



A Better Facet of Dynamic Information Flow Control

Minh Ngo, Nataliia Bielova, Cormac Flanagan, Tamara Rezk, Alejandro Russo, Thomas Schmitz

► To cite this version:

Minh Ngo, Nataliia Bielova, Cormac Flanagan, Tamara Rezk, Alejandro Russo, et al.. A Better Facet of Dynamic Information Flow Control. WWW '18 Companion: The 2018 Web Conference Companion, Apr 2018, Lyon, France. pp.1-9. hal-01723723

HAL Id: hal-01723723

<https://inria.hal.science/hal-01723723>

Submitted on 5 Mar 2018

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

A Better Facet of Dynamic Information Flow Control

Minh Ngo
INRIA, France
nguyen-nhat-minh.ngo@inria.fr

Nataliia Bielova
INRIA, France
nataliia.bielova@inria.fr

Cormac Flanagan
UCSC, USA
cormac@ucsc.edu

Tamara Rezk
INRIA, France
tamara.rezk@inria.fr

Alejandro Russo
Chalmers University of Technology,
Sweden
russo@chalmers.se

Thomas Schmitz
UCSC, USA
tschmitz@ucsc.edu

ABSTRACT

Multiple Facets (MF) is a dynamic enforcement mechanism which has proved to be a good fit for implementing information flow security for JavaScript. It relies on multi executing the program, once per each security level or view, to achieve soundness. By looking inside programs, MF encodes the views to reduce the number of needed multi-executions.

In this work, we extend Multiple Facets in three directions. First, we propose a new version of MF for arbitrary lattices, called Generalised Multiple Facets, or GMF. GMF strictly generalizes MF, which was originally proposed for a specific lattice of principals. Second, we propose a new optimization on top of GMF that further reduces the number of executions. Third, we strengthen the security guarantees provided by Multiple Facets by proposing a termination sensitive version that eliminates covert channels due to termination.

KEYWORDS

Multiple Facets; Dynamic Information Flow Control; Secure Multi-Execution; Noninterference

1 INTRODUCTION

JavaScript has become the de facto programming language of the Web. Web browsers daily execute thousands of JavaScript lines which usually have access to confidential information, for example cookies that mark that the user in a web session is authenticated. It is not surprising that JavaScript is a common target for attacks. While browsers deploy security measures in the form of access control (e.g., SOP and CSP), they are insufficient [12, 17, 30] to protect confidentiality of data.

Information flow control (IFC) is a promising technology which provides a systematic solution to handle unintentional or malicious leaks of confidential information. Recently, dynamic IFC analyses have received a lot of attention [1–3, 5, 7, 9, 10, 14, 26, 33], due, in part, to its applicability to JavaScript—where static analyses are rather an awkward fit [29].

In order to scale, a suitable IFC technique for the web not only needs to be dynamic but also needs to reduce to the minimum the modifications required to existing JavaScript code. In this light, an interesting dynamic IFC technique which fulfills both of these requirements consists in executing several copies of a program: one execution per each security level or view. In that manner, each copy of the program (view) depends only on information observable to the corresponding security level, where no leaks are therefore

possible. Secure Multi Execution (SME) [14] and Multiple Facets (MF) [3] are two techniques based on this idea.

Both techniques have been proved to be a good fit for information flow security in the web since they have been successfully implemented as extensions of the Firefox browser [13, 33].

Although both SME and MF are based on multi-executions, they present important differences [7]. On one hand, SME is black-box [24], i.e., it is a mechanism that does not *look inside* programs but rather change the semantics of inputs and outputs to ensure security. For a moment, we assume a scenario where security levels are simply sets of principals (e.g., web origins) which denote those authorities with confidentiality concerns over data. In such a scenario, SME needs to spawn one execution for any possible set of principals—where the number of executions grows exponentially with respect to the number of principals! Instead, MF [3] is designed to reduce the number of multi-executions and the memory footprint of SME. It does so by inspecting programs code and multi-executing instructions and multiplexing memory only when needed. While MF is more resource-friendly than SME, SME provides stronger security guarantees when it comes to leaks via abnormal termination [7].

Our broad goal is to augment the efficiency of techniques based on MF and SME to general cases. In particular, we discovered that MF might sometimes spawn more multi-executions than SME—something that is counter-intuitive when considering the purpose of MF (see Section 2). Our first contribution consists on a novel technique to further reduce the number of multi-executions (and memory footprint) of MF. Our second contribution is to generalize MF to work for *arbitrary finite lattices* (see Section 3) rather than being restricted to the security lattice of principals as in the original proposal [3]. This becomes useful when, for instance, a program depends on 5 security levels. In such case, as stated originally, MF will need to encode them by using (at least) 3 principals ($2^3 > 5$), and thus execute the program $2^3 = 8$ times, while SME will execute it only 5 times (one per security level). Finally, we combine MF and SME into a single new dynamic IFC mechanism in order to provide security guarantees as strong as SME (i.e., termination sensitive non-interference) while avoiding multi-executions as much as our optimized version of MF allows it. All proofs can be found in [23].

2 BACKGROUND ON SME AND MF

In this section, we discuss how on one hand, the underpinning mechanism in MF reduces the number of executions compared to SME, and on the other hand, may run more multi executions

$$\begin{array}{c}
\text{SKIP} \frac{}{(\text{skip}, \mu) \Downarrow \mu} \quad \text{ASSIGN} \frac{v = \mu(e)}{(x := e, \mu) \Downarrow \mu[x \mapsto v]} \\
\text{IF} \frac{\mu(e) = v \quad (P_v, \mu) \Downarrow \mu'}{(\text{if } e \text{ then } P_{\text{tt}} \text{ else } P_{\text{ff}}, \mu) \Downarrow \mu'} \quad \text{SEQ} \frac{(P_1, \mu) \Downarrow \mu' \quad (P_2, \mu') \Downarrow \mu''}{(P_1; P_2, \mu) \Downarrow \mu''} \\
\text{WHILE} \frac{(\text{if } e \text{ then } P; \text{while } e \text{ do } P \text{ else skip}, \mu) \Downarrow \mu'}{(\text{while } e \text{ do } P, \mu) \Downarrow \mu'}
\end{array}$$

Figure 1: Language semantics

than SME because of the security lattice based on principals. Our goal here is partly pedagogical and partly to motivate and provide intuition on the optimization proposed in Section 4.

Language and Semantics To investigate the foundation of multiple facets, we use a simple, deterministic while language. Its syntax includes programs P , variables x , expressions e , and values v . We use the symbol $\oplus$ for binary expression operators. A value is either an integer value or a boolean value.

(programs) $P ::= \text{skip} \mid x := e \mid \text{if } e \text{ then } P_1 \text{ else } P_2 \mid \text{while } e \text{ do } P \mid P_1; P_2$
(expressions) $e ::= v \mid x \mid e \oplus e$

Figure 1 presents standard big-step semantics of the language. Memories μ map variables to values; we overload the notation of memory and use $\mu(e)$ as the evaluation function for expression e in memory μ , where $\mu(v) = v$ and $\mu(e_1 \oplus e_2) = \mu(e_1) \oplus \mu(e_2)$. We write $(P, \mu) \Downarrow \mu'$ to mean that the evaluation of program P on memory μ terminates with memory μ' . We use $\mu[x \mapsto v]$ for the memory μ' where $\mu'(y) = \mu(y)$ if $y \neq x$, and $\mu'(y) = v$ if $y = x$.

MF may use fewer resources than SME SME [14] multi executes programs, in a blackbox manner, as many times as security levels in a lattice. Let's define an SME memory as a function that maps each variable to an array of values, one value per security level. For the sake of simplicity, let's consider first a security lattice with only two elements H and L where $H \not\sqsubseteq L$ is the only disallowed flow. Thus, an SME memory $\hat{\mu}$ maps variables to an array of 2 (possibly different) values: one corresponding to the H view and one corresponding to the L view. Let's denote such array of values as $\langle v_1 : v_2 \rangle$, where v_1 is a private, H , view and v_2 is a public, L , view. Assume that $H(\hat{\mu})$ (resp. $L(\hat{\mu})$) is a memory in the standard semantics, obtained by projection of $\hat{\mu}$, mapping variables to single values of the high view (resp. low view). Then, the SME monitoring rule¹ for such a language can be given by the relation $\Downarrow_{\text{SME-TINI}}$ as follows:

$$\text{SME-TINI} \frac{(P, H(\hat{\mu})) \Downarrow \mu_1 \quad (P, L(\hat{\mu})) \Downarrow \mu_2}{(P, \hat{\mu}) \Downarrow_{\text{SME-TINI}} \mu_1 \odot \mu_2}$$

where $\odot$ combines two normal memories into a SME memory in such a way that $H(\mu_1 \odot \mu_2) = \mu_1$ and $L(\mu_1 \odot \mu_2) = \mu_2$. The SME mechanism will blindly execute the program as many times as possible views (or positions of the array) may exist.

Consider a program $h := l$ where initial views for variables l and h are given by: $\hat{\mu}(h) = \langle 1 : 0 \rangle$ and $\hat{\mu}(l) = \langle 1 : 1 \rangle$. In SME,

¹We give here the termination insensitive version of SME.

using the SME-TINI rule, the assignment will be executed twice: once with $H(\hat{\mu}) = [h \mapsto 1, l \mapsto 1]$ for the high view and once with $L(\hat{\mu}) = [h \mapsto 0, l \mapsto 1]$ for the low view. After execution, the final SME memory will map h to $\langle 1 : 1 \rangle$. One way to reduce the number of executions is to exploit the knowledge that the high and the low view for variable l are equal, i.e., $H(\hat{\mu})(l) = L(\hat{\mu})(l)$. Since the semantics is deterministic, there is no need to execute the program twice. We can use this knowledge by specialising SME at the granularity of commands and include the following assignment rule:

$$\text{SME-OPTIM} \frac{H(\hat{\mu})(e) = L(\hat{\mu})(e) \quad (x := e, L(\hat{\mu})) \Downarrow \mu}{(x := e, \hat{\mu}) \Downarrow_{\text{SME}} \hat{\mu}[x \mapsto \langle \mu(x), \mu(x) \rangle]}$$

Notice that this SME optimization requires to *look inside* the shape of the program to evaluate if expression e of an assignment satisfies the hypothesis.

In general, in order to reduce the number of executions using the multi-execution technique of SME-TINI, it is sufficient to (i) identify in an SME memory which values in the array of values are equal and (ii) remember which values correspond to which views. MF uses the multi-execution technique, implements (i) and (ii) and hence, reduces the number of executions. MF encodes values in SME memories (arrays with as many positions as lattice elements) as ordered binary trees, where the order is given by the elements of the lattice. For example, for a SME memory where $\hat{\mu}(h) = \langle 1 : 0 : 0 : 0 \rangle$ for a lattice of 4 elements with top element $\top$, an equivalent MF memory encodes this array as $\langle \top ? 1 : 0 \rangle$ with the meaning that 1 is the view for $\top$ and 0 for the rest. Every execution that depends on that value, will multi execute twice instead of 4 times as in SME.

Moreover, MF further uses the view information provided by the encoding in order to multi execute less in case of branching commands. For example, for SME-TINI with SME memory $\hat{\mu}(h) = \langle 1 : 0 : 0 : 0 \rangle$ the program:

```

1: if h = 0 then
2:   h := h + 1

```

executes 4 times (where the assignment at line 2 executes 3 times).

Using the MF memory encoding $\hat{\mu}(h) = \langle \top ? 1 : 0 \rangle$, MF remembers that at line 2 there is no possible observation for the view $\top$ (because for view $\top$ the value of h is 1 so it doesn't take the then branch). Hence, the assignment $h := h + 1$ only executes once with a memory where h is 0 (the view of variable h corresponding to the 3 levels which are not $\top$).

For a program $h := l$, where $\hat{\mu}(l)$ is $\langle 1 : 1 \rangle$ in SME, MF keeps only the value 1: a single value represents the fact that all views can observe the same value. Thus the assignment $h := l$ executes once (and all future executions dependent on h will also be reduced).

Hence when encoding of an SME memory can be reduced effectively, multi executions are reduced accordingly. As shown in the following sections, preservation of MF memories encoding through execution requires: to represent arrays of values as trees called faceted values and to evaluate expressions depending on faceted values. In particular, the definition of the evaluation of

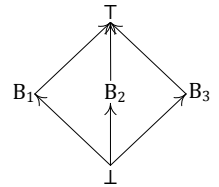


Figure 2: Lattice $\langle \mathcal{L}_B, \sqsubseteq \rangle$

expressions on faceted values depends highly on the shape of expressions and their values according to different views, and thus is contradictory to the blackbox property of a monitor.

MF may run more multi executions than SME Original MF has one limitation with respect to SME: it was designed only for a security lattice of principals: for n principals, such a lattice contains 2^n security levels. The following Ad Exchange platform [35] example demonstrates that MF may be less efficient than SME in practice, when the security lattice is not based on principals.

Example 2.1. An Ad Exchange platform needs to put an advertisement on a publisher's website. For that, it implements a Real-time Bidding (RTB) system [36], where advertisers can bid for the space on the publisher's website to get their ad published. The system receives as input all the bid offers from bidders and sorts them. According to the RTB algorithm, the second best offer wins.

We present the lattice of 5 elements for this example in Fig. 2. For simplicity, we consider only 3 bidders called B_1 , B_2 , and B_3 , an Ad Exchange ($\top$ level) which is able to see all the bids, and a public view $\perp$. Because MF is designed for a principal lattice, to encode 5 security levels, it uses 3 principals k_1 , k_2 , and k_3 , and create a lattice of $8 = 2^3$ levels, and thus has a potential to run some parts of the program 8 times, while SME always executes the program 5 times.

We consider one test that naively checks the order of bid offers and decides the winner. The encoding of the lattice is: $\top = \{k_1, k_2, k_3\}$, $B_i = \{k_i\}$, and $\perp = \emptyset$.

- 1: *winner* := 0;
- 2: *test* := ($x_1 \leq x_2$) and ($x_2 \leq x_3$);
- 3: **if** *test* **then** *winner* := 2 **else** **skip**

The bid values from bidders are $x_1 = \langle k_1 ? 10 : 0 \rangle$, $x_2 = \langle k_2 ? 5 : 0 \rangle$, and $x_3 = \langle k_3 ? 7 : 0 \rangle$. Thus, the resulting value of *test* at line 2 is

$$\langle k_1 ? \langle k_2 ? \langle k_3 ? \text{ff} : \text{ff} \rangle : \langle k_3 ? \text{ff} : \text{ff} \rangle \rangle : \langle k_2 ? \langle k_3 ? \text{tt} : \text{ff} \rangle : \langle k_3 ? \text{tt} : \text{tt} \rangle \rangle.$$

Therefore, the original MF executes the if instruction 8 times with 3 useless executions for levels $\{k_1, k_2\}$, $\{k_2, k_3\}$, and $\{k_1, k_3\}$.

Moreover, because different views of a variable may contain the same values, MF may execute the same statement several times. For example, in the execution described above, original MF executes the then branch 3 times, while it only needs to run once since the three executions for the then branch can be merged into one.

3 MF FOR ARBITRARY SECURITY LATTICE

We present an extension to the original Multiple Facets mechanism [3] for an arbitrary security lattice $\langle \mathcal{L}, \sqsubseteq \rangle$, which we call Generalised Multiple Facets mechanism, or *GMF*. Similarly to Multiple Facets, GMF operates over a *faceted memory* $\hat{\mu}$ that maps variables to simple values or faceted values. A *faceted value* is of the form $\langle l ? V_1 : V_2 \rangle$ where $l \in \mathcal{L}$ is a security level, and V_i can be either a faceted value or a simple value. The first facet V_1 of $\langle l ? V_1 : V_2 \rangle$ is called *private*, and visible to the observers at security level l or higher levels in the lattice; the second facet V_2 is called *public*, and visible to security levels that are lower or incomparable to l . We use V as a meta-variable for faceted values or simple values. Every evaluation in GMF (see Fig. 6) is marked with a set of security levels pc , for which the current computation is visible.

3.1 Expression evaluation

$$\begin{aligned} \hat{\mu}^{pc}(v) &= v \\ \hat{\mu}^{pc}(x) &= \ominus^{pc}(\hat{\mu}(x)) \\ \hat{\mu}^{pc}(e_1 \oplus e_2) &= \hat{\mu}^{pc}(e_1) \oplus^{pc} \hat{\mu}^{pc}(e_2) \\ \ominus^{pc}(v) &= v \\ \ominus^{pc}(\langle l ? V_1 : V_2 \rangle) &= \begin{cases} \ominus^{pc}(V_1) & \text{if } l \leq pc \\ \ominus^{pc}(V_2) & \text{if } l \not\leq pc \\ \langle l ? \ominus^{pc_1}(V_1) : \ominus^{pc_2}(V_2) \rangle & \text{otherwise} \end{cases} \\ v_1 \oplus^{pc} v_2 &= v_1 \oplus v_2 \\ v \oplus^{pc} \langle l ? V_1 : V_2 \rangle &= \begin{cases} v \oplus^{pc} V_1 & \text{if } l \leq pc \\ v \oplus^{pc} V_2 & \text{if } l \not\leq pc \\ \langle l ? (v \oplus^{pc_1} V_1) : (v \oplus^{pc_2} V_2) \rangle & \text{otherwise} \end{cases} \\ \langle l ? V_1 : V_2 \rangle \oplus^{pc} V &= \begin{cases} V_1 \oplus^{pc} V & \text{if } l \leq pc \\ V_2 \oplus^{pc} V & \text{if } l \not\leq pc \\ \langle l ? (V_1 \oplus^{pc_1} V) : (V_2 \oplus^{pc_2} V) \rangle & \text{otherwise} \end{cases} \end{aligned}$$

where $pc_1 = \{l' \in pc \mid l \sqsubseteq l'\}$ and $pc_2 = pc \setminus pc_1$.

Figure 4: Expression evaluation

By $\hat{\mu}^{pc}(e)$ we denote the evaluation of expression e in faceted memory $\hat{\mu}$ with set of security levels pc . The definition of $\hat{\mu}^{pc}(e)$ is presented in Fig. 4. For example, consider the evaluation of x when the faceted value x in memory $\hat{\mu}$ is $\langle l ? V_1 : V_2 \rangle$. To define which facet is useful given a pc , we consider the following cases:

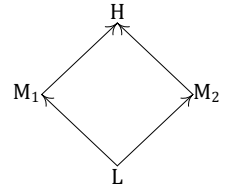


Figure 3: Lattice $\langle \mathcal{L}, \sqsubseteq \rangle$

- All the levels in pc are greater than or equal to l , denoted $l \leq pc$ (i.e. $\forall l' \in pc. l \sqsubseteq l'$): the evaluation can use the private facet V_1 because the public facet V_2 is anyway not useful for every level in this pc .
- All the levels in pc are lower than or incomparable to l , denoted $l \not\leq pc$ (i.e. $\forall l' \in pc. l \not\sqsubseteq l'$): the evaluation can only use the public facet V_2 because V_2 is a facet visible to any view that is lower than or incomparable to l .
- Otherwise, we say that l and pc are *incomparable* and denote it by $l \parallel pc$ (i.e. $\exists l', l'' \in pc. l \sqsubseteq l' \wedge l \not\sqsubseteq l''$): we first evaluate V_1 with $pc_1 = \{l' \in pc \mid l \sqsubseteq l'\}$ – the set of all levels in pc which are greater than or equal to l . Then, we evaluate V_2 with $pc_2 = pc \setminus pc_1$ which is the set of all levels in pc which are lower than or incomparable to l . Finally, we combine the two results in a new faceted value.

To evaluate a variable x , we use a special unary operator $\ominus^{pc}(\hat{\mu}(x))$, which returns the value that is visible to all the levels in the pc . Let's consider the case of $\ominus^{pc}(\langle l ? V_1 : V_2 \rangle)$. Notice that, if pc and l are incomparable, meaning that there are some levels in pc that are higher than or equal to l and other levels in pc that are lower than or incomparable to l , denoted by $l \parallel pc$, then the evaluation returns the faceted value $\langle l ? \ominus^{pc_1}(V_1) : \ominus^{pc_2}(V_2) \rangle$. The form of the result of $\hat{\mu}^{pc}(e)$ is described in Lemma 3.1.

LEMMA 3.1. *If $\hat{\mu}^{pc}(e) = \langle l ? V_1 : V_2 \rangle$, then $l \parallel pc$.*

$$\mu \uparrow_{\Gamma}^{def}(x) = \begin{cases} \mu(x) & \text{if } \Gamma(x) = \text{glb}(\mathcal{L}), \\ \langle \Gamma(x) ? \mu(x) : def(x) \rangle & \text{otherwise.} \end{cases}$$

$$\hat{\mu}|_{\Gamma}(x) = l(\hat{\mu})(x) \text{ where } l = \Gamma(x)$$

Figure 5: Functions for faceted and normal memories.

Example 3.2 (Expression evaluation). Consider the lattice $\langle \mathcal{L}, \sqsubseteq \rangle$ from Fig. 3, and the evaluation of $x + y$ in $\hat{\mu}$, where $\hat{\mu}(x) = \langle M_1 ? 10 : 0 \rangle$ and $\hat{\mu}(y) = \langle M_2 ? 5 : 0 \rangle$.

Suppose that $pc = \{M_1, H\}$. Since all the levels in pc are higher than or equal to M_1 , the evaluation of x returns $\hat{\mu}^{pc}(x) = 10$. Since pc and M_2 are incomparable, the evaluation of y returns $\hat{\mu}^{pc}(y) = \langle M_2 ? 5 : 0 \rangle$. Next, the evaluation of $10 +^{pc} \langle M_2 ? 5 : 0 \rangle$ is split into two: one uses a facet visible to M_2 (and hence H), and another one uses a public facet that will be visible to M_1 .

$$\begin{aligned} \hat{\mu}^{pc}(x + y) &= \hat{\mu}^{pc}(x) +^{pc} \hat{\mu}^{pc}(y) \\ &= \ominus^{pc}(\langle M_1 ? 10 : 0 \rangle) +^{pc} \ominus^{pc}(\langle M_2 ? 5 : 0 \rangle) \\ &= 10 +^{pc} \langle M_2 ? 5 : 0 \rangle = \langle M_2 ? 10 + \{H\} 5 : 10 + \{M_1\} 0 \rangle \\ &= \langle M_2 ? 15 : 10 \rangle \end{aligned}$$

3.2 Semantics

We abuse the notation and use l as a *projection function* on simple values, faceted values and faceted memories. For any V , $l(V)$ returns the value in V which is visible to users at level l . For any $\hat{\mu}$, $l(\hat{\mu})$ returns the memory in $\hat{\mu}$ which is visible to users at level l .

$$l(v) = v \quad l(\langle l_1 ? V_1 : V_2 \rangle) = \begin{cases} l(V_1) & \text{if } l_1 \sqsubseteq l, \\ l(V_2) & \text{otherwise.} \end{cases}$$

$$l(\hat{\mu})(x) = l(\hat{\mu}(x))$$

The projection function l is used in the definition of $\hat{\mu}|_{\Gamma}$ function that converts a faceted memory to a simple memory (see Fig. 5).

The semantics of GMF is defined in Fig. 6 as a big-step evaluation relation $\Gamma \vdash (P, \mu) \Downarrow_{GMF} \mu'$, where program P is executed in a memory μ and a security environment Γ that maps variables to security levels in a given security lattice $\langle \mathcal{L}, \sqsubseteq \rangle$.

The main rule GMF first constructs a faceted memory from the standard memory using the transformation $\mu \uparrow_{\Gamma}^{def}$ from Fig. 5, where $\text{glb}(\mathcal{L})$ is the greatest lower bound of $\mathcal{L}$. The resulting faceted memory keeps original value of each variable x in a private facet, and adds default values (defined by *def* function) in a public facet. In a special case when the level of x is the smallest level in a lattice, we keep only a simple value $\mu(x)$ that is visible to all security levels. We then evaluate the program with the constructed faceted memory and $pc = \mathcal{L}$. The resulting faceted memory is transformed back to a normal memory by using the projection function $\hat{\mu}|_{\Gamma}$.

The semantics rules for skip, sequence and while loop are straightforward. The GAssign rule uses a faceted evaluation $\hat{\mu}^{pc}(e)$ defined in Section 3.1.

Before describing the semantics of if instruction, we first define several auxiliary functions. Let $\text{dom}(\hat{\mu})$ be the domain of $\hat{\mu}$ and y be a fresh variable, i.e. $y \notin \text{dom}(\hat{\mu})$. By $\hat{\mu} \uplus (y \mapsto V)$ we denote a new memory $\hat{\mu}'$, such that $\text{dom}(\hat{\mu}') = \text{dom}(\hat{\mu}) \cup \{y\}$, $\hat{\mu}'(y) = V$ and for all $x \in \text{dom}(\hat{\mu})$, $\hat{\mu}'(x) = \hat{\mu}(x)$. By $\hat{\mu} \otimes^l y$, we remove y from

$$\text{GMF} \frac{(P, \mu \uparrow_{\Gamma}^{def}) \Downarrow_{\mathcal{L}} \hat{\mu}'}{\Gamma \vdash (P, \mu) \Downarrow_{GMF} \hat{\mu}'|_{\Gamma}}$$

$$\begin{aligned} \text{GSKIP} & \frac{}{(\text{skip}, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}} & \text{GASSIGN} & \frac{}{(x := e, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}[x \mapsto \hat{\mu}^{pc}(e)]} \\ \text{GSEQ} & \frac{(P_1, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}' \quad (P_2, \hat{\mu}') \Downarrow_G^{pc} \hat{\mu}''}{(P_1; P_2, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}''} \\ \text{GIF-C} & \frac{\hat{\mu}^{pc}(e) = v \quad (P_v, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}'}{(\text{if } e \text{ then } P_{\text{tt}} \text{ else } P_{\text{ff}}, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}'} \\ \text{GIF-S} & \frac{\begin{array}{l} \hat{\mu}^{pc}(e) = \langle l ? V_1 : V_2 \rangle \quad pc_1 = \{l' \in pc \mid l \sqsubseteq l'\} \\ pc_2 = pc \setminus pc_1 \quad \hat{\mu}_1 = \hat{\mu} \uplus (y \mapsto V_1) \quad \hat{\mu}_2 = \hat{\mu} \uplus (y \mapsto V_2) \\ P' = \text{if } y \text{ then } P_1 \text{ else } P_2 \quad (P', \hat{\mu}_1) \Downarrow_G^{pc_1} \hat{\mu}'_1 \quad (P', \hat{\mu}_2) \Downarrow_G^{pc_2} \hat{\mu}'_2 \end{array}}{(\text{if } e \text{ then } P_1 \text{ else } P_2, \hat{\mu}) \Downarrow_G^{pc} (\hat{\mu}'_1 \otimes^l \hat{\mu}'_2 \otimes^l y)} \\ \text{GWHILE} & \frac{(\text{if } e \text{ then } P; \text{while } e \text{ do } P \text{ else skip}, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}'}{(\text{while } e \text{ do } P, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}'} \end{aligned}$$

where $\hat{\mu}_1 \otimes^l \hat{\mu}_2(x) = \llbracket \langle l ? \hat{\mu}_1(x) : \hat{\mu}_2(x) \rangle \rrbracket$

Figure 6: Multiple facets for arbitrary security lattice

the domain of $\hat{\mu}$, that is, $\hat{\mu} \otimes^l y$ constructs a new memory $\hat{\mu}'$, where $\text{dom}(\hat{\mu}') = \text{dom}(\hat{\mu}) \setminus \{y\}$ and for all $x \neq y$, $\hat{\mu}'(x) = \hat{\mu}(x)$.

Consider the evaluation of the if instruction **if** e **then** P_1 **else** P_2 with $\hat{\mu}$ and pc . If e is evaluated to a constant value (**tt** or **ff**), then only P_{tt} or P_{ff} is evaluated (see rule GIF-C).

When e is evaluated to a faceted value $\langle l ? V_1 : V_2 \rangle$, we construct a new program **if** y **then** P_1 **else** P_2 , where y is a fresh variable. From Lemma 3.1, we have that $l \sqsubseteq pc$, and hence $pc_1 = \{l' \in pc \mid l \sqsubseteq l'\}$ and $pc_2 = pc \setminus pc_1$ are non-empty. In this case, we run the new program **if** y **then** P_1 **else** P_2 twice: once with the "higher view" than l , i.e., with $pc_1 = \{l' \in pc \mid l \sqsubseteq l'\}$ and y set to a private facet V_1 , and another time with "lower or incomparable view" than l , i.e. with $pc_2 = pc \setminus pc_1$ and y set to a public facet V_2 . We then combine the resulting memories using the $\otimes^l$ operator. The combination of faceted memories is based on the fact that when pc is split into pc_1 and pc_2 in the GIF-S rule, all levels in pc_1 is larger than or equal to l , and all levels in pc_2 is smaller than or incomparable to l .

Notice that the form of a faceted value constructed by combining values can be reduced. For example, a faceted value of the form $\langle H ? \langle M_1 ? V_{11} : V_{12} \rangle : V_2 \rangle$ can be reduced to $\langle H ? V_{11} : V_2 \rangle$ because $M_1 \sqsubseteq H$ and the projection of the original value at any level is either V_{11} or V_2 . We use the optimisation on the constructed faceted values from Fig. 7.

Therefore, in the GIF-S rule after the evaluation of P' in two contexts, we combine the resulting faceted memories $\hat{\mu}'_1 \otimes^l y$ and $\hat{\mu}'_2 \otimes^l y$ and apply an optimisation operator $\llbracket \cdot \rrbracket$ for each newly constructed faceted value. The correctness of $\llbracket \cdot \rrbracket$ used to optimize faceted values is proven in Lemma 3.3.

LEMMA 3.3. For all l , all V , it follows that $l(V) = l(\llbracket V \rrbracket)$.

$$\llbracket v \rrbracket = v$$

$$\llbracket \langle l ? V_1 : V_2 \rangle \rrbracket = \begin{cases} \llbracket V \rrbracket & \text{if } V_1 = V_2, \\ \llbracket \langle l ? V_{11} : V_{22} \rangle \rrbracket & \text{elseif } l_1 \sqsubseteq l, l \sqsubseteq l_2, V_1 = \langle l_1 ? V_{11} : V_{12} \rangle, \\ & V_2 = \langle l_2 ? V_{21} : V_{22} \rangle, \\ \llbracket \langle l ? V_{11} : V_2 \rangle \rrbracket & \text{elseif } l_1 \sqsubseteq l, V_1 = \langle l_1 ? V_{11} : V_{12} \rangle, \\ \llbracket \langle l ? V_1 : V_{22} \rangle \rrbracket & \text{elseif } l \sqsubseteq l_2, V_2 = \langle l_2 ? V_{21} : V_{22} \rangle, \\ \langle l ? \llbracket V_1 \rrbracket : \llbracket V_2 \rrbracket \rangle & \text{otherwise.} \end{cases}$$

Figure 7: Optimisation of a faceted value.

Example 3.4 (Evaluation of if instruction). Consider the security lattice $\langle \mathcal{L}_\circ, \sqsubseteq \rangle$ from Fig. 3 and the evaluation of the following program P with $pc = \mathcal{L}_\circ$ and $\hat{\mu}$, where $\hat{\mu}(x) = \langle M_1 ? \langle H ? \text{tt} : \text{ff} \rangle : \text{tt} \rangle$.

1: **if** x **then** $z := 10$ **else** $z := 5$

The evaluation follows the Gif-S rule since $M_1 \parallel \mathcal{L}_\circ$. We construct $P' = \text{if } y_1 \text{ then } P_1 \text{ else } P_2$ and first evaluate P' with $pc_1 = \{l' \in pc \mid M_1 \sqsubseteq l'\} = \{M_1, H\}$ and $\hat{\mu}_1 = \hat{\mu} \uplus (y \mapsto \langle H ? \text{tt} : \text{ff} \rangle)$, and then evaluate P' with $pc_2 = pc \setminus pc_1 = \{M_2, L\}$, $\hat{\mu}_2 = \hat{\mu} \uplus (y \mapsto \text{tt})$.

Since $pc_1 = \{H, M_1\}$ and $\hat{\mu}^{pc}(y) = \langle H ? \text{tt} : \text{ff} \rangle$, the evaluation of P' with pc_1 and $\hat{\mu}_1$ is split again to two evaluations: one with $P'' = \text{if } t \text{ then } P_1 \text{ else } P_2$, $pc_{11} = \{H\}$, and $\hat{\mu}_{11} = \hat{\mu}_1 \uplus (t \mapsto \text{tt})$; and the other one with P'' , $pc_{12} = \{M_1\}$, and $\hat{\mu}_{12} = \hat{\mu}_1 \uplus (t \mapsto \text{ff})$.

The evaluation of P'' with pc_{11} and with pc_{12} follow the Gif-C rule and we get two faceted memories $\hat{\mu}'_{11}$ and $\hat{\mu}'_{12}$, where $\hat{\mu}'_{11}(z) = 10$ and $\hat{\mu}'_{12}(z) = 5$. Then, $\hat{\mu}'_{11} \setminus t$ and $\hat{\mu}'_{12} \setminus t$ are combined and we get $\hat{\mu}'_1$, where $\hat{\mu}'_1(z) = \langle H ? 10 : 5 \rangle$.

The evaluation of P_2 with pc_2 follows the Gif-C rule and the result is $\hat{\mu}'_2$, where $\hat{\mu}'_2(z) = 10$. At this point, $\hat{\mu}'_1 \setminus y_1$ and $\hat{\mu}'_2 \setminus y_1$ are combined and the result is $\hat{\mu}'$, where $\hat{\mu}'(z) = \langle M_1 ? \langle H ? 10 : 5 \rangle : 10 \rangle$.

Example 3.5 (Evaluation with the GMF rule). Consider the lattice $\langle \mathcal{L}_\circ, \sqsubseteq \rangle$ from Fig. 3 and program P from Example 3.4 with one more instruction $x := x_1 > x_2$. Suppose that $\Gamma(x_1) = M_1$, $\Gamma(x_2) = H$, $\Gamma(z) = H$, $\mu(x_1) = 10$, $\mu(x_2) = 5$, the default values for x_1 and x_2 are respectively 100 and 20². Let $\hat{\mu} = \mu \uparrow_{\Gamma}^{def}$. It follows that $\hat{\mu}(x_1) = \langle M_1 ? 10 : 100 \rangle$ and $\hat{\mu}(x_2) = \langle H ? 5 : 20 \rangle$.

1: $x := x_1 > x_2$

2: **if** x **then** $z := 10$ **else** $z := 5$

Following GMF rule, the program is evaluated with $pc = \mathcal{L}_\circ = \{H, M_1, M_2, L\}$. For the assignment instruction, the value of x is updated to $\hat{\mu}^{pc}(x_1 > x_2) = \langle M_1 ? \langle H ? \text{tt} : \text{ff} \rangle : \text{tt} \rangle$. The rest of the evaluation is described in Example 3.4, and the resultant faceted memory is $\hat{\mu}'$, where $\hat{\mu}'(z) = \langle M_1 ? \langle H ? 10 : 5 \rangle : 10 \rangle$.

The memory after the application of rule GMF is $\mu' = \hat{\mu}'|_{\Gamma}$. Since $\Gamma(z) = H$, the value of z is $\mu'(z) = H(\langle M_1 ? \langle H ? 10 : 5 \rangle : 10 \rangle) = 10$.

3.3 Equivalence to SME-TINI and Security Guarantee

SME-TINI. The semantics of SME-TINI, termination-insensitive version of SME, for an arbitrary security lattice is presented below, where $\vec{\mu}$ is a vector that maps levels to normal memories; $\mu \uplus^l \Gamma$ constructs a memory where values of variables at levels that are not

²The values and default values for x_1 and x_2 are chosen so that the value of x after the evaluation of the assignment instruction is $\langle M_1 ? \langle H ? \text{tt} : \text{ff} \rangle : \text{tt} \rangle$.

visible to l are replaced by default values; $\odot_{\Gamma}(\vec{\mu})(x) \triangleq \vec{\mu}[\Gamma(x)](x)$ constructs a memory by combining all memories in $\vec{\mu}$; and def is a function mapping variables to default values.

$$\text{SME-TINI} \frac{\forall l \in \mathcal{L} : (P, \mu \uplus^l \Gamma) \Downarrow \vec{\mu}[l]}{\Gamma \vdash (P, \mu) \Downarrow_{\text{SME-TINI}} \odot_{\Gamma}(\vec{\mu})}$$

$$\mu \uplus^l \Gamma \triangleq \begin{cases} def(x) & \text{if } \Gamma(x) \not\sqsubseteq l, \\ \mu(x) & \text{if } \Gamma(x) \sqsubseteq l. \end{cases}$$

We now prove that SME-TINI enforces *termination-insensitive noninterference (TINI)*. Two memories μ and μ' are *equivalent at* l w.r.t. Γ (denoted by $\mu \equiv_l^{\Gamma} \mu'$) iff for all x , $\Gamma(x) \sqsubseteq l \implies \mu(x) = \mu'(x)$. When Γ is clear from the context, $\mu \equiv_l^{\Gamma} \mu'$ is written as $\mu \equiv_l \mu'$.

Definition 3.6 (TINI). An enforcement mechanism A is *termination insensitive non-interferent (TINI)* if for all security environments Γ , programs P , and memories μ_1 , and μ_2 , we have

$$\mu_1 \equiv_l \mu_2 \wedge \Gamma \vdash (P, \mu_1) \Downarrow_A \mu'_1 \wedge \Gamma \vdash (P, \mu_2) \Downarrow_A \mu'_2 \implies \mu'_1 \equiv_l \mu'_2.$$

THEOREM 3.7. *SME-TINI is TINI.*

Equivalence to SME-TINI. To prove the equivalence between GMF and SME-TINI, we formally define the semantic equivalence of two mechanisms.

Definition 3.8. Two enforcement mechanisms A and B are *equivalent* if for any Γ , P and μ , we have that $\Gamma \vdash (P, \mu) \Downarrow_A \mu'$ iff $\Gamma \vdash (P, \mu) \Downarrow_B \mu'$.

We next establish the relation between the execution with GMF semantics and the execution with the standard semantics.

LEMMA 3.9. $(P, \hat{\mu}) \Downarrow_G^{pc} \hat{\mu}'$ iff $(P, l(\hat{\mu})) \Downarrow l(\hat{\mu}')$ for all $l \in pc$.

Thanks to Lemma 3.9, we now prove the equivalence of GMF and SME-TINI.

THEOREM 3.10. *GMF and SME-TINI are equivalent.*

As a consequence, we have that GMF is TINI.

REMARK 3.1. *MF [3] is constructed for a set of principals. When the set $\mathbf{P}$ of principals is fixed, we can use GMF to encode MF: we construct the lattice $\langle 2^{\mathbf{P}}, \subseteq \rangle$, where each element is a set of principals; we prove that GMF for $\langle 2^{\mathbf{P}}, \subseteq \rangle$ and MF for $\mathbf{P}$ are equivalent [23].*

4 OPTIMIZING GMF

In Section 3, we presented the semantics of Generalised Multiple Facets (GMF) for arbitrary lattice and have proven it to be equivalent to SME-TINI. However, GMF from Fig. 6 can be further optimised and avoid repeating evaluations of the same commands. The following example demonstrates the sub-optimality of GMF.

Example 4.1 (GMF is not optimal). We consider the below program from Example 2.1. The lattice is $\langle \mathcal{L}_B, \sqsubseteq \rangle$ from Fig. 2.

1: *winner* := 0;

2: *test* := $(x_1 \leq x_2)$ and $(x_2 \leq x_3)$;

3: **if** *test* **then** *winner* := 2 **else** *skip*

Suppose that the bid offers of B_1 , B_2 , and B_3 are respectively 10, 5, and 7, and the default values for B_i are 0. W.r.t. this setting, the initial faceted memory is $\hat{\mu}$, where $\hat{\mu}(x_1) = \langle B_1 ? 10 : 0 \rangle$, $\hat{\mu}(x_2) = \langle B_2 ? 5 : 0 \rangle$,

and $\hat{\mu}(x_3) = \langle B_3 ? 7 : 0 \rangle$. We consider the execution of the program with GMF.

After line 2, $test = \langle B_1 ? \langle B_2 ? ff : ff \rangle : \langle B_2 ? ff : \langle B_3 ? tt : tt \rangle \rangle$. Following the semantics of GMF, the assignment instruction $winner := 2$ is evaluated twice with $pc_{B_3} = \{B_3\}$, and $pc_{\perp} = \{\perp\}$; the **skip** instruction is evaluated three times with $pc_{\top} = \{\top\}$, $pc_{B_1} = \{B_1\}$, and $pc_{B_2} = \{B_2\}$.

The main idea of our optimisation lays in reducing the number of sub-evaluations and hence the number of faceted memory combinations. For Example 4.1, we propose a mechanism that merges the evaluations corresponding to pc_{B_3} and $pc_{\perp}$ into one evaluation with $pc_1 = \{B_3, \perp\}$. This simplification is possible since $test$ denotes the same value (i.e., tt) under pc_{B_3} and $pc_{\perp}$. Similarly, our simplification merges the evaluations corresponding to $pc_{\top}$, pc_{B_1} , and pc_{B_2} , where $test$ denotes ff , into one evaluation with $pc_2 = \{\top, B_1, B_2\}$, and thus evaluates each branch of the if command only once.

In this section, we propose semantics of *optimized GMF* (OGMF) that reduces the number of sub-evaluations, and hence is more resource-friendly than GMF.

4.1 Semantics

The ideas behind the OGMF rule, and the rules for skip, assignment, sequence, and while instructions are similar to the corresponding ones of GMF. The functions $\mu \uparrow_{\Gamma}^{def}(x)$ and $\hat{\mu}|_{\Gamma}(x)$ are defined in Fig. 5. We now explain the semantic rules for the conditional instruction.

Consider evaluation of the program **if** e **then** P_1 **else** P_2 with pc and memory $\hat{\mu}$, and $\hat{\mu}^{pc}(e) = V$. In order to evaluate each branch of the conditional only once, we split the pc in two subsets: in the first subset pc_1 the visible value of V is true, and in the remaining subset pc_2 , V is false. We now have three distinct cases.

If $pc_1 = pc$, meaning that for all levels in pc , the visible value of V is true, then P_1 is evaluated (rule OIf-T). If $pc_2 = pc$, then for all levels in pc , the visible value of V is false, and only P_2 is evaluated (rule OIf-F). Finally, when pc is split in non-empty pc_1 and pc_2 , then both P_1 and P_2 are evaluated, and their results ($\hat{\mu}'_1$ and $\hat{\mu}'_2$) are combined by $\hat{\mu}'_1 \oplus^{pc_1, pc_2} \hat{\mu}'_2$ (rule OIf-S) to a new faceted memory. The intuition behind this combination is that the projection of $\hat{\mu}'_1 \oplus^{pc_1, pc_2} \hat{\mu}'_2$ at $l \in pc_1$ is taken from the evaluation of P_1 and its projection at $l \in pc_2$ is taken from the evaluation of P_2 .

In the definition of combination of memories for OGMF (bottom of Fig. 8), we distinguish two cases. If for some variable x , its value in both faceted memories is the same, ($\hat{\mu}'_1(x) = \hat{\mu}'_2(x)$), then we do not need to construct a new faceted value. Instead, we optimize the current value using the optimisation operator from Fig. 7.

If the values of x in $\hat{\mu}'_1(x)$ and $\hat{\mu}'_2(x)$ are different, then we construct a new faceted value $V = \mathbb{F}(V_1, V_2, pc_1, pc_2)$ and apply further optimisation on the resulting value V using a new optimisation operator that takes into account a faceted value and the current pc : $\llbracket V, pc \rrbracket$ optimizes the form of V and is described in Fig. 9. We show an example of such optimisation in Example 4.4.

To combine two faceted memories, we first construct a new faceted value by using $\mathbb{F}(V_1, V_2, pc_1, pc_2)$:

$$\mathbb{F}(V_1, V_2, pc_1, pc_2) = \langle \text{List}(pc_1 \cup pc_2), V_1, V_2, pc_1, pc_2 \rangle$$

$$\begin{array}{c}
\text{OGMF} \frac{\boxed{\frac{(P, \mu \uparrow_{\Gamma}^{def}) \downarrow_O^{\mathcal{L}} \hat{\mu}'}{\Gamma \vdash (P, \mu) \Downarrow_{OGMF} \hat{\mu}' | \Gamma}}}{\Gamma \vdash (P, \mu) \Downarrow_{OGMF} \hat{\mu}' | \Gamma} \\
\\
\text{OAssign} \frac{}{(x := e, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}[x \mapsto \hat{\mu}^{pc}(e)]} \quad \text{OSkip} \frac{}{(\text{skip}, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}} \\
\\
\text{OSeq} \frac{(P_1, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}' \quad (P_2, \hat{\mu}') \downarrow_O^{pc} \hat{\mu}''}{(P_1; P_2, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}''} \\
\\
\text{OIf-T} \frac{\hat{\mu}^{pc}(e) = V \quad pc_1 = \{l \in pc \mid l(V) = tt\} \quad pc_1 = pc \quad (P_1, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}'}{(\text{if } e \text{ then } P_1 \text{ else } P_2, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}'} \\
\\
\text{OIf-F} \frac{\hat{\mu}^{pc}(e) = V \quad pc_1 = \{l \in pc \mid l(V) = tt\} \quad pc_2 = pc \setminus pc_1 \quad pc_2 = pc \quad (P_2, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}'}{(\text{if } e \text{ then } P_1 \text{ else } P_2, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}'} \\
\\
\text{OIf-S} \frac{\hat{\mu}^{pc}(e) = V \quad pc_1 = \{l \in pc \mid l(V) = tt\} \quad pc_2 = pc \setminus pc_1 \quad pc_1 \neq \emptyset \quad pc_2 \neq \emptyset \quad (P_1, \hat{\mu}) \downarrow_O^{pc_1} \hat{\mu}'_1 \quad (P_2, \hat{\mu}) \downarrow_O^{pc_2} \hat{\mu}'_2}{(\text{if } e \text{ then } P_1 \text{ else } P_2, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}'_1 \oplus^{pc_1, pc_2} \hat{\mu}'_2} \\
\\
\text{OWHile} \frac{P' = \text{if } e \text{ then } P; \text{while } e \text{ do } P \text{ else skip} \quad (P', \hat{\mu}) \downarrow_O^{pc} \hat{\mu}'}{(\text{while } e \text{ do } P, \hat{\mu}) \downarrow_O^{pc} \hat{\mu}'} \\
\\
(\hat{\mu}'_1 \oplus^{pc_1, pc_2} \hat{\mu}'_2)(x) = \begin{cases} \llbracket \hat{\mu}'_1(x) \rrbracket & \text{if } \hat{\mu}'_1(x) = \hat{\mu}'_2(x), \\ \llbracket \mathbb{F}(\hat{\mu}'_1(x), \hat{\mu}'_2(x), pc_1, pc_2), pc_1 \cup pc_2) \rrbracket & \text{otherwise.} \end{cases}
\end{array}$$

Figure 8: Optimized multiple facets for arbitrary lattice

where $List(S)$ is a list of security levels from a set S , such that if l appears before l' in $List(S)$ then $l \not\sqsubseteq l'$. If the relation $\sqsubseteq$ in a given security lattice is not a total order, we can transform it into a total order $\sqsubseteq_T$ provided that $\sqsubseteq$ is a finite partial order. We can then view $List(S)$ as a list such that for any l and l' in this list, if l appears before l' , then $l \sqsubseteq_T l'$.

The definition of $\mathbb{F}(V_1, V_2, pc_1, pc_2)$ uses the following operator that creates a faceted value based on an ordered list of security levels L , two faceted values, pc_1 and pc_2 :

$$\langle \langle L, V_1, V_2, pc_1, pc_2 \rangle \rangle = \begin{cases} l(V_1) & \text{if } L = l, l \in pc_1, \\ l(V_2) & \text{if } L = l, l \in pc_2, \\ \langle l ? l(V_1) : \langle T, V_1, V_2, pc_1, pc_2 \rangle \rangle & \text{if } L = l.T, T \neq [], l \in pc_1, \\ \langle l ? l(V_2) : \langle T, V_1, V_2, pc_1, pc_2 \rangle \rangle & \text{if } L = l.T, T \neq [], l \in pc_2. \end{cases}$$

Notice that the form of the faceted value created by $\mathbb{F}(V_1, V_2, pc_1, pc_2)$ may be suboptimal.

Example 4.2 (Faceted value construction). Suppose that $V_1 = 2$, $V_2 = 0$, $pc_1 = \{B_3, \perp\}$, $pc_2 = \{\top, B_1, B_2\}$, $List(pc_1 \cup pc_2)$ is $\top, B_1, B_2, B_3, \perp$, and the lattice $\langle \mathcal{L}_B, \sqsubseteq \rangle$ is from Fig. 2.

Following the definition of combination of faceted memories, we have $\mathbb{F}(2, 0, pc_1, pc_2) = \langle \top ? 0 : \langle B_1 ? 0 : \langle B_2 ? 0 : \langle B_3 ? 2 : 2 \rangle \rangle \rangle$. This value can be further reduced to $\langle B_1 ? 0 : \langle B_2 ? 0 : 2 \rangle \rangle$.

We therefore define an optimisation function $\llbracket V, pc \rrbracket$ that further optimises the result V of a $\mathbb{F}()$ function. The optimisation uses the observation that faceted value returned by $\mathbb{F}()$ has the form of $\langle l ? v : V' \rangle$, where V' is either a simple value or a faceted value³.

The function $\llbracket V, pc \rrbracket$ is defined in Fig. 9. If V is of the form $\langle l ? v : v' \rangle$, then the optimisation is straightforward. We now consider the case when V is of the form $\langle l ? v : \langle l' ? v' : V' \rangle \rangle$. For demonstration, consider the lattice $\langle \mathcal{L}_B, \sqsubseteq \rangle$ from Fig. 2.

If the faceted value V is of the form $\langle \top ? v : \langle B_1 ? v : V' \rangle \rangle$ (formally, $l' \sqsubseteq l$ and $v = v'$), then it can be reduced to $\llbracket \langle B_1 ? v : V' \rangle, pc' \rrbracket$ (formally, $\llbracket \langle l' ? v' : V' \rangle, pc' \rrbracket$), where $pc' = pc \setminus \{\top\}$.

If the faceted value V is of the form $\langle B_1 ? v : \langle B_2 ? v : V' \rangle \rangle$, (l and l' are incomparable and $v = v'$), and moreover for all the levels in the pc , for which either B_1 or B_2 is visible, it is guaranteed that they observe the same value v (see the definition of $\text{cond}(V, pc)$ below), then we distinguish the following two cases.

$$\text{cond}(V, pc) \triangleq V = \langle l ? v : \langle l' ? v' : V' \rangle \rangle \wedge$$

$$\forall l_1 \in pc : \text{glb}(l, l') \sqsubseteq l_1 \implies l_1(V) = v.$$

- If all levels in pc are greater than or equal to $\text{glb}(l, l')$ (i.e. $\text{glb}(l, l') \leq pc$), then V is reduced to v . For example, if $pc = \{B_1, B_2, B_3\}$, $\text{glb}(B_1, B_2) = \perp$, then $\text{glb}(B_1, B_2) \leq pc$, and thanks to the $\text{cond}(V, pc)$ we know that $B_1(V) = B_2(V) = B_3(V) = v$, then we can reduce such faceted value to simply v because value V' is not useful for such pc .
- If only some levels in pc are greater than or equal to $\text{glb}(l, l')$ (i.e. $\text{glb}(l, l') \not\leq pc$), then V is reduced to $\langle \text{glb}(l, l') ? v : V'' \rangle$ and this value is reduced further recursively. Consider that we add one more security level L to the lattice $\langle \mathcal{L}_B, \sqsubseteq \rangle$ such that $L \sqsubseteq \perp$. If $pc = \{B_1, B_2, L\}$, $\text{glb}(B_1, B_2) = \perp$, then $\text{glb}(B_1, B_2) \not\leq pc$ because $\perp \not\sqsubseteq L$. We then construct a set of security levels S from pc , which are higher or equal than $\text{glb}(l, l')$, and therefore the view on V from all these levels is v (because $\text{cond}(V, pc)$ holds). In our example, $S = \{B_1, B_2\}$, and we construct a new faceted value $V'' = \langle \langle \{L\}, V' \rangle \rangle = L(V')$. We then define a new $pc' = (pc \setminus S) \cup \{\text{glb}(l, l')\} = \{L, \perp\}$, and we need to keep $\text{glb}(l, l')$ in pc' because we must ensure that all the levels present in the new faceted value are also present in pc . Therefore, the reduced faceted value for our example is $\llbracket \langle \perp ? v : L(V') \rangle, \{\perp, L\} \rrbracket$.

Finally, if none of the above conditions hold then we recursively reduce the facet $\langle l' ? v' : V' \rangle$.

The correctness of $\hat{\mu}_1 \oplus^{pc_1, pc_2} \hat{\mu}_2$ in the OIf-S rule is proven in Lemma 4.3.

LEMMA 4.3. *For all levels l , variables x , sets of security levels pc_1 and pc_2 , and memories $\hat{\mu}_1$ and $\hat{\mu}_2$,*

- *if $l \in pc_1$, then $l(\hat{\mu}_1 \oplus^{pc_1, pc_2} \hat{\mu}_2)(x) = l(\hat{\mu}_1)(x)$,*
- *if $l \in pc_2$, then $l(\hat{\mu}_1 \oplus^{pc_1, pc_2} \hat{\mu}_2)(x) = l(\hat{\mu}_2)(x)$.*

Example 4.4 (Optimisation of faceted value). Consider a faceted value $\langle \top ? 0 : \langle B_1 ? 0 : \langle B_2 ? 0 : \langle B_3 ? 2 : 2 \rangle \rangle \rangle \rangle$ and $pc = \{\top, B_1, B_2, B_3, \perp\}$ from Example 4.2. We show how this value is optimised with our optimisation function $\llbracket \cdot, \cdot \rrbracket$:

$$\begin{aligned} & \llbracket \langle \top ? 0 : \langle B_1 ? 0 : \langle B_2 ? 0 : \langle B_3 ? 2 : 2 \rangle \rangle \rangle \rangle, \{\top, B_1, B_2, B_3, \perp\} \rrbracket = \\ & = \llbracket \langle B_1 ? 0 : \langle B_2 ? 0 : \langle B_3 ? 2 : 2 \rangle \rangle \rangle, \{B_1, B_2, B_3, \perp\} \rrbracket = \\ & = \langle B_1 ? 0 : \llbracket \langle B_2 ? 0 : \langle B_3 ? 2 : 2 \rangle \rangle, \{B_2, B_3, \perp\} \rrbracket \rangle = \\ & = \langle B_1 ? 0 : \langle B_2 ? 0 : \llbracket \langle B_3 ? 2 : 2 \rangle, \{B_3, \perp\} \rrbracket \rangle = \langle B_1 ? 0 : \langle B_2 ? 0 : 2 \rangle \rangle \end{aligned}$$

Example 4.5 (OGMF is more resource-friendly than GMF). Consider the program from Example 4.1. To show optimisation of OGMF, we evaluate it with $pc = \{\top, B_1, B_2, B_3, \perp\}$ and $\hat{\mu}$, where $\hat{\mu}(x_1) = \langle B_1 ? 10 : 0 \rangle$, $\hat{\mu}(x_2) = \langle B_2 ? 5 : 0 \rangle$, and $\hat{\mu}(x_3) = \langle B_3 ? 7 : 0 \rangle$. After the execution of the instruction at line 2, the faceted memory is $\hat{\mu}' = \hat{\mu}[\text{winner} \mapsto 0, \text{test} \mapsto V]$, where $V = \langle B_1 ? \langle B_2 ? \text{ff} : \text{ff} \rangle : \langle B_2 ? \text{ff} : \langle B_3 ? \text{tt} : \text{tt} \rangle \rangle \rangle$. We consider the execution of the if instruction.

For levels $pc_1 = \{B_3, \perp\}$, the evaluation of test is $\text{tt} : B_3(\hat{\mu}^{pc}(\text{test})) = \perp(\hat{\mu}^{pc}(\text{test})) = \text{tt}$. Moreover, $pc_1 \neq pc$, therefore, the rule OIf-S applies. The evaluation of the program is split to two: the first evaluation is with $P_1 = \text{winner} := 2$ and $pc_1 = \{B_3, \perp\}$; and the second evaluation is with $P_2 = \text{skip}$ and $pc_2 = \{\top, B_1, B_2\}$. Each branch of the conditional will be evaluated only once.

The evaluation of P_1 with pc_1 terminates with $\hat{\mu}_1''(\text{winner}) = 2$. The evaluation of P_2 with pc_2 terminates with $\hat{\mu}_2''(\text{winner}) = 0$. These two faceted memories are combined to $\hat{\mu}''$, where $\hat{\mu}''(\text{winner}) = \langle B_1 ? 0 : \langle B_2 ? 0 : 2 \rangle \rangle$. The construction of this faceted memory is presented in Examples 4.2 and 4.4.

In the example above, OGMF has only two sub-evaluations, while GMF has five, moreover OGMF combines faceted memories once, while GMF combines them four times. Therefore, OGMF is more resource-friendly than GMF.

4.2 Equivalence to SME-TINI and Security Guarantee

We first establish the relation between the standard semantics and the semantics of OGMF.

LEMMA 4.6. $(P, \hat{\mu}) \Downarrow_O^{pc} \hat{\mu}'$ if and only if $(P, l(\hat{\mu})) \Downarrow l(\hat{\mu}')$ for all $l \in pc$.

We now can prove the semantic equivalence result for OGMF and SME-TINI.

THEOREM 4.7. *OGMF and SME-TINI are equivalent.*

As a consequence, OGMF and GMF are equivalent even though OGMF is optimized. In addition, OGMF is TINI.

5 A TERMINATION SENSITIVE VERSION OF MULTIPLE FACETS

A *termination sensitive* model assumes that an attacker can observe termination of evaluations. In [19], the model is explained further: an attacker at level l can observe the termination of evaluations at level l and lower. In the case of GMF and OGMF, an evaluation marked with pc is an evaluation at l if $l \in pc$. Notice that an evaluation is at more than one level whenever pc is not a singleton.

As illustrated by Example 5.1, GMF and OGMF do not prevent the influence of private data at higher levels to the termination of the evaluations at lower levels. In other words, GMF and OGMF do not prevent leakage on termination channel [19].

Example 5.1. Suppose that $\mathcal{L} = \{L, H\}$, where $L \sqsubseteq H$. We look at the evaluation of **if** x **then** (**while** tt **do** **skip**) **else** **skip** with

³The function $\mathbb{F}()$ cannot return a simple value since it is called on non-empty pc_1 and pc_2 .

$$\begin{aligned}
\llbracket \langle l ? v : v' \rangle, pc \rrbracket &= \begin{cases} v & \text{if } v = v' \\ \langle l ? v : v' \rangle & \text{otherwise.} \end{cases} \\
\llbracket \langle l ? v : \langle l' ? v' : V' \rangle \rangle, pc \rrbracket &= \begin{cases} \llbracket \langle l' ? v : V' \rangle, pc' \rrbracket & \text{if } l' \sqsubseteq l, v = v', \text{ where } pc' = pc \setminus \{l\}, \\ v & \text{if } l \parallel l', v = v', \text{ cond}(V, pc) \text{ and } \text{glb}(l, l') \leq pc, \\ \llbracket \langle \text{glb}(l, l') ? v : V'' \rangle, pc' \rrbracket & \text{if } l \parallel l', v = v', \text{ cond}(V, pc) \text{ and } \text{glb}(l, l') \parallel pc, \text{ where } pc' = (pc \setminus S) \cup \{\text{glb}(l, l')\}, \\ & S = \{l_1 \in pc \mid \text{glb}(l, l') \sqsubseteq l_1\}, \text{ and } V'' = \langle \text{List}(pc \setminus S), V' \rangle, \\ \langle l ? v : \llbracket \langle l' ? v' : V' \rangle, pc' \rrbracket \rangle & \text{otherwise, where } pc' = pc \setminus \{l\}. \end{cases} \\
\langle\langle l, V \rangle\rangle &= \begin{cases} l(V) & \text{if } l = l, \\ \langle l ? l(V) : \langle\langle T, V \rangle\rangle \rangle & \text{if } l = l.T, T \neq []. \end{cases}
\end{aligned}$$

Figure 9: Definition of $\llbracket V, pc \rrbracket$, and optimisation of a faceted value V with respect to the set of security levels pc .

$pc = \mathcal{L}$ and $\hat{\mu}(x) = \langle H ? \text{tt} : \text{ff} \rangle$. When GMF or OGMF is used, the evaluation is split into two: one is with $pc_1 = \{H\}$, the other one is with $pc_2 = \{L\}$. The evaluation with pc_2 converges, while the evaluation with pc_1 diverges since its executing program is **while tt do skip**. Therefore, the evaluation of the whole program with $pc = \{L, H\}$ also diverges and hence, to an attacker at L , the evaluation at L diverges. However, if the program is evaluated with $\hat{\mu}'(x) = \langle H ? \text{ff} : \text{ff} \rangle$, to the attacker at L , the evaluation at L converges. Based on observations on those two evaluations, an attacker at L can gain insight about the high facet of x . In other words, GMF and OGMF do not prevent the influence of data at H to the termination of the evaluation at L .

Therefore, we propose *Termination Sensitive Multiple Facets* (TSMF), a version of MF that takes into account the termination sensitive model. TSMF is a generalization of a version of MF presented in [8, Appendix A]. The basic idea of TSMF is that when an if instruction is evaluated, TSMF performs a bounded evaluation of the instruction by using OGMF. If the OGMF evaluation does not terminate within the given time bound, then the instruction is evaluated instead using SME semantics with a low-prio scheduler [14]. The security guarantees offered by TSMF are the same as SME with the same low-prio scheduler [19]. The semantics of TSMF and the proofs about its security guarantees can be found in [23].

6 RELATED WORK

SME. Devriese and Piessens introduce the idea of Secure Multi-Execution [14]. Since then, many researchers have developed different aspects of this approach. Close to our work, Kashyap et al. [19] discuss how schedulers might affect security guarantees (i.e., TSNI and TINI) based on the chosen scheduler and the lattice ordering. They show several schedulers and classify them according to the strength of security guarantees and according to fairness properties. This work complement theirs by providing a similar analysis but for an interplay of MF and SME semantics. SME [14] has many implementations: as a library in Haskell [18], as an experimental web browser based on Firefox [13], as a static program transformation for both Python and JavaScript [4], and as an adaptation to reactive systems [6]. In the work above, SME preserves the semantics of secure programs up to interleaving of events. To remedy that, Zanarini et al. [37] carefully leverage SME to design a precise monitor which exactly preserves semantics of secure programs *up to termination*. Several other works [10, 26, 33] expand SME and introduce

declassification. In this work, we focus on semantics guarantees up to interleaving of events—as in the SME original formulation.

MF. Austin and Flanagan introduce MF semantics [3]—a technique often referred as an optimization for SME. However, as shown by Bielova and Rezk [7], they do not provide the same security guarantees (i.e., TINI vs. TSNI) and differ in their treatment of default values. This work provides yet another look into a comparison between both techniques to show their differences, while introducing novel value-based optimizations to MF. Another work by the same authors [9] compare and contrast five dynamic techniques, including MF and SME, to mainly reason about the preservation of semantics of secure programs, a property known as *transparency*. In this work, we show that GMF and OGMF enjoy the same transparency guarantees as SME-TINI (Theorems 3.10 and 4.7).

Tools. Most information flow control tools provide TINI, e.g., Jif [21], FlowCaml [25], Laminar [27], Paragon [11], and JSFlow [16]. Similarly, termination leaks are often ignored in security tools coming from the operating system research community, e.g., Asbestos [15], HiStar [38], and Flume [20]. A few exceptions to this trend are the security libraries LIO [31] and MAC [28], which provide TSNI for concurrent programs.

Decentralized label models. The decentralized label model (DLM), allows one to express the interests of mutually-distrusting principals without a central authority [22]. The set of labels forms a pre-order where the order relationship does not require to know all the points in the relationship to determine the result of comparing two labels—bearing in mind that there might be an infinite number of labels due to the dynamic creation of principals at runtime. In a similar spirit, DC-labels [32, 34] provides a decentralized label format which allows one to express rich policies dictated by mutually-distrusting principals as propositional logic formulas (without negation). In this work, we require to know all the points in the chosen lattice in order to optimize MF as shown by OGMF. Extending our techniques to DLM or DC-labels is an interesting direction for future work.

7 CONCLUSION AND PERSPECTIVES

This work contributes to develop techniques to secure programs using dynamic information flow—a promising approach to secure existing JavaScript code. We specially focus on proposing a technique that achieves a smaller number of executions than MF (and hence smaller memory footprint) without diminishing security guarantees. We further extend our MF-based technique to work

with arbitrary finite lattices (GMF) based on the observation that off-the-shelf lattices with principals are not always the most convenient ones to use. Knowing all the points in the lattice allows for further optimizations: spawning multi-executions could be done on a value-based basis (OGMF) rather than on security levels—as in original MF. Finally, we propose a hybrid approach which present an interesting balance between the number of executions and security guarantees: it behaves as OGMF as long as it can and switches to SME when termination leaks could occur (TSMF). In other words, TSMF prioritizes resource usage as long as there are no risks for termination leaks. We expect that these insights will help inform future development of multi-execution-based techniques. In fact, an intriguing question is what it would take for our optimizations (or future ones) to work on potentially infinite lattices like the DLM or DC-labels—an interesting direction for future work.

ACKNOWLEDGMENTS

We would like to thank anonymous reviewers for feedback that helped to improve this paper. This research has been partially supported by the ANR projects AJACS ANR-14-CE28-0008, CISC ANR-17-CE25-0014-01, the National Science Foundation under grants CCF-1337278 and CCF-1421016, and Swedish research agencies Vetenskapsrådet and SSF Cyber Security projects WebSec: Securing Web-driven Systems and Octopi: Secure Programming for the Internet of Things.

REFERENCES

- [1] Thomas H. Austin and Cormac Flanagan. 2009. Efficient Purely-dynamic Information Flow Analysis. In *Proc. of PLAS 2009 (PLAS '09)*. 113–124.
- [2] Thomas H. Austin and Cormac Flanagan. 2010. Permissive Dynamic Information Flow Analysis. In *Proc. of PLAS 2010 (PLAS '10)*. 1–12.
- [3] Thomas H. Austin and Cormac Flanagan. 2012. Multiple Facets for Dynamic Information Flow. In *Proc. of POPL 2012 (POPL '12)*. 165–178.
- [4] Gilles Barthe, Juan Manuel Crespo, Dominique Devriese, Frank Piessens, and Exequiel Rivas. 2012. Secure multi-execution through static program transformation. In *Formal Techniques for Distributed Systems*. Springer, 186–202.
- [5] Abhishek Bichhawat, Vineet Rajani, Deepak Garg, and Christian Hammer. 2014. Generalizing Permissive-Upgrade in Dynamic Information Flow Analysis. In *Proc. of PLAS 2014 (PLAS '14)*. 15–24.
- [6] N. Bielova, D. Devriese, F. Massacci, and F. Piessens. 2011. Reactive non-interference for a Browser model. In *Proc. of NSS 2011*. 97–104.
- [7] Natalia Bielova and Tamara Rezk. 2016. Spot the Difference: Secure Multi-execution and Multiple Facets. In *Proc. of ESORICS 2016*. 501–519.
- [8] Natalia Bielova and Tamara Rezk. 2016. *Spot the Difference: Secure Multi-Execution and Multiple Facets*. Technical Report. <https://goo.gl/b7yoQ9>.
- [9] Natalia Bielova and Tamara Rezk. 2016. A Taxonomy of Information Flow Monitors. In *Proc. of POST 2016*. 46–67.
- [10] Iulia Boloşteanu and Deepak Garg. 2016. Asymmetric Secure Multi-execution with Declassification. In *Proc. of POST 2016*. 24–45.
- [11] Niklas Broberg, Bart van Delft, and David Sands. 2013. Paragon for Practical Programming with Information-Flow Control. In *Proc. of APLAS 2013 (LNCS)*, Vol. 8301. Springer, 217–232.
- [12] Stefano Calzavara, Alvise Rabitti, and Michele Bugliesi. 2016. Content Security Problems?: Evaluating the Effectiveness of Content Security Policy in the Wild. In *Proc. of CCS 2016 (CCS '16)*. 1365–1375.
- [13] Willem De Groef, Dominique Devriese, Nick Nikiforakis, and Frank Piessens. 2012. FlowFox: a web browser with flexible and precise information flow control. In *Proc. of CCS 2012*. ACM, 748–759.
- [14] Dominique Devriese and Frank Piessens. 2010. Noninterference Through Secure Multi-execution. In *Proc. of IEEE SP 2010 (SP '10)*. 109–124.
- [15] Petros Efstathopoulos, Maxwell Krohn, Steve VanDeBogart, Cliff Frey, David Ziegler, Eddie Kohler, David Mazières, Frans Kaashoek, and Robert Morris. 2005. Labels and event processes in the Asbestos operating system. In *Proc. of SOSP 2005 (SOSP)*. ACM.
- [16] D. Hedin, A. Birgisson, L. Bello, and A. Sabelfeld. 2014. JSFlow: Tracking information flow in JavaScript and its APIs. In *Proc. of SAC 2014*. ACM.
- [17] Collin Jackson and Adam Barth. 2008. Beware of Finer-Grained Origins. In *Web 2.0 Security and Privacy (W2SP'08)*.
- [18] Mauro Jaskelioff and Alejandro Russo. 2011. Secure multi-execution in Haskell. In *International Andrei Ershov Memorial Conference on Perspectives of System Informatics*. Springer, 170–178.
- [19] Vineeth Kashyap, Ben Wiedermann, and Ben Hardekopf. 2011. Timing- and Termination-Sensitive Secure Information Flow: Exploring a New Approach. In *Proc. of IEEE SP 2011 (SP '11)*. 413–428.
- [20] Maxwell Krohn, Alexander Yip, Micah Brodsky, Natan Cliffer, M. Frans Kaashoek, Eddie Kohler, and Robert Morris. 2007. Information Flow Control for Standard OS Abstractions. In *Proc. of SOSP 2007 (SOSP)*.
- [21] Andrew C. Myers. 1999. JFlow: Practical mostly-static information flow control. In *Proc. of POPL 1999*. ACM, 228–241.
- [22] Andrew C. Myers and Barbara Liskov. 2000. Protecting privacy using the decentralized label model. *ACM Transactions on Software Engineering and Methodology (TOSEM)* 9, 4 (2000), 410–442.
- [23] Minh Ngo, Natalia Bielova, Cormac Flanagan, Tamara Rezk, Alejandro Russo, and Thomas Schmitz. 2017. A Better Facet of Dynamic Information Flow Control. (2017). <https://goo.gl/Y2SEnw>.
- [24] Minh Ngo, Fabio Massacci, Dimitar Milushev, and Frank Piessens. 2015. Runtime Enforcement of Security Policies on Black Box Reactive Programs. In *Proc. of POPL 2015*.
- [25] F. Pottier and V. Simonet. 2002. Information Flow Inference for ML. In *ACM Symp. on Principles of Programming Languages*. 319–330.
- [26] Willard Rafnsson and Andrei Sabelfeld. 2013. Secure Multi-execution: Fine-Grained, Declassification-Aware, and Transparent. In *Proc. of CSF 2013*. 33–48.
- [27] Indrajit Roy, Donald E. Porter, Michael D. Bond, Kathryn S. McKinley, and Emmett Witchel. 2009. Laminar: Practical Fine-grained Decentralized Information Flow Control. In *Proc. of PLDI 2009 (PLDI)*. ACM.
- [28] Alejandro Russo. 2015. Functional Pearl: Two Can Keep a Secret, if One of Them Uses Haskell. In *Proc. of ICFP 2015 (ICFP)*. ACM.
- [29] José Fragoso Santos, Thomas Jensen, Tamara Rezk, and Alan Schmitt. 2015. Hybrid Typing of Secure Information Flow in a JavaScript-Like Language. In *Proc. of TGC 2015*. 63–78.
- [30] Dolière Francis Some, Natalia Bielova, and Tamara Rezk. 2017. On the Content Security Policy Violations Due to the Same-Origin Policy. In *Proc. of WWW 2017 (WWW '17)*. 877–886.
- [31] Deian Stefan, Alejandro Russo, Pablo Buiras, Amit Levy, John C Mitchell, and David Mazieres. 2012. Addressing covert termination and timing channels in concurrent information flow systems. In *Proc. of ICFP 2012*, Vol. 47. ACM, 201–214.
- [32] D. Stefan, A. Russo, D. Mazières, and J. C. Mitchell. 2011. Disjunction Category Labels. In *Proc. of NordSec 2011*. Springer-Verlag.
- [33] Mathy Vanhoef, Willem De Groef, Dominique Devriese, Frank Piessens, and Tamara Rezk. 2014. Stateful Declassification Policies for Event-Driven Programs. In *Proc. of CSF 2014 (CSF '14)*. 293–307.
- [34] Lucas Waye, Pablo Buiras, Dan King, Stephen Chong, and Alejandro Russo. 2015. It's My Privilege: Controlling Downgrading in DC-Labels. In *International Workshop on Security and Trust Management*.
- [35] Wikipedia. 2017. Ad exchange. (2017). https://en.wikipedia.org/wiki/Ad_exchange. Checked on Nov 08, 2017.
- [36] Wikipedia. 2017. Real-time bidding. (2017). https://en.wikipedia.org/wiki/Real-time_bidding. Checked on Nov 08, 2017.
- [37] Dante Zaranini, Mauro Jaskelioff, and Alejandro Russo. 2013. Precise enforcement of confidentiality for reactive systems. In *Proc. of CSF 2013*. IEEE, 18–32.
- [38] Nikolai Zeldovich, Silas Boyd-Wickizer, Eddie Kohler, and David Mazières. 2006. Making information flow explicit in HiStar. In *USENIX Symp. on Operating Systems Design and Implementation*. USENIX.