

Detecting speculative leaks with compositional semantics

XAVER FABIAN, CISPA Helmholtz Center for Information Security, Germany and University of Trento, Italy

MARCO GUARNIERI, IMDEA Software Institute, Spain

BORIS KÖPF, Azure Research, Microsoft, UK

JOSE F. MORALES, IMDEA Software Institute, Spain

MARCO PATRIGNANI, University of Trento, Italy

JAN REINEKE, Saarland University, Germany

ANDRES SANCHEZ, Amazon, Spain

Speculative execution enhances processor performance by predicting intermediate results and executing instructions based on these predictions. However, incorrect predictions can lead to security vulnerabilities, as speculative instructions leave traces in microarchitectural components that attackers can exploit. This is demonstrated by the family of Spectre attacks. Unfortunately, existing countermeasures to these attacks lack a formal security characterization, making it difficult to verify their effectiveness.

In this paper, we propose a novel framework for detecting information flows introduced by speculative execution and reasoning about software defenses. The theoretical foundation of our approach is *speculative non-interference* (SNI), a novel semantic notion of security against speculative execution attacks. SNI relates information leakage observed under a standard non-speculative semantics to leakage arising under semantics that explicitly model speculative execution. To capture their combined effects, we extend our framework with a mechanism to safely compose multiple speculative semantics, each focussing on a single aspect of speculation. This allows us to analyze the complex interactions and resulting leaks that can arise when multiple speculative mechanisms operate together. On the practical side, we develop Spectector, a symbolic analysis tool that uses our compositional framework and leverages SMT solvers to detect vulnerabilities and verify program security with respect to multiple speculation mechanisms. We demonstrate the effectiveness of Spectector through evaluations on standard security benchmarks and new vulnerability scenarios.

CCS Concepts: • Security and privacy → Formal security models; Systems security.

Additional Key Words and Phrases: Spectre; Speculative Execution; Speculative information flows; Speculative non-interference

ACM Reference Format:

Xaver Fabian, Marco Guarnieri, Boris Köpf, Jose F. Morales, Marco Patrignani, Jan Reineke, and Andres Sanchez. 2026. Detecting speculative leaks with compositional semantics. 1, 1 (July 2026), 60 pages. <https://doi.org/10.1145/nnnnnnn.nnnnnnn>

1 INTRODUCTION

Speculative execution avoids pipeline stalls by predicting the results of intermediate computations and by speculatively executing instructions based on such predictions. Whenever a prediction turns out to be incorrect, the processor squashes

Authors' addresses: Xaver Fabian, CISPA Helmholtz Center for Information Security, Saarbrücken, Germany, xaver.fabian@cispa.de and University of Trento, Trento, Italy; Marco Guarnieri, IMDEA Software Institute, Madrid, Spain, marco.guarnieri@imdea.org; Boris Köpf, Azure Research, Microsoft, UK; Jose F. Morales, IMDEA Software Institute, Madrid, Spain; Marco Patrignani, University of Trento, Trento, Italy, marco.patrignani@unitn.it; Jan Reineke, Saarland University, Saarbrücken, Germany; Andres Sanchez, Amazon, Madrid, Spain.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

© 2026 Association for Computing Machinery.

Manuscript submitted to ACM

Manuscript submitted to ACM

Listing 1. Code vulnerable to branch speculation.

```

1  if (y < size)
2  temp &= B[A[y] * 512];

```

Listing 2. Code vulnerable to store speculation.

```

1  p = &secret;
2  p = &public;
3  temp = B[*p * 512];

```

Listing 3. Speculative leak arising from speculation over branch and store instructions combined.

```

1  x = 0;
2  p = &secret;
3  p = &public;
4  if (x != 0)
5      temp &= A[*p];

```

Fig. 1. Code snippets vulnerable to different kinds of speculation attacks.

the speculatively executed instructions, thereby rolling back their effects on the architectural state, which consists of registers, flags, and memory. However, the execution of speculative instructions leaves footprints in a CPU’s microarchitectural components (like caches, predictors, and internal buffers) that may persist even after these instructions have been squashed. As shown by Spectre [49] and follow-up attacks [10, 15, 50, 53, 74, 78, 79, 81], these microarchitectural side effects can be exploited to compromise the security of programs and to leak information about speculatively accessed data.

Speculative Leaks. Modern CPUs employ a wide range of speculation mechanisms (branch predictors, memory disambiguators, etc.) that are used to speculate over different kinds of instructions and intermediate results, such as conditional branches [49], indirect jumps [49], store and load operations [43], return instructions [8, 50], and load addresses and loaded results [46, 47, 69]. All these speculation mechanisms can be exploited by attackers to leak data.

As an example, the code in Listing 1 is vulnerable to a Spectre-PHT attack [49], which exploits speculation over branch instructions. Whenever the branch prediction mispredicts the outcome of the condition $y > \text{size}$ on Line 1, this results in speculatively executing Line 2, which leaks the out-of-bounds value pointed by $A[y]$ through the data cache. As another example, Listing 2 depicts a code snippet vulnerable to a Spectre-STL attack [43], which exploits speculation over memory disambiguation checks. Whenever the memory write on Line 2 is predicted to have a different address than the memory read $*p$ on Line 3, the content of `secret` is leaked to an attacker. Essentially, the memory write in Line 2 is bypassed because of speculation.

The majority of well-known attacks [7, 43, 49, 50, 53] only exploit leaks introduced by individual speculation mechanisms, e.g., branch predictors in Spectre-PHT (Listing 1) and memory disambiguators in Spectre-STL (Listing 2). However, some speculative leaks only arise due to the interaction of multiple mechanisms. For example, the code in Listing 3 can speculatively leak the value of `&secret` in Line 5 whenever (1) the memory write to `p` in Line 3 is predicted to have a different address than the memory read $*p$ on Line 5, and (2) the branch instruction on Line 4 is mispredicted as taken. This leak, therefore, arises from the *combination* of two speculation mechanisms—branch prediction and memory disambiguation prediction—and *it cannot be detected* when considering the two mechanisms in isolation.

Limitations of existing leak-detection approaches. Since the discovery of Spectre [49], a number of approaches have been proposed for reasoning about the security of programs against speculative leaks [20, 25, 36, 39, 42, 44, 60, 61, 63, 76, 77, 80].

These approaches employ a wide array of techniques for reasoning about leaks (e.g., symbolic execution [20, 25, 36, 39, 76, 77], abstract interpretation [80], testing [42, 44, 60, 61, 63]), but they all share a critical limitation. They support only *fixed* speculation mechanisms, such as branch prediction [36, 38, 75–77] and memory disambiguation prediction [20, 25, 65], which are *hard-coded*, e.g., in the underlying formal models. As a result, extending them to reason about a new

speculation mechanism is often complicated (e.g., it might require modifying the underlying formal model and updating the corresponding security proofs), since none of these approaches has been designed in a *compositional* manner.

As shown by code snippets like Listing 3, sound reasoning about speculative leaks in a program requires accounting for *all* speculation mechanisms, since ignoring some mechanisms might lead to missed leaks. Unfortunately, (1) we lack a precise understanding of the speculation mechanisms implemented in current CPUs, as demonstrated time and again by attacks discovering previously unknown speculation mechanisms, and (2) new CPU generations often implement refined and improved speculation mechanisms to improve performance [46, 47]. Since we cannot develop an analysis that would account for future microarchitectures, approaches for reasoning about speculative leaks must be *compositional* and easy to *extend* whenever a new speculation mechanism is discovered. However, we currently lack a precise characterization of security against speculative leaks and an associated program analysis framework that is *extensible and compositional* in the underlying speculation mechanisms.

Our approach: In this paper, we develop a novel, principled approach for reasoning about leaks introduced by speculatively executed instructions and for reasoning about software defenses against Spectre-style attacks. Our approach is backed by a semantic characterization of security against speculative leaks and it comes with an algorithm, based on symbolic execution, for proving the absence of leaks. Crucially, our approach is *compositional*. It allows specifying individual operational semantics—each one capturing the effects of a different speculation mechanism—and combining them in a single composed semantics, which captures leaks arising from the interactions of all these mechanisms, thereby leading to simpler formalizations. Additionally, these semantics can be directly integrated in our verification algorithm, whose soundness proof is compositional and defined in terms of the component semantics, thereby maximizing proof reuse. Next, we describe our contributions in more detail:

(1) **Modeling speculation:** We propose a generalized template for modeling the effects of speculative execution as an operational semantics. Our template extends an architectural (non-speculative) semantics (Section 3) to a *speculative semantics* (Section 4) that captures the effects of speculatively executed instructions. In a nutshell, the speculative semantics follow mispredicted paths for a bounded number of steps before backtracking and restarting the architectural execution, while relying on a *prediction oracle* to model the speculation mechanism. Note that this template is flexible enough to capture both control-flow and data-flow speculation. Following this template, we introduce five speculative semantics (denoted $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$, and $\mathcal{L}_{SLs}$) which capture speculation over branches, jumps, stores, and return instructions (Section 6). This is the most comprehensive collection of individual speculative semantics in a single framework to date.

(2) **Speculative non-interference:** We propose *speculative non-interference* (Section 5), a novel semantic notion of security against speculative execution attacks. Speculative non-interference is based on comparing a program with respect to two different semantics: (a) a standard, *non-speculative semantics*, which we use as a proxy for the intended program behavior, and (b) a *speculative semantics* that captures the effects of speculation induced by one or more speculation mechanisms, which captures the effect of speculatively executed instructions. In a nutshell, speculative non-interference requires that *speculatively executed instructions do not leak more information into the microarchitectural state than what the intended behavior does*, i.e., than what is leaked by the standard, non-speculative semantics.

To capture “leakage into the microarchitectural state”, we consider an observer of the program execution that sees the locations of memory accesses and jump targets. This observer model is commonly used for characterizing “side-channel free” or “constant-time” code [4, 57] in the absence of detailed models of the microarchitecture. Under this observer model, an adversary may distinguish two initial program states if they yield different traces of memory locations and

jump targets. *Speculative non-interference* (SNI) requires that two initial program states can be distinguished under the speculative semantics only if they can also be distinguished under the standard, non-speculative semantics.

The speculative semantics, and hence SNI, depends on the decisions taken by the prediction oracle. We show that one can abstract away from the specific oracle by considering a worst-case oracle that *always mispredicts*. SNI w.r.t. this always-mispredict oracle implies SNI w.r.t. a large class of prediction strategies. This allows reasoning about speculative leaks while ignoring the details of the specific prediction strategy, which are often undocumented.

(3) **Composition framework:** We propose a framework for composing speculative semantics that capture speculation due to different mechanisms (Section 7). The framework allows specifying individual semantics for each speculation mechanism and combining them into a composed semantics. The combination yields a single operational semantics which can be used to reason about leaks involving the different kinds of speculation it comprises (as in Listing 3). We also formalize the key properties of our framework: if the individual semantics fulfill some (expected) safety conditions (which we prove for all the semantics we combine), then the composed semantics is well-formed and can be used to reason about leakage in a sound manner. That is, whenever a program is leaky under one of the individual semantics then it is leaky under the composed semantics and, vice versa, whenever a program is SNI under the composed semantics, then it is SNI w.r.t. all individual semantics. We apply the composition framework to all our individual semantics, thereby obtaining 18 combined speculative semantics (Section 8).

(4) **Checking speculative non-interference:** We propose SPECTECTOR (Section 9), an algorithm to automatically prove that programs satisfy SNI w.r.t. any speculative semantics (or composition) in our framework. Given a program p , SPECTECTOR uses symbolic execution with respect to the speculative semantics to derive a concise representation of the traces of memory accesses and jump targets during execution along all possible program paths. Based on this representation, SPECTECTOR creates a symbolic formula capturing that, whenever two initial program states produce the same (instruction and data) memory access patterns in the standard semantics, they also produce the same access patterns in the speculative semantics. Validity of this formula for each program path implies SNI.

(5) **Implementation and evaluation:** We implement a prototype of SPECTECTOR (Section 10), with a front-end parsing (a subset of) x86 assembly and a back-end that uses the Z3 SMT solver to perform symbolic execution against all our speculative semantics (and their combinations) and to determine whether programs contain speculative leaks. We validate SPECTECTOR on both existing microbenchmarks (for speculation on branches, indirect jumps, store and return instructions) and on new ones (for speculative leaks arising from combinations of speculation mechanisms) that we introduce. These microbenchmarks contain both leaky programs as well as programs patched with well-known compiler-level countermeasures against Spectre. Using SPECTECTOR, we successfully (1) detect all leaks pointed out in [48], (2) detect novel, subtle leaks that are out of the scope of existing approaches that check for known vulnerable code patterns [77], (3) detect leaks arising from speculation over multiple mechanisms that are not detected by existing symbolic approaches [14], and (4) identify cases where compilers unnecessarily inject countermeasures, i.e., opportunities for optimization without sacrificing security.

Scope of this paper. This paper provides a unified, extended, and updated version of the results presented in two conference papers by Guarnieri et al. [38] and Fabian et al. [31]. In addition to the original results from [31, 38], (1) the paper introduces a general template for speculative operational semantics that generalizes existing formalizations, (2) it adds new speculative semantics for straight-line speculation and for speculation over indirect jumps (and their combinations), and (3) it provides additional technical details about the speculative semantics as well as about composition proofs. The paper is accompanied by an updated version of SPECTECTOR (with an updated evaluation) that

Listing 4. SPECTRE-PHT- Assembly code

```

1  mov    size, %rax
2  mov    y, %rbx
3  cmp    %rbx, %rax
4  jbe    END
5  mov    A(%rbx), %rax
6  shl    $9, %rax
7  mov    B(%rax), %rax
8  and    %rax, temp

```

supports all new speculative semantics and combinations presented in the paper. Overall, we see this paper as a unified reference for the modeling of speculative leaks at program level and for the SPECTECTOR program analysis tool.

Additional materials. The SPECTECTOR program analysis tool is open-source and available at [37]. Given the number of different speculative semantics studied in this paper, we leave the full formalization of all semantics, technicalities, and proof details to the associated technical report, which is available at [30].

2 ILLUSTRATIVE EXAMPLE

To illustrate our approach, we show how SPECTECTOR applies to the SPECTRE-PHT example [49] shown in Listing 1 using a speculative semantics capturing branch speculation.

SPECTRE-PHT. The program checks whether the index stored in the variable `y` is less than the size of the array `A`, stored in the variable `size`. If that is the case, the program retrieves `A[y]`, amplifies it with a multiple (here: 512) of the cache line size, and uses the result as an address for accessing the array `B`.

If `size` is not cached, evaluating the branch condition requires traditional processors to wait until `size` is fetched from main memory. Modern processors instead speculate on the condition’s outcome and continue the computation. Hence, the memory accesses in line 2 may be executed even if `y ≥ size`.

When `size` becomes available, the processor checks whether the speculated branch is the correct one. If it is not, it rolls back the architectural (i.e., ISA) state’s changes and executes the correct branch. However, the speculatively executed memory accesses leave a footprint in the microarchitectural state, in particular in the cache, which enables an adversary to retrieve `A[y]`, even for `y ≥ size`, by probing the array `B`.

Detecting leaks with SPECTECTOR. SPECTECTOR automatically detects leaks introduced by speculatively executed instructions, or proves their absence. Specifically, SPECTECTOR detects a leak whenever executing the program under the speculative semantics, which captures that the execution can go down a mispredicted path for a bounded number of steps, leaks more information into the microarchitectural state than executing the program under a non-speculative semantics.

To illustrate how SPECTECTOR operates, we consider the x86 assembly¹ translation of Listing 1’s program (cf. Listing 4). SPECTECTOR performs symbolic execution with respect to the speculative semantics to derive a concise representation of the concrete traces of memory accesses and program counter values along each path of the program. These symbolic traces capture the program’s effect on the microarchitectural state.

¹We use a simplified AT&T syntax without operand sizes

Listing 5. SPECTRE-PHT- Assembly code with speculative load hardening. CLANG inserted instructions 3, 6, 9, and 11.

```

1  mov    size, %rax
2  mov    y, %rbx
3  mov    $0, %rdx
4  cmp    %rbx, %rax
5  jbe    END
6  cmovbe $-1, %rdx
7  mov    A(%rbx), %rax
8  shl    $9, %rax
9  or     %rdx, %rax
10 mov    B(%rax), %rax
11 or     %rdx, %rax
12 and    %rax, temp

```

During speculative execution, the speculatively executed parts are determined by the predictions of the branch predictor. As shown in Section 6.1, leakage due to speculative execution is maximized under a branch predictor that mispredicts every branch. The code in Listing 4 yields two symbolic traces w.r.t. the speculative semantics that mispredicts every branch:²

$$\text{start} \cdot \text{rlb} \cdot \bar{\tau} \quad \text{when } y < \text{size} \quad (1)$$

$$\text{start} \cdot \bar{\tau} \cdot \text{rlb} \quad \text{when } y \geq \text{size} \quad (2)$$

where $\bar{\tau} = \text{load}(A + y) \cdot \text{load}(B + A[y] * 512)$. Here, the argument of `load` is visible to the observer, while `start` and `rlb` denote the start and the end of a misspeculated execution. The traces of the *non-speculative* semantics are obtained from those of the speculative semantics by removing all observations in between `start` and `rlb`.

Trace 1 shows that whenever `y` is in bounds (i.e., $y < \text{size}$) the observations of the speculative semantics and the non-speculative semantics coincide (i.e. they are both $\bar{\tau}$). In contrast, Trace 2 shows that whenever $y \geq \text{size}$, the speculative execution generates observations $\bar{\tau}$ that depend on `A[y]` whose value is not visible in the non-speculative execution. This is flagged as a leak by SPECTECTOR.

Proving Security with SPECTECTOR. The CLANG 7.0.0 C++ compiler implements a countermeasure, called speculative load hardening [19], that applies conditional masks to addresses to prevent leaks into the microarchitectural state. Listing 5 depicts the protected output of CLANG on the program from Listing 1.

The symbolic execution of the speculative semantics produces, as before, Trace 1 and Trace 2, but with

$$\bar{\tau} = \text{load}(A + y) \cdot \text{load}(B + (A[y] * 512) | \text{mask}),$$

where $\text{mask} = \text{ite}(y < \text{size}, 0x0, 0xFF \dots FF)$ corresponds to the conditional move in line 6 and $|$ is a bitwise-or operator. Here, $\text{ite}(y < \text{size}, 0x0, 0xFF \dots FF)$ is a symbolic if-then-else expression evaluating to `0x0` if $y < \text{size}$ and to `0xFF ... FF` otherwise.

The analysis of Trace 1 is as before. For Trace 2, however, SPECTECTOR determines (via a query to Z3 [26]) that, for all $y \geq \text{size}$ there is exactly *one* observation that the adversary can make during the speculative execution, namely $\text{load}(A + y) \cdot \text{load}(B + 0xFF \dots FF)$, from which it concludes that no information leaks into the microarchitectural state, i.e., the countermeasure is effective in securing the program.

²For simplicity of presentation, the example traces capture only loads but not the program counter.

(Programs) $p := n : i \mid p_1; p_2$ (Functions) $\mathcal{F} := \emptyset \mid \mathcal{F}; f \mapsto n$
 (Expressions) $e := n \mid x \mid \ominus e \mid e_1 \otimes e_2$ (Registers) $x \in \text{Regs}$ (Values) $n, l \in \text{Vals} = \mathbb{N} \cup \{\perp\}$
 (Instructions) $i := \text{skip} \mid x \leftarrow e \mid \text{load } x, e \mid \text{store } x, e \mid \text{jmp } e \mid \text{beqz } x, l \mid x \xleftarrow{e'} e \mid \text{spbarr} \mid \text{call } f \mid \text{ret}$

Fig. 2. Syntax of the μASM language

3 μASM LANGUAGE

In this section, we present μASM , a simple assembly-style language serving as the basis for our framework. We start by describing the attacker model we consider (Section 3.1). Then, we present the syntax (Section 3.2) and the base non-speculative semantics (Section 3.3) of μASM .

3.1 Observations and Attacker Model

We adopt a commonly-used attacker model [3, 20, 25, 35, 36, 38, 64, 75]: a passive attacker observing the execution of a program through events τ . These events, which we call *observations*, model timing leaks through cache and control flow while abstracting away low-level microarchitectural details. Specifically, we consider an attacker that observes the program counter and the locations of memory accesses during computation. This attacker model is commonly used to formalize timing side-channel free code [4, 57], without requiring microarchitectural models. In particular, it captures through data and instruction caches without requiring an explicit cache model.

$$\text{Obs} ::= \text{load } n \mid \text{store } n \mid \text{pc } n \mid \text{call } f \mid \text{ret } n \quad \tau ::= \varepsilon \mid \text{Obs} \quad \bar{\tau} ::= \emptyset \mid \bar{\tau} \cdot \tau$$

The $\text{store } n$ and $\text{load } n$ events denote read and write accesses to memory location n , so they capture leaks through the data cache. In contrast, $\text{pc } n$, $\text{call } f$, and $\text{ret } n$ events record the control-flow of the program, thereby capturing leaks through, for instance, the instruction cache. An observation τ is either an event Obs or the empty observation ε . Traces $\bar{\tau}$ are sequences of observations; we indicate sequences of elements $[e_1; \dots; e_n]$ as $\bar{e}$, and adding an element e to $\bar{e}$ as $\bar{e} \cdot e$.

3.2 Syntax of μASM

μASM is an assembly-like language whose syntax is presented in Figure 2. In μASM , programs p are sequences of mappings from natural numbers n (i.e., the instruction address or label) to instructions i .

Instructions i include skipping **skip**, register assignments $x \leftarrow e$, loads **load** x, e , stores **store** x, e , indirect jumps **jmp** e , conditional branches **beqz** x, l , conditional assignments $x \xleftarrow{e'} e$, speculation barriers **spbarr**, calls **call** f , and returns **ret**. Speculation barriers (**spbarr**) and conditional assignments ($x \xleftarrow{e'} e$) are both commonly used to implement Spectre countermeasures. In particular, speculation barriers stop speculation outright, whereas conditional assignments are often used to replace branching instructions.

Instructions can refer to expressions e , which are constructed by combining registers and values with unary $\ominus$ and binary $\otimes$ operators. Registers come from the set Regs , containing register identifiers and designated registers **pc** and **sp** modelling the program counter and stack pointer respectively, whereas values come from the set Vals , which includes natural numbers and $\perp$. Note that we use $\perp$ to denote the termination of the program.

The function map $\mathcal{F}$ is a partial map from function names f to instruction numbers n .

We say that a program is *well-formed* if (1) it does not contain duplicate labels, (2) it contains an instruction labelled with 0, i.e., the initial instruction, (3) every function name f occurring in a `call  $f$` instruction is in the domain of $\mathcal{F}$ and (4) it does not contain branch instructions of the form $n : \text{beqz } x, n + 1$. Note that the last constraint ensures that all the outcomes of branch instructions are always reflected in the value of the program counter.

Example 1 (SPECTRE-PHT Example). The SPECTRE-PHT example from Listing 1 can be expressed in μASM as follows:

```

0 :  $x \leftarrow y < \text{size}$ 
1 :  $\text{beqz } x, \perp$ 
2 :  $\text{load } z, A + y$ 
3 :  $z \leftarrow z * 512$ 
4 :  $\text{load } w, B + z$ 
5 :  $\text{temp} \leftarrow \text{temp} \& w$ 

```

Here, registers `y`, `size`, and `temp` store the respective variables. Similarly, registers `A` and `B` store the memory addresses of the first elements of the arrays `A` and `B`.

In the following, we use instruction keywords to denote the set of all instructions of a given type. For instance, `beqz` is the set of all branch instructions, i.e., $\text{beqz} = \{\text{beqz } x, l \mid x \in \text{Regs} \wedge l \in \text{Vals}\}$.

3.3 Non-speculative Semantics of μASM

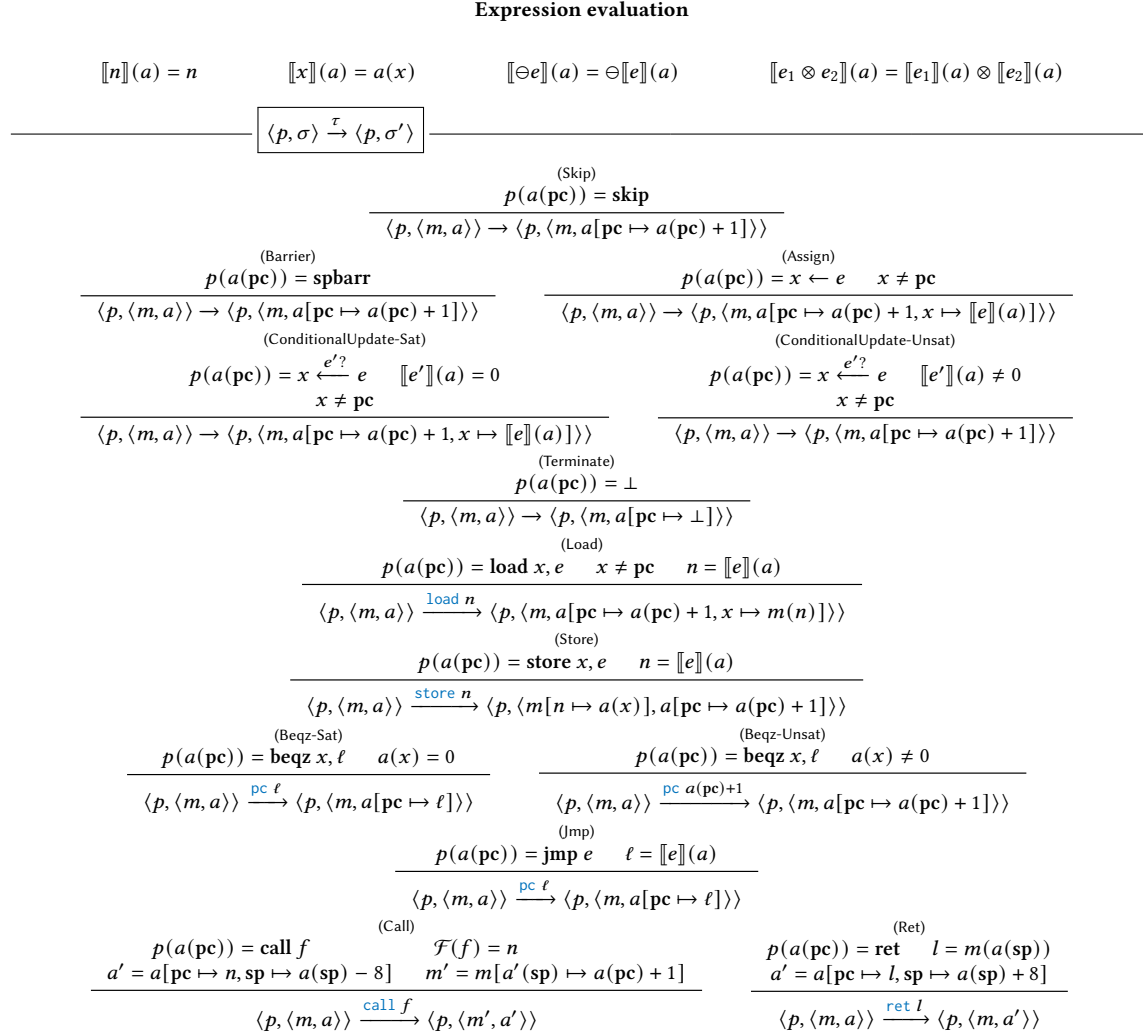
We now introduce the standard, non-speculative semantics of μASM program, which models their execution on a platform without speculation. For this, we describe next the small-step operational non-speculative semantics $\rightarrow$.

The judgment for this semantics is $\langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle$ and it reads: “a program state $\langle p, \sigma \rangle$ steps to a new program state $\langle p, \sigma' \rangle$ producing observation τ ”. Program states $\langle p, \sigma \rangle \in \Omega$ consist of a program p and a configuration σ . The program p is used to look up the current instruction, whereas the configuration $\sigma = \langle m, a \rangle$ records the state of the memory m and register file a . Memories map addresses (which are natural numbers) to values, whereas register files map register identifiers to values.

$$\begin{array}{ll}
 \text{Configurations } \sigma ::= \langle m, a \rangle & \text{Prog. States } \Omega ::= \langle p, \sigma \rangle \\
 \text{RegisterFile } a ::= \emptyset \mid r; x \mapsto v & \text{Registers } x \in \text{Regs} \\
 \text{Memory } m ::= \emptyset \mid m; n \mapsto v & \text{where } n \in \mathbb{N}
 \end{array}$$

Figure 3 presents the rules defining the non-speculative semantics $\rightarrow$. The rules rely on the evaluation of expressions (indicated as $\llbracket e \rrbracket(a) = n$) where expression e is evaluated to value n under register file a . In the rules, $a[x \mapsto n]$, where $x \in \text{Regs} \cup \mathbb{N}$ and $n \in \text{Vals}$, denotes the update of a map (memory or registers), whereas $a(x)$ denotes reading from a map. Finally, $\sigma(x)$, where $x \in \text{Regs}$ and $\sigma = \langle m, a \rangle$, denotes $a(x)$.

Most of the rules of the semantics in Figure 3 are standard; we describe selected rules below. Branch instructions emit observations recording the outcome of the branch (Rule `Beqz-Sat`, Rule `Beqz-Unsat`), while memory operations emit observations recording the accessed memory (Rule `Store`, Rule `Load`). A call to function f is a jump to the function’s starting line number n , as indicated by the function map $\mathcal{F}$. A call stores the return address on the stack at the value of


 Fig. 3. The non-speculative small-step semantics of μAsm .

the stack pointer sp and decreases sp (Rule **Call**). A return does the inverse: it looks up the return address via the stack pointer sp and then increases the stack pointer (Rule **Ret**).

Figure 4 presents the rules for deriving the trace of observations associated with an execution according to the non-speculative semantics $\rightarrow$. We denote by $(\langle p, \sigma \rangle) \Downarrow_{\text{NS}} \bar{\tau}$ the trace $\bar{\tau}$ associated with executing the program p starting from an initial configuration σ until reaching a final configuration σ' (Rule **NS-Trace**) using the reflexive-transitive-closure $\Downarrow$ of the non-speculative semantics $\rightarrow$ (Rule **NS-Reflection** and Rule **NS-Single**). To simplify our notation, we introduce the shorthand $\text{Beh}_{\text{NS}}(p, \sigma)$ to denote the specific trace $\bar{\tau}$ generated by this execution. Finally, the *non-speculative behaviour* $\text{Beh}_{\text{NS}}(p)$ of a program p is the set of all traces generated by all possible initial states for program p (Rule **NS-Beh**).

$$\begin{array}{c}
\boxed{(\langle p, \sigma \rangle) \Downarrow_{NS} \bar{\tau}} \\
\hline
\begin{array}{cc}
\text{(NS-Reflection)} & \text{(NS-Single)} \\
\frac{}{\langle p, \sigma \rangle \Downarrow_{\varepsilon} \langle p, \sigma \rangle} & \frac{\langle p, \sigma \rangle \Downarrow_{\bar{\tau}} \langle p, \sigma'' \rangle \quad \langle p, \sigma'' \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle}{\langle p, \sigma \rangle \Downarrow_{\bar{\tau} \cdot \tau} \langle p, \sigma' \rangle} \\
\text{(NS-Trace)} & \text{(NS-Beh)} \\
\frac{\exists \sigma' \quad \sigma' \in \text{FinalConf} \quad \langle p, \sigma \rangle \Downarrow_{\bar{\tau}} \langle p, \sigma' \rangle}{(\langle p, \sigma \rangle) \Downarrow_{NS} \bar{\tau}} & \frac{}{Beh_{NS}(p) = \{\bar{\tau} \mid \exists \sigma \in \text{InitConf}. (\langle p, \sigma \rangle) \Downarrow_{NS} \bar{\tau}\}}
\end{array}
\end{array}$$

Fig. 4. The Reflexive-transitive-closure of $\rightarrow$ and behaviour of μAsm programs.

4 MODELLING SPECULATIVE EXECUTION

In this section, we propose a general approach for modeling, at program-level, the effects of speculative execution induced by different kinds of speculation mechanisms. For this, we extend a non-speculative (i.e., architectural) semantics to a *speculative semantics* that is equipped with a *prediction oracle* (which captures how predictions are made) used to explore mispredicted paths. For this, we first informally explain the core concepts behind our model (Section 4.1). Then, we formalize prediction oracles and prediction histories (Section 4.2). We conclude by presenting the states and the evaluation relation associated with the (speculative) oracle semantics (Section 4.3).

We remark that here we present only a general template for defining the oracle semantics. Such template needs to be instantiated to capture specific speculation mechanisms, and we do so in Section 6 for a large class of mechanisms. In the following, we indicate a speculative semantics as $\Downarrow_x$, where the metavariable x is used as a placeholder indicating the type of speculation mechanism.

4.1 High-level Description

The core idea behind our approach is to extend a non-speculative semantics—like the one from Section 3—to a *speculative semantics* that captures the effect of speculatively executed instructions. For this, our speculative semantics is parametric in two components: (1) a set $\text{instr}_{\hat{\Delta}}$ of μAsm instructions that might trigger speculation, and (2) a *prediction oracle* O capturing the prediction strategy associated with the modeled speculation mechanism.

In a nutshell, our speculative semantics works as follows. All instructions not in $\text{instr}_{\hat{\Delta}}$ are executed as in the standard non-speculative semantics. In contrast, whenever an instruction in $\text{instr}_{\hat{\Delta}}$ is reached, the oracle O is queried to obtain the next predicted state, which records the effects of the prediction. For instance, for branch prediction, the program counter in the predicted state is updated to decide which of the two branches to execute speculatively.

To enable a subsequent rollback in case of a misprediction, a snapshot of the current program state is taken, before starting a *speculative transaction*. In this speculative transaction, the program is executed speculatively starting from the predicted state for a bounded number w of computation steps, called the *speculative window*. To abstract from the complexity of estimating the actual speculative window (which depends on the details of a CPU's microarchitecture), in our model the speculation window w is also provided by the prediction oracle.

At the end of a speculative transaction, the correctness of the prediction is evaluated: (1) If the prediction was *correct*, the transaction is committed and the computation continues using the current configuration. (2) In contrast, if the prediction was *incorrect*, the transaction is aborted, the original configuration is restored, and the computation continues on the correct branch.

In the following sections, we formalize the behavior intuitively described above in the general structure of the *speculative oracle semantics*, whereas in Section 6 we will present instantiations of the oracle semantics (and $\text{instr}_{\mathbb{Q}}$) for different speculation mechanisms.

4.2 Prediction Oracles

In our semantics, the prediction oracle serves two distinct purposes: (1) predicting the next program state (i.e., modeling how the speculation mechanism works), and (2) determining the length of the speculative transactions (i.e., determining for how long the speculative execution lasts). Our notion of prediction oracle is general enough to account for both control-flow speculation, such as branch prediction, and data-flow speculation, such as value prediction.

A *prediction oracle* O is a partial function that takes as input a program p , a prediction history h , and the current program state σ such that $p(\sigma(\text{pc}))$ is an instruction triggering speculation, i.e., $p(\sigma(\text{pc})) \in \text{instr}_{\mathbb{Q}}$. The prediction oracle returns as output a triple $\langle \delta, w, \tau \rangle$, where $\delta \in \text{Conf}$ is a partial configuration indicating the predicted values, $w \in \mathbb{N}$ represents the speculative transaction's length, and τ is a (potentially empty) observation. One can then join the returned partial configuration δ with the current configuration σ , indicated as $\sigma \uplus \delta$, to obtain the next speculative configuration. Informally, $\sigma \uplus \delta$ simply updates σ with the values in δ at the same places (registers, memory locations, etc). We say that a prediction oracle O has *speculative window at most* w if the length of the transactions generated by its predictions is at most w , i.e., for all programs p , histories h , and configurations σ , $O(p, h, \sigma) = \langle \delta, w', \tau \rangle$, for some δ, τ such that $w' \leq w$.

Taking into account the prediction history enables us to capture history-based predictors, a general class of predictors that base their decisions on past behaviors. Formally, a *prediction history* is a sequence of triples $\langle \ell, id, \delta \rangle$, where $\ell \in \text{Vals}$ is the label of an instruction related to speculation, $id \in \mathbb{N}$ is the unique identifier of the transaction that triggered speculation, and δ are the predicted values.

Example 2 (Control-flow Prediction). The “backward taken forward not taken” (BTFNT) branch predictor, implemented in early CPUs [40], predicts the branch as taken if the target instruction address is lower than the program counter. It can be formalized by the *BTFNT* oracle below, for a fixed speculative window w , as follows:

$$\text{BTFNT}(p, h, \sigma) = \langle [\text{pc} \mapsto \ell''], w, \text{pc } \ell'' \rangle, \text{ where } \ell'' = \min(\sigma(\text{pc}) + 1, \ell') \text{ and } p(\sigma(\text{pc})) = \text{beqz } x, \ell'$$

BTFNT speculatively updates the program counter to the minimum of the branch target ℓ' and the branch-not-taken target (i.e., $\sigma(\text{pc}) + 1$). Note that the *BTFNT* oracle is static, that is, it is independent of the prediction history h . Dynamic branch predictors, such as simple 2-bit predictors and more complex correlating or tournament predictors [40], can also be formalized using prediction oracles.

Example 3 (Data-flow Prediction). Oleksenko et al. [61] discovered that certain Intel CPUs speculate on division operations, called Zero Dividend Injection (ZDI), thereby resulting in a restricted form of value prediction. We model this behavior with the following oracle, which predicts that the upper bits of a 128-bit by 64-bit division operation are 0:

$$\text{ZDI}(p, h, \sigma) = \langle \delta, w, \epsilon \rangle, \text{ where } p(\sigma(\text{pc})) = k \leftarrow (x : y) \div z \text{ and } \delta = [k \mapsto (0 : y) \div z].$$

Above, $(x : y)$ denotes the 128-bit value obtained by concatenating registers x and y whereas $\div$ is the division operator.

4.3 Oracle Semantics

Next, we formalize the speculative oracle semantics. We start by formalizing the states over which the semantics operates and continue by presenting the evaluation relation given a prediction oracle O .

Speculative States. The speculative oracle semantics operates on speculative states X_x , each one consisting of a stack of speculative instances Ψ_x . The state is a stack for two reasons. First, speculation can be nested: a new (mis)prediction may occur while the processor is already executing speculatively, pushing a further transaction on top of the current one. Second, the stack records, for each ongoing transaction, the architectural state from which that speculation started, so that the effects of a mispredicted transaction can be discarded and execution resumed from that state upon rollback

Each instance Ψ_x contains the program p , a counter ctr that uniquely identifies the speculative transaction associated with the instance, a configuration σ , the prediction history h , and the remaining speculation window n describing the number of instructions that can still be executed speculatively (or $\perp$ when no speculation is happening). Depending on the specific speculation that is modelled, additional data might be tracked (indicated with $\dots$ below as well as in the evaluation rules).

$$\text{Spec. States } X_x ::= \bar{\Psi}_x \qquad \text{Spec. Instances } \Psi_x ::= \langle p, ctr, \sigma, h, n, \dots \rangle$$

Evaluation rules. We describe the speculative oracle semantics, given a prediction oracle O_x , with the relation $\rightsquigarrow_x^{O_x} \subseteq X_x \times \text{Obs}^* \times X_x$, which describes how speculative states are modified along the computation while producing observations capturing leaks. To denote the start and end of speculative transactions, we extend the set of observations Obs to include three new observations $\text{start}_x n$, $\text{rlb}_x n$, and $\text{commit}_x n$ marking the start and end (with a commit or a rollback) of speculative transaction with id n respectively.

$$\text{Obs} ::= \dots \mid \text{start}_x n \mid \text{rlb}_x n \mid \text{commit}_x n$$

Next, we describe in detail the rules formalizing the oracle semantics, which capture the intuition from Section 4.1 and are depicted in Figure 5. Whenever none of the speculative instances in the current state has a speculative window of 0 or is stuck (denoted by $\text{canStep}(\bar{\Psi}_x)$), the computation proceeds by doing one step for the instance at the top of the stack (Rule **O-Context**) and by updating the speculative window of the other instances. In particular, if the instruction to be executed in the topmost instance is a speculative barrier, then all speculative windows are set to 0 ($\bar{\Psi}_x'' = \text{zeroes}(\bar{\Psi}_x)$), otherwise all windows are decremented by 1 ($\bar{\Psi}_x'' = \text{decr}(\bar{\Psi}_x)$).

In the topmost speculative instance, instructions that do not trigger speculation (i.e., they are not in $\text{instr}_{\hat{\Delta}}$) are executed following the non-speculative semantics while also decrementing the current speculative window by 1 (Rule **O-NoSpec**). Speculation barriers are handled similarly except for setting the speculative window in the topmost instance to 0 (Rule **O-barr**). Finally, Rule **O-Spec** handles instructions that trigger speculation (i.e., those that belong to $\text{instr}_{\hat{\Delta}}$). In this case, the oracle is queried to obtain the prediction δ and the next speculative window w , and a new speculative instance (starting from the predicted state $\sigma \uplus \delta$ and with an updated prediction history h') is pushed on top of the stack to start the new speculative transaction. This fresh speculative instance is decorated with the original program state σ and with the predicted values δ , which will later on be used to determine whether the transaction should be committed or rolled back. The rule also appends the observation $\text{start}_x n$ to the trace to record the beginning of a speculative transaction.

A speculative transaction must be finalized (committed or rolled back) whenever its speculation window reaches 0 or the execution cannot proceed further. We refer to the latter case as the instance being ‘stuck’, formally indicated by

$$\begin{array}{c}
\boxed{\bar{\Psi}_x \xrightarrow{\tau}^{O_x} \bar{\Psi}'_x} \\
\text{(O-Context)} \\
\frac{\Psi_x \xrightarrow{\tau}^{O_x} \bar{\Psi}'_x \quad \text{canStep}(\bar{\Psi}_x \cdot \Psi_x)}{\text{if } p(\Psi_x \cdot \sigma(\text{pc})) = \text{spbarr} \text{ then } \bar{\Psi}_x'' = \text{zeroes}(\bar{\Psi}_x) \text{ else } \bar{\Psi}_x'' = \text{decr}(\bar{\Psi}_x)} \\
\frac{\bar{\Psi}_x \cdot \Psi_x \xrightarrow{\tau}^{O_x} \bar{\Psi}_x'' \cdot \bar{\Psi}'_x}{\text{(O-Rollback)}} \\
\frac{\langle p, \sigma'' \rangle \xrightarrow{\tau} \langle p, \sigma''' \rangle \quad (\sigma''' \neq \sigma'' \uplus \delta \text{ or } p, \sigma' \vdash \text{fin})}{\bar{\Psi}_x \cdot \langle p, \text{ctr}, \sigma, h, n, \dots \rangle \cdot \langle p, \text{ctr}', \sigma', h', 0, \dots \rangle \langle \sigma'', \delta \rangle \cdot \bar{\Psi}_x' \xrightarrow{\text{rlb}_x \text{ ctr} \cdot \tau}^{O_x} \bar{\Psi}_x \cdot \langle p, \text{ctr}', \sigma''', h', n-1, \dots \rangle} \\
\text{(O-Commit)} \\
\frac{\langle p, \sigma'' \rangle \xrightarrow{\tau} \langle p, \sigma''' \rangle \quad (\sigma''' = \sigma'' \uplus \delta \text{ or } p, \sigma' \vdash \text{fin})}{\bar{\Psi}_x \cdot \langle p, \text{ctr}, \sigma, h, n, \dots \rangle \cdot \langle p, \text{ctr}', \sigma', h', 0, \dots \rangle \langle \sigma'', \delta \rangle \cdot \bar{\Psi}_x' \xrightarrow{\text{commit}_x \text{ ctr} \cdot \tau}^{O_x} \bar{\Psi}_x \cdot \langle p, \text{ctr}', \sigma', h', n, \dots \rangle \cdot \bar{\Psi}_x'} \\
\boxed{\Psi_x \xrightarrow{\tau}^{O_x} \bar{\Psi}'_x} \\
\text{(O-NoSpec)} \quad \text{(O-barr)} \\
\frac{p(\sigma(\text{pc})) \notin \text{instr}_{\hat{\Pi}} \cup \{\text{spbarr}\} \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle}{\langle p, \text{ctr}, \sigma, h, n+1, \dots \rangle \xrightarrow{\tau}^{O_x} \langle p, \text{ctr}, \sigma', h, n, \dots \rangle} \quad \frac{p(\sigma(\text{pc})) = \text{spbarr} \quad \sigma \xrightarrow{\tau} \sigma'}{\langle p, \text{ctr}, \sigma, h, n, \dots \rangle \xrightarrow{\tau}^{O_x} \langle p, \text{ctr}, \sigma', h, 0, \dots \rangle} \\
\text{(O-Spec)} \\
\frac{p(\sigma(\text{pc})) \in \text{instr}_{\hat{\Pi}} \quad O(p, n, h, \sigma) = \langle \delta, w, \tau \rangle \quad h' = h \cdot \langle \sigma(\text{pc}), n+1, \delta \rangle \quad \bar{\tau} = \text{start}_x \text{ ctr} \cdot \tau}{\langle p, \text{ctr}, \sigma, h, n+1, \dots \rangle \xrightarrow{\bar{\tau}}^{O_x} \langle p, \text{ctr}, \sigma, h, n+1, \dots \rangle \cdot \langle p, \text{ctr}+1, \sigma \uplus \delta, h', w, \dots \rangle \langle \sigma, \delta \rangle}
\end{array}$$

Fig. 5. Evaluation rules for the speculative oracle semantics $\xrightarrow{\tau}^{O_x}$ given an oracle O_x

the judgment $p, \sigma' \vdash \text{fin}$. This judgment holds whenever the current instruction is undefined, i.e., $p(\sigma(\text{pc})) = \perp$ (see Rule **Terminate**). To determine whether to commit or rollback, the semantics inspects the snapshot state σ and prediction δ that decorate the instance $\Psi_x^{\langle \sigma, \delta \rangle}$ whose window has reached 0.

- If the prediction was *correct* (i.e., doing one step from σ results indeed in the predicted state $\sigma \uplus \delta$), the transaction is committed and the computation continues using the current configuration (Rule **O-Commit**).
- If the prediction was *incorrect* (i.e., doing one step from σ produces a state different from the predicted state $\sigma \uplus \delta$), the transaction is aborted and the computation continues on the correct branch (Rule **O-Rollback**).

Thus, committing and rolling back speculative transactions happen along the stack of states. Rolling back deletes all the instances above the rolled back instance, whereas committing updates the configuration, the counter, the prediction history h and additional data tracked by the semantics of the instance below and the committed instance is deleted. Rule **O-Commit** and Rule **O-Rollback** also record the end of speculative transactions on the trace with the `commitx n` and `rlbx n` observations respectively.

Analogously to the non-speculative case, the behaviour $\text{Beh}_x^O(p)$ of a program p under the oracle semantics is the set of all traces generated from an initial state until termination.

Non-Speculative Consistency. Fundamentally, a valid speculative semantics must preserve the program's architectural behaviour. Speculation should only introduce observational side-effects during transient execution, without altering

the program's non-speculative behaviour. We formalize this correctness requirement in Property 1, stating that the standard non-speculative behaviour can be exactly recovered from the oracle behaviour by applying the non-speculative projection $\downarrow_{ns}$, which (1) removes from the trace τ any sub-trace enclosed between $\text{start}_x n$ and $\text{rlb}_x n$ for some n , and then (2) drops all remaining $\text{start}_x n$ and $\text{commit}_x n$ observations. In Section 6, we prove that all our speculative semantics instances satisfy Property 1.

Property 1 (NS Consistency Oracle). $\text{Beh}_{NS}(p) = \text{Beh}_x^O(p) \downarrow_{ns}$

5 SPECULATIVE NON-INTERFERENCE

In this section, we introduce *speculative non-interference* (SNI, Section 5.1), a semantic notion characterizing the leaks introduced by speculatively-executed instructions. Next, we introduce the notion of *always-mispredict speculative semantics* (Section 5.2), that is, a speculative semantics that facilitates reasoning about leaks w.r.t. *any* prediction oracle.

5.1 Speculative Non-Interference

Speculative non-interference (SNI) is a semantic notion of security characterizing those information leaks that are introduced by speculative execution. Intuitively, SNI requires that programs do not leak more information under the speculative semantics than what is leaked under the non-speculative semantics.

SNI is parametric in a policy ϕ and in the speculative semantics x , which models how the program executes. The policy ϕ describes which parts of the program are public/low information, i.e., those that are known by an adversary. Formally, a security policy P is a finite subset of $\text{Regs} \cup \mathbb{N}$ specifying the low register identifiers and memory addresses. Two configurations σ, σ' are called *low-equivalent* for a policy ϕ , written $\sigma \sim_\phi \sigma'$, if they agree on all register and memory locations in ϕ .

Example 4 (Example Policy for Example 1). A policy P for the program from Example 1 may state that the content of the registers y , size , A , and B is non-sensitive, i.e., $P = \{y, \text{size}, A, B\}$.

Policies need not be manually specified but can in principle be inferred from the context in which a piece of code executes, e.g., whether a variable is reachable from public input or not.

A program p satisfies SNI (Definition 1) for a speculative semantics x if any pair of low-equivalent initial configurations σ and σ' that generate the same observations under the non-speculative semantics also generate the same observations under the speculative semantics.

Definition 1 (Speculative Non-Interference). *Program p satisfies SNI (denoted $p \vdash_x \text{SNI}$) for a speculative semantics x if for all σ, σ' , if $\sigma \sim_\phi \sigma'$ and $\text{Beh}_{NS}(p, \sigma) = \text{Beh}_{NS}(p, \sigma')$ then $\text{Beh}_x(p, \sigma) = \text{Beh}_x(p, \sigma')$.*

SNI is parametric in the speculative semantics x . For the oracle semantics, which additionally depends on a prediction oracle O , we make the oracle explicit and write $p \vdash_x^O \text{SNI}$. The always-mispredict semantics, introduced in Section 5.2, does not depend on an oracle; for it we keep the plain notation $p \vdash_x \text{SNI}$.

Speculative non-interference is a variant of non-interference. While non-interference compares what is leaked by a program with a policy specifying the allowed leaks, speculative non-interference compares the program leakage under two semantics, the non-speculative and the speculative one. The security policy and the non-speculative semantics, together, specify what the program may leak under the speculative semantics.³

³Conceptually, the non-speculative semantics induces declassification assertions for the speculative semantics [66].

Example 5 (SNI for Example 1). The program p from Example 1 does not satisfy speculative non-interference for the BTFNT oracle from Example 2 and the policy P from Example 4. Consider two initial configurations $\sigma := \langle m, a \rangle$, $\sigma' := \langle m', a' \rangle$ that agree on the values of y , size , A , and B but disagree on the value of $B[A[y] * 512]$. Say, for instance, that $m(a(A) + a(y)) = 0$ and $m'(a'(A) + a'(y)) = 1$. Additionally, assume that $y \geq \text{size}$.

Executing the program under the non-speculative semantics produces the trace $\text{pc} \perp$ when starting from σ and σ' . Moreover, the two initial configurations are indistinguishable with respect to the policy P . However, executing p under the speculative semantics produces two distinct traces:

$$\begin{aligned}\tau &= \text{start } 0 \cdot \text{pc } 3 \cdot \text{load } v_1 \cdot \text{load } (a'(B) + 0) \cdot \text{rlb } 0 \cdot \text{pc } \perp \\ \tau' &= \text{start } 0 \cdot \text{pc } 3 \cdot \text{load } v_1 \cdot \text{load } (a'(B) + 1) \cdot \text{rlb } 0 \cdot \text{pc } \perp\end{aligned}$$

where $v_1 = a(A) + a(y) = a'(A) + a'(y)$. Therefore, p does not satisfy speculative non-interference.

Note that $p \vdash_x^O$ SNI requires a specific oracle. We restrict attention to oracles whose predictions do not themselves depend on secrets.

Definition 2 (Non-Interfering oracle). A prediction oracle O is non-interfering w.r.t. a policy ϕ if, for every program p and history h and for all configurations $\sigma \sim_\phi \sigma'$, $O(p, h, \sigma) = O(p, h, \sigma')$.

Intuitively, a non-interfering oracle bases its predictions only on the public part of the configuration. This costs no generality in the speculation mechanisms we can model: prediction in real hardware is driven by the program's execution history and by microarchitectural state, never by secret architectural data. The only oracles excluded are physically unrealizable ones that inspect secrets to decide whether or how to speculate. Such an oracle would itself be a secret-dependent control-flow channel. For the remainder of the paper we therefore assume that all prediction oracles are non-interfering.

5.2 Always-Mispredict (AM) Semantics

The oracle speculative semantics from Section 4.3 and, as a result, SNI are parametric in the prediction oracle O . Often, however, it is desirable to obtain security guarantees w.r.t. *any* prediction oracle, since the details about speculation mechanisms might differ between different CPUs or may even be unknown. To this end, we introduce a variant of the speculative semantics, which we call the *always-mispredict speculative semantics*, that facilitates reasoning about leaks w.r.t. any prediction oracle. We formally define the AM semantics as a speculative template rather than a concrete mechanism. This template abstracts the core logic of exploring mispredicted paths and serves as a unified foundation for the specific instantiations defined in Section 6.

For simplicity, consider the case of branch prediction. In this case, leakage due to speculative execution is maximized under a predictor that mispredicts every time. This intuition holds true unless speculative transactions are nested, where a correct prediction of a nested branch sometimes yields more leakage than a misprediction.

Example 6. Consider the following variation of the SPECTRE-PHT example [49] from Figure 1, and assume that the function `benign()` runs for longer than the speculative window and does not leak any information.

Then, under a branch predictor that mispredicts every branch, the speculative transaction corresponding to the outer branch will be rolled back before reaching line 4. On the other hand, given a correct prediction of the inner branch, line 4 would be reached and a speculative leak would be present.

$$\begin{array}{c}
\boxed{\bar{\Phi}_x \xrightarrow{\tau} \bar{\Phi}'_x} \\
\hline
\begin{array}{c}
\text{(AM-NoSpec)} \quad \frac{p(\sigma(\text{pc})) \notin \text{instr}_{\bar{\Delta}} \cup \{\text{spbarr}\} \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle}{\bar{\Phi}_x \cdot \langle p, \text{ctr}, \sigma, n+1, \dots \rangle \xrightarrow{\tau} \bar{\Phi}_x \cdot \langle p, \text{ctr}, \sigma', n, \dots \rangle} \\
\text{(AM-barr)} \quad \frac{p(\sigma(\text{pc})) = \text{spbarr} \quad \sigma \xrightarrow{\tau} \sigma'}{\bar{\Phi}_x \cdot \langle p, \text{ctr}, \sigma, n, \dots \rangle \xrightarrow{\tau} \bar{\Phi}_x \cdot \langle p, \text{ctr}, \sigma', 0, \dots \rangle} \\
\text{(AM-Spec)} \quad \frac{p(\sigma(\text{pc})) = \text{instr}_{\bar{\Delta}} \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle \quad j = \min(\omega, n) \quad \bar{\tau} = \tau \cdot \text{start}_x \text{ ctr} \dots}{\bar{\Phi}_x \cdot \langle p, \text{ctr}, \sigma, n+1 \rangle \xrightarrow{\bar{\tau}} \bar{\Phi}_x \cdot \langle p, \text{ctr}, \sigma', n \rangle \cdot \Phi'_x} \\
\text{(AM-Rollback)} \quad \frac{\Phi'_x.n = 0 \text{ or } p, \Phi'_x.\sigma \vdash \text{fin}}{\bar{\Phi}_x \cdot \langle p, \text{ctr}, \sigma, n \rangle \cdot \Phi'_x \xrightarrow{\text{rlb}_x \text{ ctr}} \bar{\Phi}_x \cdot \langle p, \Phi'_x.\text{ctr}', \sigma, n \rangle}
\end{array}
\end{array}$$

Fig. 6. Always-mispredict speculative semantics

```

1  if (y < size)
2    if (y-1 < size)
3      benign();
4    temp &= B[A[y] * 512];

```

A simple but inefficient approach to deal with this challenge would be to consider both cases, correct and incorrect predictions, upon every branch. This, however, would result in an exponential explosion of the number of paths to consider. Furthermore, it would not be applicable to speculation mechanisms where there might be more than one incorrect prediction, e.g.,

indirect branch speculation or value speculation.

Intuition. To address these issues, we introduce the *always-mispredict speculative semantics* that differs from the oracle speculative semantics in three key ways:

- (1) It mispredicts every time, hence its name. For every instruction that might trigger speculation (i.e., the instruction belongs to $\text{instr}_{\bar{\Delta}}$), the always-mispredict semantics first speculatively explores all possible *wrong* paths for a bounded number of steps and then continues with the correct one. Thus, the semantics is not parametric in the prediction oracle.
- (2) It initializes the length of every *non-nested* transaction to w , and the length of every *nested* transaction to the remaining length of its enclosing transaction, decremented by 1.
- (3) Upon executing instructions, only the remaining length of the innermost transaction is decremented.

The consequence of these modifications is that nested transactions do not reduce the number of steps that the semantics may explore the correct path for, after the nested transactions have been rolled back. In Example 6, after rolling back the nested speculative transaction, the outer transaction continues as if the nested branch had been correctly predicted in the first place, and thus the speculative leak in line 4 is reached.

Formalization. The state Σ_x of the AM semantics is a stack of speculative instances Φ_x . As shown in all rules in Figure 6, reductions in the always-mispredict semantics happen only on top of the stack. Each instance Φ_x contains the program p , a counter ctr that identifies the speculative transaction, a configuration σ , and the remaining speculation window n describing the number of instructions that can still be executed speculatively (or $\perp$ when no speculation is happening). Depending on the specific speculation that is modelled, additional data is tracked (indicated with $\dots$

in the rules). Throughout the paper, we fix the maximal speculation window, i.e., the maximum number of speculative instructions, to a global constant ω .

$$\text{Spec. States } \Sigma_x ::= \bar{\Phi}_x \qquad \text{Spec. Instances } \Phi_x ::= \langle p, ctr, \sigma, n, \dots \rangle$$

Modifications (1)–(3) are captured in the four rules given in Figure 6. The judgement for the AM semantics is: $\Sigma_x \xrightarrow{\tau} \mathcal{J}_x \Sigma'_x$. If the instruction is not related to speculation and it is not a speculation barrier, then the speculative instance on the top of the stack is updated according to the non-speculative semantics $\rightarrow$ (Rule **AM-NoSpec**). In contrast, whenever the current instruction is a speculation barrier **spbarr**, the remaining speculation window of the topmost instance is set to 0.

Whenever speculation starts, one or more speculative instances are pushed on top of the stack (Rule **AM-Spec**) and when speculation ends, the speculative instance is then popped (Rule **AM-Rollback**). Finally, we note that how exactly the new speculative instances are created w.r.t. the auxiliary data depends on the specific speculative semantics, as we show in Section 6. This is reflected in Rule **AM-Spec** not specifying how the new speculative instances $\bar{\Phi}'_x$ are defined. In all cases, however, the new instances $\bar{\Phi}'_x$ are pushed with remaining speculation window $j = \min(\omega, n)$; only the auxiliary data is instance-specific.

The always-mispredict behaviour $Beh_x^\omega(p)$ of a program p is the set of all traces generated from an initial state until termination using the reflexive-transitive closure of $\xrightarrow{\tau} \mathcal{J}_x$. Crucially, it is possible to connect the AM semantics and the non-speculative semantics through the non-speculative projection $\upharpoonright_{ns}$, which removes from a trace $\bar{\tau}$ all events related to speculation. We lift the projection $\upharpoonright_{ns}$ to the behaviour of a program p in the natural way.

We require that any instantiation of this template (e.g., for branches or stores) preserves the non-speculative behaviour of the program (Similar to the Oracle Semantics in Section 4.3). We formalize this as a correctness requirement that all our instances must satisfy:

Property 2 (NS Consistency AM). $Beh_{NS}(p) = Beh_x^\omega(p) \upharpoonright_{ns}$

The goal of the AM semantics is to derive security guarantees that are independent of the choice of a prediction oracle. Property 3 formalizes this intuition by precisely connecting the oracle and the AM semantics of any speculative semantics $\mathcal{J}_x$. In particular, Property 3 states that checking SNI w.r.t. the AM semantics is sufficient to obtain security guarantees w.r.t. *all* prediction oracles. That is, if a program is SNI w.r.t. the always-mispredict semantics, then it is SNI irrespectively of the choice of prediction oracle. Similarly, if a program violates SNI w.r.t. the always-mispredict semantics, then there is one prediction oracle for which SNI is violated. As we show in Section 6, this holds for all instances studied in this paper.

Property 3 (Oracle Overapproximation). $p \vdash_x \text{SNI} \text{ iff } \forall \mathcal{O}. p \vdash_x^{\mathcal{O}} \text{SNI}$

6 INSTANCES OF SPECULATIVE SEMANTICS

Here we present specific instances of speculative semantics modelling the effect of speculative execution over branch instructions (Section 6.1), **store** instructions (Section 6.2), **ret** instructions (Section 6.3, Section 6.4) and indirect jump instructions (Section 6.5) using the semantics templates presented in Section 4.

Before detailing these specific instances, we must establish the properties that well-formed semantics must satisfy in our framework. We summarise these properties in a single definition:

Definition 3 (Well-Formed Speculative Semantics). *A speculative semantics $\mathcal{J}_x$ is well-formed (denoted $\vdash \mathcal{J}_x$ WFSS) if:*

Manuscript submitted to ACM

- *Oracle Overapproximation*: $p \vdash_x \text{SNI} \text{ iff } \forall O. p \vdash_x^O \text{SNI}$
- *NS Consistency*: $\text{Beh}_x^\omega(p) \upharpoonright_{ns} = \text{Beh}_{NS}(p) = \text{Beh}_x^O(p) \upharpoonright_{ns}$
- *Symbolic Consistency*: $\text{Beh}_x^\omega(p) = \mu(\text{Beh}_x^S(p))$

Intuitively, a well-formed speculative semantics is made of three components: an AM semantics, an oracle semantics, and a symbolic AM semantics.

First, the AM semantics must overapproximate the oracle semantics (for any oracle), guaranteeing that it is sufficient to check a program p for SNI w.r.t. the AM semantics. Next, both the AM and the Oracle semantics must preserve the non-speculative behaviour of a program p . Applying the non-speculative projection ($\upharpoonright_{ns}$) to their traces exactly recovers the standard non-speculative behaviour. Thus, we can execute p only once to get the (non-)speculative behaviour of that program run.

Finally, to enable automated verification, we require the Symbolic AM semantics. Both the symbolic semantics and the concretization function μ (which maps symbolic values to concrete ones) are defined in Section 9, as they are used only by our analysis tool. The Symbolic Consistency property mandates that the concrete AM traces exactly match the concretization of the symbolic traces, where $\mu(\text{Beh}_x^S(p))$ conceptually denotes the instantiation of the symbolic traces with concretizations that satisfy the collected path conditions.

For conciseness' sake, here we only present the concrete AM semantics in full detail and refer to the technical report for the full details of the oracle/symbolic semantics [30], and defer the detailed discussion of how the symbolic semantics is utilized in practice to the implementation of SPECTECTOR in Section 9. For each of the speculative semantics, we prove that they satisfy all validity requirements, that is, that they are well-formed speculative semantics according to Definition 3.

6.1 $\mathcal{L}_B$: Speculation on Branch Instructions

Modern hardware uses a branch predictor to predict the outcome of branching decisions since this speeds up program execution. However, mispredictions can be steered and exploited by attackers leading to SPECTRE-PHT attacks [49].

6.1.1 The AM Semantics. At every branch instruction, the always-mispredict semantics first speculatively executes the wrong branch for a fixed number of steps and then continues with the correct one. As a result, this semantics is deterministic and agnostic to implementation details of the branch predictor.

The state Σ_B of the AM semantics is a stack of speculative instances Φ_B where reductions happen only on top of the stack. Each instance Φ_B contains the program p , a counter ctr that uniquely identifies the speculation instance, a configuration σ , and the remaining speculation window n describing the number of instructions that can still be executed speculatively (or $\perp$ when no speculation is happening). In this semantics, we have $\text{instr}_{\mathcal{A}} = \text{beqz}$ since branches are the source of speculation.

$$\text{Spec. States } \Sigma_B ::= \bar{\Phi}_B$$

$$\text{Spec. Instances } \Phi_B ::= \langle p, ctr, \sigma, n \rangle$$

The judgement for the AM semantics is: $\Sigma_B \xRightarrow{\tau} \Sigma'_B$.

$$\boxed{\bar{\Phi}_B \xRightarrow{\tau} \bar{\Phi}'_B}$$

$$\begin{array}{c}
 \text{(B:AM-NoSpec)} \\
 \frac{p(\sigma(\text{pc})) \notin \text{beqz} \cup Z_B \cup \{\text{spbarr}\} \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle}{\langle p, \text{ctr}, \sigma, n+1 \rangle \xrightarrow{\tau} \langle p, \text{ctr}, \sigma', n \rangle} \\
 \text{(B:AM-Barr)} \quad \text{(B:AM-Rollback)} \\
 \frac{p(\sigma(\text{pc})) = \text{spbarr} \quad \sigma \xrightarrow{\tau} \sigma' \quad n' = 0 \text{ or } p \text{ is stuck}}{\langle p, \text{ctr}, \sigma, n \rangle \xrightarrow{\tau} \langle p, \text{ctr}, \sigma', 0 \rangle \quad \langle p, \text{ctr}, \sigma, n \rangle \cdot \langle p, \text{ctr}', \sigma', n' \rangle \xrightarrow{\text{rlb}_B \text{ ctr}} \langle p, \text{ctr}', \sigma, n \rangle} \\
 \text{(B:AM-Spec)} \\
 \frac{p(\sigma(\text{pc})) = \text{beqz } x, \ell \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle \quad j = \min(\omega, n) \quad \sigma'' = \sigma[\text{pc} \mapsto l'] \quad \bar{\tau} = \tau \cdot \text{start}_B \text{ ctr} \cdot \text{pc } l}{l' = \begin{cases} \sigma(\text{pc}) + 1 & \text{if } \sigma'(\text{pc}) = l \\ l & \text{if } \sigma'(\text{pc}) \neq l \end{cases} \quad \langle p, \text{ctr}, \sigma, n+1 \rangle \xrightarrow{\bar{\tau}} \langle p, \text{ctr}, \sigma', n \rangle \cdot \langle p, \text{ctr} + 1, \sigma'', j \rangle}
 \end{array}$$

As mentioned, Rule **B:AM-Spec** pushes a new speculative state with the wrong branch, followed by the state with the correct one. When speculation ends, Rule **B:AM-Rollback** pops the related state. All other instructions are handled by delegating back to the non-speculative semantics (Rule **B:AM-NoSpec**).

Rule **B:AM-NoSpec** differs slightly from what was presented before in Section 5.2: it applies to instructions that are not branch or barrier instructions **and** are not in the metaparameter Z_B (in gray). The latter is a set of instructions and is part of our composition framework (which we explain in Section 7.1). Instantiating Z_B allows us to restrict when to apply non-speculative steps in composed semantics. When we consider $\mathcal{L}_B$ in isolation, Z_B is the empty set (so, Rule **B:AM-NoSpec** applies to everything except branch and barrier instructions). However, we will instantiate Z_B in different manners when building the composed semantics. In the following, we write $\mathcal{L}_B^{Z_B}$ to stress the value of Z_B when needed but we often omit Z_B for simplicity.

The always-mispredict behaviour $\text{Beh}_B^\omega(p)$ of a program p is the set of all traces generated from an initial state until termination using the reflexive-transitive closure of $\xrightarrow{\tau} \mathcal{L}_B$.

Note that Rule **B:AM-NoSpec**, Rule **B:AM-Barr** and Rule **B:AM-Rollback** are direct instantiations of the general rules described in Section 5.2. This holds true for all the other speculative semantics that we will present, thus, we omit them.

6.1.2 Oracle Semantics. At every branch instruction, the oracle semantics queries the explicit oracle O_B for the branching decision.

Here, we summarize the key differences with the AM semantics. First, speculative instances are extended to track the branching history h , which records the outcomes of prior branch instructions. Second, when executing a **beqz** instruction, the oracle predicts the branch outcome (based on the branching history h) and a new speculative instance is pushed on top of the stack (Rule **B:Spec**). Finally, whenever the speculation window of an instance anywhere on the stack reaches 0, the execution needs to be rolled back or committed.

$$\Psi_B \xrightarrow{\tau} \Psi'_B$$

$$\begin{array}{c}
\text{(B:NoBranch)} \\
\frac{p(\sigma(\text{pc})) \notin \cup Z \quad \sigma \xrightarrow{\tau} \sigma'}{\langle p, \text{ctr}, \sigma, h, n+1 \rangle \xrightarrow{\tau}_{\text{B}}^O \langle p, \text{ctr}, \sigma', h, n \rangle} \\
\text{(B:Spec)} \\
\frac{
\begin{array}{l}
p(\sigma(\text{pc})) = \text{beqz } x, \ell \quad \sigma \xrightarrow{\tau} \sigma' \quad O(p, n, h, \sigma) = (\delta, \omega, \text{pc } \ell'') \\
\delta = [\text{pc} \mapsto \ell''] \quad h' = h \cdot \langle \sigma(\text{pc}), \text{ctr}, \delta \rangle \\
\bar{\tau} = \tau \cdot \text{start}_{\text{B}} \text{ctr} \cdot \text{pc } \ell''
\end{array}
}{\langle p, \text{ctr}, \sigma, h, n+1 \rangle \xrightarrow{\bar{\tau}}_{\text{B}}^O \langle p, \text{ctr}, \sigma', h', n \rangle \cdot \langle p, \text{ctr}+1, \sigma \uplus \delta, h', \omega \rangle^{\langle \sigma, \delta \rangle}}
\end{array}$$

As in the oracle-semantics template, rolling back deletes all the instances above the rolled back instance, whereas committing updates the configuration, the counter and the branching history h of the instance below and the committed instance is deleted. These rules are not shown since they are direct instantiations of the general oracle rules presented in Section 4.3.

6.2 $\mathcal{L}_S$: Speculation on Store Instructions

Modern processors write **stores** to main memory asynchronously to reduce delays caused by the memory subsystem. Processors employ a *Store Queue* where not-yet-committed **store** instructions are stored before being permanently written to memory. When executing a **load** instruction, the processor first inspects the store queue for a matching memory address. If there is a match, the value is retrieved from the store queue (called *store-to-load forwarding*), and otherwise, the memory request is issued to the memory subsystem. To speed up computation, processors employ memory disambiguation predictors to predict if the memory addresses of loads and stores match. Since the prediction can be incorrect, processors may speculatively bypass a **store** instruction in the store queue, leading to a **load** instruction retrieving a stale value [43].

Example 7 (SPECTRE-STL). Consider the code in Listing 6, which is the μASM translation of Listing 2:

Listing 6. SPECTRE-STL in μASM .

```

1  store secret, p
2  store public, p
3  load eax, p
4  load edx, B + eax

```

Assume that the **store** instructions in Lines 1 and 2 are still in the *store queue* and not yet committed to main memory. A misprediction of the memory disambiguator for the **load** instruction in Line 3 causes it to bypass the **store** instruction in Line 2 and retrieve the value from the stale **store** instruction in Line 1. The speculative access of the memory is then leaked into the microarchitectural state by the array access into **B** in Line 4.

Here, we present the speculative AM semantics (Section 6.2.1) and the corresponding oracle semantics (Section 6.2.2).

6.2.1 Speculative Semantics. The overall structure of the $\mathcal{L}_S$ semantics is similar to that of $\mathcal{L}_B$: speculative execution is modelled using a stack of speculative states, instructions that do not start speculative transactions are executed by delegating back to the non-speculative semantics, and speculative transactions are rolled back whenever the speculative window reaches 0. The key difference between $\mathcal{L}_S$ and $\mathcal{L}_B$ is the differing source of speculation $\text{instr}_{\mathcal{L}}$. Here $\text{instr}_{\mathcal{L}} = \text{store}$ for $\mathcal{L}_S$ instead of $\text{instr}_{\mathcal{L}} = \text{beqz}$ as in $\mathcal{L}_B$.

The states used in $\mathcal{L}_S$ are similar to those of $\mathcal{L}_B$:

$$\begin{array}{ll}
\text{Spec. States } \Sigma_S ::= \bar{\Phi}_S & \text{Spec. Instance } \Phi_S ::= \langle p, \text{ctr}, \sigma, n \rangle
\end{array}$$

We also add a **bypass** n observation denoting that the **store** instruction at program counter n was speculatively bypassed.

$$Obs_S ::= Obs \mid \text{bypass } n$$

The judgement $\Sigma_S \xrightarrow{\tau} \Sigma'_S$ describes how Σ_S steps to Σ'_S emitting observation τ . As in $\mathcal{L}_B$, reductions only happen on top of the stack.

$$\boxed{\overline{\Phi}_S \xrightarrow{\tau} \overline{\Phi}'_S}$$

(S:AM-Spec)

$$\frac{\begin{array}{l} p(\sigma(\text{pc})) = \text{store } x, e \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle \quad j = \min(\omega, n) \\ \sigma'' = \sigma[\text{pc} \mapsto \sigma(\text{pc}) + 1] \quad \tau' = \tau \cdot \text{bypass } \sigma(\text{pc}) \cdot \text{start}_S \text{ ctr} \end{array}}{\langle p, \text{ctr}, \sigma, n+1 \rangle \xrightarrow{\tau'} \langle p, \text{ctr}, \sigma', n \rangle \cdot \langle p, \text{ctr} + 1, \sigma'', j \rangle}$$

To model the effect of bypassing a **store** instruction, Rule S:AM-Spec bypasses the **store** instruction by increasing the program counter without updating the memory and starts a new speculative transaction by pushing a new speculative instance on top of the state. A **load** instruction loading from the same memory location as the bypassed **store** instruction, therefore, retrieves a stale value. Our semantics models store-to-load forwarding as all-or-nothing: a **load** reads either the bypassed (stale) value or the stored value. We do not capture partial or size-mismatched forwarding, where only part of a store is forwarded.

The set $Beh_S^\omega(p)$ contains all traces generated from an initial state until termination using the reflexive-transitive closure of $\xrightarrow{\tau}$.

6.2.2 Oracle Semantics. Instead of bypassing **store**, the oracle semantics employs an oracle $\mathcal{O}$ that decides if the **store** instruction should be speculatively bypassed or not. As before, the behaviour $Beh_S^{\mathcal{O}}(p)$ of a program p is the set of all traces starting from an initial state until termination using the reflexive-transitive closure of the oracle semantics.

6.3 $\mathcal{L}_R$: Speculation on Return Instructions

The return-stack buffer (RSB) is a small stack the CPU uses to save return addresses upon **call** instructions. These saved return addresses are speculatively used when the function returns because accessing the RSB is faster than looking up the return address on the stack (stored in main memory). This works well because return addresses rarely change during function execution. However, mispredictions can be exploited by an attacker [50, 53].

Example 8 (Return Speculation Vulnerability). Consider the example in Listing 7 and recall that register **sp** is used to find return addresses saved on the stack.

Each function call pushes a return address on the stack and decrements the **sp** register. After reaching the function *Manip_Stack*, the **sp** register is incremented (line 2). Thus, **sp** points to the previous return address on the stack (i.e., line 11), and the non-speculative execution continues in *Main* and terminates. However, the return address of the call in line 5 is line 6 and it is on top of the RSB. Thus, the CPU speculatively executes lines 6–7 and leaks the secret.

This section describes the AM semantics (Section 6.3.1) and the oracle semantics (Section 6.3.2). Then, it discusses formalising different implementations of the RSB in the CPU (Section 6.3.3).

6.3.1 Speculative Semantics. Unlike before, the state of $\mathcal{L}_R$ contains a model of the RSB which is used to retrieve return addresses instead of relying on the stack.

Listing 7. A program exploiting RSB speculation.

```

1  Manip_Stack:
2      sp ← sp + 8
3      ret
4  Speculate:
5      call Manip_Stack
6      load eax, secret
7      load edx, eax
8      ret
9  Main:
10     call Speculate
11     skip

```

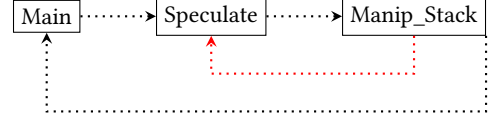


Fig. 7. Control Flow of the vulnerable program in Listing 7. A red arrow indicates a speculative control flow that happens because of misprediction with the RSB.

Thus, speculative instances of $\mathcal{L}_R$ are extended with an additional entry $\mathbb{R}$ for tracking the RSB, whose size is limited by a global constant $\mathbb{R}_{size}$ denoting the maximal RSB size. A speculative instance Φ_R now consists of the program p , the counter ctr , the configuration σ , the speculation window ω , and the RSB $\mathbb{R}$. As before, a state Σ_R is a stack of speculative instances $\bar{\Phi}_R$.

Spec. States $\Sigma_R ::= \bar{\Phi}_R$

Spec. Instance $\Phi_R ::= \langle p, ctr, \sigma, \mathbb{R}, n \rangle$

As before, in $\Sigma_R \xrightarrow{\tau} \mathcal{L}_R \Sigma_R$ reductions happen on the top of the stack and in this semantics we have $instr_{\mathcal{L}} = \text{call} \cup \text{ret}$ since both **call** and **ret** instructions interact with the RSB and **ret** instructions are the source of speculation.

$$\begin{array}{c}
 \hline
 \boxed{\bar{\Phi}_R \xrightarrow{\tau} \mathcal{L}_R \bar{\Phi}'_R} \\
 \hline
 \text{(R:AM-Spec)} \quad \begin{array}{l} p(\sigma(\text{pc})) = \text{ret} \quad \sigma = \langle m, a \rangle \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle \\ \mathbb{R} = \mathbb{R}' \cdot l \quad j = \min(\omega, n) \quad l \neq m(a(\text{sp})) \\ \sigma'' = \sigma[\text{pc} \mapsto l, \text{sp} \mapsto a(\text{sp}) + 8] \quad \bar{\tau} = \tau \cdot \text{start}_R \text{ctr} \cdot \text{ret } l \end{array} \\
 \hline
 \langle p, ctr, \sigma, \mathbb{R}, n+1 \rangle \xrightarrow{\bar{\tau}} \mathcal{L}_R \langle p, ctr, \sigma', \mathbb{R}', n \rangle \cdot \langle p, ctr+1, \sigma'', \mathbb{R}', j \rangle
 \end{array}
 \quad
 \begin{array}{c}
 \text{(R:AM-Call)} \\
 \hline
 \begin{array}{l} p(\sigma(\text{pc})) = \text{call } f \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle \\ \mathbb{R}' = \mathbb{R} \cdot \langle a(\text{pc}) + 1 \rangle \quad |\mathbb{R}| < \mathbb{R}_{size} \end{array} \\
 \hline
 \langle p, ctr, \sigma, \mathbb{R}, n+1 \rangle \xrightarrow{\tau} \mathcal{L}_R \langle p, ctr, \sigma', \mathbb{R}', n \rangle
 \end{array}$$

During **call** instructions (Rule **R:AM-Call**), the return address is pushed on top of the RSB (if there is space available) and during **ret** instructions, the return address stored on the RSB is used if the entry on top of the RSB is different from the one stored on the stack (Rule **R:AM-Spec**). Then, the rule creates a new speculative instance that uses the return address from the RSB $\mathbb{R}$. Note that speculation only happens when the return address from the RSB differs from the one on the stack (stored in $m(a(\text{sp}))$).

Here, we overview how our semantics behaves with empty and full RSB. Whenever the RSB is empty, executing a **ret** instruction does not cause speculation and we return to the address pointed by **sp**. In contrast, whenever the RSB is full, executing a **call** instruction does not add entries to the RSB, i.e., we model an *acyclic* RSB.⁴

⁴We follow the way AMD processors handle this kind of speculation [53].

The behaviour $Beh_R^\omega(p)$ is the set of all traces generated from an initial state until termination using $\approx_R$.

6.3.2 Oracle Semantics. Unlike before, the oracle cannot decide the outcome of the **ret** instruction, because the CPU always uses the return address stored in the RSB (if there is one) and it does not speculate otherwise [20]. The only thing the oracle decides here is the size of the speculation window ω .

6.3.3 Different Behaviours of Empty and Full RSBs. Modern CPUs use different RSB implementations that differ in the way they handle underflows and overflows, i.e., when the RSB is empty or full [53]. For example, cyclic RSB implementations overwrite old entries when the RSB is full. Alternatively, CPUs can fallback to other predictors (like the indirect branch predictor) to predict return addresses whenever the RSB is empty.

In our model, the RSB is not cyclic and there is no speculation when the RSB is empty (Rule **R:AM-Ret-Empty**).

$$\frac{\text{(R:AM-Ret-Empty)} \quad p(\sigma(\mathbf{pc})) = \mathbf{ret} \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle}{\langle p, ctr, \sigma, \emptyset, n+1 \rangle \approx_R^\tau \langle p, ctr, \sigma', \emptyset, n \rangle}$$

We remark that extending $\approx_R$ to support different RSBs implementations can be done with minimal effort.

6.4 $\approx_{SLS}$: Straightline Speculation

Modern CPUs use the Branch Target Buffer (BTB) to assist with branch prediction. The BTB is indexed by possible jump instructions and predicts the next program counter. However, it also stores information about the type of branch (i.e. no branch, direct branch, indirect branch, or return) encountered at that location. Thus, it can happen that the BTB correctly predicts the location of a branch but mispredicts the type of the branch. For example, a **ret** instruction can be mispredicted to have the type no branch which in turn means that the CPU speculatively executes code right past the **ret** instruction. This kind of speculation is called straight-line speculation (SLS) [5, 8]. Note that a non-branch instruction can also be predicted as a branch which results in ghost jumps. Here, however, we focus only on straight-line speculation.

Example 9 (Straightline Speculation Vulnerability). Consider the example in Listing 8:

Listing 8. Code vulnerable to straightline speculation.

```

1  ret
2  load eax, p
3  load edx, B + eax

```

Assume that **p** is an attacker-controlled value during the execution. After the execution of the **ret** instruction, the BTB of the processor predicts a no branch for the **ret** instruction. Thus, the processor speculatively bypasses the **ret** instruction and executes the following **load** instruction in Line 2. The speculative access of the memory is then leaked into the microarchitectural state by the array access into **B** in Line 3.

This section describes the AM semantics (Section 6.4.1) and the oracle semantics (Section 6.4.2).

6.4.1 Speculative Semantics.

$$\overline{\Phi}_{SLS} \approx_{SLS}^\tau \overline{\Phi}'_{SLS}$$

$$\begin{array}{c}
 \text{(SLS:AM-Spec)} \\
 \frac{p(\sigma(\text{pc})) = \text{ret} \quad \sigma \xrightarrow{\tau} \sigma' \quad j = \min(\omega, n)}{\sigma'' = \sigma[\text{pc} \mapsto \sigma(\text{pc}) + 1] \quad \tau' = \tau \cdot \text{bypass } \sigma(\text{pc}) \cdot \text{start}_{\text{SLS}} \text{ ctr}} \\
 \langle p, \text{ctr}, \sigma, n + 1 \rangle \xrightarrow{\tau'}_{\text{SLS}} \langle p, \text{ctr}, \sigma', n \rangle \cdot \langle p, \text{ctr} + 1, \sigma'', j \rangle
 \end{array}$$

Judgement $\Sigma_{\text{SLS}} \xrightarrow{\tau}_{\text{SLS}} \Sigma'_{\text{SLS}}$ describes how Σ_{SLS} steps to Σ'_{SLS} emitting observation τ . As in $\mathcal{L}_{\text{B}}$, reductions only happen on top of the stack. Here we have $\text{instr}_{\text{B}} = \text{ret}$ because **ret** instructions are the source of speculation.

To model the effect of bypassing a **ret** instruction, Rule **SLS:AM-Spec** bypasses the **ret** instruction by increasing the program counter instead of using the return address to update the program counter and starts a new speculative transaction by pushing a new speculative instance on top of the state. This is in contrast to speculation on **ret** in $\mathcal{L}_{\text{R}}$, which uses the RSB to update the program counter accordingly.

The set $\text{Beh}_{\text{SLS}}^{\omega}(p)$ contains all traces generated from an initial state until termination using the reflexive-transitive closure of $\xrightarrow{\tau}_{\text{SLS}}$.

6.4.2 Oracle Semantics. Instead of bypassing every **ret** instruction, the oracle semantics employs an oracle O that decides if the **ret** instruction should be speculatively bypassed or not. As before, the behaviour $\text{Beh}_{\text{SLS}}^O(p)$ of a program p is the set of all traces starting from an initial state until termination using the reflexive-transitive closure of the oracle semantics.

6.5 $\mathcal{L}_{\text{J}}$: Speculation on (Indirect) Jump Instructions

Modern processors use an indirect branch predictor to predict the outcome of indirect branches. Indirect branches are branches where the outcome is not known until the execution of the program. For example, **jmp r1**, which jumps to the location specified by register **r1**. Instead of waiting for the value of **r1** to be available, the processor predicts the jump target based on the branch target buffer. However, attackers can poison the content of the branch target buffer, which allows the attacker to speculatively divert the control flow of the program [49].

Example 10 (A Program Exploiting Jump Speculation). Consider the example in Listing 9 implementing a small jump table. In Line 2 until Line 7 the program assigns the jump target to **r1** depending on the value of **x**. Next, the indirect jump in Line 8 executes and non-speculative execution continues at either **J1** or **J2**. In both cases, execution terminates by jumping to the end in Line 17. By exploiting jump speculation, an attacker could guide the indirect jump in Line 8 to Line 15 thus leaking the contents of the private variable **p** in Line 16.

This section describes the AM semantics (Section 6.5.1), the oracle semantics (Section 6.5.2).

6.5.1 Speculative Semantics. The attacker can inject any address as the new jump target, making analysis especially tricky. Consider a hypothetical speculation rule capturing the speculative behaviour of indirect jumps (note, this is not the version we use, which is presented below):

$$\begin{array}{c}
 \text{(J:AM-Spec (Non-Final-Version))} \\
 \frac{p(\sigma(\text{pc})) = \text{jmp } e \quad \sigma \xrightarrow{\tau} \sigma' \quad \exists x \subset e. x \in \text{Regs} \quad j = \min(\omega, n)}{S = \{l \mid l \in p \setminus \{\sigma'.\text{pc}\}\} \quad \langle p, \text{ctr}, \sigma, j \rangle \vdash^S \bar{\Phi}_j} \\
 \langle p, \text{ctr}, \sigma, n + 1 \rangle \xrightarrow{\tau}_{\text{J}} \langle p, \text{ctr}, \sigma, n \rangle \cdot \bar{\Phi}_j
 \end{array}$$

Listing 9. Code vulnerable to indirect jump speculation.

```

1  Main:
2      beqz x, L1
3      r1 <- J1
4      jmp End
5  L1:
6      r1 <- J2
7  End:
8      jmp r1
9  J1:
10     x <- 0
11     jmp Fin
12 J2:
13     x <- 1
14     jmp Fin
15     load eax, p
16     load edx, B + eax
17  Fin:

```

Listing 10. Example addition of endbr instructions.

```

1  Main:
2      beqz x, L1
3      r1 <- J1
4      jmp End
5  L1:
6      r1 <- J2
7  End:
8      jmp r1
9  J1:
10     endbr
11     x <- 0
12     jmp Fin
13 J2:
14     endbr
15     x <- 1
16     jmp Fin
17     load eax, p
18     load edx, B + p
19  Fin:

```

Fig. 8. Example program vulnerable to jump speculation (Listing 9) and a version restricting the jump speculation using **endbr** instructions (Listing 10).

$$\begin{array}{c}
\text{(Helper-Base)} \\
\hline
\langle p, ctr, \sigma, n \rangle \vdash^{\emptyset} \emptyset
\end{array}
\quad
\begin{array}{c}
\text{(Helper-Ind)} \\
\hline
\frac{l \in S \quad \langle p, ctr + 1, \sigma, n \rangle \vdash^{S \setminus l} \bar{\Phi}_j}{\langle p, ctr, \sigma, n \rangle \vdash^S \langle p, ctr + 1, \sigma[\text{pc} \mapsto l], n \rangle \cdot \bar{\Phi}_j}
\end{array}$$

Rule **J:AM-Spec (Non-Final-Version)** allows the execution to speculatively jump to any location and we would need to create a speculative instance for **each** of these locations. Here, Rule **Helper-Ind** and Rule **Helper-Base** are helpers creating the speculative transactions and we use $x \subset e$ to denote the subexpressions of e , thereby ensuring that only indirect jumps are used for speculation.

However, creating a speculative instance for each location in the program makes the analysis of the program infeasible because of the amount of states that need to be explored. Furthermore, against such a model, almost any “interesting” program would be considered insecure since speculation is, in practice, unrestricted.

One of the techniques employed by modern CPUs to restrict the scope of indirect-jump speculation is (hardware) control-flow-integrity (CFI) [1]. CFI leverages tags to mark valid jump targets in the code and ensures that the code can only jump to these tagged targets. There are different ways tagging and enforcement mechanisms can be implemented (see Burow et al. [17] for a comprehensive survey). We focus on the hardware implementations of CFI of Intel (Intel-CET [68]) and ARM (branch target identification [9]) because they are available on new hardware and constrain *speculative* control flow: when enabled, a predicted indirect-branch target that does not begin with an **endbr** instruction limits or halts transient execution [68].

Both implementations add a new instruction **endbr** (in the case of ARM this instruction is called **bti**) that marks the legal targets of indirect jumps in a program and enables coarse-grained forward edge control flow integrity. Enforcement is done by a state machine in hardware that ensures only valid jumps are allowed. Furthermore, all available standard

compilers (i.e. GCC, CLANG) already emit these **endbr** instructions and these instructions are handled as no-ops if the current hardware does not support CFI, making it backwards-compatible. Thus, we similarly add a **endbr** instruction to μASM , which marks legal targets for indirect jumps in the program and otherwise behaves like a **skip** instruction (Rule **Endbr**). This restricts the scope of indirect-jump speculation and allows for a tractable analysis.

$$\text{(Instructions)} \ i := \dots \mid \text{endbr} \quad \frac{\text{(Endbr)} \quad p(a(\text{pc})) = \text{endbr}}{\langle p, \langle m, a \rangle \rangle \rightarrow \langle p, \langle m, a[\text{pc} \mapsto a(\text{pc}) + 1] \rangle \rangle}$$

First, we keep track of all possible allowed jump targets in the program. We collect the instruction number for each **endbr** instruction into a labelset L defined as follows: $L_p = \{i \mid \forall i \in \mathbb{N}. p(i) = \text{endbr}\}$. We will omit p in L_p if it is clear from context. One possible way to modify the program in Listing 9 with **endbr** instruction is the program in Listing 10. Then, the labelset is defined as: $L = \{\text{Line 10}, \text{Line 14}\}$. Next, we modify the rule for non-speculative jumps by requiring that the jump target is part of the allowed jump targets L (Rule **Jump-Mod**):

$$\frac{\text{(Jump-Mod)} \quad \begin{array}{l} p(a(\text{pc})) = \text{jmp } e \quad \ell = \llbracket e \rrbracket(a) \\ \exists x \subset e. x \in \text{Regs} \quad l \in L \end{array}}{\langle p, \langle m, a \rangle \rangle \xrightarrow{\text{pc } \ell} \langle p, \langle m, a[\text{pc} \mapsto \ell] \rangle \rangle}$$

Note that for all other semantics, we can define the set of possible jump targets to be all instruction labels in the program to recover the original rule for non-speculative jumps.

Next, we define the new speculative semantics:

$$\Phi_j \xrightarrow{\tau} \mathcal{J}_j \Phi_j'$$

$$\frac{\text{(J:AM-Spec)} \quad \begin{array}{l} p(\sigma(\text{pc})) = \text{jmp } x \quad \sigma \xrightarrow{\tau} \sigma' \quad j = \min(\omega, n) \\ \exists x \subset e. x \in \text{Regs} \quad \tau' = \tau \cdot \text{pc } \sigma'(\text{pc}) \cdot \text{start}_j \text{ ctr} \\ \langle p, \text{ctr}, \sigma, j \rangle \vdash^{L \setminus \{\sigma'.\text{pc}\}} \bar{\Phi}_j \end{array}}{\langle p, \text{ctr}, \sigma, n+1 \rangle \xrightarrow{\tau'} \mathcal{J}_j \langle p, \text{ctr}, \sigma', n \rangle \cdot \bar{\Phi}_j}$$

We replace the unmitigated Rule **J:AM-Spec (Non-Final-Version)** with the constrained Rule **J:AM-Spec**. The main difference lies in the definition of the allowed jump target set. The set S in the unmitigated rule allows jumps to arbitrary addresses in the program. In the new rule this is restricted to L , allowing jumps only to targets designated by **endbr** instructions. When an indirect branch is encountered, Rule **J:AM-Spec** creates a speculative instance for each target in L , excluding the correct non-speculative target. Thus, the rule creates exactly $|L| - 1$ speculative instances. Consider the example in Listing 10 where $L = \{\text{Line 10}, \text{Line 14}\}$. Because Rule **J:AM-Spec** can only speculatively jump to targets in L , execution cannot speculatively reach Line 17, successfully preventing the leakage of the private variable **p**. Our $\mathcal{J}_j$ semantics models speculation over indirect jumps; we do not model BTB-driven speculation triggered by call instructions.

Crucially, because this rule creates a set of instances rather than a single misprediction (diverging from the standard single-target template in Section 5.2), it introduces an asymmetry between start and rollback observations: A single **start_j ctr** observation is now associated with multiple rollback observations (exactly $|L| - 1$ many).⁵ However, because the speculative instances are pushed onto the execution stack sequentially, their execution is strictly nested. The resulting trace exhibits the following structure, where inner transactions are rolled back before the outer transaction

⁵We could have decided to add matching **start_j ctr** observations to Rule **J:AM-Spec** to balance the **start_j i** and the **rlb_j i**.

concludes:

$$\text{start}_j \text{ctr} \cdots \text{rlb}_j \text{ctr} + |L| - 2 \cdots \text{rlb}_j \text{ctr}$$

This nesting guarantees the correctness of the non-speculative projection function $\upharpoonright_{ns}$. Since $\upharpoonright_{ns}$ erases the trace segment between a start event and its matching rollback (here, the final $\text{rlb}_j \text{ctr}$), all intermediate speculative events—including the nested rollbacks—are correctly removed from the trace.

6.5.2 Oracle Semantics. Instead of predicting all of the possible indirect `jmp` targets, the oracle chooses one of the possible jump targets in L and continues only with this one choice. As before, the behaviour $\text{Beh}_j^O(p)$ of a program p is the set of all traces starting from an initial state until termination using the reflexive-transitive closure of the oracle semantics.

6.6 Safety of Semantics

We conclude this section by proving the core properties satisfied by all semantics. Theorem 1 characterizes that all the speculative semantics presented in the paper are well-formed speculative semantics, i.e., (1) their always-mispredict version over-approximates the corresponding oracle version, (2) they are consistent under non-speculative trace projection, and (3) their symbolic version is consistent with the corresponding non-symbolic version.

THEOREM 1 (WELL-FORMED SPECULATIVE SEMANTICS). *The following statements hold:*

- $(\mathcal{L}_B \text{ is WFSS}) \vdash \mathcal{L}_B \text{ WFSS}$
- $(\mathcal{L}_S \text{ is WFSS}) \vdash \mathcal{L}_S \text{ WFSS}$
- $(\mathcal{L}_R \text{ is WFSS}) \vdash \mathcal{L}_R \text{ WFSS}$
- $(\mathcal{L}_{SLS} \text{ is WFSS}) \vdash \mathcal{L}_{SLS} \text{ WFSS}$
- $(\mathcal{L}_j \text{ is WFSS}) \vdash \mathcal{L}_j \text{ WFSS}$

7 A FRAMEWORK FOR COMPOSING SPECULATIVE SEMANTICS

Although the instances of speculative semantics presented in Section 6 each capture one specific aspect of speculation (e.g., $\mathcal{L}_B$ captures branch speculation), they do not capture the vulnerability in Listing 3 (restated here for clarity) as the traces of Example 11 show.

Example 11 (SNI for Listing 11). The traces generated are:

$$\begin{aligned} \bar{\tau}_B^1 &= \bar{\tau}_B^2 := \text{store } p \cdot \text{store } p \cdot \text{start}_B 0 \cdot \text{load } p \cdot \text{load } A + \text{public} \cdot \text{rlb}_B 0 \cdot \text{pc } 9 \\ \bar{\tau}_S^1 &= \bar{\tau}_S^2 := \dots \cdot \text{store } p \cdot \text{start}_S 1 \cdot \text{bypass } 1 \cdot \text{pc } \perp \cdot \text{rlb}_S 1 \cdot \text{pc } \perp \end{aligned}$$

The program in Listing 11 seems secure since there is no secret value leaked in the speculative transaction; thus the program satisfies SNI for $\mathcal{L}_B$ and $\mathcal{L}_S$ in isolation. However, this program speculatively leaks when considering speculation over `beqz` and `store` instructions, but we would need a combined semantics to detect this vulnerability.

The vulnerability only appears when the branch predictor (Section 6.1) and the memory disambiguator (Section 6.2) are used *together*. Intuitively, we know that CPUs use all the speculation mechanisms described here (and many others as well) at the same time. Thus, we should not only focus on these individual speculation mechanisms in *isolation* but we need to look at their *combinations* as well. That is, we need a way to compose the different semantics into new semantics that can reason about these “combined” leaks.

Listing 11. $\mathcal{L}_{B+S}$ example.

```

1  x = 0;
2  p = &secret;
3  p = &public;
4  if (x != 0)
5      temp &= A[*p];

```

This section presents a novel, general framework for composing two speculative semantics x and y , each one capturing the effects of a single speculation mechanism, to allow for speculation from both mechanisms x and y . The semantics x and y are also called the *source* semantics of the composition. Next, we first introduce the new composed semantics, which consists of an always-mispredict semantics, an oracle semantics, and a symbolic semantics (Section 7.1). Then, we present the notion of *well-formed* composition which we use to study the properties of composed

semantics (Section 7.2).

New Notation. The states Σ_{xy} , instances Φ_{xy} , and the trace model Obs_{xy} are defined as the union of the source parts. Furthermore, we define a projection function $\downarrow_{xy}$ and two projections $\downarrow_{xy}^x$ and $\downarrow_{xy}^y$ that return the first and second projection of the pair from $\downarrow_{xy}$. These functions are lifted to states by applying them pointwise:

$$\begin{aligned}
 Obs_{xy} &:= Obs_x \cup Obs_y & \Phi_{xy} &:= \Phi_x \cup \Phi_y & \Sigma_{xy} &:= \Sigma_x \cup \Sigma_y \\
 \downarrow_{xy} &: \Phi_{xy} \mapsto (\Phi_x, \Phi_y) & \downarrow_{xy}^x &: \Phi_{xy} \mapsto \Phi_x & \downarrow_{xy}^y &: \Phi_{xy} \mapsto \Phi_y
 \end{aligned}$$

For example, the Φ_{S+R} state resulting from the union of Φ_S and Φ_R states (from Section 6.2.1 and Section 6.3.1 respectively) is $\langle p, ctr, \sigma, \mathbb{R}, n \rangle$, as it contains all common elements (the program p , the counter ctr , the state σ , and the speculation count n) plus the return stack buffer $\mathbb{R}$ from Φ_R only. Taking the $\cdot \downarrow_{S+R}^S$ of a Φ_{S+R} state returns the Φ_S subpart, i.e., all but the return stack buffer.

We overload $\downarrow_{xy}^x$ and $\downarrow_{xy}^y$ to also work on traces $\bar{\tau}$. The projection $\tau \downarrow_{xy}^x$ deletes all speculative transactions (marked by start_y *id* and rlb_y *id*) not generated by the source semantics x . The definition of $\downarrow_{xy}^y$ is similar by replacing x with y :

$$\begin{aligned}
 \varepsilon \downarrow_{xy}^x &= \varepsilon & (\tau \cdot \bar{\tau}) \downarrow_{xy}^x &= \tau \cdot (\bar{\tau}) \downarrow_{xy}^x \\
 (\text{start}_y \text{ id} \cdots \text{rlb}_y \text{ id} \cdot \bar{\tau}) \downarrow_{xy}^x &= \bar{\tau} \downarrow_{xy}^x
 \end{aligned}$$

We indicate source semantics for x and y as $\mathcal{L}_x$ and $\mathcal{L}_y$ respectively and use $\mathcal{L}_{xy}$ to indicate the composed semantics.

7.1 Combined Speculative Semantics

The combined semantics delegates back to the source semantics of x and y to model the effects of both speculation mechanisms (modeled by x and y). This is captured in the two core rules below:

$$\begin{array}{c}
 \text{(AM-x-Step)} \\
 \frac{\Phi_{xy} \downarrow_{xy}^x \xrightarrow{\tau \downarrow_x^{Z_{xy} \downarrow_{xy}^x}} \bar{\Phi}'_{xy} \downarrow_{xy}^x}{\Phi_{xy} \xrightarrow{\tau \downarrow_{xy}^{Z_{xy}}} \bar{\Phi}'_{xy}}
 \end{array}
 \quad
 \begin{array}{c}
 \text{(AM-y-Step)} \\
 \frac{\Phi_{xy} \downarrow_{xy}^y \xrightarrow{\tau \downarrow_y^{Z_{xy} \downarrow_{xy}^y}} \bar{\Phi}'_{xy} \downarrow_{xy}^y}{\Phi_{xy} \xrightarrow{\tau \downarrow_{xy}^{Z_{xy}}} \bar{\Phi}'_{xy}}
 \end{array}$$

The combined semantics does a step by either delegating back to the x source semantics (Rule **AM-x-Step**) or to the y one (Rule **AM-y-Step**).⁶ The rules rely on metaparameter Z_{xy} , which is a pair of two metaparameters $Z_{xy} := (Z_x, Z_y)$ — one for x and one for y . We overload the projections $\downarrow_{xy}^x$ and $\downarrow_{xy}^y$ to extract the corresponding metaparameter from Z_{xy} , e.g., $Z_{xy} \downarrow_{xy}^x = Z_x$.

⁶To simplify notation, we omit that the $\Phi_x \setminus \Phi_y$ parts of state Φ_{xy} in x-step (similar $\Phi_y \setminus \Phi_x$ in y-step) do not change between Φ_{xy} and $\bar{\Phi}'_{xy}$.

The role of Z is central to making the composed semantics work as expected. It restricts how the combined semantics delegates execution to the components to ensure that the correct rule is applied.

With $Z = (\emptyset, \emptyset)$, consider the execution of the **beqz** instruction in Line 4 in Listing 11. The combined semantics $\mathcal{L}_{B+S}$ can use Rule **AM-x-Step** to delegate back to $\mathcal{L}_B$ for **beqz** instructions, creating a new speculative transaction (Rule **B:AM-Spec**). However, $\mathcal{L}_{B+S}$ can also use Rule **AM-y-Step**, because **beqz** instructions are also handled by $\mathcal{L}_S$. Unfortunately, this does not start speculation, which happens only on **store** instructions (Rule **S:AM-Spec**).

Intuitively, $\mathcal{L}_{B+S}$ should delegate back to $\mathcal{L}_B$, so Rule **AM-y-Step** should not be applicable. This can be obtained by instantiating $Z_{B+S} = (\text{store}, \text{beqz})$, so that its projections are $Z_B = \text{store}$ and $Z_S = \text{beqz}$. Now, $\mathcal{L}_{B+S}$ can only apply Rule **AM-x-Step** on the **beqz** of Line 4, because Z_S ensures that $\mathcal{L}_S$ cannot execute **beqz** instructions, as depicted in the full rule for $\mathcal{L}_S^{\text{beqz}}$ below (where we indicate the instructions derived from $Z_S = \text{beqz}$ in blue):

$$\frac{\begin{array}{c} \text{(S:AM-NoSpec)} \\ p(\sigma(\text{pc})) \notin \text{store} \cup \text{beqz} \quad \langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle \end{array}}{\langle p, \text{ctr}, \sigma, n+1 \rangle \xrightarrow{\tau}^{\text{beqz}}_{\mathcal{L}_S} \langle p, \text{ctr}, \sigma', n \rangle}$$

Having clarified the intuition behind the semantics, we can define the behaviour $\text{Beh}_{xy}^\omega()$ as the set of all traces generated from initial states until termination using $\mathcal{L}_{xy}$.

7.1.1 Oracle Combination. Instead of using one oracle, the combination uses a pair of oracles, one from each source semantics. When delegating back to either source, the correct oracle of the source is handed over to the source semantics.

7.1.2 Symbolic Combination. Instead of using the AM semantics for delegation, the combined symbolic semantics $\mathcal{L}_{xy}^S$ uses the symbolic source semantics for delegation. Furthermore, the new notation (union, projections) is lifted to the symbolic combination to create the symbolic states Σ_{xy}^S . The behaviour $\text{Beh}_{xy}^S(p)$ of program p is the set of all traces generated using the symbolic semantics.

7.2 Properties of Composition

We now illustrate the benefits of our composition framework. For this, we first introduce a notion of well-formed composition (Section 7.2.1), which intuitively tells when a combined semantics “makes sense”. Then, we show that for well-formed compositions, if the source semantics are WFSS, so is the combined semantics (Section 7.2.2). Since we proved this property for *any* well-formed composition in our framework, all (well-formed) compositions we present in Section 8 are WFSS *for free*. This proof reuse and extensibility is our framework’s key advantage over having ad-hoc semantics combining multiple speculation mechanisms, which requires one to manually prove the WFSS results we instead obtain for free.

7.2.1 Well-formed Compositions. The well-formedness conditions for the composition in Definition 4 ensures that the delegation between the source semantics is done properly. Note that these conditions are the *minimal* set of assumptions that let us derive WFSS of the combined semantics for free:

Definition 4 (Well-Formed Composition). A composition $\mathcal{L}_{xy}$ of two speculative semantics $\mathcal{L}_x$ and $\mathcal{L}_y$ is well-formed, written $\vdash \mathcal{L}_{xy} : \text{WFC}$, if:

- (1) (Confluence) Whenever $\Sigma_{xy} \xrightarrow{\tau} \mathcal{L}_{xy} \Sigma'_{xy}$ and $\Sigma_{xy} \xrightarrow{\tau} \mathcal{L}_{xy} \Sigma''_{xy}$, then $\Sigma'_{xy} = \Sigma''_{xy}$.
- (2) (Projection preservation) For all p , $\text{Beh}_x^\omega(p) = \text{Beh}_{xy}^\omega((p)) \upharpoonright_{xy}^x$ and $\text{Beh}_y^\omega(p) = \text{Beh}_{xy}^\omega((p)) \upharpoonright_{xy}^y$.
- (3) (Relation preservation) If $\Sigma_{xy} \approx_{xy} X_{xy}$ and $\Sigma_{xy} \xrightarrow{\tau} \mathcal{L}_{xy}^* \Sigma'_{xy}$ then $X_{xy} \xrightarrow{\tau}^{O_{xy}*} X'_{xy}$ and $\Sigma'_{xy} \approx_{xy} X'_{xy}$.

(4) (Symbolic preservation) If $\Sigma_{xy}^S \xrightarrow{\tau_S} \mathcal{L}_{xy}^S \Sigma_{xy}'$ and $\mu \Sigma_{xy}^S = \Sigma_{xy}$, then there is Σ_{xy}' s.t. $\Sigma_{xy} \xrightarrow{\mu \tau_S} \mathcal{L}_{xy} \Sigma_{xy}'$ and $\mu \Sigma_{xy}' = \Sigma_{xy}'$.

Next, we explain the well-formedness conditions:

- Confluence (point 1) ensures that the non-determinism of the combined semantics (that non-deterministically delegates back to its sources) is not harmful. Consider the assignment in Line 1 in Listing 3. $\mathcal{L}_{B+S}$ can delegate to either $\mathcal{L}_B$ or $\mathcal{L}_S$ to reduce the assignment. If the combined semantics is *confluent*, then it does not matter which source rule executes the assignment in Line 1 in Listing 3, the semantics reaches the same state either way.

- Projection preservation (point 2) ensures that the combined semantics is not hiding or forgetting traces of its sources. Any observable emitted by a source semantics must be propagated to the combined one, this is also the reason why Obs_{xy} is defined as the union of the source Obs .

- To explain relation preservation (point 3), we need to mention a technical detail: the state relation (denoted $\approx_{xy}$ and defined in our technical report) between the AM states (Σ_{xy}) and the Oracle ones (X_{xy}). Intuitively, two states are related if they are the same or if one is waiting on a speculation of the other to end. Then, point (3) ensures that whenever we start from related states ($\Sigma_{xy} \approx_{xy} X_{xy}$) and we do one or more steps of the AM composed semantics ($\Sigma_{xy} \xrightarrow{\tau} \mathcal{L}_{xy}^* \Sigma_{xy}'$), then we can *always* find a related state ($\Sigma_{xy}' \approx_{xy} X_{xy}'$) that is reachable by performing one or more steps of the composed oracle semantics ($X_{xy} \xrightarrow{\tau_{xy}^{O_{xy}}} X_{xy}'$). This fact is used when proving that SNI of a program under the composed AM semantics implies SNI under the composed oracle semantics (point 1 of Definition 3). Thus, it is not important for the AM and the Oracle semantics to produce the same traces, just that the two AM traces and the two Oracle traces are pairwise equivalent – which follows from the state relation.

- Finally, symbolic preservation (point 4) ensures that any step of the always-mispredict composed semantics corresponds to the concretization of a step of the symbolic composed semantics (and vice versa⁷). Note that proving symbolic preservation is almost trivial whenever both source semantics enjoy the same property (like our semantics $\mathcal{L}_B, \mathcal{L}_J, \mathcal{L}_S, \mathcal{L}_R$ and $\mathcal{L}_{SLS}$).

7.2.2 WFSS Preservation. The key result of this section is that well-formed compositions whose sources are well-formed speculative semantics (WFSS) are also WFSS (Theorem 2). Note that our proof of Theorem 2, available in the companion technical report [30], holds for *any* well-formed composition in our framework and, therefore, it applies *for free* to all the compositions in Section 8.

THEOREM 2 ($\mathcal{L}_{xy}$ IS WFSS). *If $\vdash \mathcal{L}_x$ WFSS and $\vdash \mathcal{L}_y$ WFSS and $\vdash \mathcal{L}_{xy} : WFC$, then $\vdash \mathcal{L}_{xy}$ WFSS.*

As a corollary of Theorem 2, we obtain that the security of well-formed compositions is related to the security of their components (Theorem 3). In particular, whenever a program is insecure w.r.t. one of the components, then it is insecure w.r.t. the composed semantics. Dually, if a program is secure w.r.t. the composed semantics, then it is secure w.r.t. the single components. Note, however, that there are programs that are secure for the single components but insecure w.r.t. the composed semantics like Listing 3.

THEOREM 3 (COMBINED SNI PRESERVATION). *Whenever $\vdash \mathcal{L}_{xy} : WFC$ holds:*

- If $p \not\vdash_x$ SNI or $p \not\vdash_y$ SNI, then $p \not\vdash_{xy}$ SNI.
- If $p \vdash_{xy}$ SNI, then $p \vdash_x$ SNI and $p \vdash_y$ SNI.

⁷For space reasons, Definition 4 only reports one direction (with a simplified notation).

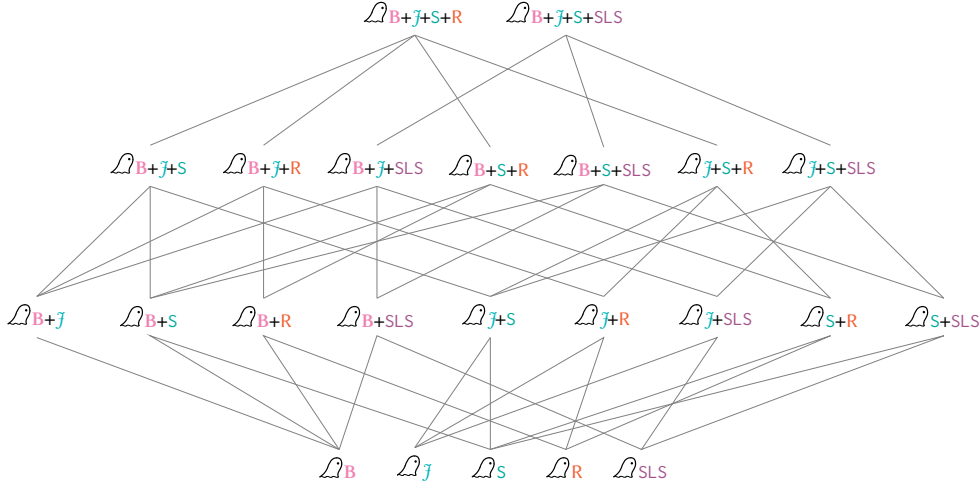


Fig. 9. Partial order of the different combinations that we instantiate using our framework

These results have an immediate practical impact on SPECTECTOR: (1) the analysis of SPECTECTOR relies on the (symbolic) speculative semantics, (2) the source semantics $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$ are *WFSS*, (3) well-formed compositions are also *WFSS*, and (4) the composition of $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$ are well-formed. So, SPECTECTOR equipped with any combination of the $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$ produces sound results, i.e., whenever the tool proves that a program is leak-free then the program satisfies SNI. In the next section, we describe all the compositions and prove they are well-formed (this implies that they are *WFSS* thanks to Theorem 2).

8 INSTANTIATING OUR FRAMEWORK

We instantiate our combination framework with all possible combinations of $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$, yielding 18 combinations depicted in Figure 9. Combinations of $\mathcal{L}_{SLS}$ and $\mathcal{L}_R$ are impossible because they speculate on the same instructions (`ret`). Thus, we cannot set the metaparameter Z in a sensible way (more details in Section 11). Since we cannot describe all of the combinations in detail, we focus on two representative combinations: $\mathcal{L}_{S+R}$ (Section 8.1) and $\mathcal{L}_{B+J+S+R}$ (Section 8.2); the other combinations can be instantiated in a similar way.

For each of these, we overview the combined AM semantics using examples and we prove that the combined semantics is well-formed, i.e., it satisfies Definition 4.

Full details and well-formedness proofs of all other combinations are available in the companion technical report [30].

8.1 $\mathcal{L}_{S+R}$ Composition

To combine semantics using our framework, we need to define the states, observations, and metaparameter Z_{S+R} for the composed semantics $\mathcal{L}_{S+R}$. The combined state Σ_{S+R} is the union of the states Σ_S and Σ_R ; thus it contains the RSB $\mathbb{R}$ as well.

$$\text{Spec. States } \Sigma_{S+R} ::= \overline{\Phi}_{S+R} \quad \text{Spec. Instance } \Phi_{S+R} ::= \langle p, ctr, \sigma, \mathbb{R}, n \rangle$$

Listing 12. $\mathcal{L}_{S+R}$ example.

1613
1614
1615
1616
1617
1618
1619
1620
1621
1622
1623
1624
1625
1626
1627

```

1  Manip_Stack:
2      sp ← sp + 8
3      ret
4  Speculate:
5      call Manip_Stack
6      store secret, p
7      store pub, p
8      load eax, p
9      load edi, eax
10     ret
11 Main:
12     call Speculate
13     skip

```

The union Obs_{S+R} of the trace models Obs_S and Obs_R is defined as:

$$Obs_{S+R} ::= \text{start}_S n \mid \text{start}_R n \mid \text{rlb}_S n \mid \text{rlb}_R n \mid \text{bypass } n \mid \dots$$

To define the metaparameter Z_{S+R} , we need to identify the instructions that are related with speculative execution for each component semantics. For $\mathcal{L}_S$, the only instruction associated with speculative execution is **store**, since the semantics can only speculatively bypass stores. For $\mathcal{L}_R$, even though the semantics speculates only over **ret** instructions, **call** instructions also affect speculative execution since $\mathcal{L}_R$ pushes return addresses onto the RSB $\mathbb{R}$ when executing calls. Therefore, we set the metaparameter Z_{S+R} to $(\text{call} \cup \text{ret}, \text{store})$. This ensures that in $\mathcal{L}_{S+R}$, **store** instructions are only executed by delegating back to $\mathcal{L}_S^{\text{call} \cup \text{ret}}$ whereas **call** and **ret** instructions are only executed by delegating back to $\mathcal{L}_R^{\text{store}}$.

Theorem 4 states the combination of $\mathcal{L}_S$ and $\mathcal{L}_R$ described above is well-formed. Given that $\mathcal{L}_S$ and $\mathcal{L}_R$ are WFSS (Theorem 1), we can derive “for free” that $\mathcal{L}_{S+R}$ is WFSS (Theorem 2).

THEOREM 4 ($\mathcal{L}_{S+R}$ IS WELL-FORMED). $\vdash \mathcal{L}_{S+R} : WFC$

Listing 12 presents a program that contains a leak that can be detected only by $\mathcal{L}_{S+R}$ but not by its components $\mathcal{L}_S$ and $\mathcal{L}_R$. Execution starts on Line 12 by calling the function *Speculate* and it continues at Line 5. Next, the function *Manip_Stack* is called and the stack pointer *sp* is incremented (Line 2). This modifies the return address of the function *Manip_Stack* to now point to Line 13 (the return address of the **call** to *Speculate*). Under $\mathcal{L}_R$, mispredicting the return address of *Manip_Stack* using the RSB leads to continuing the execution at Line 6. However, the **store** instructions in Line 7 overwrites the *secret* value stored in Line 6. Then, the **load** instructions in Line 8 and Line 9 emit only public values. As a result, no *secret* is leaked and speculation ends. Similarly, under $\mathcal{L}_S$, speculation over store bypasses has no effect in Listing 12 because the **store** instruction in Line 6 is never reached and function *Manip_Stack* returns to Line 13. Therefore, the leak is missed under $\mathcal{L}_S$ and $\mathcal{L}_R$, i.e., Listing 12 $\vdash_S$ SNI and Listing 12 $\vdash_R$ SNI.

However, under the combined semantics $\mathcal{L}_{S+R}$, the **store** instruction on Line 7 is now speculatively bypassed and when returning from function *Manip_Stack* the execution speculatively continues from Line 8. Now, the **load** instructions are executed and the *secret* is leaked, as shown in the traces below. Since *secret* is a high value, there are low-equivalent configurations σ^1, σ^2 that differ in the value of *secret*. Thus, there are two traces that differs in the observation **load secret** (highlighted in gray). Hence, the program is not secure under the combined semantics, i.e.,

Listing 12 $\mathcal{K}_{S+R}$ SNI.

```
 $\tau_{S+R}^2 \neq \tau_{S+R}^1 \stackrel{\text{def}}{=} \text{call Speculate} \dots \text{start}_R 0 \dots \text{start}_S 1 \dots \text{r1bs} 1 \dots \text{start}_S 2 \cdot \text{bypass} 7 \cdot \text{load } p \cdot \text{load secret} \dots$ 
```

The relation between the source semantics and their composition is visualised in Figure 10, which shows the insecure programs (with respect to SNI) detected under the different semantics. The combined semantics encompasses all vulnerable programs of $\mathcal{L}_S$ and $\mathcal{L}_R$ and additional programs like Listing 12. These additional programs are the reason why the semantics $\mathcal{L}_{S+R}$ is “stronger than the sum of its parts” $\mathcal{L}_S$ and $\mathcal{L}_R$.

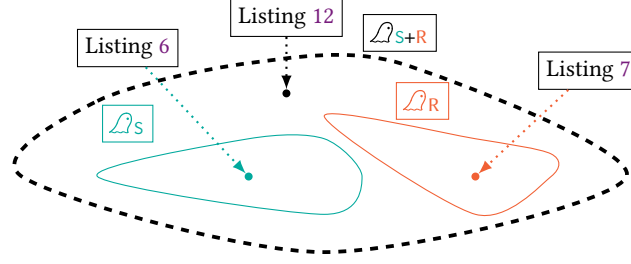


Fig. 10. Relating $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{S+R}$ w.r.t. SNI.

8.2 $\mathcal{L}_{B+J+S+R}$ Composition

We conclude this section by combining four semantics $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$ and $\mathcal{L}_R$ (as mentioned, a combination with five semantics is not possible because $\mathcal{L}_{SLS}$ and $\mathcal{L}_R$ speculate on the same instruction). Our framework (Section 7) allows us to only combine a pair of source semantics into a combined one. For simplicity, we present $\mathcal{L}_{B+J+S+R}$ as a direct combination of the four source semantics (technically, we obtain $\mathcal{L}_{B+J+S+R}$ by combining $\mathcal{L}_{B+J}$ with $\mathcal{L}_{S+R}$).

The metaparameter $Z_{B+J+S+R}$ (which we represent as a quadruple of values) is :

$$Z_B \stackrel{\text{def}}{=} (\text{call} \cup \text{ret} \cup \text{store} \cup \text{jmp})$$

$$Z_J \stackrel{\text{def}}{=} (\text{call} \cup \text{ret} \cup \text{store} \cup \text{beqz})$$

$$Z_S \stackrel{\text{def}}{=} (\text{call} \cup \text{ret} \cup \text{beqz} \cup \text{jmp})$$

$$Z_R \stackrel{\text{def}}{=} (\text{beqz} \cup \text{store} \cup \text{jmp})$$

$$Z_{B+J+S+R} \stackrel{\text{def}}{=} (Z_B, Z_J, Z_S, Z_R)$$

As a result, the combined semantics $\mathcal{L}_{B+J+S+R}$ can only delegate to the corresponding speculative semantics for the appropriate speculation sources.

As stated in Theorem 5, $\mathcal{L}_{B+J+S+R}$ is well-formed and we derive that $\mathcal{L}_{B+J+S+R}$ is WFSS through Theorem 2.

THEOREM 5 ($\mathcal{L}_{B+J+S+R}$ IS WELL-FORMED). $\vdash \mathcal{L}_{B+J+S+R} : \text{WFC}$

Figure 11 depicts a leaky program that can be detected only under $\mathcal{L}_{B+J+S+R}$, since the program satisfies SNI under $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$ and $\mathcal{L}_R$. Under $\mathcal{L}_{B+J+S+R}$, the `store` instruction in Line 18 is bypassed. Therefore, when returning from the `Manip_Stack` function, the program mispredicts the return address and speculatively returns to Line 6. Here, the indirect jump in Line 7 mispredicts and execution continues speculatively in Line 9. Finally, the `beqz` instruction in Line 11 is mispredicted and the `load` instructions are executed, which leaks the secret value.

<pre> 1717 1 Manip_Stack: 1718 2 sp <- sp + 8 1719 3 ret 1720 4 Speculate: 1721 5 call Manip_Stack 1722 6 r1 <- L2 1723 7 jmp r1 1724 8 L1: 1725 9 endbr 1726 10 x <- 0 1727 11 beqz x, L2 1728 12 load eax, p 1729 13 load edi, eax </pre>	<pre> 14 L2: 15 ret 16 Main: 17 store secret, p 18 store pub, p 19 call Speculate </pre>
--	--

Fig. 11. $\mathcal{D}_{B+J+S+R}$ example.

The resulting traces, differing in the value of *secret* exposed by the `load secret` observation (highlighted in gray), are given below:

$$\begin{aligned}
 \tau_{B+J+S+R}^2 \neq \tau_{B+J+S+R}^1 \stackrel{\text{def}}{=} & \dots \text{start}_S 1 \cdot \text{bypass } 18 \cdot \text{call Speculate} \cdot \text{call Manip_Stack} \cdot \text{start}_R 2 \cdot \text{ret } 6 \cdot \text{start}_J 3 \\
 & \cdot \text{pc } 9 \cdot \text{start}_B 4 \cdot \text{pc } 12 \cdot \text{load } p \cdot \text{load secret} \cdot \text{rlb}_B 4 \cdot \text{rlb}_J 3 \cdot \text{rlb}_R 2 \cdot \text{rlb}_S 1 \dots
 \end{aligned}$$

Thus, the program is not secure, i.e., Figure 11 $\mu_{B+J+S+R}$ SNL.

9 DETECTING SPECULATIVE INFORMATION FLOWS : THE SPECTECTOR ALGORITHM

This section presents SPECTECTOR, a program analysis for detecting speculative leaks or proving their absence. The core idea behind SPECTECTOR is to formally compare a program's behavior under speculative execution against its standard, non-speculative execution. A leak is identified if the speculative execution reveals information that is not revealed by the non-speculative execution. To make this comparison comprehensive, analyzing every possible concrete execution is infeasible. Therefore, Spectector leverages symbolic execution to represent all possible program behaviors concisely as a set of symbolic traces. Our approach requires two key formalisms:

- (1) A symbolic non-speculative semantics to establish the program's intended, baseline behavior (Section 9.1).
- (2) A symbolic architectural model (AM) semantics to model the processor's speculative behavior (Section 9.2).

By analyzing the symbolic traces generated from both semantics with an SMT solver, SPECTECTOR can identify discrepancies that correspond to memory or control-flow leaks. If no such discrepancies are found, the program is proven secure against the modeled speculative behaviors (Section 9.3). Finally, we discuss the implementation of SPECTECTOR within the CIAO logic programming system [41] (Section 9.4).

9.1 Symbolic Non-Speculative Semantics for μASM

To establish a ground truth for our analysis, we first define the program's intended behavior using a symbolic non-speculative semantics, denoted by the relation $\langle p, \sigma^S \rangle \xrightarrow{\tau_S} \langle p, \sigma^{S'} \rangle$.

This semantics lifts the concrete non-speculative semantics $\langle p, \sigma \rangle \xrightarrow{\tau} \langle p, \sigma' \rangle$ to operate over symbolic values instead of concrete ones. We now detail the key extensions.

Symbolic program states consist of a program p and a symbolic configuration $\sigma^S ::= \langle sm, sa \rangle$.

Concrete memories m are replaced with symbolic memories sm , modeled as symbolic memories using the standard theory of arrays [16]. We model memory updates as triples of the form $\mathbf{write}(sm, se, se')$, which update the symbolic memory sm by assigning the symbolic value se' to the symbolic location se . Memory reads $\mathbf{read}(sm, se)$ retrieve the value at symbolic memory location se in sm . $\mathbf{read}(sm, se)$ and $\mathbf{write}(sm, se, se')$ are used in Rule **Load-Symb** and Rule **Store-Symb** respectively, which is the only difference to the concrete non-speculative semantics. Concrete register assignments a are replaced with symbolic register assignments sa , which are functions mapping registers to symbolic expressions.

Symbolic expressions represent computations over symbolic values. A *symbolic expression* se is a concrete value $n \in \mathbf{Vals}$, a symbolic value $s \in \mathbf{SymbVals}$, an if-then-else expression $\mathbf{ite}(se, se', se'')$, or the application of a unary $\ominus$ or a binary operator $\otimes$.

$$se := n \mid s \mid \mathbf{ite}(se, se', se'') \mid \ominus se \mid se \otimes se'$$

Most operational rules of $\rightarrow_S$ are straightforward extensions of the concrete semantics, and we omit them. The primary differences w.r.t. the concrete μASM semantics arise when handling memory operations and control-flow statements with symbolic values, which are formalized in the operational semantics rules below.

$$\begin{array}{c}
 \hline
 \sigma^S \xrightarrow{\tau}_S \sigma^{S'} \\
 \hline
 \begin{array}{l}
 \text{(Load-Symb)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{load } x, e \quad x \neq \mathbf{pc} \quad se = \llbracket e \rrbracket(sa) \quad se' = \mathbf{read}(sm, se)}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{load } se}_S \langle p, sm, sa[\mathbf{pc} \mapsto sa(\mathbf{pc}) + 1, x \mapsto se'] \rangle} \\
 \text{(Store-Symb)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{store } x, e \quad se = \llbracket e \rrbracket(sa) \quad sm' = \mathbf{write}(sm, se, sa(x))}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{store } se}_S \langle p, sm', sa[\mathbf{pc} \mapsto a(\mathbf{pc}) + 1] \rangle} \\
 \text{(Beqz-Conc-Sat)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{beqz } x, \ell \quad sa(x) = 0}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{pc } \ell \cdot \mathbf{symPc}(\top)}_S \langle p, sm, sa[\mathbf{pc} \mapsto \ell] \rangle} \\
 \text{(Beqz-Conc-Unsat)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{beqz } x, \ell \quad sa(x) \neq 0}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{pc } sa(\mathbf{pc}) + 1 \cdot \mathbf{symPc}(\top)}_S \langle p, sm, sa[\mathbf{pc} \mapsto a(\mathbf{pc}) + 1] \rangle} \\
 \text{(Beqz-Symb-Sat)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{beqz } x, \ell \quad sa(x) \notin \mathbf{Vals}}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{pc } \ell \cdot \mathbf{symPc}(sa(x)=0)}_S \langle p, sm, sa[\mathbf{pc} \mapsto \ell] \rangle} \\
 \text{(Beqz-Symb-Unsat)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{beqz } x, \ell \quad sa(x) \notin \mathbf{Vals}}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{pc } sa(\mathbf{pc}) + 1 \cdot \mathbf{symPc}(sa(x) \neq 0)}_S \langle p, sm, sa[\mathbf{pc} \mapsto sa(\mathbf{pc}) + 1] \rangle} \\
 \text{(Jmp-Conc)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{jmp } e \quad \ell = \llbracket e \rrbracket(sa) \quad \ell \in \mathbf{Vals}}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{pc } \ell \cdot \mathbf{symPc}(\top)}_S \langle p, sm, sa[\mathbf{pc} \mapsto \ell] \rangle} \\
 \text{(Jmp-Symb)} \\
 \frac{p(sa(\mathbf{pc})) = \mathbf{jmp } e \quad \llbracket e \rrbracket(sa) \notin \mathbf{Vals} \quad \ell \in \mathbf{Vals}}{\langle p, sm, sa \rangle \xrightarrow{\mathbf{pc } \ell \cdot \mathbf{symPc}(\llbracket e \rrbracket(sa)=\ell)}_S \langle p, sm, sa[\mathbf{pc} \mapsto \ell] \rangle}
 \end{array}
 \end{array}$$

When symbolically executing a program, we may produce observations whose values are symbolic. To account for this, we introduce symbolic observations of the form $\mathbf{load } se$ and $\mathbf{store } se$ for symbolic load and store instructions.

Furthermore, because of the symbolic execution, we need to keep track if certain paths in our program are feasible. We encode this information using $\mathbf{symPc}(se)$ observations, which record the choices made at branches and jumps. We can see these observations in action in Rule **Beqz-Symb-Unsat** and Rule **Beqz-Symb-Sat** for branches. Rule **Jmp-Conc** and Rule **Jmp-Symb** record them analogously for indirect jumps. Because x cannot be evaluated to a value, the symbolic semantics decides if the branch is taken or not and records this decision in the observation $\mathbf{symPc}(sa(x) = 0)$. Without these observations, we could later concretize x to a value unequal to 0 even though the branch was not taken, making the path infeasible in a concrete program execution. To explore all paths of the program, our symbolic tool will negate these

symbolic branching observations to explore the other branches as well. The path condition $pthCnd(\overline{\tau_S}) = \bigwedge_{\text{symPc}(se) \in \tau} se$ of trace τ is the conjunction of all symbolic branching conditions in τ .

The value of a symbolic expression se depends on a *valuation function* $\mu : \text{SymbVals} \rightarrow \text{Vals}$ mapping symbolic values to concrete ones. The evaluation $\mu(se)$ is also standard. We write $\mu \models se$ to denote that symbolic expression se is *satisfiable* for a valuation μ , i.e., $\mu(se) \neq 0$. Every valuation that satisfies a symbolic run's path condition maps the symbolic run to a concrete one. Finally, we write $\models se$ to denote that there exists a valuation μ such that $\mu \models se$.

The *symbolic non-speculative behaviour* $Beh_{NS}^S(p)$ of a program p is the set of all traces generated by all possible initial states for program p .

Theorem 6 connects the symbolic and concrete non-speculative semantics via the concretization function μ :

THEOREM 6 (NS: SYMBOLIC CONSISTENCY). $Beh_{NS}(p) = \mu(Beh_{NS}^S(p))$.

This allows us to use the symbolic non-speculative semantics in our tool SPECTECTOR with confidence that it behaves equivalently to the concrete non-speculative semantics.

9.2 Symbolic Always-Mispredict Semantics

To enable automated verification, SPECTECTOR implements a symbolic always-mispredict semantics. Structurally, this semantics perfectly mirrors the concrete always-mispredict semantics introduced earlier, with the sole distinction that it operates over symbolic expressions and relies on the symbolic non-speculative step relation $\rightarrow_S$ rather than its concrete counterpart. Thus, we omit details related to this semantics and refer to the technical report [30].

We denote by $\mu((\langle p, \sigma^S \rangle) \Downarrow_x^\omega \overline{\tau_S})$ the set $\{(\langle p, \sigma^S \rangle) \Downarrow_x^\omega (\mu(\overline{\tau_S})) \mid \mu \models pthCnd(\overline{\tau_S})\}$ and lift it to $Beh_x^S(p)$. Here $(\langle p, \sigma^S \rangle) \Downarrow_x^\omega \overline{\tau_S}$ denotes the run of the program p with initial configuration σ^S generating $\overline{\tau_S}$ for the symbolic always-mispredict semantics x .

Symbolic Consistency. Fundamentally, a valid symbolic analysis must faithfully capture all possible concrete executions of a program. The symbolic always-mispredict semantic should explore the exact same set of execution paths and generate the same observations as its concrete counterpart, merely operating over symbolic values and constraints. We formalize this correctness requirement in Property 4, stating that the concrete always-mispredict behaviour can be exactly recovered from the symbolic behaviour by applying the concretization function.

Property 4 (Symbolic Consistency AM). $Beh_x^\omega(p) = \mu(Beh_x^S(p))$

Because we have proven that all our speculative semantics and their combinations satisfy the requirements of our framework (i.e., they are well-formed speculative semantics), they inherently satisfy Property 4. This allows us to use all of the symbolic semantics in our tool SPECTECTOR with confidence that they behave equivalently to the concrete semantics.

Example 12 (Symbolic Trace for Example 1). Executing the program from Example 1 under the symbolic speculative semantics $\mathcal{L}_B$ with speculative window 2 yields the following two symbolic traces:

$$\begin{aligned} \tau_1 &:= \text{symPc}(y < \text{size}) \cdot \text{start}_B 0 \cdot \text{pc } 2 \cdot \text{pc } 10 \cdot \text{rlb}_B 0 \cdot \text{pc } 3 \cdot \text{load } A + y \cdot \text{load } B + \text{read}(sm, (A + y)) * 512 \\ \tau_2 &:= \text{symPc}(y \geq \text{size}) \cdot \text{start}_B 0 \cdot \text{pc } 3 \cdot \text{load } A + y \cdot \text{load } B + \text{read}(sm, (A + y)) * 512 \cdot \text{rlb}_B 0 \cdot \text{pc } 2 \cdot \text{pc } 10 \end{aligned}$$

Algorithm 1 SPECTECTOR program analysis**Input:** A program p , a security policy P , a speculative window $\omega \in \mathbb{N}$.**Output:** SECURE if p satisfies speculative non-interference with respect to the policy P and speculative semantics $\mathcal{Q}_x$; INSECURE otherwise

```

1: procedure SPECTECTOR( $p, P, \omega$ )
2:   for each symbolic run  $\bar{\tau} \in \text{Beh}_x^S(p)$  do
3:     if MEMLEAK( $\bar{\tau}, P$ )  $\vee$  CTRLLEAK( $\bar{\tau}, P$ ) then
4:       return INSECURE
5:   return SECURE

6: procedure MEMLEAK( $\bar{\tau}, P$ )
7:    $\psi \leftarrow \text{pthCnd}(\bar{\tau})_{1 \wedge 2} \wedge \text{polEqv}(P) \wedge$ 
    $\text{obsEqv}(\bar{\tau} \upharpoonright_{ns}) \wedge \neg \text{obsEqv}(\bar{\tau} \upharpoonright_{se})$ 
8:   return SATISFIABLE( $\psi$ )

9: procedure CTRLLEAK( $\bar{\tau}, P$ )
10:  for each prefix  $v \cdot \text{symPc}(se)$  of  $\bar{\tau} \upharpoonright_{se}$  do
11:     $\psi \leftarrow \text{pthCnd}(\bar{\tau} \upharpoonright_{ns} \cdot v)_{1 \wedge 2} \wedge \text{polEqv}(P) \wedge$ 
     $\text{obsEqv}(\bar{\tau} \upharpoonright_{ns}) \wedge \neg \text{sameSymbPc}(se)$ 
12:    if SATISFIABLE( $\psi$ ) then
13:      return  $\top$ 
14:  return  $\perp$ 

```

9.3 The Spectector Algorithm

The SPECTECTOR program analysis approach is presented in Algorithm 1. It relies on two procedures: MEMLEAK and CTRLLEAK, to detect leaks resulting from memory and control-flow instructions, respectively. We start by discussing the SPECTECTOR algorithm and next explain the MEMLEAK and CTRLLEAK procedures.

SPECTECTOR takes as input a program p , a policy P specifying the non-sensitive information, and a speculative window ω . The algorithm iterates over all symbolic runs produced by the symbolic always-mispredict speculative semantics (lines 2-4). For each trace $\bar{\tau} \in \text{Beh}_x^S(p)$, the algorithm checks whether $\bar{\tau}$ speculatively leaks information through memory accesses or control-flow instructions. If this is the case, then SPECTECTOR has found a witness of a speculative leak and it reports p as INSECURE. If none of the traces contains speculative leaks, the algorithm terminates returning SECURE (line 5).

Detecting leaks caused by memory accesses. The procedure MEMLEAK takes as input a trace $\bar{\tau}$ and a policy P , and it determines whether $\bar{\tau}$ leaks information through symbolic **load** and **store** observations. The check is expressed as a satisfiability check of a constraint ψ . The construction of ψ is inspired by self-composition [13], which reduces reasoning about *pairs* of program runs to reasoning about single runs by replacing each symbolic variable x with two copies x_1 and x_2 . We lift the subscript notation to symbolic expressions.

The constraint ψ is the conjunction of four formulas:

- $\text{pthCnd}(\tau)_{1 \wedge 2}$ stands for $\text{pthCnd}(\bar{\tau})_1 \wedge \text{pthCnd}(\bar{\tau})_2$, which ensures that both runs follow the path associated with $\bar{\tau}$.
- $\text{polEqv}(P)$ introduces constraints $x_1 = x_2$ for each register $x \in P$ and $\text{read}(sm_1, n) = \text{read}(sm_2, n)$ for each memory location $n \in P$, which ensure that both runs agree on all non-sensitive inputs.

• $obsEqv(\bar{\tau}|_{ns})$ introduces a constraint $se_1 = se_2$ for each **load** se or **store** se in $\tau|_{ns}$, which ensures that the non-speculative observations associated with memory accesses are the same in both runs.

• $\neg obsEqv(\bar{\tau}|_{se})$ ensures that at least one speculative observation associated with memory accesses differs among the two runs.

If ψ is satisfiable, there are two P -indistinguishable configurations that produce the same non-speculative traces (since $pthCnd(\bar{\tau})_{1\wedge 2} \wedge polEqv(P) \wedge obsEqv(\bar{\tau}|_{ns})$ is satisfied) and whose speculative traces differ in a memory access observation (since $\neg obsEqv(\bar{\tau}|_{se})$ is satisfied), i.e. a violation of SNI.

Detecting leaks caused by control-flow instructions. To detect leaks caused by control-flow instructions, CTRLLEAK checks whether there are two traces in τ 's concretization that agree on the outcomes of all non-speculative branch, jump and return instructions, while differing in the outcome of at least one speculatively executed branch, jump or return instruction.

In addition to $pthCnd(\bar{\tau})$, $obsEqv(\bar{\tau})$, and $polEqv(P)$, the procedure relies on the function $sameSymbPc(se)$ that introduces the constraint $se_1 \leftrightarrow se_2$ ensuring that se is satisfied in one concretization iff it is satisfied in the other.

CTRLLEAK checks, for each prefix $v \cdot \text{symPc}(se)$ in $\bar{\tau}$'s speculative projection $\bar{\tau}|_{se}$, the satisfiability of the conjunction of $pthCnd(\bar{\tau}|_{ns} \cdot v)_{1\wedge 2}$, $polEqv(P)$, $obsEqv(\bar{\tau}|_{ns})$, and $\neg sameSymbPc(se)$. Whenever the formula is satisfiable, there are two P -indistinguishable configurations that produce the same non-speculative traces, but whose speculative traces differ on control-flow observations, i.e. a violation of SNI.

Example 13 (Example Application of SPECTECTOR). Consider the trace from Example 12:

$\bar{\tau} := \text{symPc}(y \geq \text{size}) \cdot \text{start}_B 0 \cdot \text{pc } 3 \cdot \text{load } A + y \cdot \text{load } B + \text{read}(sm, (A + y)) * 512 \cdot \text{rlb}_B 0 \cdot \text{pc } 2 \cdot \text{pc } 10$

MEMLEAK detects a leak caused by the observation **load** $B + \text{read}(sm, (A + y)) * 512$. Specifically, it detects that there are distinct symbolic valuations that agree on the non-speculative observations but disagree on the value of **load** $B + \text{read}(sm, (A + y)) * 512$. That is, the observation depends on sensitive information that is not disclosed by $\bar{\tau}$'s non-speculative projection.

Soundness and Completeness. Theorem 7 states that SPECTECTOR deems secure only speculatively non-interferent programs, and all detected leaks are actual violations of SNI.

THEOREM 7 (CORRECTNESS SPECTECTOR). *If SPECTECTOR(p, P, ω) terminates, then SPECTECTOR(p, P, ω) = SECURE iff the program p satisfies speculative non-interference w.r.t. the policy P and all prediction oracles $\mathcal{O}$ with speculative window at most ω .*

The theorem follows from Oracle Overapproximation and Symbolic Consistency of the speculative semantics. Importantly, both conditions are part of our well-formed speculative semantics definition (Definition 3). Since $\mathcal{Q}_B$, $\mathcal{Q}_J$, $\mathcal{Q}_S$, $\mathcal{Q}_R$ and $\mathcal{Q}_{SL}$ are well-formed speculative semantics $\vdash \mathcal{Q}_X$ WFSS and our combinations are well-formed (Definition 4), we have that all the combinations are WFSS as well. Thus, we get Theorem 7 for free for all the combinations.

Theorem 7 characterizes the SPECTECTOR *algorithm*, under the assumptions that concolic exploration is exhaustive, the SMT solver resolves every query, and the procedure terminates. Whenever our *implementation* returns a verdict, it respects these guarantees and never reports an incorrect result. On programs with loops or too many paths, or when the SMT solver cannot resolve a query within its time budget, SPECTECTOR reports a timeout instead. Because the SMT theories we use are decidable, such a timeout reflects resource limits rather than a fundamental obstacle.

9.4 SPECTECTOR Implementation

We implement our approach in our tool SPECTECTOR, which is available at [37]. The tool, which is implemented on top of the CIAO logic programming system [41], consists of three components: a front end that translates x86 assembly programs into μ ASM, a core engine implementing Algorithm 1, and a back end handling SMT queries.

x86 Front End. The front end translates AT&T/GAS and Intel-style assembly files into μ ASM⁸. It currently supports over 120 instructions: data movement instructions (**mov**, etc.), logical, arithmetic, and comparison instructions (**xor**, **add**, **cmp**, etc.), branching and jumping instructions (**jae**, **jmp**, etc.), conditional moves (**cmovae**, etc.), stack manipulation (**push**, **pop**, etc.), and function calls⁹ (**call**, **ret**).

It currently does not support privileged x86 instructions, e.g., for handling model specific registers and virtual memory. Further, it does not support sub-registers (like **eax**, **ah**, and **al**) and unaligned memory accesses, i.e., we assume that only 64-bit words are read/written at each address without overlaps. Finally, the translation currently maps symbolic address names to μ ASM instruction addresses, limiting arithmetic on code addresses.

Core Engine. The core engine implements Algorithm 1. It relies on a concolic approach to implement symbolic execution that performs a depth-first exploration of the symbolic runs. Starting from a concrete initial configuration, the engine executes the program under the symbolic always-mispredict speculative semantics while keeping track of the symbolic configuration and path condition. It discovers new runs by iteratively negating the last (not previously negated) conjunct in the path condition until it finds a new initial configuration, which is then used to re-execute the program concolically. In our current implementation, indirect jumps and returns are *not* included in the path conditions, and thus new symbolic runs and corresponding inputs are only discovered based on negated branch conditions.¹⁰ This process is interleaved with the MEMLEAK and CTRLLEAK checks and iterates until a leak is found or all paths have been explored.

SMT Back End. The Z3 SMT solver [26] acts as a back end for checking satisfiability and finding models of symbolic expressions using the BITVECTOR and ARRAY theories, which are used to model registers and memory. The implementation currently does not rely on incremental solving, since it was less efficient than one-shot solving for the selected theories.

Implementation of the Speculative Semantics and the Combinations in SPECTECTOR. We implemented all our semantics (the symbolic versions of $\mathcal{L}_B$, $\mathcal{L}_S$, $\mathcal{L}_J$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$ plus all 18 compositions from Section 8) in SPECTECTOR. The implementation of compositions closely follows the structure of our framework. As in Section 8, selecting one of the composed semantics in SPECTECTOR sets the metaparameter Z , which is used to delegate back to the correct individual semantics.

10 EVALUATION

This section reports on two case studies in which we apply SPECTECTOR to analyze the security of programs. The goals of the first case study are (1) to determine whether speculative non-interference realistically captures speculative leaks and (2) to assess SPECTECTOR's precision. Therefore, we analyze various examples targeting the different Spectre versions we want to capture with our speculative semantics (Section 10.1).

The goal of our second case study is to investigate if the combined speculative semantics allow us to capture stronger attacks that are not captured by individual semantics. Thus, we create code snippets that are only exploitable using a combination of different Spectre attacks and check if SPECTECTOR using the combined semantics can detect leaks (Section 10.2).

⁸This translation is currently unverified.

⁹We model the so-called "near calls", where the callee is in the same code segment as the caller.

¹⁰We plan to remove this limitation in a future release of our tool.

All code snippets as well as all scripts to reproduce our results are available in our public repository at <https://spectector.github.io/>.

10.1 Case Study: Detecting Leaks w.r.t. Individual Speculative Semantics

Using SPECTECTOR, we analyze a corpus of 41 microbenchmarks containing speculative leaks generated by different speculation mechanisms in isolation (Section 6). With these experiments, we aim to show that speculative non-interference and our individual speculative semantics can correctly identify speculative leaks associated with speculation over branches, indirect jumps, store-bypasses and return instructions.

10.1.1 Benchmarks. We start from 41 snippets of code containing leaks resulting from speculation over `beqz`, `store`, `load`, `jmp`, and `ret` instructions (and their combinations). We compile these snippets using various compilers, compiler options, and mitigations, thereby obtaining a corpus of 291 x64 assembly programs, which we analyse with SPECTECTOR. Next, we describe next in detail the composition of our benchmark corpus:

- **Spectre-PHT:** 15 snippets are variants of the Spectre-PHT vulnerability by Kocher [48]. For each of the 15 snippets, we analyse the assembly programs obtained using different compilers and compiler options. For compilers, we rely on three state-of-the-art compilers: Microsoft VISUAL C++ versions v19.15.26732.1 and v19.20.27317.96, Intel Icc v19.0.0.117, and CLANG v7.0.0 and compile each snippet using two different *optimization levels* (`-O0` and `-O2`) and three *mitigation levels*: (a) UNP: we compile without any SPECTRE mitigations. (b) FEN: we compile with automated injection of speculation barriers.¹¹ (c) SLH: we compile using speculative load hardening.¹² Compiling each of the 15 examples from [48] with each of the 3 compilers, each of the 2 optimization levels, and each of the 2-3 mitigation levels, yields a corpus of 240 x64 assembly programs. For each program, we specify a security policy that flags as “low” all registers and memory locations that can either be controlled by the adversary or can be assumed to be public. This includes variables `y` and `size`, and the base addresses of the arrays `A` and `B` as well as the stack pointer.

- **Spectre-BTB:** 2 snippets are variants of the SPECTRE-BTB vulnerability. They exploit speculation over indirect jump instructions and were created by ourselves. For each snippet, we also analyze manually patched versions obtained by inserting `LFENCES` after every `ENDBR` instruction and automatically patched versions using the `retpoline` [45] countermeasure.¹³ This results in a corpus of 6 x64 assembly programs.

To make the analysis tractable, we use the `endbr` instruction for CFI (see Section 6.5). Note that these `endbr` instructions were automatically added by the compiler.

- **Spectre-STL:** 13 snippets are variants of the Spectre-STL vulnerability. They exploit speculation over memory disambiguation, and they have been used as benchmarks in prior work [25, 65]. For each snippet, we also analyze a patched version where a manually inserted `LFENCE` instruction stops speculation over store-bypasses and prevents the leak. This results in a corpus of 26 x64 assembly programs.

- **Spectre-RSB:** 5 snippets are variants of the Spectre-RSB vulnerability. They exploit speculation over return instructions, and they are obtained from the `safeside` [34] and `transientfail` [18] projects¹⁴. For each snippet, we also analyze manually patched versions obtained by (1) inserting `LFENCES` after call instructions (i.e., at the instruction

¹¹Fences are supported by CLANG with the flag `-x86-speculative-load-hardening-lfence`, by Icc with `-mconditionalbranch=all-fix`, and by VISUAL C++ with `/Qspectre`.

¹²Speculative load hardening is supported by CLANG with the flag `-x86-speculative-load-hardening`.

¹³Retpoline is supported by Gcc with the flag `-mindirect-branch=thunk`.

¹⁴Out of the three Spectre-RSB examples from `safeside` [34], we analyze the only one that works against an acyclic RSB like the one supported by `ℒR`. Programs `ca_ip`, `ca_oop`, and `sa_ip` from `transientfail` [18] rely on concurrent execution. Since SPECTECTOR does not support concurrency, we hardcode the worst-case interleaving in terms of speculative leakage in our benchmark.

address where **ret** speculatively returns), and (2) using the modified retpoline defense proposed in [53, Section 6.1]. This results in a corpus of 15 x64 assembly programs.

- **Spectre-SLS**: 2 snippets are variants of the Spectre-SLS vulnerability. They exploit speculation over return instructions and were created by ourselves.

For each snippets, we also analyze a manually patched version obtained by inserting `lfences` after every return instruction, resulting in a corpus of 4 x64 assembly programs.¹⁵

10.1.2 Experimental Setup. The benchmarks for **Spectre-PHT** are compiled as described in the section above. The benchmarks for **Spectre-STL**, **Spectre-RSB**, and **Spectre-SLS** are implemented in C and compiled with Gcc 11.1.0 and we manually inserted `lfence`/modified retpoline countermeasures in the patched versions. The benchmarks for **Spectre-BTI** are also implemented in C and `lfences` were inserted manually while retpoline was added automatically by the compiler into the patched programs.

The benchmarks for **Spectre-PHT** were run on a Linux machine (kernel 4.9.0-8- amd64) with Debian 9.0, a Xeon Gold 6154 CPU, and 64 GB of RAM. All other experiments were run on a laptop with a Dual Core Intel Core i5-7200U CPU and 8GB of RAM.

The speculation window ω bounds how many instructions may execute speculatively before a misprediction resolves. We use $\omega = 200$ by default, motivated by typical sizes of the reorder buffer [28], which limits the lengths of speculative transactions in modern microarchitectures.

Ex.	VCC						ICC				CLANG					
	UNP		FEN 19.15		FEN 19.20		UNP		FEN		UNP		FEN		SLH	
	-O0	-O2	-O0	-O2	-O0	-O2	-O0	-O2	-O0	-O2	-O0	-O2	-O0	-O2	-O0	-O2
01	○	○	●	●	●	●	○	○	●	●	○	○	●	●	●	●
02	○	○	●	●	●	●	○	○	●	●	○	○	●	●	●	●
03	○	○	●	○	●	●	○	○	●	●	○	○	●	●	●	●
04	○	○	○	○	●	●	○	○	●	●	○	○	●	●	●	●
05	○	○	●	○	●	○	○	○	●	●	○	○	●	●	●	●
06	○	○	○	○	○	○	○	○	●	●	○	○	●	●	●	●
07	○	○	○	○	○	○	○	○	●	●	○	○	●	●	●	●
08	○	●	○	●	○	●	○	●	●	●	○	●	●	●	●	●
09	○	○	○	○	○	○	○	○	●	●	○	○	●	●	●	●
10	○	○	○	○	○	○	○	○	●	●	○	○	●	●	●	○
11	○	○	○	○	○	○	○	○	●	●	○	○	●	●	●	●
12	○	○	○	○	●	●	○	○	●	●	○	○	●	●	●	●
13	○	○	○	○	○	○	○	○	●	●	○	○	●	●	●	●
14	○	○	○	○	●	●	○	○	●	●	○	○	●	●	●	●
15	○	○	○	○	○	○	○	○	●	●	○	○	●	●	○	●

Fig. 12. Analysis of Kocher’s examples [48], compiled with different compilers and options. For each of the 15 examples, we analyzed the unpatched version (denoted by UNP), the version patched with speculation barriers (denoted by FEN), and the version patched using speculative load hardening (denoted by SLH). Programs have been compiled without optimizations (-O0) or with compiler optimizations (-O2) using the compilers VISUAL C++ (two versions), ICC, and CLANG. ○ denotes that SPECTECTOR detects a speculative leak, whereas ● indicates that SPECTECTOR proves the program secure.

¹⁵We note that compilers like CLANG have the option `mharden-sls-all` to automatically protect the code. However, this countermeasure inserts an `int3` instruction after every return on x86. Since interrupts and exceptions are not handled by SPECTECTOR, they cannot be translated to μ ASM. An `lfence` behaves similarly in this case.

10.1.3 *Results for Spectre-PHT.* Figure 12 depicts the results of applying SPECTECTOR to the 240 **Spectre-PHT** examples. We highlight the following findings:

- SPECTECTOR detects the speculative leaks in almost all unprotected programs, for all compilers (see the UNP columns). The exception is Example #8, which uses a conditional expression instead of the if statement of Listing 1:

```
1  temp &= B[A[y<size?(y+1):0]*512];
```

At optimization level -O0, this is translated to a (vulnerable) branch instruction by all compilers, and at level -O2 to a (safe) conditional move, thus closing the leak. See Appendix A.1 for the corresponding CLANG assembly.

- The CLANG and Intel ICC compilers defensively insert fences after each branch instruction, and SPECTECTOR can prove security for all cases (see the FEN columns for CLANG and ICC). In Example #8 with options -O2 and FEN, ICC inserts an **lfence** instruction, even though the baseline relies on a conditional move, see line 10 below. This **lfence** is unnecessary according to our semantics, but may close leaks on processors that speculate over conditional moves.

```
1  mov    y, %rdi
2  lea    1(%rdi), %rdx
3  mov    size, %rax
4  xor    %rcx, %rcx
5  cmp    %rax, %rdi
6  cmovb  %rdx, %rcx
7  mov    temp, %r8b
8  mov    A(%rcx), %rsi
9  shl    $9, %rsi
10 lfence
11 and    B(%rsi), %r8b
12 mov    %r8b, temp
```

- For the VISUAL C++ compiler, SPECTECTOR automatically detects all leaks pointed out in [48] (see the FEN 19.15 -O2 column for Vcc). Our analysis differs from Kocher’s only on Example #8, where the compiler v19.15.26732.1 introduces a safe conditional move, as explained above. Moreover, without compiler optimizations (which is not considered in [48]), SPECTECTOR establishes the security of Examples #3 and #5 (see the FEN 19.15 -O0 column). The latest VCC compiler additionally mitigates the leaks in Examples #4, #12, and #14 (see the FEN 19.20 column).

- SPECTECTOR can prove the security of speculative load hardening in Clang (see the SLH column for CLANG), except for Example #10 with -O2 and Example #15 with -O0. We stress that this is a translation-validation result: SPECTECTOR certifies the security of individual hardened programs, not of the hardening pass itself. Indeed, SPECTECTOR flags two hardened programs that remain insecure, matching the known result that Clang’s speculative load hardening is not secure in general [82].

Example 10 with Speculative Load Hardening. Example #10 differs from Listing 1 in that it leaks sensitive information into the microarchitectural state by conditionally reading the content of B[0], depending on the value of A[y].

```
1  if (y < size)
2      if (A[y] == k)
```

```

3      temp &= B[0];

```

SPECTECTOR proves the security of the program produced with CLANG -O0 and speculative load hardening. However, at optimization level -O2, CLANG outputs the following code that SPECTECTOR reports as insecure.

```

1  mov     size, %rdx
2  mov     y, %rbx
3  mov     $0, %rax
4  cmp     %rbx, %rdx
5  jbe     END
6  cmovbe  $-1, %rax
7  or      %rax, %rbx
8  mov     k, %rcx
9  cmp     %rcx, A(%rbx)
10 jne     END
11 cmovne  $-1, %rax
12 mov     B, %rcx
13 and     %rcx, temp
14 jmp     END

```

The reason for this is that CLANG masks only the register `%rbx` that contains the index of the memory access `A[y]`, cf. lines 6–7. However, it does *not* mask the value that is read from `A[y]`. As a result, the comparison at line 9 speculatively leaks (via the jump target) whether the content of `A[0xFF...FF]` is `k`. SPECTECTOR detects this subtle leak and flags a violation of speculative non-interference.

While this example nicely illustrates the scope of SPECTECTOR, it is likely not a problem in practice: First, the leak may be mitigated by how data dependencies are handled in modern out-of-order CPUs. Specifically, the conditional move in line 6 relies on the comparison in Line 4. If executing the conditional leak effectively terminates speculation, the reported leak is spurious. Second, the leak can be mitigated at the OS-level by ensuring that `0xFF...FF` is not mapped in the page tables, or that the value of `A[0xFF...FF]` does not contain any secret.¹⁶ Such contextual information can be expressed with policies (see Section 5.1) to improve the precision of the analysis.

Results for Spectre-BTB Figure 13c reports the analysis of the programs in the **Spectre-BTB** benchmark. Using the $\mathcal{L}_j$ semantics, SPECTECTOR detected leaks (i.e., violations of SNI) in all unpatched programs. Furthermore, SPECTECTOR proved that the manually patched programs using `lfence` and the automatically patched programs using the `retpoline` countermeasure satisfy SNI, i.e., they are free of speculative leaks.

10.1.4 Results for Spectre-STL. Figure 13a reports the results of analysing the programs in the **Spectre-STL** benchmark.¹⁷ Using the $\mathcal{L}_S$ semantics, SPECTECTOR successfully detected leaks (i.e., violations of SNI) in all unpatched programs, except programs 03, 09, and 12 which do not contain speculative leaks (consistently with other analysis results [25, 65]). Observe that Binsec/Haunted [25] flags program 13 as secure since the program can *only* speculatively

¹⁶Personal Communication with C. Marinas

¹⁷We had to slightly modify programs 02, 05, and 06 due to limitations of SPECTECTOR's x86 front-end when dealing with global values (programs 05 and 06) and 32-bit addressing (program 02). We had to limit the speculation window, due to vanilla SPECTECTOR's limitations in symbolic execution, when analyzing program 09, which contains a loop.

Test case		$\mathcal{Q}_S$	
		None	Fence
case01	(I)	○	●
case02	(I)	○	●
case03	(S)	●	●
case04	(I)	○	●
case05	(I)	○	●
case06	(I)	○	●
case07	(I)	○	●
case08	(I)	○	●
case09	(S)	●	●
case10	(I)	○	●
case11	(I)	○	●
case12	(S)	●	●
case13	(I)	○	●

(a) Results for the Spectre-STL programs under the $\mathcal{Q}_S$ semantics against unpatched programs (column “None”) and programs patched with 1fence (column “Fence”).

Test case		$\mathcal{Q}_J$		
		None	Fence	Retpoline
caseJ01	(I)	○	●	●
caseJ02	(I)	○	●	●

(c) Results for the Spectre-BTB programs under the $\mathcal{Q}_J$ semantics against unpatched programs (column “None”), programs patched with 1fence (column “Fence”), and programs patched with the retpoline defense proposed in [45] (column “Retpoline”).

Test case		$\mathcal{Q}_R$		
		None	Fence	Mod. Retpoline
<i>ret2spec_c_d</i>	(I)	○	●	●
<i>ca_ip</i>	(I)	○	●	●
<i>ca_oop</i>	(I)	○	●	●
<i>sa_ip</i>	(I)	○	●	●
<i>sa_oop</i>	(I)	○	●	●

(b) Results for the Spectre-RSB programs under the $\mathcal{Q}_R$ semantics against unpatched programs (column “None”), programs patched with 1fence (column “Fence”), and programs patched with the modified retpoline defense proposed in [53, S6.1] (column “Mod. Retpoline”).

Test case		$\mathcal{Q}_{SLS}$	
		None	Fence
caseSLS01	(I)	○	●
caseSLS02	(I)	○	●

(d) Results for the Spectre-SLS programs under the $\mathcal{Q}_{SLS}$ semantics against unpatched programs (column “None”), programs patched with 1fence (column “Fence”).

Fig. 13. Result of the analysis of our benchmarks for $\mathcal{Q}_J$, $\mathcal{Q}_S$, $\mathcal{Q}_R$ and $\mathcal{Q}_{SLS}$. For each program, ○ denotes that SPECTECTOR finds a violation of SNI under the corresponding semantics, whereas ● denotes that SPECTECTOR proves the program secure under the semantics. Next to each program, we report if the program is Secure or Insecure in its unpatched version.

leak initial values from the stack, which Binsec/Haunted treats as public by default [14]. Since we assume initial memory values to be secret (like Ponce de León and Kinder [65]), SPECTECTOR correctly detected the leak in program 13. SPECTECTOR also successfully proved that all patched programs (where an 1fence is added between store instructions) satisfy SNI and are free of speculative leaks.

10.1.5 Results for Spectre-RSB. Figure 13b reports the analysis results on the **Spectre-RSB** programs. Using $\mathcal{Q}_R$, SPECTECTOR successfully detected leaks in all unpatched programs. Moreover, SPECTECTOR successfully proved that the patched programs where a 1fence instruction is added after every call satisfy SNI, i.e., they are free of speculative leaks. SPECTECTOR also successfully proved secure the programs patched using the modified retpoline defense proposed by Maisuradze and Rossow [53], which replaces return instructions with a construct that traps the speculation in an infinite loop.

10.1.6 Results for Spectre-SLS. Figure 13d reports the analysis results on the **Spectre-SLS** programs. Using $\mathcal{Q}_{SLS}$, SPECTECTOR successfully detected leaks in all unpatched programs and proved the security (i.e. SNI) of the patched programs where an LFENCE instruction was added after every return instruction.

10.2 Case Study: Detecting Leaks w.r.t. Combined Speculative Semantics

Using SPECTECTOR, we analyze a corpus of 18 microbenchmarks containing speculative leaks generated by a combination of speculation mechanisms (for all composed semantics from Section 8). In particular, we consider one microbenchmark for each of all possible combinations of the $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$ semantics, i.e., 18 combined semantics (as mentioned, combinations containing both $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$ are not possible). With these experiments, we aim to show that our combined semantics can detect novel leaks that are otherwise undetectable when considering single speculation mechanisms in isolation.

10.2.1 Benchmarks. For each of the 18 combined semantics, we manually crafted an example μ ASM program that has a speculative leak under that specific combination. For each program, we also analyze a manually patched version where `lfence` instructions prevent speculative leaks. We refer to this benchmark, consisting of 36 μ ASM programs, as **Spectre-Comb**.

10.2.2 Experimental Setup. For each μ ASM program in **Spectre-Comb**, we run SPECTECTOR with all individual speculative semantics as well as all their combinations. All experiments were run on a laptop with a Dual Core Intel Core i5-7200U CPU and 8GB of RAM.

10.2.3 Results for Spectre-Comb. Table 1 reports the results of our analysis on the **Spectre-Comb** programs, which involve leaks arising from a combination of multiple speculation mechanisms.

SPECTECTOR equipped with the individual semantics $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$ is not able to detect the speculative leaks in any of the 18 programs and, therefore, proves them secure. This is expected since the programs contain leaks that arise from a combination of semantics.

SPECTECTOR can successfully identify leaks in all programs when using the semantics corresponding to that example. For example, **comb25** is proven insecure by the respective semantics $\mathcal{L}_{S+R}$.

Furthermore, as Figure 9 suggests, all 'stronger' semantics (higher in the combination lattice) also detect the vulnerability in the programs made for a 'weaker' semantics (lower in the combination lattice) and the 'weaker' ones do not detect the vulnerability of programs made for 'stronger' semantics. For example, **comb25** is also proven to be insecure by SPECTECTOR when using $\mathcal{L}_{B+J+R}$, $\mathcal{L}_{J+S+R}$, $\mathcal{L}_{B+J+S+R}$, while the program **comb1245** can only be proven insecure by $\mathcal{L}_{B+J+S+R}$.

Semantics that are in no particular order in the lattice like $\mathcal{L}_{B+J+S+R}$ and $\mathcal{L}_{B+J+S+SLS}$ are also not able to detect the vulnerability in **comb1246** and **comb1245** respectively. This is expected since they each miss one speculative mechanism necessary to detect the vulnerability.

Because there is no combination of all our source speculative semantics $\mathcal{L}_B$, $\mathcal{L}_J$, $\mathcal{L}_S$, $\mathcal{L}_R$ and $\mathcal{L}_{SLS}$, there is also no unique combination that can find the leaks in all programs in **Spectre-Comb**. However, using $\mathcal{L}_{B+J+S+R}$ and $\mathcal{L}_{B+J+S+SLS}$, SPECTECTOR is able to successfully detect leaks in all programs.

We also analyzed programs manually patched with `lfence` statements. The examples are not shown here for brevity and can be found with the other code snippets at [37]. As before, SPECTECTOR successfully proves the security of all patched programs. Even for leaks that arise from multiple speculation mechanisms, it is often sufficient to insert a single `lfence` to secure the entire program, e.g., an `lfence` after the `beqz` instruction in **comb15** is enough to make the program SNI with respect to $\mathcal{L}_{B+S+R}$.

Program	$\mathcal{L}_B$	$\mathcal{L}_J$	$\mathcal{L}_S$	$\mathcal{L}_R$	$\mathcal{L}_{SLS}$	$\mathcal{L}_{B+J}$	$\mathcal{L}_{B+S}$	$\mathcal{L}_{B+R}$	$\mathcal{L}_{B+SLS}$	$\mathcal{L}_{J+S}$	$\mathcal{L}_{J+R}$	$\mathcal{L}_{J+SLS}$	$\mathcal{L}_{S+R}$	$\mathcal{L}_{S+SLS}$
comb12	•	•	•	•	•	◦	•	•	•	•	•	•	•	•
comb14	•	•	•	•	•	◦	◦	•	•	•	•	•	•	•
comb15	•	•	•	•	•	•	•	◦	•	•	•	•	•	•
comb16	•	•	•	•	•	•	•	•	◦	•	•	•	•	•
comb24	•	•	•	•	•	•	•	•	•	◦	•	•	•	•
comb25	•	•	•	•	•	•	•	•	•	•	◦	•	•	•
comb26	•	•	•	•	•	•	•	•	•	•	•	◦	•	•
comb45	•	•	•	•	•	•	•	•	•	•	•	•	◦	•
comb46	•	•	•	•	•	•	•	•	•	•	•	•	•	◦
comb124	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb125	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb126	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb145	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb146	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb245	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb246	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb1245	•	•	•	•	•	•	•	•	•	•	•	•	•	•
comb1246	•	•	•	•	•	•	•	•	•	•	•	•	•	•

Program	$\mathcal{L}_{B+J+S}$	$\mathcal{L}_{B+J+R}$	$\mathcal{L}_{B+J+SLS}$	$\mathcal{L}_{B+S+R}$	$\mathcal{L}_{B+S+SLS}$	$\mathcal{L}_{J+S+R}$	$\mathcal{L}_{J+S+SLS}$	$\mathcal{L}_{B+J+S+R}$	$\mathcal{L}_{B+J+S+SLS}$
comb12	◦	◦	◦	•	•	•	•	◦	◦
comb14	◦	•	•	◦	◦	•	•	◦	◦
comb15	•	◦	•	◦	•	•	•	◦	•
comb16	•	•	◦	•	◦	•	•	•	◦
comb24	◦	•	•	•	•	◦	◦	◦	◦
comb25	•	◦	•	•	•	◦	•	◦	•
comb26	•	•	◦	•	•	•	◦	•	◦
comb45	•	•	•	◦	•	◦	•	◦	•
comb46	•	•	•	•	◦	•	◦	•	◦
comb124	◦	•	•	•	•	•	•	◦	◦
comb125	•	◦	•	•	•	•	•	◦	•
comb126	•	•	◦	•	•	•	•	•	◦
comb145	•	•	•	◦	•	•	•	◦	•
comb146	•	•	•	•	◦	•	•	•	◦
comb245	•	•	•	•	•	◦	•	◦	•
comb246	•	•	•	•	•	•	◦	•	◦
comb1245	•	•	•	•	•	•	•	◦	•
comb1246	•	•	•	•	•	•	•	•	◦

Table 1. Results of the analysis. For each program, ◦ denotes that SPECTECTOR finds a violation of SNI whereas • denotes that SPECTECTOR proves the program secure under the corresponding semantics. The different programs were devised to be vulnerable under a specific combination. The numbers correspond to that specific version 1 = B, 2 = J, 4 = S, 5 = R and 6 = SLS. For example, comb125 was devised to be vulnerable under $\mathcal{L}_{B+J+R}$.

11 DISCUSSION

11.1 Scope of the Models

Lifting the results of the security analysis for our speculative semantics to real-world CPUs is only possible to the extent that these semantics capture the information flows in the target system. Thus, SPECTECTOR’s result may incorrectly classify programs as secure (if our semantics do not capture information flows happening in real-world CPUs) or insecure (if our semantics admit speculations that are impossible on real systems).

We capture “leakage into the microarchitectural state” using the relatively powerful observer of the program execution that sees the location of memory accesses and the jump targets. This observer could be replaced by a weaker one, which accounts for more detailed models of a CPU’s memory hierarchy, and SPECTECTOR could be adapted accordingly, e.g. by adopting the cache models from CacheAudit [29]. We believe, however, that highly detailed models are not actually desirable for several reasons: (a) they encourage brittle designs that break under small changes to the model, (b) they have to be adapted frequently, and (c) they are hard to understand and reason about for compiler developers and hardware engineers. The “constant-time” observer model adopted in this paper has proven to offer a good tradeoff between precision and robustness [4, 57].

11.2 Other Speculation Mechanisms

There are many speculation mechanisms beyond those modeled in the speculative semantics $\mathcal{L}_B, \mathcal{L}_S, \mathcal{L}_J, \mathcal{L}_R$ and $\mathcal{L}_{SLS}$ from Section 6:

- CPUs speculate over **ret** instructions in different ways. For instance, there are many different ways of implementing return stack buffers (e.g., cyclic versus acyclic RSBs [53] or RSBs that fall back to indirect branch prediction [78]). This kind of speculation can be modeled by modifying the Rule **R:AM-Spec** rule in $\mathcal{L}_R$.
- Many proposals for value prediction over different kinds of instructions exist [51, 56, 69]. While naive speculative semantics might have to explore *all* possible values as prediction, semantics that model specific prediction mechanisms might restrict the set of predicted values (thereby leading to a more tractable analysis).

We expect that most of these mechanisms can be modeled as speculative semantics satisfying our well-formedness conditions. Hence, they could work with our composition framework.

11.3 Limitations of Composition

Our composition framework has two main limitations:

(1) The metaparameter Z is expressed in terms of μ ASM instructions, i.e., the smallest unit of computation in our framework. Since Z restricts how the composed semantics delegates execution to its sources, this limit the expressiveness of composed semantics. For instance, $\mathcal{L}_{S+R}$ cannot speculate over the *implicit store* writing the return address to the stack that happens as part of **call** instructions.

(2) Our framework does not support combinations where a single instruction performs speculation-relevant changes in both source semantics. For instance, consider a combination of $\mathcal{L}_R$ with $\mathcal{L}_{SLS}$. Here, both semantics start different speculative transactions on executing **ret** instructions. However, instantiating Z as $(\emptyset, \emptyset)$, which enables both speculations, violates the confluence well-formedness condition for the composed semantics, whereas setting $Z = (x, y)$ so that only one of x and y is **ret** would only capture one of the two speculation mechanisms.

The first limitation is inherent to the instruction granularity of Z , and we leave it to future work. The second can be lifted with a moderate extension, which we sketch here. When a single instruction (e.g., **ret**) triggers speculation in both

source semantics, the composed semantics—rather than delegating to one source and blocking the other via Z —creates two speculative instances at that instruction, one per source semantics, and pushes both onto the speculation stack. Since a speculative instance is popped when its window ends, execution resumes with the second instance, so both mechanisms are explored. The same construction applies to any pair of mechanisms acting on a single instruction, such as the load-address (LAP [47]) and load-value (LVP [46]) predictors of recent Apple processors. We leave a full formalization to future work.

11.4 Scalability

SPECTECTOR is a proof-of-concept aimed at detecting speculative leaks and their combinations rather than at analyzing large code bases, and three factors currently limit its scalability. First, the front end supports only a fraction of the x64 ISA (cf. Section 9.4); extending it to the full ISA is conceptually straightforward but engineering-intensive. Second, SPECTECTOR relies on manually specified policies partitioning configurations into public and secret parts; precise manual policies do not scale to large code bases, and automatic policy inference from the calling context is not yet supported. Third, SPECTECTOR is based on symbolic execution and therefore inherits its path explosion. That said, checking SNI itself adds little on top of symbolic execution. Although SNI is a relational property—and relational analyses are known to scale poorly—SPECTECTOR only ever compares executions that follow the *same* symbolic path: because the program counter is observable, speculative non-interference never requires relating two executions that disagree on their path conditions. The cost of the analysis is thus dominated by the underlying symbolic execution rather than by the relational SNI check itself.

11.5 Different Flavours of Non-Interference

Here, we discuss the relation of SNI with other well-known notions of security against side-channel leaks.

11.5.1 General Non-Interference. We start by introducing General Non-Interference (GNI), a notion of security capturing that secrets do not leak through side channels. This notion is adapted from the hardware-software contracts framework of Guarnieri et al. [36] into our setting.

Definition 5 (General Non-Interference (GNI)). *Program p satisfies GNI (denoted $p \vdash_x \text{GNI}$) for a semantics x iff for all σ, σ' , if $\sigma \sim_\phi \sigma'$ then $\text{Beh}_x^\omega(p, \sigma) = \text{Beh}_x^\omega(p, \sigma')$.*

In a nutshell, a program p satisfies GNI for a given semantics x if any pair of low-equivalent initial configurations σ and σ' results in the same observations, i.e., $\text{Beh}_x^\omega(p, \sigma) = \text{Beh}_x^\omega(p, \sigma')$.

Observe that GNI is an *absolute security property* [21], that is, it ensures that *all* observations are the same for any two executions starting from low-equivalent configurations. In contrast, SNI is a *relative security property* [21]. That is, it ensures that speculatively executed instructions (and, thus, speculative observations) do not leak more information than what is already leaked by the program’s non-speculative execution. Naturally, satisfying GNI implies satisfying SNI for the same semantics x :

Corollary 1 (GNI implies SNI). *If $p \vdash_x \text{GNI}$ then $p \vdash_x \text{SNI}$.*

11.5.2 Relation with constant-time variants. A common defense against several side-channel timing attacks is the constant-time programming discipline [4, 57]. This programming discipline requires that a program’s control flow and memory access patterns are strictly independent of sensitive data.

Following [4], constant-time programming can be modeled as a non-interference property by requiring that low-equivalent executions result in the same architectural control-flow and in the same sequence of memory accesses. That is, we can instantiate constant-time security by requiring the equivalence of the non-speculative trace, as indicated next:

Definition 6 (Non-Speculative Constant-Time (CT)). *Program p satisfies CT (denoted $p \vdash \text{SeqCT}$) iff for all σ, σ' , if $\sigma \sim_{\phi} \sigma'$ then $\text{Beh}_{NS}(p, \sigma) = \text{Beh}_{NS}(p, \sigma')$.*

Observe that CT is an instantiation of GNI w.r.t. the non-speculative semantics, i.e., $p \vdash \text{SeqCT} \Leftrightarrow p \vdash_{NS} \text{GNI}$. Note also that we named Definition 6 as “Sequential Constant-Time” rather than the more standard “Constant-Time” to stress that CT traditionally applies only to the architectural, non-speculative semantics.

To account for the effects of speculatively executed instructions, several works [3, 12, 20, 25, 75] proposed security conditions that also restrict leaks through speculatively executed instructions. That is, all these works enforce variants of GNI, differing in the underlying speculative semantics or in what observations are part of the trace. For instance, in Cauligi et al. [20] the low-projection of the final configuration is also part of the trace (in addition to the usual observations associated with the constant-time attacker) to ensure the absence of explicit leaks.

Crucially, we can precisely connect CT, SNI, and GNI. That is, general non-interference w.r.t. a speculative semantics x is exactly the conjunction of constant-time and speculative non-interference w.r.t. x [36, Proposition 4]:

Proposition 1 (GNI Decomposition). *Given a policy ϕ and a program p : $p \vdash \text{SeqCT} \wedge p \vdash_x \text{SNI} \Leftrightarrow p \vdash_x \text{GNI}$*

Proposition 1 allows us to decompose the check for GNI into two independent checks. An analysis only needs to prove that the program is constant-time (SeqCT) and that it satisfies speculative non-interference (SNI).

11.5.3 Extending SPECTECTOR with support for GNI. We conclude by discussing how SPECTECTOR can be extended to check GNI. As stated above, SeqCT can be obtained by instantiating GNI with the non-speculative semantics, so this extension can be used also to check the classic constant-time property.

Recall that SPECTECTOR algorithm from Section 9.1 searches for pairs of execution traces that have equivalent non-speculative behavior but differing speculative observations.

To support GNI, we need to modify MEMLEAK and CTRLLEAK as indicated in Figure 14. Specifically, we enforce that the entire symbolic trace $\bar{\tau}$ can only produce identical concrete observations for all low-equivalent inputs.

For this, we modify MEMLEAK to check the satisfiability of the following formula:

$$pthCnd(\bar{\tau})_{1\wedge 2} \wedge polEqv(P) \wedge \neg obsEqv(\bar{\tau}).$$

In particular, $pthCnd(\bar{\tau})_{1\wedge 2}$ ensures that the same symbolic path is followed, $polEqv(P)$ ensures that the initial states are low-equivalent, and $\neg obsEqv(\bar{\tau})$ asserts that at least one observation (either architectural or speculative) associated with memory accesses differ among the two runs. That is, this formula is satisfiable only if there are two concrete traces associated with the symbolic trace $\bar{\tau}$ where memory operation accesses a different address.

The CTRLLEAK procedure is modified analogously to ensure that all control-flow observations (architectural or speculative) are identical between traces.

Thus, with minor modifications to the SPECTECTOR algorithm, we are able to check for GNI-style properties for all the semantics discussed in the paper. The full algorithm for checking GNI can be found in Appendix B.

2549			
2550	Algorithm 2	MEMLEAK	for GNI
2551	1:	procedure	$\text{MEMLEAK}(\bar{\tau}, P)$
2552	2:	$\psi \leftarrow \text{pthCnd}(\bar{\tau})_{1 \wedge 2} \wedge \text{polEqv}(P) \wedge$	
2553		$\neg \text{obsEqv}(\bar{\tau})$	
2554	3:	return	$\text{SATISFIABLE}(\psi)$
2555			
2556			
2557			
2558			
2559			
2560			
2561			
2562			
2563			
2564			
2565			
2566			
2567			
2568			
2569			
2570			
2571			
2572			
2573			
2574			
2575			
2576			
2577			
2578			
2579			
2580			
2581			
2582			
2583			
2584			
2585			
2586			
2587			
2588			
2589			
2590			
2591			
2592			
2593			
2594			
2595			
2596			
2597			
2598			
2599			
2600			

Fig. 14. The modified versions of MEMLEAK and CTRLLEAK for GNI

12 RELATED WORK

Speculative Execution Attacks: After Spectre [49] has been disclosed to the public in 2018, researchers have identified many other speculative execution attacks [10, 15, 50, 53, 79, 81]. These attacks differ in the exploited speculation sources [43, 50, 53], the covert channels [2, 52, 67, 72, 73] used, or the target platforms [23]. We refer the reader to Canella et al. [18] for a survey of existing attacks.

Security Properties for Speculative Leaks: Researchers have proposed many program-level properties for security against speculative leaks, which can be classified in three main groups [21]:

(1) Non-interference definitions ensure the security of speculative *and* non-speculative instructions. For instance, speculative constant-time [3, 12, 20, 25, 75] extends the constant-time security condition to account also for transient instructions.

(2) Relative non-interference definitions [22, 35, 39, 70] ensure that transient instructions do not leak more information than what is leaked by non-transient instructions. The notion of Speculative Non-Interference (SNI) introduced in this paper falls into this category, as it restricts the information leaked by speculatively executed instructions relative to the program's non-speculative behavior.

(3) Definitions that formalise security as a safety property [64, 65], which may over-approximate definitions from the groups above.

Finally, we highlight the relationship between Speculative Constant-Time (SCT) and our properties SNI and GNI. As shown by Cauligi et al. [20] if the program is sequentially constant-time, SNI, and the resulting architectural states are low-equivalent, then the program is also SCT. In relation to GNI, SCT enforces the exact same trace equivalence but imposes the additional constraint that final architectural states must be low-equivalent. We can verify SCT on top of our SPECTECTOR-GNI check by explicitly exposing the final architectural state as a terminal observation within the execution trace.

Operational Semantics for Speculative Leaks: In the last few years, there has been a growing interest in developing formal models and principled program analyses for detecting leaks caused by speculatively executed instructions. We refer the reader to [21] for a comprehensive survey. In the following, we discuss the approaches that are particularly relevant to our paper.

Some approaches explicitly model microarchitectural components like multiple pipeline stages, caches, and branch predictors. For instance, KLEESpectre [76] and SpecuSym [39] consider a semantics that explicitly models the cache, enabling reasoning about the cache content. McIlroy et al. [55] go a step further and model a multi-stage pipeline with

explicit cache and branch predictor. Their semantics can only model speculation over branch instructions since it lacks store-forwarding or RSB.

Cauligi et al. [20]’s semantics model speculation over branch instructions, store-bypasses, and return instructions. Differently from our semantics, their 3-stage pipeline semantics explicitly models several microarchitectural components like a reorder buffer and an RSB. Their tool detects violations of speculative constant-time induced by speculation over branch instructions and store-bypasses.

Barthe et al. [12] extend the Jasmin [3] cryptographic verification framework to reason about speculative constant-time and supports speculation over store-bypasses and branch instructions. Shivakumar et al. [71] equip Jasmin with an information flow control type system that enforces SCT. However, only against SPECTRE-PHT attacks.

Binsec/Haunted [25] detect violations of speculative constant-time due to speculation over store-bypasses and branch instructions. For this, they explicitly model the store buffer, which $\mathcal{L}_S$ abstracts away. Blade [75] uses a JIT-step semantics that translates high-level commands (from a WHILE-language) to low-level machine instructions while tracking control-flow predictions. This allows for source-level reasoning and their semantics captures speculation over branch instructions.

While several of these models support multiple speculation mechanisms, these mechanisms are *hard-coded* and no existing approach provides a composition framework or extensible ways of extending the main theoretical results to new mechanisms “for free”. Moreover, while we could have used other semantics as a basis for our framework, this would have resulted in more difficult proofs (since semantics like the one in [20] are significantly more complex than ours).

Axiomatic Semantics for Speculative Leaks: A few approaches formalise the effects of speculatively executed instructions using axiomatic semantics inspired by work on weak memory models. For instance, Colvin and Winter [24] and Disselkoe et al. [27] capture the effects of branch speculation but both lack program analyses.

Ponce de León and Kinder [65] illustrate how one can model leaks resulting from speculation over branch instructions and store-bypasses using the CAT modeling language for memory consistency, and they present a bounded model checking analysis for detecting speculative leaks. Interestingly, they talk about composing several of their semantics [65, SIV.F], which should allow them to detect vulnerabilities like Listing 3 (which we detect under $\mathcal{L}_{B+S}$). However, they do not formally characterize compositions and, therefore, they cannot derive interesting results “for free” about the composed semantics (like we do in Theorem 2). Moreover, even though they state that composability is an advantage of axiomatic models, our framework (and tool implementation) shows that composability can be done with operational semantics as well.

Secure Compilation for Speculative Leaks: In addition to program analyses like the ones described above, researchers have proposed compiler passes to prevent and mitigate speculative leaks. For instance, Patrignani and Guarnieri [64] and Fabian et al. [32] study the security of compiler-level countermeasures implemented in major compilers using Secure Compilation theory, whereas Vassena et al. [75] propose a compiler pass for automatically patching speculative leaks using a type system to track speculative leaks. Shivakumar et al. [70] show that speculative execution can leak secrets even at declassification sites and propose selective speculative load hardening (selSLH), a compiler-level countermeasure they prove sound via a relative non-interference property. In a similar vein, Ultimate SLH [?] strengthens SLH to cover a broader class of speculative leaks.

Mosier et al. [58] propose a set of compiler passes that, together with hardware support (e.g. Intel CET-IBT, a shadow stack for return addresses) offers protection against SPECTRE-PHT, SPECTRE-BTB, SPECTRE-STL, SPECTRE-RSB, and predictive store forwarding [6]. We note that Mosier et al. [58, Section 3.3] showed a counterexample that our

semantics $\mathcal{L}_S$ does not catch. When devising our speculative semantics, we encountered a trade-off: keeping the related SPECTECTOR implementation tractable (albeit imprecise) or making the semantics precise (but not implementable due to state explosion of the tool). We decided to model the speculation in that way to keep the analysis tractable. Furthermore, our case study in Section 10.1 was used by other detection tools as well [25, 65] and our semantics $\mathcal{L}_S$ detects all vulnerabilities there. Thus, we think this is a reasonable trade-off between precision and performance.

Detecting Leaks through Testing: The Revizor testing tool [42, 60–62] and SpecFuzz [63] use fuzz testing to find vulnerabilities caused by speculation in CPUs. Similarly, SpecDoctor by Hur et al. [44] uses differential fuzzing to fuzz for transient vulnerabilities on the RTL level. Marinaro et al. [54] extend Scam-V [59] to optimize software mitigations for SPECTRE-PHT using a relational testing approach.

AMuLeT [33] adapts Revizor to the design phase. AMuLeT instruments microarchitectural simulators (e.g., gem5) to check secure speculation countermeasures against leakage contracts. LMTTest [11] provides a framework for testing cryptographic implementations against parameterized leakage models (defined in a DSL called LmSpec). It generates test cases to detect secret-dependent leaks in the presence of proposed microarchitectural optimizations.

Unlike SPECTECTOR, these approaches cannot detect the absence of leaks.

13 CONCLUSION

This paper presented an approach to automatically detect speculative leaks in programs. We introduced speculative non-interference, the first semantic notion of security against speculative execution attacks, and we defined multiple speculative semantics to capture 5 classes of Spectre attacks. Next, we defined a general framework to reason about the composition of different speculative semantics and instantiated the framework with our speculative semantics. Our framework yields safety of the composed semantics (almost) for free, given the safety of its parts. We utilize our theoretical development and implement a program analysis tool called SPECTECTOR that automatically detects speculative leaks or proves their absence, for our speculative semantics and their combinations.

ACKNOWLEDGMENTS

This work is supported by the Spanish Ministry of Science, Innovation, and University under the Ramón y Cajal grant RYC2021-032614-I; the Spanish Ministry of Science, Innovation, and University under the project PID2022-142290OB-I00 ESPADA; the Spanish Ministry of Science, Innovation, and University under the Europa Excelencia project EUR2025-164828 SINTRAZAS; the Spanish Ministry of Science, Innovation, and University under the project CEX2024-001471-M; the European Union under the Horizon Europe projects 101230068 PRIMULA and 101020415 SafeSecS.

REFERENCES

- [1] Martín Abadi, Mihai Budiu, Úlfar Erlingsson, and Jay Ligatti. 2005. Control-Flow Integrity. In *Proceedings of the 12th ACM Conference on Computer and Communications Security* (Alexandria, VA, USA) (CCS '05). Association for Computing Machinery, New York, NY, USA, 340–353. <https://doi.org/10.1145/1102120.1102165>
- [2] Alejandro Cabrera Aldaya, Billy Bob Brumley, Sohaib ul Hassan, Cesar Pereida García, and Nicola Taveri. 2019. Port Contention for Fun and Profit. In *2019 IEEE Symposium on Security and Privacy (SP)*. 870–887. <https://doi.org/10.1109/SP.2019.00066>
- [3] José Bacelar Almeida, Manuel Barbosa, Gilles Barthe, Arthur Blot, Benjamin Grégoire, Vincent Laporte, Tiago Oliveira, Hugo Pacheco, Benedikt Schmidt, and Pierre-Yves Strub. 2017. Jasmin: High-Assurance and High-Speed Cryptography. In *Proceedings of the 24th ACM SIGSAC Conference on Computer and Communications Security (CCS '17)*. ACM.
- [4] José Bacelar Almeida, Manuel Barbosa, Gilles Barthe, François Dupressoir, and Michael Emmi. 2016. Verifying constant-time implementations. In *Proceedings of the 25th USENIX Conference on Security Symposium (SEC'16)*. USENIX, USA.
- [5] AMD. 2020. Whitepaper Straight-line Speculation. <https://www.amd.com/system/files/documents/technical-guidance-for-mitigating-branch-type-confusion.pdf>.
- [6] AMD. 2021. Security analysis of AMD predictive store forwarding. <https://www.amd.com/system/files/documents/security-analysis-predictive-store-forwarding.pdf>. Accessed: 2024-03-11.
- [7] arm. 2020. Straight-line Speculation. <https://developer.arm.com/documentation/102825/0100/?lang=en>. Accessed: 2022-09-05.
- [8] ARM. 2020. Whitepaper Straight-line Speculation. <https://developer.arm.com/documentation/102825/0100/>.
- [9] ARM. 2021. Arm Armv9-A A64 Instruction Set Architecture. <https://developer.arm.com/documentation/100076/0100/A64-Instruction-Set-Reference/A64-General-Instructions/BTI?>
- [10] Enrico Barberis, Pietro Frigo, Marius Muench, Herbert Bos, and Cristiano Giuffrida. 2022. Branch history injection: On the effectiveness of hardware mitigations against cross-privilege Spectre-v2 attacks. In *Proceedings of the 31st USENIX Security Symposium (USENIX Security '22)*. USENIX Association.
- [11] Gilles Barthe, Marcel Böhme, Sunjay Cauligi, Chitchanok Chuengsatiansup, Daniel Genkin, Marco Guarnieri, David Mateos Romero, Peter Schwabe, David Wu, and Yuval Yarom. 2024. Testing Side-channel Security of Cryptographic Implementations against Future Microarchitectures. In *Proceedings of the 2024 ACM SIGSAC Conference on Computer and Communications Security (CCS '24)*. ACM. <https://doi.org/10.1145/3658644.3670319>
- [12] Gilles Barthe, Sunjay Cauligi, Benjamin Grégoire, Adrien Koutsos, Kevin Liao, Tiago Oliveira, Swarn Priya, Tamara Rezk, and Peter Schwabe. 2021. High-Assurance Cryptography in the Spectre Era. In *Proceedings of the 42nd IEEE Symposium on Security and Privacy (S&P '21)*. IEEE.
- [13] Gilles Barthe, Pedro R D'argenio, and Tamara Rezk. 2004. Secure information flow by self-composition. In *Proceedings of the 17th IEEE Computer Security Foundations Workshop (CSF '04)*. IEEE.
- [14] Binsec/Haunted Benchmark. 2021. Result of case_13. https://github.com/binsec/haunted_bench/issues/2.
- [15] Atri Bhattacharyya, Alexandra Sandulescu, Matthias Neugschwandtner, Alessandro Sorniotti, Babak Falsafi, Mathias Payer, and Anil Kurmus. 2019. SMOtherSpectre: Exploiting Speculative Execution through Port Contention. In *Proceedings of the 26th ACM SIGSAC Conference on Computer and Communications Security (CCS '19)*. ACM.
- [16] Aaron R. Bradley and Zohar Manna. 2007. *The Calculus of Computation: Decision Procedures with Applications to Verification*. Springer.
- [17] Nathan Burow, Scott A. Carr, Joseph Nash, Per Larsen, Michael Franz, Stefan Brunthaler, and Mathias Payer. 2017. Control-Flow Integrity: Precision, Security, and Performance. *ACM Comput. Surv.* 50, 1 (apr 2017). <https://doi.org/10.1145/3054924>
- [18] Claudio Canella, Jo Van Bulck, Michael Schwarz, Moritz Lipp, Benjamin von Berg, Philipp Ortner, Frank Piessens, Dmitry Evtvushkin, and Daniel Gruss. 2019. A Systematic Evaluation of Transient Execution Attacks and Defenses. In *Proceedings of the 28th USENIX Security Symposium (USENIX Security '19)*. USENIX Association.
- [19] C. Carruth. 2018. Speculative load hardening. <https://lvm.org/docs/SpeculativeLoadHardening.html>
- [20] Sunjay Cauligi, Craig Disselkoen, Klaus v. Gleissenthall, Dean Tullsen, Deian Stefan, Tamara Rezk, and Gilles Barthe. 2020. Constant-Time Foundations for the New Spectre Era. In *Proceedings of the 41st ACM SIGPLAN Conference on Programming Language Design and Implementation (PLDI '20)*. ACM.
- [21] Sunjay Cauligi, Craig Disselkoen, Daniel Moghimi, Gilles Barthe, and Deian Stefan. 2022. SoK: Practical Foundations for Software Spectre Defenses. In *Proceedings of the 43rd IEEE Symposium on Security and Privacy (S&P '22)*. IEEE.
- [22] Kevin Cheang, Cameron Rasmussen, Sanjit Seshia, and Pramod Subramanyan. 2019. A Formal Approach to Secure Speculation. In *Proceedings of the 32nd IEEE Computer Security Foundations Symposium (CSF '19)*. IEEE.
- [23] Guoxing Chen, Sanchuan Chen, Yuan Xiao, Yinqian Zhang, Zhiqiang Lin, and Ten H. Lai. 2019. Stealing Intel Secrets from SGX Enclaves via Speculative Execution. In *Proceedings of the 4th IEEE European Symposium on Security and Privacy (EuroS&P '19)*. IEEE.
- [24] Robert J. Colvin and Kirsten Winter. 2019. An Abstract Semantics of Speculative Execution for Reasoning About Security Vulnerabilities. In *Proceedings of the 19th Refinement Workshop (Refine '19)*. Springer.
- [25] Lesly-Ann Daniel, Sébastien Bardin, and Tamara Rezk. 2021. Hunting the Haunter – Efficient relational symbolic execution for Spectre with Haunted ReSe. In *Proceedings of the 28th Annual Network and Distributed System Security Symposium (NDSS '21)*. The Internet Society.
- [26] Leonardo de Moura and Nikolaj Bjørner. 2008. Z3: An Efficient SMT Solver. In *Tools and Algorithms for the Construction and Analysis of Systems*, C. R. Ramakrishnan and Jakob Rehof (Eds.). Springer Berlin Heidelberg, Berlin, Heidelberg, 337–340.

- [27] Craig Disselkoe, Radha Jagadeesan, Alan Jeffrey, and James Riely. 2019. The Code That Never Ran: Modeling Attacks on Speculative Evaluation. In *Proceedings of the 40th IEEE Symposium on Security and Privacy (S&P '19)*. IEEE.
- [28] Jack Doweck, Wen-Fu Kao, Allen Kuan-yu Lu, Julius Mandelblat, Anirudha Rahatekar, Lihu Rappoport, Efraim Rotem, Ahmad Yasin, and Adi Yoaz. 2017. Inside 6th-Generation Intel Core: New Microarchitecture Code-Named Skylake. *IEEE Micro* (2017). <https://doi.org/10.1109/MM.2017.38>
- [29] Goran Doychev, Boris Köpf, Laurent Mauborgne, and Jan Reineke. 2015. CacheAudit: A Tool for the Static Analysis of Cache Side Channels. *ACM Trans. Inf. Syst. Secur.* 18 (2015).
- [30] Xaver Fabian, Marco Guarnieri, Boris Köpf, José F. Morales, Marco Patrignani, Jan Reineke, and Andrés Sánchez. 2025. [Technical Report]. Supplementary Material. Submitted to ACM TOPLAS with this manuscript.
- [31] Xaver Fabian, Marco Patrignani, and Marco Guarnieri. 2022. Automatic Detection of Speculative Execution Combinations. In *Proceedings of the 29th ACM Conference on Computer and Communications Security (CCS 2022)*. ACM.
- [32] Xaver Fabian, Marco Patrignani, and Michael Backes. 2025. Do You Even Lift? Strengthening Compiler Security Guarantees against Spectre Attacks (POPL'25). ACM, New York, NY, USA. <https://doi.org/10.1145/3704867>
- [33] Bo Fu, Leo Tenenbaum, David Adler, Assaf Klein, Arpit Gogia, Alaa R. Alameldeen, Marco Guarnieri, Mark Silberstein, Oleksii Oleksenko, and Gururaj Saileshwar. 2025. AMuLeT: Automated Design-Time Testing of Secure Speculation Countermeasures. In *Proceedings of the 30th ACM International Conference on Architectural Support for Programming Languages and Operating Systems, Volume 2 (ASPLOS '25)*. ACM. <https://doi.org/10.1145/3676641.3716247>
- [34] Google. 2019. SafeSide. <https://github.com/google/safeside>
- [35] Roberto Guanciale, Musard Balliu, and Mads Dam. 2020. InSpectre: Breaking and Fixing Microarchitectural Vulnerabilities by Formal Analysis. In *Proceedings of the 27th ACM SIGSAC Conference on Computer and Communications Security (CCS '20)*. ACM.
- [36] Marco Guarnieri, Boris Köpf, Jan Reineke, and Pepe Vila. 2021. Hardware-Software Contracts for Secure Speculation. In *Proceedings of the 42nd IEEE Symposium on Security and Privacy (S&P '21)*. IEEE.
- [37] Marco Guarnieri, Boris Köpf, José F. Morales, Jan Reineke, and Andrés Sánchez. 2020. Spectector – Automatic detection of speculative information flows. <https://spectector.github.io/>
- [38] Marco Guarnieri, Boris Köpf, José F. Morales, Jan Reineke, and Andrés Sánchez. 2020. Spectector: Principled Detection of Speculative Information Flows. In *Proceedings of the 41st IEEE Symposium on Security and Privacy (S&P '20)*.
- [39] Shengjian Guo, Yueqi Chen, Peng Li, Yueqiang Cheng, Huibo Wang, Meng Wu, and Zhiqiang Zuo. 2020. SpecuSym: Speculative Symbolic Execution for Cache Timing Leak Detection. In *Proceedings of the 42nd ACM/IEEE International Conference on Software Engineering (ICSE '20)*. ACM.
- [40] John L. Hennessy and David A. Patterson. 2011. *Computer Architecture, Fifth Edition: A Quantitative Approach* (5th ed.). Morgan Kaufmann Publishers Inc., San Francisco, CA, USA.
- [41] Manuel V. Hermenegildo, Francisco Bueno, Manuel Carro, Pedro López-García, Edison Mera, José F. Morales, and German Puebla. 2012. An overview of Ciao and its design philosophy. *Theory and Practice of Logic Programming* 12, 1-2 (2012), 219–252.
- [42] Jana Hofmann, Emanuele Vannacci, Cedric Fournet, Boris Köpf, and Oleksii Oleksenko. 2023. Speculation at Fault: Modeling and Testing Microarchitectural Leakage of CPU Exceptions. In *32nd USENIX Security Symposium (USENIX Security 23)*. USENIX Association. <https://www.usenix.org/conference/usenixsecurity23/presentation/hofmann>
- [43] J. Horn. 2018. Speculative execution, variant 4: Speculative store bypass. <https://bugs.chromium.org/p/project-zero/issues/detail?id=1528>. Accessed: 2021-04-11.
- [44] Jaewon Hur, Suhwan Song, Sunwoo Kim, and Byoungyoung Lee. 2022. SpecDoctor: Differential Fuzz Testing to Find Transient Execution Vulnerabilities. In *Proceedings of the 2022 ACM SIGSAC Conference on Computer and Communications Security (Los Angeles, CA, USA) (CCS '22)*. Association for Computing Machinery, New York, NY, USA, 1473–1487. <https://doi.org/10.1145/3548606.3560578>
- [45] Intel. 2018. Retpoline: A Branch Target Injection Mitigation. <https://www.intel.com/content/dam/develop/external/us/en/documents/retpoline-a-branch-target-injection-mitigation.pdf>
- [46] Jason Kim, Jalen Chuang, Daniel Genkin, and Yuval Yarom. 2025. FLOP: Breaking the Apple M3 CPU via False Load Output Predictions. In *USENIX Security*.
- [47] Jason Kim, Daniel Genkin, and Yuval Yarom. 2025. SLAP: Data Speculation Attacks via Load Address Prediction on Apple Silicon. In *S&P*.
- [48] P. Kocher. 2018. Spectre Mitigations in Microsoft's C/C++ Compiler. <https://www.paulkocher.com/doc/MicrosoftCompilerSpectreMitigation.html>. Accessed: 2021-04-11.
- [49] Paul Kocher, Jann Horn, Anders Fogh, Daniel Genkin, Daniel Gruss, Werner Haas, Mike Hamburg, Moritz Lipp, Stefan Mangard, Thomas Prescher, Michael Schwarz, and Yuval Yarom. 2019. Spectre Attacks: Exploiting Speculative Execution. In *Proceedings of the 40th IEEE Symposium on Security and Privacy (S&P '19)*.
- [50] Esmail Mohammadian Koruyeh, Khaled N. Khasawneh, Chengyu Song, and Nael Abu-Ghazaleh. 2018. Spectre Returns! Speculation Attacks Using the Return Stack Buffer. In *Proceedings of the 12th USENIX Workshop on Offensive Technologies (WOOT'18)*. USENIX Association.
- [51] Mikko H. Lipasti, Christopher B. Wilkerson, and John Paul Shen. 1996. Value locality and load value prediction. In *Proceedings of the 7th International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS '96)*. ACM.
- [52] Moritz Lipp, Andreas Kogler, David Oswald, Michael Schwarz, Catherine Easdon, Claudio Canella, and Daniel Gruss. 2021. PLATYPUS: Software-based Power Side-Channel Attacks on x86. In *2021 IEEE Symposium on Security and Privacy (SP)*. 355–371. <https://doi.org/10.1109/SP40001.2021.00063>

- [53] Giorgi Maisuradze and Christian Rossow. 2018. Ret2spec: Speculative Execution Using Return Stack Buffers. In *Proceedings of the 25th ACM SIGSAC Conference on Computer and Communications Security (CCS '18)*. ACM.
- [54] Tiziano Marinaro, Pablo Buiras, Andreas Lindner, Roberto Guanciale, and Hamed Nemati. 2024. Beyond Over-Protection: A Targeted Approach to Spectre Mitigation and Performance Optimization. In *Proceedings of the 19th ACM Asia Conference on Computer and Communications Security (ASIA CCS '24)*. Association for Computing Machinery, New York, NY, USA. <https://doi.org/10.1145/3634737.3637651>
- [55] Ross McIlroy, Jaroslav Sevcik, Tobias Tebbi, Ben L. Titzer, and Toon Verwaest. 2019. Spectre is here to stay: An analysis of side-channels and speculative execution. (2019). arXiv:1902.05178
- [56] Sparsh Mittal. 2017. A survey of value prediction techniques for leveraging value locality. *Concurrency and computation: practice and experience* (2017).
- [57] David Molnar, Matt Piotrowski, David Schultz, and David Wagner. 2005. The program counter security model: automatic detection and removal of control-flow side channel attacks. In *Proceedings of the 8th International Conference on Information Security and Cryptology (Seoul, Korea) (ICISC'05)*. Springer.
- [58] N. Mosier, H. Nemati, J. C. Mitchell, and C. Trippel. 2024. Serberus: Protecting Cryptographic Code from Spectres at Compile-Time. In *2024 IEEE Symposium on Security and Privacy (SP)*. IEEE Computer Society, Los Alamitos, CA, USA, 48–48. <https://doi.org/10.1109/SP54263.2024.00048>
- [59] Hamed Nemati, Pablo Buiras, Andreas Lindner, Roberto Guanciale, and Swen Jacobs. 2020. Validation of Abstract Side-Channel Models for Computer Architectures. In *Computer Aided Verification: 32nd International Conference, CAV 2020, Los Angeles, CA, USA, July 21–24, 2020, Proceedings, Part I*. Springer-Verlag, Berlin, Heidelberg. https://doi.org/10.1007/978-3-030-53288-8_12
- [60] Oleksii Oleksenko, Christof Fetzner, Boris Köpf, and Mark Silberstein. 2022. Revizor: Testing Black-Box CPUs against Speculation Contracts. In *Proceedings of the 27th ACM International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS '22)*. ACM.
- [61] Oleksii Oleksenko, Marco Guarnieri, Boris Köpf, and Mark Silberstein. 2023. Hide and Seek with Spectres: Efficient discovery of speculative information leaks with random testing. In *Proceedings of the 44th IEEE Symposium on Security and Privacy (S&P 2023)*. IEEE.
- [62] Oleksii Oleksenko, Flavien Solt, Cédric Fournet, Jana Hofmann, Boris Köpf, and Stavros Volos. 2026. Enter, Exit, Page Fault, Leak: Testing Isolation Boundaries for Microarchitectural Leaks. In *Proceedings of the 47th IEEE Symposium on Security and Privacy (S&P)*. <https://arxiv.org/abs/2507.06039> To appear.
- [63] Oleksii Oleksenko, Bohdan Trach, Mark Silberstein, and Christof Fetzner. 2020. SpecFuzz: Bringing Spectre-type vulnerabilities to the surface. In *29th USENIX Security Symposium (USENIX Security 20)*. USENIX Association, 1481–1498. <https://www.usenix.org/conference/usenixsecurity20/presentation/oleksenko>
- [64] Marco Patrignani and Marco Guarnieri. 2021. Exorcising Spectres with Secure Compilers. In *Proceedings of the 28th ACM Conference on Computer and Communications Security (CCS '21)*. ACM.
- [65] Hernán Ponce de León and Johannes Kinder. 2022. Cats vs. Spectre: An Axiomatic Approach to Modeling Speculative Execution Attacks. In *Proceedings of the 43rd IEEE Symposium on Security and Privacy (S&P '22)*. IEEE.
- [66] A. Sabelfeld and D. Sands. 2005. Dimensions and principles of declassification. In *18th IEEE Computer Security Foundations Workshop (CSFW'05)*. 255–269. <https://doi.org/10.1109/CSFW.2005.15>
- [67] Michael Schwarz, Martin Schwarzl, Moritz Lipp, Jon Masters, and Daniel Gruss. 2019. NetSpectre: Read Arbitrary Memory over Network. In *Proceedings of the 24th European Symposium on Research in Computer Security (ESORICS '19)*. Springer.
- [68] Vedvyas Shanhogue, Deepak Gupta, and Ravi Sahita. 2019. Security Analysis of Processor Instruction Set Architecture for Enforcing Control-Flow Integrity. In *Proceedings of the 8th International Workshop on Hardware and Architectural Support for Security and Privacy (HASP '19)*. ACM.
- [69] Rami Sheikh, Harold W. Cain, and Raguram Damodaran. 2017. Load value prediction via path-based address prediction: Avoiding mispredictions due to conflicting stores. In *Proceedings of the 50th Annual IEEE/ACM International Symposium on Microarchitecture (MICRO '17)*. ACM.
- [70] Basavesh Ammanaghatta Shivakumar, Jack Barnes, Gilles Barthe, Sunjay Cauligi, Chitchanok Chuengsatiansup, Daniel Genkin, Sioli O'Connell, Peter Schwabe, Rui Qi Sim, and Yuval Yarom. 2023. Spectre Declassified: Reading from the Right Place at the Wrong Time. In *Proceedings of the 44th IEEE Symposium on Security and Privacy (S&P '23)*. IEEE.
- [71] Basavesh Ammanaghatta Shivakumar, Gilles Barthe, Benjamin Gregoire, Vincent Laporte, Tiago Oliveira, Swarn Priya, Peter Schwabe, and Lucas Tabary-Maujean. 2023. Typing High-Speed Cryptography against Spectre v1. In *2023 IEEE Symposium on Security and Privacy (SP)*.
- [72] Julian Stecklina and Thomas Prescher. 2018. LazyFP: Leaking FPU Register State using Microarchitectural Side-Channels. *CoRR* (2018). arXiv:1806.07480
- [73] Caroline Trippel, Daniel Lustig, and Margaret Martonosi. 2018. MeltdownPrime and SpectrePrime: Automatically-Synthesized Attacks Exploiting Invalidation-Based Coherence Protocols. *CoRR* (2018). arXiv:1802.03802
- [74] Daniel Trujillo, Johannes Wikner, and Kaveh Razavi. 2023. Inception: Exposing New Attack Surfaces with Training in Transient Execution. In *32nd USENIX Security Symposium (USENIX Security 23)*. USENIX Association, Anaheim, CA. <https://www.usenix.org/conference/usenixsecurity23/presentation/trujillo>
- [75] Marco Vassena, Craig Disselkoen, Klaus von Gleissenthall, Sunjay Cauligi, Rami Gökhan Kıcı, Ranjit Jhala, Dean Tullsen, and Deian Stefan. 2021. Automatically Eliminating Speculative Leaks from Cryptographic Code with Blade. *Proceedings of the ACM on Programming Languages* 5, POPL (2021).

- [76] Guanhua Wang, Sudipta Chattopadhyay, Arnab Kumar Biswas, Tulika Mitra, and Abhik Roychoudhury. 2020. KLEESpectre: Detecting Information Leakage through Speculative Cache Attacks via Symbolic Execution. *ACM Transactions on Software Engineering and Methodology* 29, 3 (2020).
- [77] Guanhua Wang, Sudipta Chattopadhyay, Ivan Gotovchits, Tulika Mitra, and Abhik Roychoudhury. 2021. oo7: Low-Overhead Defense Against Spectre Attacks via Program Analysis. *IEEE Transactions on Software Engineering* 47, 11 (2021).
- [78] Johannes Wikner and Kaveh Razavi. 2022. RETBLEED: Arbitrary Speculative Code Execution with Return Instructions. In *Proceedings of the 31st USENIX Security Symposium (USENIX Security '22)*. USENIX Association.
- [79] Johannes Wikner, Daniël Trujillo, and Kaveh Razavi. 2023. Phantom: Exploiting Decoder-detectable Mispredictions. In *MICRO*. Paper=https://comsec.ethz.ch/wp-content/files/phantom_micro23.pdfURL=<https://comsec.ethz.ch/research/microarch/inception>
- [80] Meng Wu and Chao Wang. 2019. Abstract interpretation under speculative execution. In *Proceedings of the 40th ACM SIGPLAN Conference on Programming Language Design and Implementation (PLDI' 19)*. ACM, New York, NY, USA. <https://doi.org/10.1145/3314221.3314647>
- [81] Tao Zhang, Kenneth Koltermann, and Dmitry Evtvushkin. 2020. Exploring Branch Predictors for Constructing Transient Execution Trojans. In *Proceedings of the 25th International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS '20)*. ACM.
- [82] Zhiyuan Zhang, Gilles Barthe, Chitchanok Chuengsatiansup, Peter Schwabe, and Yuval Yarom. 2023. Ultimate SLH: Taking Speculative Load Hardening to the Next Level. In *Proceedings of the 32nd USENIX Security Symposium*. USENIX.

A CODE FROM CASE STUDIES

A.1 Example #8

In Example #8, the bounds check of Listing 1 is implemented using a conditional operator:

```
1  temp  &=  B[A[y<size?(y+1):0]*512];
```

When compiling the example without countermeasures or optimizations, the conditional operator is translated to a branch instruction (cf. line 4), which is a source of speculation. Hence, the resulting program contains a speculative leak, which SPECTECTOR correctly detects.

```
1  mov    size, %rcx
2  mov    y, %rax
3  cmp    %rcx, %rax
4  jae    .L1
5  add    $1, %rax
6  jmp    .L2
7
8  xor    %rax, %rax
9  jmp    .L2
10
11 mov    A(%rax), %rax
12 shl    $9, %rax
13 mov    B(%rax), %rax
14 mov    temp, %rcx
15 and    %rax, %rcx
16 mov    %rcx, temp
```

In the UNP -O2 mode, the conditional operator is translated as a conditional move (cf. line 6), for which SPECTECTOR can prove security.

```
1  mov    size, %rax
2  mov    y, %rdx
3  xor    %rcx, %rcx
4  cmp    %rdx, %rax
5  lea    1(%rdx), %rax
6  cmova  %rax, %rcx
7  mov    A(%rcx), %rax
8  shl    $9, %rax
9  mov    B(%rax), %rax
10 and    %rax, temp
```

A.2 Example #15 in SLH mode

Here, the adversary provides the input via the pointer `*y`:

```

1  if (*y < size)
2      temp &= B[A[*y] * 512];

```

In the `-O0 SLH` mode, CLANG hardens the address used for performing the memory access `A[*y]` in lines 8–12, but not the resulting value, which is stored in the register `%cx`. However, the value stored in `%cx` is used to perform a second memory access at line 14. An adversary can exploit the second memory access to speculatively leak the content of `A[0xFF...FF]`. In our experiments, SPECTECTOR correctly detected such leak.

```

1  mov    $0, %rax
2  mov    y, %rdx
3  mov    (%rdx), %rsi
4  mov    size, %rdx
5  cmp    %rdx, %rsi
6  jae    END
7  cmovae $-1, %rax
8  mov    y, %rcx
9  mov    (%rcx), %rcx
10 mov    %rax, %rdx
11 or     %rcx, %rdx
12 mov    A(%rdx), %rcx
13 shl    $9, %rcx
14 mov    B(%rcx), %rcx
15 mov    temp, %rdx
16 and    %rcx, %rdx
17 mov    %rdx, temp

```

In contrast, when Example #15 is compiled with the `-O2` flag, CLANG correctly hardens `A[*y]`'s result (cf. line 10). This prevents information from flowing into the microarchitectural state during speculative execution. Indeed, SPECTECTOR proves that the program satisfies speculative non-interference.

```

1  mov    $0, %rax
2  mov    y, %rdx
3  mov    (%rdx), %rdx
4  mov    size, %rsi
5  cmp    %rsi, %rdx
6  jae    END
7  cmovae $-1, %rax
8  mov    A(%rdx), %rcx
9  shl    $9, %rcx
10 or     %rax, %rcx

```

```

3017 11  mov    B(%rcx), %rcx
3018 12  or     %rax, %rcx
3019 13  and    %rcx, temp
3020

```

B VARIANTS OF SPECTECTOR

In Section 11.5 we discussed the variation of the SPECTECTOR algorithm to check for GNI. Here, we will present the full algorithm as well as the algorithm for checking sequential constant-time:

Algorithm 4 SPECTECTOR SeqCT

Input: A program p , a security policy P , a speculative window $\omega \in \mathbb{N}$.

Output: SECURE if p satisfies SeqCT with respect to the policy P ; INSECURE otherwise

```

1: procedure SPECTECTOR( $p, P, \omega$ )
2:   for each symbolic run  $\bar{\tau} \in Beh_x^S(p)$  do
3:     if MEMLEAK( $\bar{\tau}, P$ )  $\vee$  CTRLLEAK( $\bar{\tau}, P$ ) then
4:       return INSECURE
5:   return SECURE

6: procedure MEMLEAK( $\bar{\tau}, P$ )
7:    $\psi \leftarrow pthCnd(\bar{\tau})_{1\wedge 2} \wedge polEqv(P) \wedge$ 
       $\neg obsEqv(\bar{\tau} \upharpoonright_{ns})$ 
8:   return SATISFIABLE( $\psi$ )

9: procedure CTRLLEAK( $\bar{\tau}, P$ )
10:  for each prefix  $v \cdot symPc(se)$  of  $\bar{\tau} \upharpoonright_{ns}$  do
11:     $\psi \leftarrow pthCnd(\bar{\tau} \upharpoonright_{ns} \cdot v)_{1\wedge 2} \wedge polEqv(P) \wedge$ 
       $\neg sameSymbPc(se)$ 
12:    if SATISFIABLE( $\psi$ ) then
13:      return  $\top$ 
14:  return  $\perp$ 

```

Algorithm 5 SPECTECTOR GNI

Input: A program p , a security policy P , a speculative window $\omega \in \mathbb{N}$.

Output: SECURE if p satisfies general non-interference with respect to the policy P and speculative semantics $\mathcal{Q}_x$;
INSECURE otherwise

```

1: procedure SPECTECTOR( $p, P, \omega$ )
2:   for each symbolic run  $\bar{\tau} \in \text{Beh}_x^S(p)$  do
3:     if MEMLEAK( $\bar{\tau}, P$ )  $\vee$  CTRLLEAK( $\bar{\tau}, P$ ) then
4:       return INSECURE
5:   return SECURE

6: procedure MEMLEAK( $\bar{\tau}, P$ )
7:    $\psi \leftarrow \text{pthCnd}(\bar{\tau})_{1\wedge 2} \wedge \text{polEqv}(P) \wedge$ 
       $\neg \text{obsEqv}(\bar{\tau})$ 
8:   return SATISFIABLE( $\psi$ )

9: procedure CTRLLEAK( $\bar{\tau}, P$ )
10:  for each prefix  $v \cdot \text{symPc}(se)$  of  $\bar{\tau}$  do
11:     $\psi \leftarrow \text{pthCnd}(\bar{\tau}|_{ns} \cdot v)_{1\wedge 2} \wedge \text{polEqv}(P)$ 
       $\wedge \neg \text{sameSymbPc}(se)$ 
12:    if SATISFIABLE( $\psi$ ) then
13:      return  $\top$ 
14:  return  $\perp$ 

```
