessary CHC background is given in Section 3. The two layers of our

validation approach are presented in [Section 4](#). ATHENA is detailed in [Section 5](#). The evaluation is discussed in [Section 6](#). Finally, our conclusions are laid out in [Section 7](#).

2 Related Work

As logic solvers became more powerful they were quickly adopted as the back-end reasoning engines of many verification tools. The need to validate the answers from these solvers arose soon after, with the complexity of the validation increasing hand-in-hand with the expressiveness of the underlying formalism. In this section we provide an overview of related witness validation approaches, targeting SAT and SMT solving as well as other contexts. For details on SAT and SMT solving we refer readers to the book of D. Kroening and O. Strichman [\[44\]](#).

2.1 Validation of SAT Witnesses

In line with its relative simplicity, witness validation in the context of Boolean satisfiability was the first to be investigated. Validating a satisfying model is an easy task: one simply substitutes the variables of the formula with their values from the model and checks if the resulting Boolean expression over constants *true* and *false* simplifies to *true*. The validation of unsatisfiability proofs, however, is far from trivial, even in such a restricted domain. Many proof formats have been proposed, offering different trade-offs between proof compactness and checking efficiency. Initial formats were based on resolution [\[56\]](#) and clausal proofs [\[35\]](#), with resolution asymmetric tautology (RAT) [\[37\]](#) being a base for many recent developments, e.g., deletion RAT (DRAT) [\[60\]](#), linear RAT (LRAT) [\[21\]](#), and flexible RAT (FRAT) [\[4\]](#). The production of proofs in the DRAT format has been a requirement in the SAT competition since its 2014 edition, with DRAT-TRIM [\[60\]](#) being the standard proof checker for proofs following this format.

2.2 Validation of SMT Witnesses

In regards to satisfiability modulo theories, witness validation is complicated by the presence of theories and quantifiers. No standard way of representing SMT models currently exists, with a consistent push by the organisers of the SMT competition having been made in recent years for the adoption of a unified format in line with the SMT-LIB standard [\[7\]](#). A separate, experimental, model validation track has been established and has seen a steady increase in the SMT-LIB logics supported, with PySMT [\[30\]](#) and DOLMEN [\[18\]](#) used as validating tools. Despite recent advances, model validation is still restricted to quantifier-free formulas. For unsatisfiability proofs, different formats, often attached to a specific solver, have been proposed. The ALETHE format [\[55\]](#) was initially supported by VERiT, but has since also been integrated into cvc5's proof production. cvc5 also caters for proofs based on the logical framework with side conditions (LFSC) [\[57\]](#), with LFSC support preceding ALETHE's integration and dating back to the CVC3 version of the tool. SMTINTERPOL [\[39\]](#), OPENSMT [\[48\]](#), and Z3 [\[22\]](#) also support their own, unnamed, proof formats. Each format has one or more associate tools that can consume the proofs produced, with said tools being either interactive or automated. In the interactive side, proof assistants discharge some verification conditions to external logic solvers as a way to increase their level of automation, with the proofs produced providing new theorems to be checked by the proof assistant's internal engine, as it has been done with Coq [\[3, 25\]](#) and ISABELLE/HOL [\[6, 12, 14, 28\]](#). When it comes to automated checkers, their goal is mainly to serve as independent lightweight validators, with potential to be integrated into tools such as model checkers. Automated checkers are available for a variety of formats [\[2, 39, 48, 57, 60\]](#). As of the time of writing, no proof format is enforced by the SMT competition, with an experimental track being available as a way to showcase the existing formats.

2.3 Validation of Other Types of Witnesses

In addition to logic solving, witness validation is also pursued in other contexts. A good example of this is the use of validation in the annual competition on software verification [9]. Software verification witnesses are different from those used by logic solvers, being categorized as either correctness or violation witnesses, with their own formats³ and limitations [10]. The tool that maybe best illustrates usage of witness is Korn [26], a participant in the software verification competition that relies on Horn solvers as its back-end and produces witnesses for its reasoning about C programs' properties from the witnesses produced by the underlying solvers.

3 Constrained Horn Clauses

The constrained Horn clauses formalism has been proposed as a unified, purely logic-based, intermediate language for reasoning about verification tasks [31]. It builds upon the success achieved with SAT and SMT, using logical constraints to capture various verification tasks, including safety, termination, and loop invariant computation, from a variety of domains. In this section the necessary CHC background is presented.

Following standard SMT terminology [8], we assume a first-order theory $\mathcal{T}$ and a set of uninterpreted predicates $\mathcal{P}$ disjoint from the signature of $\mathcal{T}$. A constrained Horn clause is a formula $\varphi \wedge P_1 \wedge \dots \wedge P_n \implies H$, where φ is an interpreted formula in the language of $\mathcal{T}$, each P_i is an application of a symbol $p \in \mathcal{P}$ to terms of $\mathcal{T}$, H is either an application of a symbol $p \in \mathcal{P}$ to terms of $\mathcal{T}$ or *false*, and all variables in the formula are implicitly universally quantified. Commonly, the antecedent and the consequent of the implication are denoted as the *body* and the *head* of the Horn clause, and φ is referred to as its *constraint*. Furthermore, a clause is usually called a *query* if its head is equal to *false* and a *fact* if no uninterpreted predicates are present in its body.

Given a set of constrained Horn clauses $\mathcal{S}$ over the uninterpreted predicates $\mathcal{P}$ and theory $\mathcal{T}$, we say that $\mathcal{S}$ is satisfiable if there exists a model $\mathcal{M}$ of $\mathcal{T}$ extended with an interpretation for all the uninterpreted predicates $\mathcal{P}$ such that all the clauses evaluate to *true* in $\mathcal{M}$, i.e., every clause evaluates to *true* in $\mathcal{M}$ for all possible instantiations of the universally quantified variables. In practice, we are interested in interpretations that are definable in the language of $\mathcal{T}$, i.e., we want a mapping of the predicates to a set of formulas in the language of $\mathcal{T}$ such that each clause from $\mathcal{S}$ is valid in $\mathcal{T}$ after the uninterpreted predicates are replaced by their interpretations. This is called *syntactic solvability*, as opposed to the more general *semantic solvability* [53]. Having LIA as our theory $\mathcal{T}$, an example of an interpretation is that of a set of (in)equalities, with the elements of the set being conjoined when replacing the predicate they are an interpretation for. The interpretations of the predicates from the discovered model serve as witnesses for the satisfiability answer. An unsatisfiability answer, on the other hand, needs to be witnessed by a derivation of *false* from the original clauses, likely via universal instantiation and resolution.

4 Validation of CHC Models

A CHC solver is a complex piece of software, often implementing sophisticated algorithms relying on decision and interpolation procedures, which allows for subtle bugs to occur and lead to incorrect answers. In addition to providing much needed stronger guarantees in regards to SAT results, model validation also ensures the correctness of the models themselves, which are commonly relied upon by verification tools to, for instance, establish inductive invariants of programs. Model validation is therefore critical for assurance not only of solvers' results, but also of all structures derived from models presented to end users.

³See <https://gitlab.com/sosy-lab/benchmarking/sv-witnesses>.

We propose a two-layered validation approach for CHC models, detailed in Figure 2; since our focus is on models, the illustration assumes the benchmark is satisfiable. The first of the two layers in the approach handles model validation via SMT solving. Following the CHC model definition laid out in Section 3, model validation can be done via a number of SMT queries which is linear w.r.t. to the number of Horn clauses present in the input. Each such query checks if a specific Horn clause is logically valid in the theory $\mathcal{T}$ after its uninterpreted predicates are substituted by their interpretations given by the model. This is done by checking if the negation of the Horn clause, augmented with the interpretations, is satisfiable, i.e., if a satisfying assignment for $\varphi \wedge P_1 \wedge \dots \wedge P_n \wedge \neg H$ exists. This check is well suited for SMT solving, with a valid model leading to all queries being unsatisfiable. An important note is that, depending on the theory $\mathcal{T}$, the query checking might be intractable for existing SMT solvers, and in some cases even undecidable.

As an example of the computation done in the first layer of our approach, consider the following CHC system, consisting of three Horn clauses and a single uninterpreted predicate *Inv*:

$$x \leq 0 \implies \text{Inv}(x)$$

$$\text{Inv}(x) \wedge x < 5 \wedge x' = x + 1 \implies \text{Inv}(x')$$

$$\text{Inv}(x) \wedge \neg(x < 10) \implies \text{false}$$

This system is satisfiable with a potential model being one that contains the interpretation $\text{Inv}(x) \equiv x \leq 5$. To validate this model we need to establish that the following three formulas are logically valid in the LIA theory:

$$x \leq 0 \implies x \leq 5$$

$$x \leq 5 \wedge x < 5 \wedge x' = x + 1 \implies x' \leq 5$$

$$x \leq 5 \wedge \neg(x < 10) \implies \text{false}$$

The validation can be done by showing that the three formulas below are all unsatisfiable in the LIA theory, which can be trivially seen for this small example as the assignments to x in each formula lead to a contradiction:

$$x \leq 0 \wedge \neg(x \leq 5)$$

$$x \leq 5 \wedge x < 5 \wedge x' = x + 1 \wedge \neg(x' \leq 5)$$

$$x \leq 5 \wedge \neg(x < 10) \wedge \neg \text{false}$$

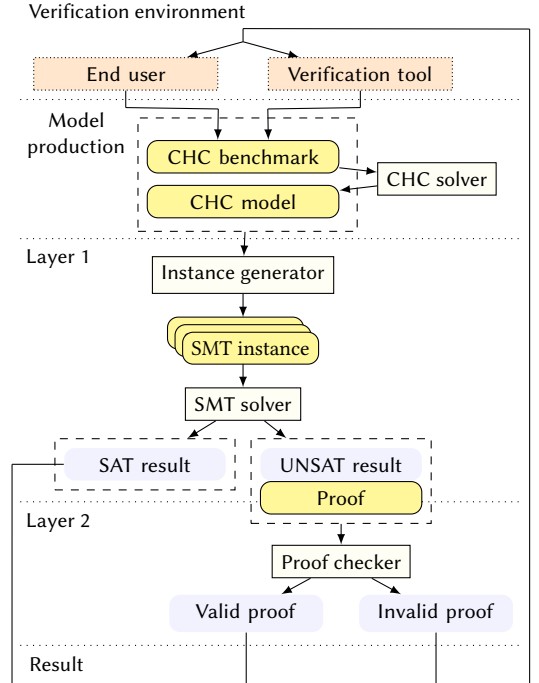


Fig. 2. Breakdown of our two-layered approach for validation of CHC satisfiability. A valid model will have all the SMT instances generated from it yield an UNSAT result backed by a valid SMT proof.

While it can be easy to validate models such as the one above, this is far from the case when dealing with real world examples. As a consequence, SMT solvers are, like their CHC counterparts, very complex tools that are susceptible to bugs. The second layer in the approach we propose tackles this issue via the validation of SMT solvers' results. A number of SMT solvers produce unsatisfiability proofs that can be independently checked. These proofs provide much needed guarantees regarding unsatisfiability results, which are at the core of the validation done in Layer 1.

input : pair (S, M) , with S being a set of CHC clauses and M being a candidate model
output : Boolean value representing whether M is a valid model for S or not

```

1 foreach clause  $C \in S$  do
2    $C_I \leftarrow \text{InstantiateClause}(C, M)$ 
3    $(\text{answer}, \text{proof}) \leftarrow \text{CheckSMTSatisfiability}(\neg C_I)$ 
4   if  $\text{answer} \neq \text{UNSAT}$  then return false
5   if not  $\text{ValidateSMTProof}(\neg C_I, \text{proof})$  then return false
6 end
7 return true

```

Algorithm 1: Validation of a CHC model.

By relying on the currently untapped power of SMT proofs we provide additional correctness guarantees for CHC model validation; a generic example of the computation done in the second layer is not possible because it is specific to the proof format and proof checker being used.

The pseudocode describing the whole validation process can be seen in [Algorithm 1](#). The `InstantiateClause` function replaces the uninterpreted predicates in the given clause by the interpretations present in the model, while the `CheckSMTSatisfiability` and `ValidateSMTProof` functions stand for calls to an SMT solver and a proof checker. Our approach is theory independent and can be applied to any CHC, with the only requirement being that a proof producing SMT solver and a proof checker are available for the theory in question. In addition to validating direct end user usage of CHC solvers, our approach can also be embedded into CHC-based verification tools, enhancing their own guarantees.

5 Implementation

To enable the practical use of our approach, with the immediate goal of ascertaining the capabilities of state-of-the-art CHC and SMT solvers, we developed the modular constrained Horn clauses model validation framework, ATHENA for short. Our framework is capable of validating CHC models via SMT solving while using different solver combinations. ATHENA also handles the production and checking of SMT proofs, for the SMT solvers with proof production capabilities. In addition, metrics such as model and proof sizes can be gathered and analysed. The framework consists of 3485 lines of Shell and Python code in total, is fully automated, and uses GNU PARALLEL [58] to achieve a large degree of parallelisation in order to better tackle the high computation cost. One important assumption is the correctness of the instance generator, i.e., the `InstantiateClause` function in [Algorithm 1](#). The generator’s small size, 197 lines of Python code, makes manual inspection a reasonable way to ensure correctness at this stage. ATHENA is open-source⁴, enabling third-parties to make full use of it, with one of our goals being to provide the groundwork for model validation at CHC-COMP.

6 Evaluation

We first describe the benchmarks and tools used, in [Section 6.1](#), and then discuss the results obtained related to CHC model validation, i.e., layer 1, in [Section 6.2](#), and SMT proof checking, i.e., layer 2, in [Section 6.3](#). We used a machine with 64 AMD EPYC 7452 processors and 256 GB of memory for the evaluation. All individual tool executions had a timeout of 30 minutes and a memory limit of 10 gigabytes. In total, our evaluation took 2392 hours of CPU time.

⁴Available at <https://github.com/usi-verification-and-security/athena>.

Table 2. Results for solving the CHC benchmarks of the LIA and LIA-Arrays tracks.

		SAT	UNSAT	Unknown	Timeout	Memout	Error
LIA-lin (300 benchmarks)	ELDARICA	130	59	0	110	1	0
	GOLEM	115	67	0	118	0	0
	SPACER	111	60	2	88	39	0
LIA-nonlin (300 benchmarks)	ELDARICA	152	92	0	55	1	0
	GOLEM	145	98	0	57	0	0
	SPACER	175	95	2	25	3	0
LIA-Arrays-lin (150 benchmarks)	ELDARICA	56	23	0	70	0	1
	GOLEM	-	-	-	-	-	-
	SPACER	29	27	0	94	0	0
LIA-Arrays-nonlin (300 benchmarks)	ELDARICA	79	66	0	153	0	2
	GOLEM	-	-	-	-	-	-
	SPACER	87	109	0	98	6	0

6.1 Benchmarks and Tools

We used all benchmarks of the two LIA and the two LIA-Arrays tracks of CHC-COMP 2024 [27]. The LIA-lin and LIA-Arrays-lin tracks consist of benchmarks containing only linear Horn clauses, and the LIA-nonlin and LIA-Arrays-nonlin tracks consist of benchmarks containing nonlinear Horn clauses. We decided to use LIA benchmarks for two reasons: first, the LIA tracks are the most traditional in CHC-COMP, being present in every edition of the competition and having the most competing solvers, and second, the LIA theory is covered by all proof producing SMT solvers available, even if for some only in its quantifier-free fragment. Regarding the LIA-Arrays benchmarks, their tracks are also very traditional, being present in the competition since its second edition and also having a significant number of competitors, and the ALIA theory is covered by three out of five of the available proof producing SMT solvers.

For model production we chose the three best performing CHC solvers in the LIA tracks of CHC-COMP 2023 for comparison, which are, in alphabetical order, ELDARICA [41], GOLEM [13], and SPACER [43]. Regarding the other competitors, LOAT [29] and THETA [59] are currently not capable of producing models, and the two competitors based on the ULTIMATE framework, ULTIMATE TREEAUTOMIZER [24] and ULTIMATE UNIHORN [34], have no clear build instructions available.

For model validation we used all SMT solvers that competed in the proof exhibition track of the 2022 SMT competition, which are, in alphabetical order, cvc5 [5], OPENSMT [17], SMTINTERPOL [20], and VERIT [15], as well as Z3 [23], which can produce proofs but did not compete in the track. To check the SMT proofs we used the fully automated checkers ALFC [52], for ALETHELF proofs produced by cvc5, CARCARE [2], for ALETHE proofs produced by cvc5 and VERIT, LFSC checker [57], for LFSC proofs produced by cvc5, SMTINTERPOL checker [39], for proofs produced by SMTINTERPOL, and TSWC [48], for proofs produced by OPENSMT; to the best of our knowledge there is currently no independent automated checker for proofs produced by Z3⁵.

6.2 Model Validation Results

To produce the CHC models we executed the selected CHC solvers with all benchmarks; the results are summarised in Table 2. All tools were executed with their default engine configurations, with

⁵For details on proof efforts involving the Z3 solver see [https://microsoft.github.io/z3guide/programming/Proof Logs](https://microsoft.github.io/z3guide/programming/Proof%20Logs).

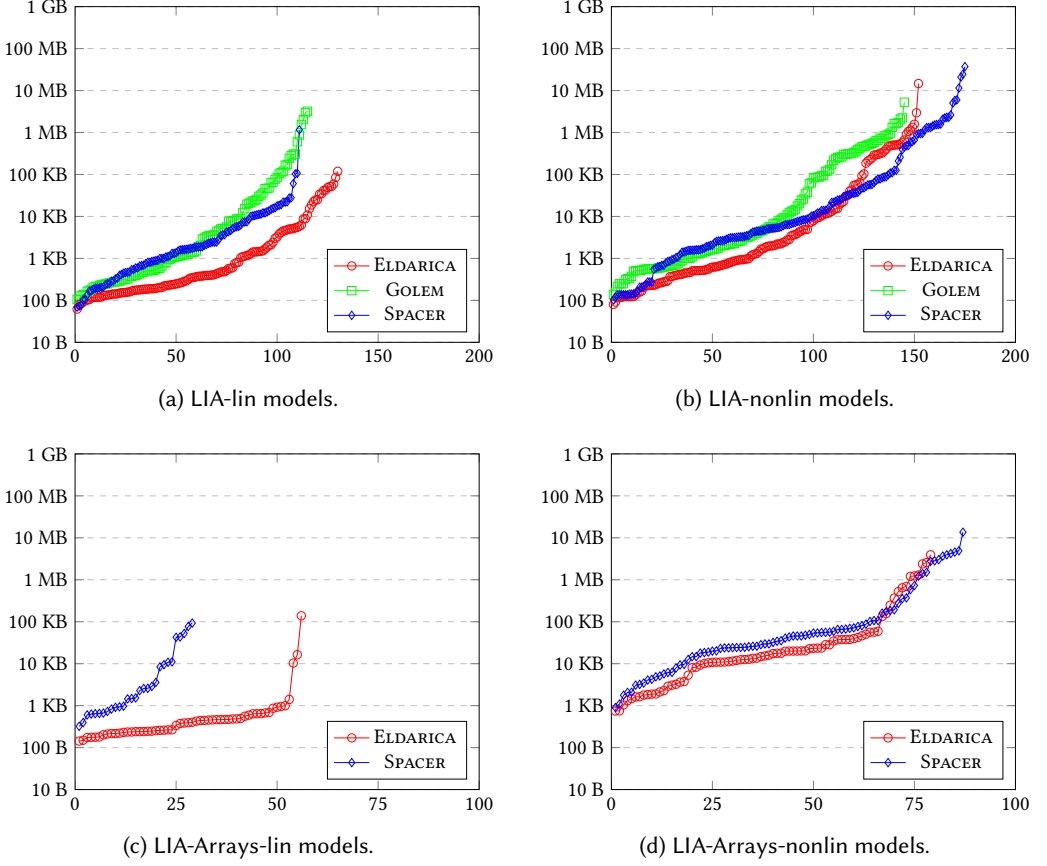


Fig. 3. Sizes of the CHC models. The models are ordered according to their size, the x-axis indicates their position in the order and the y-axis indicates their size.

GOLEM currently not supporting benchmarks containing arrays. The performance of the tools is in line with the CHC-COMP results, with an unknown result, meaning that the solver terminated within the allocated time frame but was not able to decide if the benchmark was satisfiable or not, being yielded only by SPACER in four of its executions. Three errors, i.e., tool crashes, were observed with ELDARICA, due to the presence of quantifiers in predicate interpretations computed⁶. A number of crashes were also observed with GOLEM while testing our framework⁷, as well as syntactically malformed models being produced by it⁸, with the underlying causes of both issues being addressed before the full-scale evaluation. Regarding the models' sizes (in bytes), ELDARICA's models tended to be the most compact, followed by SPACER's, with GOLEM producing most of the larger models, as can be seen in Figure 3. The last point of note is that all models produced by GOLEM are quantifier-free, while ELDARICA produced 68 quantified model, 7 from LIA-Arrays-lin benchmarks and 61 from LIA-Arrays-nonlin benchmarks, and SPACER produced 238 quantified

⁶See <https://github.com/uuverifiers/eldarica/issues/46> and [../eldarica/issues/39](https://github.com/uuverifiers/eldarica/issues/39).

⁷See <https://github.com/usi-verification-and-security/golem/issues/29>.

⁸See <https://github.com/usi-verification-and-security/golem/issues/27> and [../golem/issues/68](https://github.com/usi-verification-and-security/golem/issues/68).

Table 3. Results for solving the SMT instances generated from the LIA models.

		SAT	UNSAT	Unknown	Timeout	Memout	Error	Unsupported	
(1433 instances)	LIA-lin	CVC5	0	1426	0	1	0	6	0
		OPENSMT	0	1427	0	0	0	6	0
	ELDARICA	SMTINTERPOL	0	1427	0	0	0	6	0
		VERiT	0	1348	0	0	0	6	79
		Z3	0	1433	0	0	0	0	0
(1004 instances)	LIA-lin	CVC5	0	1004	0	0	0	0	0
		OPENSMT	0	1004	0	0	0	0	0
	GOLEM	SMTINTERPOL	0	1004	0	0	0	0	0
		VERiT	0	964	0	1	0	0	39
		Z3	0	1004	0	0	0	0	0
(1273 instances)	LIA-lin	CVC5	4	1269	0	0	0	0	0
		OPENSMT	0	647	0	0	0	0	626
	SPACER	SMTINTERPOL	4	1268	1	0	0	0	0
		VERiT	0	602	0	0	0	0	671
		Z3	4	1269	0	0	0	0	0
(13540 instances)	LIA-nonlin	CVC5	0	13529	0	0	0	11	0
		OPENSMT	0	13529	0	0	0	11	0
	ELDARICA	SMTINTERPOL	0	13529	0	0	0	11	0
		VERiT	0	13463	0	0	0	11	66
		Z3	1	13539	0	0	0	0	0
(13109 instances)	LIA-nonlin	CVC5	0	13109	0	0	0	0	0
		OPENSMT	0	13109	0	0	0	0	0
	GOLEM	SMTINTERPOL	0	13106	0	3	0	0	0
		VERiT	0	13102	0	0	7	0	0
		Z3	0	13105	0	0	4	0	0
(14716 instances)	LIA-nonlin	CVC5	39	14677	0	0	0	0	0
		OPENSMT	3	5028	0	0	0	0	9685
	SPACER	SMTINTERPOL	39	14297	146	199	35	0	0
		VERiT	3	5027	0	0	0	0	9686
		Z3	39	14675	0	2	0	0	0

models, 42 from LIA-lin benchmarks, 97 from LIA-nonlin benchmarks, 19 from LIA-Arrays-lin benchmarks, and 80 from LIA-Arrays-nonlin benchmarks.

To validate each model we executed the SMT instances generated from it with the selected SMT solvers; in this section all the reported SMT solvers' executions were done with proof production disabled. One SMT instance is generated for each Horn clause in the CHC benchmark for which the model was produced, thus many SMT instances, sometimes hundreds, can be generated for a single model. We consider the SMT instances generated for the models produced by each CHC solver, by track, as separate instance sets, thus we have three instance sets per track. The results for the LIA and LIA-Arrays tracks can be seen in [Table 3](#) and [Table 4](#), respectively; the number of SMT instances generated for each CHC solver is related to the number of models it produced, with

Table 4. Results for solving the SMT instances generated from the LIA-Arrays models.

		SAT	UNSAT	Unknown	Timeout	Memout	Error	Unsupported
LIA-Arrays-lin (880 instances)	CVC5	0	880	0	0	0	0	0
	OPENSMT	0	844	0	0	0	0	36
	ELДАРICA	0	880	0	0	0	0	0
	SMTINTERPOL	0	880	0	0	0	0	0
	VERiT	-	-	-	-	-	-	-
LIA-Arrays-lin (365 instances)	Z3	0	880	0	0	0	0	0
	CVC5	0	356	5	4	0	0	0
	OPENSMT	0	140	0	0	0	0	225
	SPACER	0	355	5	5	0	0	0
	SMTINTERPOL	0	355	5	5	0	0	0
LIA-Arrays-nonlin (8348 instances)	VERiT	-	-	-	-	-	-	-
	Z3	0	363	2	0	0	0	0
	CVC5	0	5348	106	0	0	2894	0
	OPENSMT	0	311	0	0	0	169	7868
	ELДАРICA	0	5305	86	86	10	2861	0
LIA-Arrays-nonlin (9326 instances)	SMTINTERPOL	0	5305	86	86	10	2861	0
	VERiT	-	-	-	-	-	-	-
	Z3	0	5407	31	16	9	2885	0
	CVC5	17	8577	710	20	1	0	0
	OPENSMT	0	189	0	0	0	0	9137
LIA-Arrays-nonlin (9326 instances)	SPACER	5	8416	259	638	8	0	0
	SMTINTERPOL	5	8416	259	638	8	0	0
	VERiT	-	-	-	-	-	-	-
LIA-Arrays-nonlin (9326 instances)	Z3	16	8688	583	39	0	0	0
	VERiT	-	-	-	-	-	-	-

VERiT currently not supporting instances containing arrays. The combination of related tools, e.g., pairing of SPACER with Z3 or Golem with OPENSMT, might yield better solving performance, but it weakens the strength of independent validation.

The validation results provide a useful insight into the quality of the models produced by each CHC solver. The models produced by Golem are the only ones for which no invalid result, i.e., a SAT output from the SMT solver, was observed. Both ELДАРICA and SPACER produced invalid models, although the latter to a significantly higher degree. ELДАРICA's invalid model is due to the embedding of Boolean values into arithmetic operations⁹, which leads to an error in most SMT solvers and is the cause of the errors reported in the tables, but can be solved by Z3 in some cases via unit propagation¹⁰. SPACER's invalid models are due to problematic internal transformations¹¹. While not implying that the SAT results the models are supposed to serve as witnesses for are incorrect, since the problem can be in the model production itself, this is a serious issue. Another aspect of model quality is the presence of quantifiers, which can make solving harder and is unsupported by both OPENSMT and VERiT. Two last points of note are the occurrence of errors when OPENSMT was handling instances generated from ELДАРICA and SPACER models in our preliminary experiments, which was due to a limitation in scoping in the presence of different sorts¹² that was fixed prior to the full-scale evaluation, and the small, but consistent, number of

⁹See <https://github.com/uuverifiers/eldarica/issues/51>.

¹⁰See <https://github.com/Z3Prover/z3/issues/6719>.

¹¹See <https://github.com/Z3Prover/z3/issues/6716>.

¹²See <https://github.com/usi-verification-and-security/opensmt/issues/613>.

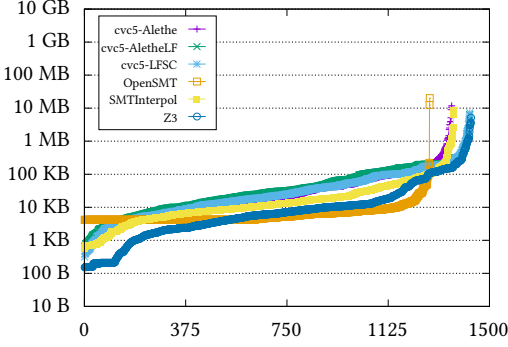
Table 5. Number of proofs produced by the selected SMT solvers; cvc5 has separate results for its three proof formats. Each column shows the amount of proofs produced from a given instance set, with the total amount of instances in each set shown below the CHC solver that produced the models for it.

	Proofs produced					
	LIA-lin			LIA-nonlin		
	ELDARICA	GOLEM	SPACER	ELDARICA	GOLEM	SPACER
	(1433)	(1004)	(1273)	(13540)	(13109)	(14716)
cvc5-ALETHE	1360	865	1029	12941	12039	9473
cvc5-ALETHELF	1427	1000	1269	13528	13106	14675
cvc5-LFSC	1427	1002	1269	13528	13109	14676
OPENSMT	1279	1004	647	12208	13109	1336
SMTINTERPOL	1368	935	1199	13244	12827	13736
VERiT	0	0	0	0	0	0
Z3	1432	1004	1269	13539	13104	14675
	LIA-Arrays-lin			LIA-Arrays-nonlin		
	ELDARICA	GOLEM	SPACER	ELDARICA	GOLEM	SPACER
	(880)	-	(365)	(8348)	-	(9326)
cvc5-ALETHE	-	-	-	-	-	-
cvc5-ALETHELF	880	-	344	3414	-	5427
cvc5-LFSC	880	-	345	4709	-	7943
OPENSMT	-	-	-	-	-	-
SMTINTERPOL	880	-	353	5216	-	8294
VERiT	-	-	-	-	-	-
Z3	880	-	364	5405	-	8694

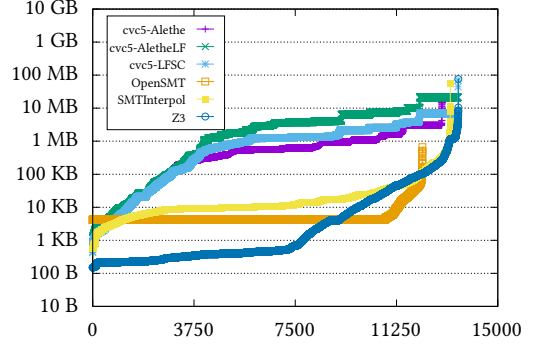
instances unsupported solely by VERiT. After a manual inspection, it was discovered that the lack of support observed with VERiT is due to the LIA tracks containing some benchmarks that, although semantically belonging to the LIA fragment of first-order logic, use operators reserved for the nonlinear integer arithmetic (NIA) logic of the SMT-LIB standard; the competition organizers were informed of this finding.

6.3 Proof Checking Results

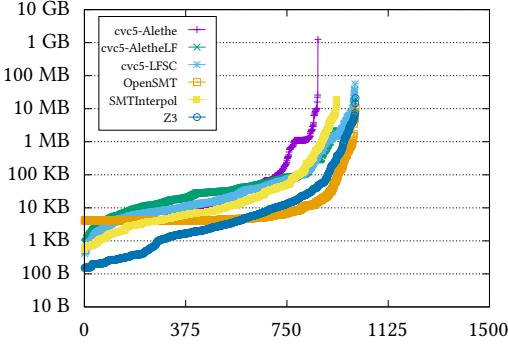
To validate the UNSAT results given by the SMT solvers we rely on the proofs produced by them. For each SMT instance generated from a CHC model we executed the selected SMT solvers in proof production mode. The number of proofs produced by each SMT solver can be seen in Table 5. In addition to VERiT not handling arrays, cvc5-ALETHE and OPENSMT cannot produce proofs for instances containing arrays. Since proof production adds an overhead to solver execution, the number of proofs produced is expected to be lower than the amount of UNSAT results reported in Table 3 and Table 4. Concretely, the combined percentage of proofs produced in relation to the previous UNSAT results, for the ten instance sets (or six if arrays are not supported), is: 83% for cvc5-ALETHE, 91% for cvc5-ALETHELF, 97% for cvc5-LFSC, 85% for OPENSMT, 97% for SMTINTERPOL, 0% for VERiT, and 100% for Z3. The reduction in performance observed is overall small for all solvers, with the main reason for the lower amount of UNSAT results in proof production mode being the errors present exclusively in this mode. The production of proofs by cvc5, particularly of ALETHE



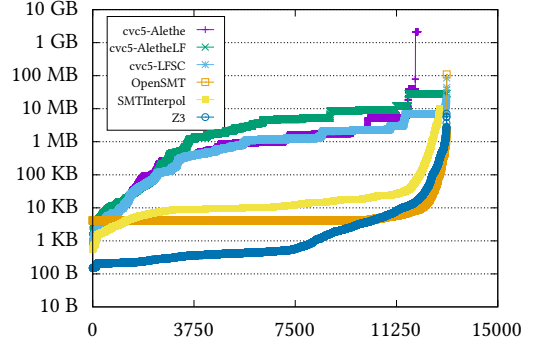
(a) LIA-lin proofs for ELDARICA models.



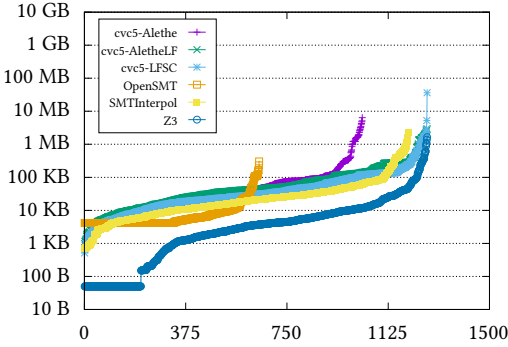
(b) LIA-nonlin proofs for ELDARICA models.



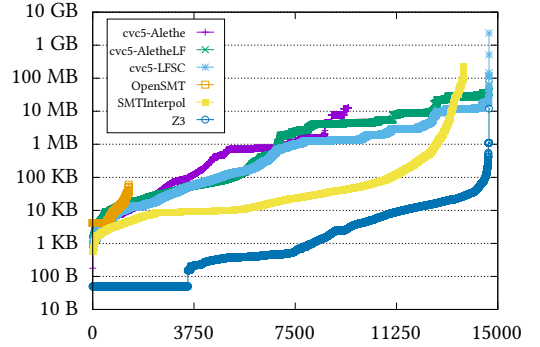
(c) LIA-lin proofs for GOLEM models.



(d) LIA-nonlin proofs for GOLEM models.



(e) LIA-lin proofs for SPACER models.

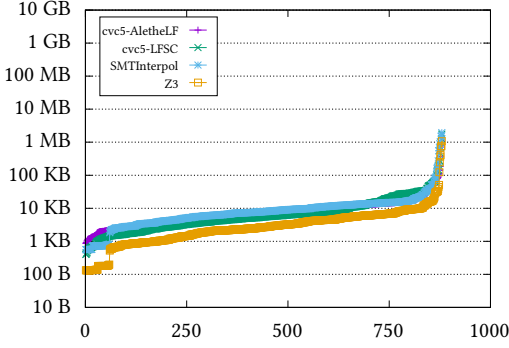


(f) LIA-nonlin proofs for SPACER models.

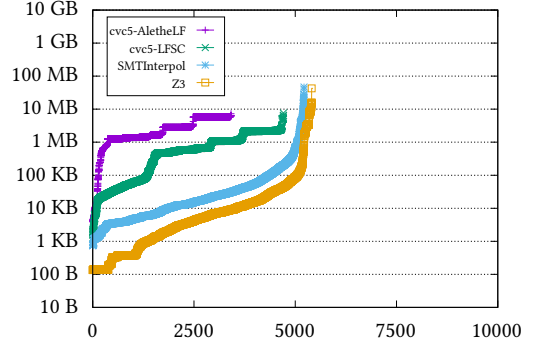
Fig. 4. Sizes of the LIA SMT proofs. The proofs are ordered according to their size, the x-axis indicates their position in the order and the y-axis indicates their size

proofs, is currently unstable¹³, while OPENSMT in proof production mode still suffers from scoping problems¹². One note is that VERiT was not able to produce any proofs, with the reason being that

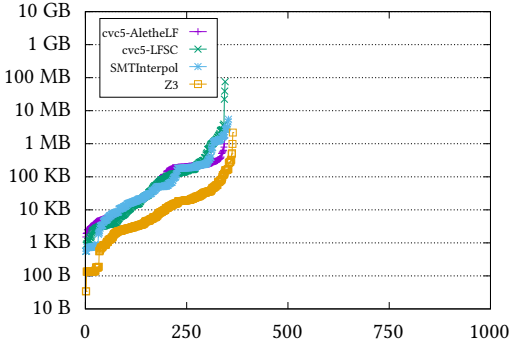
¹³See <https://github.com/cvc5/cvc5/issues/9770> and [./cvc5/issues/10836](https://github.com/cvc5/cvc5/issues/10836).



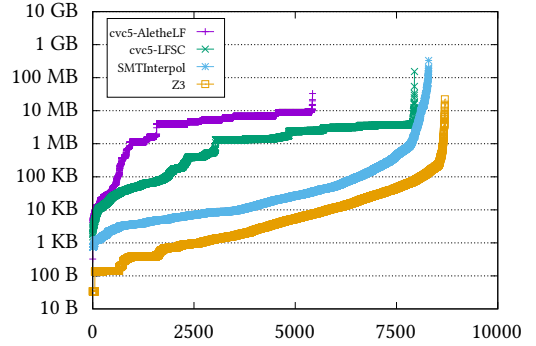
(a) LIA-Arrays-lin proofs for ELDARICA models.



(b) LIA-Arrays-nonlin proofs for ELDARICA models.



(c) LIA-Arrays-lin proofs for SPACER models.



(d) LIA-Arrays-nonlin proofs for SPACER models.

Fig. 5. Sizes of the LIA-Arrays SMT proofs. The proofs are ordered according to their size, the x-axis indicates their position in the order and the y-axis indicates their size.

the `define_fun` construct of the SMT-LIB standard, present in the models produced by all CHC solvers, is supported by `veriT` in its default configuration, but not in its proof production mode.

The proof formats used by each SMT solver can be quite different, not only in shape, but also in the amount of information stored, with the choice of finer or coarser proofs potentially having a significant effect on proof size. The sizes of all proofs produced in our evaluation can be seen in Figure 4 and Figure 5. The ranking of proof sizes between the SMT solvers depends on which CHC solver’s models the instances are generated from. A large number of Z3 proofs, all with a size of 50 B, consisted simply of `(proof (asserted false))`, showcasing how coarse proofs can be; although very compact, these extreme examples make checking essentially degenerate into solving the instance again. The single largest proof produced, by `cvc5-LFSC` from an instance generated from a SPACER LIA-nonlin model, had a size of 2.3 gigabytes, which is a good illustration of the need of compact proofs.

To check the proofs we used the available automated checkers suitable for each proof format, namely `ALFC`, `CARCARA`, and `TSWC` for the proofs produced by `cvc5-ALETHELF`, `cvc5-ALETHE`, and `OPENSMT`, and the `LFSC` and `SMTINTERPOL` checkers for the proofs produced by `cvc5-LFSC` and `SMTINTERPOL`. The results for the proofs produced for the instance sets of the LIA and LIA-Arrays tracks can be seen in Table 6 and Table 7, respectively. Overall, the five checkers were able to

Table 6. Results for checking the proofs produced by solving the SMT instances generated for the LIA benchmarks. The percentage of proofs validated in relation to their total amount is shown in parentheses.

		Valid	Invalid	Timeout	Memout	Error
LIA-lin ELDARICA	ALFC	1427 (100%)	0	0	0	0
	CARCARA	1360 (100%)	0	0	0	0
	LFSC	1427 (100%)	0	0	0	0
	SMTINTERPOL	1368 (100%)	0	0	0	0
	TSWC	1279 (100%)	0	0	0	0
LIA-lin GOLEM	ALFC	995 (99.5%)	0	1	4	0
	CARCARA	865 (100%)	0	0	0	0
	LFSC	996 (99.4%)	0	2	4	0
	SMTINTERPOL	935 (100%)	0	0	0	0
	TSWC	1004 (100%)	0	0	0	0
LIA-lin SPACER	ALFC	1269 (100%)	0	0	0	0
	CARCARA	1029 (100%)	0	0	0	0
	LFSC	1269 (100%)	0	0	0	0
	SMTINTERPOL	1199 (100%)	0	0	0	0
	TSWC	647 (100%)	0	0	0	0
LIA-nonlin ELDARICA	ALFC	13527 (99.9%)	0	0	1	0
	CARCARA	12941 (100%)	0	0	0	0
	LFSC	13517 (99.9%)	0	0	11	0
	SMTINTERPOL	13244 (100%)	0	0	0	0
	TSWC	12208 (100%)	0	0	0	0
LIA-nonlin GOLEM	ALFC	13106 (100%)	0	0	0	0
	CARCARA	11950 (99.2%)	0	0	89	0
	LFSC	13097 (99.9%)	0	0	12	0
	SMTINTERPOL	12827 (100%)	0	0	0	0
	TSWC	13108 (99.9%)	0	1	0	0
LIA-nonlin SPACER	ALFC	14658 (99.8%)	0	0	0	17
	CARCARA	9472 (99.9%)	0	0	1	0
	LFSC	14666 (99.9%)	0	1	9	0
	SMTINTERPOL	13731 (99.9%)	2	3	0	0
	TSWC	1336 (100%)	0	0	0	0

validate most of the proofs produced, with the LFSC checker being the only tool to be significantly affected by the resource constraints, specifically the memory limit of 10 gigabytes. An important discovery is that SMTINTERPOL produced 25 invalid proofs¹⁴. While not implying that the UNSAT results the proofs are supposed to validate are incorrect, since the problem can be in the proof production itself, this is a serious issue. This is combined with the 562 invalid proofs produced by cvc5-ALETHE, due to incorrect proof steps¹⁵, uncovered in our iFM’23 experiments [49]. Still in regards to our previous experiments, they also led to 44282 errors with CARCARA, when checking

¹⁴See <https://github.com/ultimate-pa/smtinterpol/issues/148>.

¹⁵See <https://github.com/cvc5/cvc5/issues/9760>.

Table 7. Results for checking the proofs produced by solving the SMT instances generated for the LIA-Arrays benchmarks. The percentage of proofs validated in relation to their total amount is shown in parentheses.

		Valid	Invalid	Timeout	Memout	Error
LIA-Arrays-lin ELDARICA	ALFC	880 (100%)	0	0	0	0
	CARCARA	-	-	-	-	-
	LFSC checker	880 (100%)	0	0	0	0
	SMTINTERPOL	880 (100%)	0	0	0	0
	TSWC	-	-	-	-	-
LIA-Arrays-lin SPACER	ALFC	343 (99.7%)	0	1	0	0
	CARCARA	-	-	-	-	-
	LFSC checker	345 (100%)	0	0	0	0
	SMTINTERPOL checker	353 (100%)	0	0	0	0
	TSWC	-	-	-	-	-
LIA-Arrays-nonlin ELDARICA	ALFC	3414 (100%)	0	0	0	0
	CARCARA	-	-	-	-	-
	LFSC checker	4695 (99.7%)	0	0	14	0
	SMTINTERPOL checker	5193 (99.5%)	23	0	0	0
	TSWC	-	-	-	-	-
LIA-Arrays-nonlin SPACER	ALFC	5416 (99.7%)	0	1	0	10
	CARCARA	-	-	-	-	-
	LFSC checker	6081 (76.5%)	0	0	1862	0
	SMTINTERPOL checker	8292 (99.9%)	0	2	0	0
	TSWC	-	-	-	-	-

proofs produced for SMT instances generated from SPACER models due to the presence of attribute annotations in models containing quantifiers¹⁶, and 3 errors with the LFSC checker, due to a type mismatch¹⁷. Lastly, 27 errors were also found with the newly integrated checker ALFC¹⁸.

With the production of a CHC model, the generation of the associated SMT instances, the solving of the instances and the production of SMT proofs, and finally the checking of the proofs, it is possible to say with a high degree of confidence that a CHC benchmark is satisfiable. Applying this whole process to our benchmarks, we obtain the results shown in Table 8.

7 Conclusions

We presented a novel two-layered approach for CHC model validation that relies on SMT proofs to provide additional correctness guarantees. The approach is supported by a modular evaluation framework, ATHENA, that allows models to be validated by many different SMT solvers and the SMT solving results to be validated by available proof checkers. A large-scale evaluation was conducted using all LIA and ALIA benchmarks from CHC-COMP 2024 to compare three CHC solvers, five SMT solvers, and five proof checkers. The results indicate that the approach is feasible in practice, with potential to benefit CHC-based verification tools; the results also highlight proof sizes as a crucial practicality factor. A final important point is that many bugs were found in the tools

¹⁶See <https://github.com/ufmg-smite/carcara/issues/12>.

¹⁷See <https://github.com/cvc5/LFSC/issues/87>.

¹⁸See <https://github.com/cvc5/alfc/issues/49>.

Table 8. Number of benchmarks shown to be satisfiable by our two-layered approach. Each column shows the amount of benchmarks shown to be satisfiable by a combination of CHC solver, SMT solver, and proof checker, with the total amount of benchmarks to be validated by each CHC solver displayed in parentheses.

	CHC benchmarks validated					
	LIA-lin			LIA-nonlin		
	ELDARICA (130)	GOLEM (115)	SPACER (111)	ELDARICA (152)	GOLEM (145)	SPACER (175)
CVC5-ALETHE	84	63	48	99	64	20
CVC5-ALETHELF	128	107	109	150	142	161
CVC5-LFSC	128	109	110	144	135	163
OPENSMT	105	115	69	112	144	48
SMTINTERPOL	113	99	93	129	119	117
VERiT	0	0	0	0	0	0
Z3	-	-	-	-	-	-
	LIA-Arrays-lin			LIA-Arrays-nonlin		
	ELDARICA	GOLEM	SPACER	ELDARICA	GOLEM	SPACER
	(56)	-	(29)	(79)	-	(87)
CVC5-ALETHE	-	-	-	-	-	-
CVC5-ALETHELF	56	-	24	12	-	23
CVC5-LFSC	56	-	24	16	-	29
OPENSMT	-	-	-	-	-	-
SMTINTERPOL	56	-	25	13	-	12
VERiT	-	-	-	-	-	-
Z3	-	-	-	-	-	-

compared, including invalid models being produced by two state-of-the-art CHC solvers, which confirms the need to provide modern verification tooling with additional correctness guarantees.

Directions for future work are threefold, namely (i) evaluating the approach with other FOL theories, (ii) embedding the approach into CHC-based verification tooling, and (iii) designing a complementary approach to validate CHC proofs. For the first direction, enhancements can be made to the framework’s implementation to cater to theories other than those already supported, with algebraic data types being a suitable candidate for this. For the second direction, the use of proof-backed model validation in CHC-based model checkers is a direct application. For the third direction, one possibility is to use the ALETHE format to represent and check CHC proofs, since it is rich enough to describe the necessary proof steps. An important unknown regarding potential ALETHE CHC proofs is the correct level of granularity, as it is unclear if coarse proofs can be efficiently checked, either by CARCARA or any future checker, or if additional burden needs to be put on the solvers to produce fine-grained proofs.

Acknowledgments

Rodrigo Otoni’s work was supported by the Swiss National Science Foundation grant 200021_197353. Martin Blicha’s work was supported by the Swiss National Science Foundation grant 200021_185031 and the Czech Science Foundation grant 23-06506S. The authors thank Fedor Gromov for his help in preparing the instance generator.

References

- [1] Leonardo Alt, Martin Blicha, Antti E. J. Hyvärinen, and Natasha Sharygina. 2022. SolCMC: Solidity Compiler’s Model Checker. In *Proceedings of the 34th International Conference on Computer Aided Verification*. 325–338.
- [2] Bruno Andreotti, Hanna Lachnitt, and Haniel Barbosa. 2023. Carcara: An Efficient Proof Checker and Elaborator for SMT Proofs in the Alethe Format. In *Proceedings of the 29th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 367–386.
- [3] Michael Armand, Germain Faure, Benjamin Grégoire, Chantal Keller, Laurent Théry, and Benjamin Werner. 2011. A Modular Integration of SAT/SMT Solvers to Coq through Proof Witnesses. In *Proceedings of the 1st International Conference on Certified Programs and Proofs*. 135–150.
- [4] Seulkee Baek, Mario Carneiro, and Marijn J. H. Heule. 2021. A Flexible Proof Format for SAT Solver-Elaborator Communication. In *Proceedings of the 27th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 59–75.
- [5] Haniel Barbosa, Clark Barrett, Martin Brain, Gereon Kremer, Hanna Lachnitt, Makai Mann, Abdalrhman Mohamed, Mudathir Mohamed, Aina Niemetz, Andres Nötzli, Alex Ozdemir, Mathias Preiner, Andrew Reynolds, Ying Sheng, Cesare Tinelli, and Yoni Zohar. 2022. CVC5: A Versatile and Industrial-Strength SMT Solver. In *Proceedings of the 28th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 415–442.
- [6] Haniel Barbosa, Jasmin Christian Blanchette, Mathias Fleury, and Pascal Fontaine. 2020. Scalable Fine-Grained Proofs for Formula Processing. *Journal of Automated Reasoning* 64, 3 (2020), 485–510.
- [7] Clark Barrett, Pascal Fontaine, and Cesare Tinelli. 2021. The SMT-LIB Standard: Version 2.6. <https://smtlib.cs.uiowa.edu/papers/smt-lib-reference-v2.6-r2021-05-12.pdf>.
- [8] Clark Barrett, Roberto Sebastiani, Sanjit Seshia, and Cesare Tinelli. 2021. Satisfiability Modulo Theories. In *Handbook of Satisfiability*, Armin Biere, Marijn J. H. Heule, Hans van Maaren, and Toby Walsh (Eds.). IOS Press.
- [9] Dirk Beyer. 2023. Competition on Software Verification and Witness Validation: SV-COMP 2023. In *Proceedings of the 29th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 495–522.
- [10] Dirk Beyer and Jan Strejček. 2022. Case Study on Verification-Witness Validators: Where We Are and Where We Go. In *Proceedings of the 29th International Symposium on Static Analysis*. 160–174.
- [11] Nikolaj Bjørner, Arie Gurfinkel, Ken McMillan, and Andrey Rybalchenko. 2015. Horn Clause Solvers for Program Verification. *Fields of Logic and Computation II. Lecture Notes in Computer Science* 9300 (2015), 24–51.
- [12] Jasmin Christian Blanchette, Sascha Böhme, Mathias Fleury, Steffen Juilf Smolka, and Albert Steckermeier. 2016. Semi-intelligible Isar Proofs from Machine-Generated Proofs. *Journal of Automated Reasoning* 56, 2 (2016), 155–200.
- [13] Martin Blicha, Konstantin Britikov, and Natasha Sharygina. 2023. The Golem Horn Solver. In *Proceedings of the 35th International Conference on Computer Aided Verification*. 209–223.
- [14] Sascha Böhme and Tjark Weber. 2010. Fast LCF-Style Proof Reconstruction for Z3. In *Proceedings of the 1st International Conference on Interactive Theorem Proving*. 179–194.
- [15] Thomas Bouton, Diego Caminha B. de Oliveira, David Déharbe, and Pascal Fontaine. 2009. veriT: An Open, Trustable and Efficient SMT-Solver. In *Proceedings of the 22nd International Conference on Automated Deduction*. 151–156.
- [16] Martin Bromberger, Jochen Hoenicke, and François Bobot. 2023. SMT-COMP 2023: Competition Report. <https://smt-workshop.cs.uiowa.edu/2023/slides/smtcomp.pdf>.
- [17] Roberto Bruttomesso, Edgar Pek, Natasha Sharygina, and Aliaksei Tsitovich. 2010. The OpenSMT Solver. In *Proceedings of the 16th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 150–153.
- [18] Guillaume Bury. 2021. Dolmen: A Validator for SMT-LIB and Much More. In *Proceedings of the 19th International Workshop on Satisfiability Modulo Theories*. 32–39.
- [19] Stefano Calzavara, Ilya Grishchenko, and Matteo Maffei. 2016. HornDroid: Practical and Sound Static Analysis of Android Applications by SMT Solving. In *Proceedings of the 1st IEEE European Symposium on Security and Privacy*. 47–62.
- [20] Jürgen Christ, Jochen Hoenicke, and Alexander Nutz. 2012. SMTInterpol: An Interpolating SMT Solver. In *Proceedings of the 19th International SPIN Workshop*. 248–254.
- [21] Luís Cruz-Filipe, Marijn J. H. Heule, Warren A. Hunt, Matt Kaufmann, and Peter Schneider-Kamp. 2017. Efficient Certified RAT Verification. In *Proceedings of the 26th International Conference on Automated Deduction*. 220–236.
- [22] Leonardo de Moura and Nikolaj Bjørner. 2008. Proofs and Refutations, and Z3. In *Proceedings of the 7th International Workshop on the Implementation of Logics*. 123–132.
- [23] Leonardo de Moura and Nikolaj Bjørner. 2008. Z3: An Efficient SMT Solver. In *Proceedings of the 14th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 337–340.
- [24] Daniel Dietsch, Matthias Heizmann, Jochen Hoenicke, Alexander Nutz, and Andreas Podelski. 2019. Ultimate TreeAutomizer (CHC-COMP Tool Description). In *Proceedings of the 6th Workshop on Horn Clauses for Verification and Synthesis*. 42–47.

- [25] Burak Ekici, Alain Mebsout, Cesare Tinelli, Chantal Keller, Guy Katz, Andrew Reynolds, and Clark Barrett. 2017. SMTCoq: A Plug-In for Integrating SMT Solvers into Coq. In *Proceedings of the 29th International Conference on Computer Aided Verification*. 126–133.
- [26] Gidon Ernst. 2023. Korn - Software Verification with Horn Clauses (Competition Contribution). In *Proceedings of the 29th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 559–564.
- [27] Gidon Ernst and José F. Morales. 2024. CHC-COMP 2024: Competition Report. <https://chc-comp.github.io/CHC-COMP2024Report-HCSV.pdf>.
- [28] Pascal Fontaine, Jean-Yves Marion, Stephan Merz, Leonor Prensa Nieto, and Alwen Tiu. 2006. Expressiveness + Automation + Soundness: Towards Combining SMT Solvers and Interactive Proof Assistants. In *Proceedings of the 12th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 167–181.
- [29] Florian Frohn and Jürgen Giesl. 2022. Proving Non-Termination and Lower Runtime Bounds with LoAT (System Description). In *Proceedings of the 11th International Joint Conference on Automated Reasoning*. 712–722.
- [30] Marco Gario and Andrea Micheli. 2015. PySMT: a Solver-agnostic Library for Fast Prototyping of SMT-based Algorithms. In *Proceedings of the 13th International Workshop on Satisfiability Modulo Theories*. 1–10.
- [31] Sergey Grebenshchikov, Nuno P. Lopes, Corneliu Popeea, and Andrey Rybalchenko. 2012. Synthesizing Software Verifiers from Proof Rules. In *Proceedings of the 33rd ACM SIGPLAN Conference on Programming Language Design and Implementation*. 405–416.
- [32] Arie Gurfinkel and Nikolaj Bjørner. 2019. The Science, Art, and Magic of Constrained Horn Clauses. In *Proceedings of the 21st International Symposium on Symbolic and Numeric Algorithms for Scientific Computing*. 6–10.
- [33] Arie Gurfinkel, Temesghen Kahsai, Anvesh Komuravelli, and Jorge A. Navas. 2015. The SeaHorn Verification Framework. In *Proceedings of the 27th International Conference on Computer Aided Verification*. 343–361.
- [34] Matthias Heizmann, Daniel Dietsch, Jochen Hoenicke, Alexander Nutz, Andreas Podelski, and Frank Schüssele. 2024. Ultimate Unihorn - Solver Description. <https://doi.org/10.4204/EPTCS.402.10>. Page 99.
- [35] Marijn J.H. Heule, Warren A. Hunt, and Nathan Wetzler. 2013. Trimming While Checking Clausal Proofs. In *Proceedings of the 13th Conference on Formal Methods in Computer-Aided Design*. 181–188.
- [36] Marijn J. H. Heule, Warren A. Hunt, Matt Kaufmann, and Nathan Wetzler. 2017. Efficient, Verified Checking of Propositional Proofs. In *Proceedings of the 8th International Conference on Interactive Theorem Proving*. 269–284.
- [37] Marijn J. H. Heule, Warren A. Hunt, and Nathan Wetzler. 2013. Verifying Refutations with Extended Resolution. In *Proceedings of the 24th International Conference on Automated Deduction*. 345–359.
- [38] C. A. R. Hoare. 1969. An Axiomatic Basis for Computer Programming. *Commun. ACM* 12, 10 (1969), 576–580.
- [39] Jochen Hoenicke and Tanja Schindler. 2022. A Simple Proof Format for SMT. In *Proceedings of the 20th International Workshop on Satisfiability Modulo Theories*. 54–70.
- [40] Hossein Hojjat, Philipp Rümmer, Pavle Subotic, and Wang Yi. 2014. Horn Clauses for Communicating Timed Systems. In *Proceedings of the 1st Workshop on Horn Clauses for Verification and Synthesis*. 39–52.
- [41] Hossein Hojjat and Philipp Rümmer. 2018. The Eldarica Horn Solver. In *Proceedings of the 18th Conference on Formal Methods in Computer-Aided Design*. 1–7.
- [42] Temesghen Kahsai, Philipp Rümmer, Huascar Sanchez, and Martin Schäfer. 2016. JayHorn: A Framework for Verifying Java programs. In *Proceedings of the 28th International Conference on Computer Aided Verification*. 352–358.
- [43] Anvesh Komuravelli, Arie Gurfinkel, and Sagar Chaki. 2016. SMT-based Model Checking for Recursive Programs. *Formal Methods in System Design* 48, 3 (2016), 175–205.
- [44] Daniel Kroening and Ofer Strichman. 2016. *Decision Procedures - An Algorithmic Point of View* (2 ed.). Springer Berlin, Heidelberg.
- [45] Ramana Kumar, Magnus O. Myreen, Michael Norrish, and Scott Owens. 2014. CakeML: A Verified Implementation of ML. In *Proceedings of the 41st ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages*. 179–191.
- [46] Peter Lammich. 2020. Efficient Verified (UN)SAT Certificate Checking. *Journal of Automated Reasoning* 64, 3 (2020), 513–532.
- [47] Yusuke Matsushita, Takeshi Tsukada, and Naoki Kobayashi. 2021. RustHorn: CHC-Based Verification for Rust Programs. *ACM Transactions on Programming Languages and Systems* 43, 4 (2021), 1–54.
- [48] Rodrigo Otoni, Martin Blicha, Patrick Eugster, Antti E. J. Hyvärinen, and Natasha Sharygina. 2021. Theory-Specific Proof Steps Witnessing Correctness of SMT Executions. In *Proceedings of the 58th ACM/IEEE Design Automation Conference*. 541–546.
- [49] Rodrigo Otoni, Martin Blicha, Patrick Eugster, and Natasha Sharygina. 2023. CHC Model Validation with Proof Guarantees. In *Proceedings of the 18th International Conference on Integrated Formal Methods*. 62–81.
- [50] Rodrigo Otoni, Matteo Marescotti, Leonardo Alt, Patrick Eugster, Antti Hyvärinen, and Natasha Sharygina. 2023. A Solicitous Approach to Smart Contract Verification. *ACM Transactions on Privacy and Security* 26, 2 (2023), 1–28.
- [51] Joseph E. Reeves, Benjamin Kiesl-Reiter, and Marijn J. H. Heule. 2023. Propositional Proof Skeletons. In *Proceedings of the 29th International Conference on Tools and Algorithms for the Construction and Analysis of Systems*. 329–347.

- [52] Andrew Reynolds and Hans-Jörg Schurr. 2024. AletheLF Checker (alfc). https://cvc5.github.io/docs/cvc5-1.1.2/proofs/output_alf.html.
- [53] Philipp Rümmer, Hossein Hojjat, and Viktor Kuncak. 2015. On Recursion-Free Horn Clauses and Craig Interpolation. *Formal Methods In System Design* 47, 1 (2015), 1–25.
- [54] Adam Sandberg Ericsson, Magnus O. Myreen, and Johannes Åman Pohjola. 2017. A Verified Generational Garbage Collector for CakeML. In *Proceedings of the 8th International Conference on Interactive Theorem Proving*. 444–461.
- [55] Hans-Jörg Schurr, Mathias Fleury, Haniel Barbosa, and Pascal Fontaine. 2021. Alethe: Towards a Generic SMT Proof Format. In *Proceedings of the 7th Workshop on ProofExchange for Theorem Proving*. 49–54.
- [56] Carsten Sinz and Armin Biere. 2006. Extended Resolution Proofs for Conjoining BDDs. In *Proceedings of the 1st International Symposium on Computer Science in Russia*. 600–611.
- [57] Aaron Stump, Duckki Oe, Andrew Reynolds, Liana Hadarean, and Cesare Tinelli. 2013. SMT Proof Checking Using a Logical Framework. *Formal Methods in System Design* 42, 1 (2013), 91–118.
- [58] Ole Tange. 2011. GNU Parallel - The Command-Line Power Tool. ;login: *The USENIX Magazine* 36, 1 (2011), 42–47.
- [59] Tamás Tóth, Ákos Hajdu, András Vörös, Zoltán Micskei, and István Majzik. 2017. Theta: a Framework for Abstraction Refinement-Based Model Checking. In *Proceedings of the 17th Conference on Formal Methods in Computer-Aided Design*. 176–179.
- [60] Nathan Wetzler, Marijn J. H. Heule, and Warren A. Hunt. 2014. DRAT-trim: Efficient Checking and Trimming Using Expressive Clausal Proofs. In *Proceedings of the 17th International Conference on Theory and Applications of Satisfiability Testing*. 422–429.

Received 1 January 20XX; revised 1 January 20XX; accepted 1 January 20XX