

Improving dynamic code analysis by code abstraction

Isabella Mastroeni

Department of Computer Science, University of Verona (Italy)

`isabella.mastroeni@univr.it`

Vincenzo Arceri

Department of Environmental Sciences, Informatics and Statistics,
Ca' Foscari University of Venice (Italy)

`vincenzo.arceri@unive.it`

In this paper, our aim is to propose a model for code abstraction, based on abstract interpretation, allowing us to improve precision of a recently proposed static analysis by abstract interpretation of dynamic languages. The problem we tackle here is that the analysis may add some spurious code to the string-to-execute abstract value and this code may need some abstract representations in order to make it analyzable. This is precisely what we propose here, where we drive the code abstraction by the analysis we have to perform.

1 Introduction

The possibility of dynamically building code instructions as the result of text manipulation is a key aspect in dynamic programming languages. In this scenario, programs can turn text, which can be built at run-time, into executable code [25]. These features are often used in code protection and tamper resistant applications, employing camouflage for escaping attack or detection [23], in malware, in mobile code, in web servers, in code compression, and in code optimization, e.g., in Just-in-Time (JIT) compilers, employing optimized run-time code generation.

While the use of dynamic code generation may simplify considerably the *art and performance of programming*, this practice is also highly dangerous, making the code prone to unexpected behaviors and malicious exploits of its dynamic vulnerabilities, such as code/object-injection attacks for privilege escalation, database corruption, and malware propagation. It is clear that more advanced and secure functionalities based on string-to-code statements could be permitted if we better master how to safely generate, analyze, debug, and deploy programs that dynamically generate and manipulate code.

There are lots of good reasons to analyze programs building strings that can be later executed as code. An interesting example is code obfuscation. Recently, several techniques have been proposed for JavaScript code obfuscation¹, meaning that also client-side code protection is becoming an increasingly important problem to be tackled by the research community and by practitioners. Hence, it is not always possible to simply ignore `eval` without accepting to lose the possibility of analyzing the rest of the program [3].

¹<https://www.daftlogic.com/projects-online-javascript-obfuscator.htm>,
<http://www.danstools.com/javascript-obfuscate/>,
<http://javascript2img.com/>,
<https://javascriptobfuscator.herokuapp.com/>,
<https://javascriptobfuscator.com/>

```

str = "x=5";
while (i < 3) {
  str = str + "5";
  i = i + 1;
}
str = str + ";"; eval(str);

```

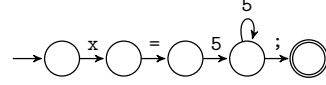


Figure 1: A s.t. $\mathcal{L}(A) = \{x = 5^n; | n > 0\}$, where 5^n means 5 repeated n times.

The context: Analyzing dynamic code. A major problem in presence of dynamic code generation is that static analysis becomes extremely hard if not impossible. This happens because program data structures, such as the control-flow graph and the system of recursive equations associated with the program in question, are themselves dynamically mutating objects. Recently [3], the problem of analyzing dynamic code has been tackled by *treating code as any other dynamic structure that can be statically analyzed by abstract interpretation, and to treat the abstract interpreter as any other program function that can be recursively called*. In particular, in [3], we provide a static analyzer architecture for a core dynamic language, containing non-removable `eval` statements, that still has some limitation in terms of precision but provides the necessary ground for studying more precise solutions to the problem. In particular,

- We have designed an automata-based string abstract domain [4] for analyzing string values during execution. Automata (FA) provide the perfect choice for abstracting strings that may be executed by `eval` since they allow us to over-approximate the set of possible values of string variables by keeping enough information for both analyzing properties of string variables that are never executed by an `eval` during computation and for extracting the potential executable sub-language.
- In order to statically analyze the code potentially executed by an `eval`, we have designed a systematic process for extracting from the (abstract) argument of `eval` (i.e., from the FA collection of its potential arguments) an over-approximation of executable code that this collection contains. Clearly, this approximation must keep a form that the analyzer can interpret.
- We designed a static analyzer for dynamic languages performing a recursive call of the interpreter on the (over-approximated) code that `eval` may execute.

The problem: Improve precision analysis by abstracting code. This analysis provides a first step towards the analysis of dynamic languages but still has some important precision loss [3]. In particular, there are particular forms of FA (which occur when the string is dynamically generated by loops) avoiding the possibility of generating a control flow graph (CFG) able to approximate the code executed by an `eval`. For instance, when the FA accepts a language such as $\{x=(5)^n; | n > 0\}$, the analysis in [3] cannot extract, from the FA, the CFG approximating the `eval` argument. In order to better explain the problem, consider the code in Fig. 1, where the value of `i` is statically unknown. In Fig. 1, we draw the automaton A representing the abstract value of `str` before the `eval` execution. The problem is that A has a cycle not involving a whole statement [3]. This situation makes the analyzer unable to build a CFG over-approximating the code potentially executed since, intuitively, such a CFG should be infinite. Indeed, only an infinite CFG could capture all the possible assignments described by the FA, namely all the assignments of any possible number formed only by 5 to the variable `x` (i.e., `x=5`; `x=55`; `x=555`; ...).

In order to make it possible to overcome this limitation, at least for a set of potential `eval` patterns, we propose to define a form of *abstract* CFG able to finitely represent a potential infinite set of CFGs, e.g., we look for a CFG representing `x=5*`.

$$\begin{aligned}
\text{Exp} \ni e &::= a \mid s \\
\text{AExp} \ni a &::= x \mid n \mid a + a \mid a - a \mid a * a \\
\text{BExp} \ni b &::= x \mid \text{true} \mid \text{false} \mid e = e \mid e > e \mid e < e \mid b \wedge b \mid \neg b \\
\text{SExp} \ni s &::= x \mid \text{"}\sigma\text{"} \mid \text{concat}(s, s) \mid \text{substr}(s, a, a) \\
\text{Comm} \ni c &::= \ell_1 \text{skip}^{\ell_2} \mid \ell_1 x := e^{\ell_2} \mid \ell_1 c; \ell_2 c^{\ell_3} \mid \ell_1 \text{if } (b) \{ \ell_2 c^{\ell_3} \} \{ \ell_4 c^{\ell_5} \}^{\ell_6} \\
&\quad \mid \ell_1 \text{while } (b) \{ \ell_2 c^{\ell_3} \}^{\ell_4} \mid \ell_1 \text{eval}(s)^{\ell_2} \\
\text{Imp} \ni P &::= \ell_1 c; \ell_2 \quad \text{where } \text{Id} \ni x \text{ (Identifiers)}, n \in \mathbb{Z}, \sigma \in \Sigma^*
\end{aligned}$$

Figure 2: Syntax of Imp

Unfortunately, things are not so easy as it may seem, since this abstract code representation has to be built in such a way that the analyzer may still be able to interpret it.

Contribution. The main contributions for tackling the problem above are:

- We first define the notion of *abstract* CFG, based on the idea of making it possible to still perform a given analysis. The idea is to leave the control structure unchanged, while approximating the edge labels (the statements to execute) to sets of labels, i.e., those sharing a fixed abstract property.
- We show how completeness of code abstraction w.r.t. the semantic observation models the possibility, for the static analyzer, of interpreting also the abstract code, and we show how we can make any code abstraction complete.
- We provide a systematic approach, based on the one proposed in [3], allowing us to analyse also the `eval` patterns described above, for which, instead, the analysis in [3] loses precision.

2 The core language: Imp

The language is quite standard (see Fig. 2²), and each statement is annotated with a label $\ell \in \text{Lab}$ (not part of the syntax) corresponding to the statement program point³.

In order to analyze a program $P \in \text{Imp}$, we need to model it by building a corresponding control flow graph [26] (CFG for short), which embeds the control structure in the graph structure and leaves in the edges (or equivalently on the nodes) only the access to states, i.e., manipulation of the states (assignments) and guards. The approach we use is quite standard, and we follow [26] for the construction of the control flow graph. For technical details see [3], here we show the construction on the example in Fig. 3, where i denotes the node corresponding to the program point ℓ_i . Note that, by construction [3], the language of the CFG edge labels is an intermediate language slightly different from the Imp grammar. Edge labels correspond to a primitive statement (i.e., an assignment or `eval`) or a boolean guard, namely they form the language Ψ generated by the grammar $l ::= x := e \mid b \mid \text{eval}(s)$.

²We use n to denote the semantic value corresponding to the syntactic symbol n .

³We suppose that there exists a function that, taken a well written program, can label it with a fresh label for each program point.

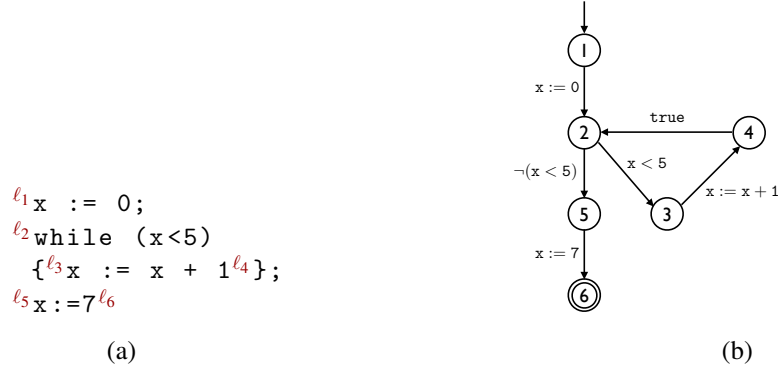


Figure 3: Example of CFG: (a) Fragment of code and (b) corresponding CFG.

Concrete semantics. The concrete semantics of our language `Imp` is intuitive and it is fully reported in [2]. Since our aim is to analyze `Imp` programs by analyzing their CFGs, we focus here only on the interpretation of CFG's labels [26]. In particular, we have to specify the semantics associated with each possible edge of the CFG. In other words, we have to formalize how each statement transforms a current state, which is represented as a store, namely as an association between identifiers and values. It is well known that static program analysis works by computing (abstract) collecting semantics, namely for each program point ℓ and for each variable x , it computes the set of values that the variable x can have in any computation at the program point ℓ . Hence, we define (collecting) memories $\mathfrak{m}$, associating with each variable a *set* of values. The basic values of `Imp` are integers, booleans and strings, hence we define the set of memories as $\mathbb{M} \stackrel{\text{def}}{=} \text{Var} \rightarrow (\wp(\mathbb{Z}) \cup \text{Bool} \cup \wp(\Sigma^*))$, ranged over the meta-variable $\mathfrak{m}$, where $\text{Bool} = \wp(\{\text{false}, \text{true}\})$. Let us denote by $\mathbb{V}$ this domain of collections of values $\wp(\mathbb{Z}) \cup \text{Bool} \cup \wp(\Sigma^*)$. The update of memory $\mathfrak{m}$ for a variable x with set of values v is denoted $\mathfrak{m}[x/v]$. The partial order $\sqsubseteq$ between memories is defined as $\mathfrak{m}_1 \sqsubseteq \mathfrak{m}_2 \Leftrightarrow \forall x \in \text{Id}. \mathfrak{m}_1(x) \subseteq \mathfrak{m}_2(x)$. Finally, lub and glb of memories are computed point-wise, i.e., $\mathfrak{m}_1 \sqcup \mathfrak{m}_2 \stackrel{\text{def}}{=} \lambda x. \mathfrak{m}_1(x) \cup \mathfrak{m}_2(x)$ and $\mathfrak{m}_1 \sqcap \mathfrak{m}_2 \stackrel{\text{def}}{=} \lambda x. \mathfrak{m}_1(x) \cap \mathfrak{m}_2(x)$. The collecting (input/output) semantics of statements $c \in \Psi$ is defined as the function $\llbracket c \rrbracket : \mathbb{M} \rightarrow \mathbb{M}$. We denote by $\llbracket \cdot \rrbracket$ the collecting semantics of expressions, defined as additive lift⁴ to sets of memories of the standard expression semantics. We abuse notation by denoting as $\llbracket \cdot \rrbracket$ also its additive lift to sets of statements.

$$\begin{aligned}
 \llbracket x := e \rrbracket \mathfrak{m} &= \mathfrak{m}[x / \llbracket e \rrbracket \mathfrak{m}] & \llbracket b \rrbracket \mathfrak{m} &= \mathfrak{m} \sqcap \sqcup \{ \mathfrak{m} \mid \llbracket b \rrbracket \mathfrak{m} = \text{true} \} \\
 \llbracket \text{eval}(s) \rrbracket \mathfrak{m} &= \llbracket (s) \mathfrak{m} \cap \text{Imp} \rrbracket \mathfrak{m}
 \end{aligned}$$

where $\cap$ is the intersection in the set of `Imp` programs. By computing the traces of application of this transfer function, starting from any possible input memory, we precisely compute the maximal trace semantics [24].

Static analysis on CFG: Semantic abstraction. It is well known that when we perform static analysis on a CFG, we interpret, on the corresponding abstract domain, all the edges, and more specifically all the labels (in Ψ) [26]. This is also a quite standard approach, but we recall it here for fixing the notation used. We suppose to abstract values on the coalesced sum [2] of the Sign abstract domain for integers,

⁴Let $f : S \rightarrow S$ be a generic function, by *additive lift* we mean its extension to sets of elements, i.e., $\forall X \subseteq S$ we define $f(X) \stackrel{\text{def}}{=} \{ f(x) \mid x \in S \}$. If $f : S \rightarrow \wp(S)$, then its lift to sets of memories is $f(X) \stackrel{\text{def}}{=} \bigcup \{ f(x) \mid x \in S \}$

of the concrete domain for booleans and of the (deterministic) finite state automata abstract domain for strings [2]⁵. Let us consider an abstraction $\rho \in uco(\mathbb{V})$ ⁶ of the values manipulated by our language, we denote by $\mathbb{M}^\rho : \text{Var} \longrightarrow \rho(\mathbb{V})$ the set of (collecting) memories, where sets of values are abstracted by ρ , ranged over $\mathbb{m}^\rho$. In the following, we abuse notation by applying ρ to memories in $\mathbb{M}$, simply by defining $\rho(\mathbb{m}) \in \mathbb{M}^\rho$ as $\rho(\mathbb{m}) : x \in \text{Var} \mapsto \rho(\mathbb{m}(x))$ ⁷. In this way, we can see abstract memories as sets of concrete memories, and therefore as particular collecting memories, i.e., $\mathbb{M}^\rho \subseteq \mathbb{M}$. Finally, we can define the abstract edge effect $\llbracket \cdot \rrbracket^\rho$ [26] telling us how to abstractly interpret each edge of the CFG:

$$\begin{aligned} \llbracket x := e \rrbracket^\rho \mathbb{m}^\rho &= \mathbb{m}^\rho[x/\rho(\llbracket e \rrbracket^\rho \mathbb{m}^\rho)] & \llbracket b \rrbracket^\rho \mathbb{m}^\rho &= \mathbb{m}^\rho \sqcap \rho(\sqcup \{ \mathbb{m} \mid \text{true} \in \llbracket b \rrbracket^\rho \mathbb{m}^\rho \}) \\ \llbracket \text{eval}(s) \rrbracket^\rho \mathbb{m}^\rho &= \llbracket \llbracket s \rrbracket^\rho \mathbb{m}^\rho \rrbracket^\rho \mathbb{m}^\rho \end{aligned}$$

where $\llbracket \cdot \rrbracket^\rho \stackrel{\text{def}}{=} \rho \circ \llbracket \cdot \rrbracket \circ \rho$. The semantics of a path in the CFG is the composition of the interpretation of each edge, and the interpretation of an edge is the interpretation, given above, of its label [26].

This is clearly, what happens when the CFG is not abstracted, namely when the edge labels are single statements. Finally, since we deal with potential abstract CFG, we have to say how we execute them, potentially on an abstract semantics. The idea is simple, since we move from executing single statements to executing sets of statements, we simply take as execution of the abstract CFG the additive lift of the single statements executions. Since the semantics is always additive⁸, in order to guarantee that everything works, also the semantic abstraction ρ must be additive. Hence, in the following of the paper we always require ρ to be additive.

3 Semantic-driven code abstraction

In this section, we study how we can model a syntactic abstraction of the CFG and which is its *relation* with the semantic abstraction, i.e., the code analysis.

Modeling code abstraction. Following the standard approach for abstracting objects, we should abstract each CFG in a set of CFGs sharing an invariant property, i.e., an equivalence class of CFGs. In particular, since we aim at abstracting code (CFG) without changing the analysis performed on the code, we choose to abstract CFG by abstracting edge labels, and by leaving unchanged the control structure of the CFG. In other words, an abstract CFG, denoted $\text{CFG}^\#$, is a pair $\langle \text{Nodes}, \text{Edges}^\# \rangle$, where we leave the nodes unchanged, while the edge labels are abstracted to sets of labels. Formally, $\text{Edges}^\# \subseteq \text{Nodes} \times \wp(\Psi) \times \text{Nodes}$, where Ψ is the CFG label language.

Given $\eta \in uco(\wp(\Psi))$, $G^\eta \stackrel{\text{def}}{=} \langle \text{Nodes}(G), \text{Edges}^\eta(G) \rangle$ is the $\text{CFG}^\#$ built from a CFG G in terms of η , where $\text{Edges}^\eta(G) \subseteq \text{Nodes}(G) \times \eta(\wp(\Psi)) \times \text{Nodes}(G)$.

As an example, consider the CFG in Fig. 3, in Fig. 4 we have the $\text{CFG}^\#$ where numerical expressions are abstracted by $\{ \mathbb{m} \mid m \in \text{Sign}(n) \}$ ⁹ (where Sign is the well-known sign abstraction $\text{Sign} \in uco(\wp(\mathbb{Z}))$) such as $\text{Sign}(\wp(\mathbb{Z})) = \{ \top, \mathbb{Z}^+, \mathbb{Z}^-, \{0\}, \emptyset \}$. For instance, $x := x + 1$ is abstracted in $x := x + \mathbb{Z}^+$ where $x + \mathbb{Z}^+ \stackrel{\text{def}}{=} \{ x + n \mid n \in \mathbb{Z}^+ \}$, being $\text{Sign}(1) = \mathbb{Z}^+$.

⁵A string static analyzer using finite state automata abstract domain has been developed and it is available in [2].

⁶For the sake of simplicity here we abuse notation by considering a unique ρ which is indeed the coalesced sum of three abstractions, one for integers, one for booleans and one for strings.

⁷For the sake of simplicity of presentation and implementation, we have considered here non-relational abstractions of data, anyway we believe that it is possible to easy extend our work to relational abstractions.

⁸A function is said to be *additive* if it commutes with least upper bound.

⁹We use n to denote the semantic value corresponding to the syntactic symbol n .

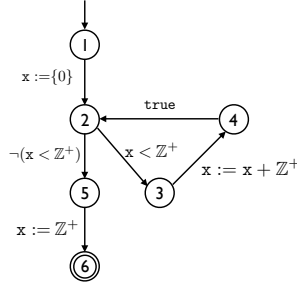


Figure 4: CFG abstracted by signs.

Abstracting code vs abstracting semantics. As previously noted, we aim at characterizing code abstractions, for dynamically generated code, for which the given analysis works precisely. Formally, let us consider the following equation:

$$\forall \mathfrak{m}^\rho \in \mathbb{M}^\rho \subseteq \mathbb{M}. \forall \varphi \in \Psi. \llbracket \eta(\varphi) \rrbracket^{\mathfrak{m}^\rho} = \llbracket \eta(\varphi) \rrbracket^{\rho \mathfrak{m}^\rho} \quad (1)$$

If this equality does not hold it means that the abstract semantic interpretation $\llbracket \cdot \rrbracket^\rho$ merges predicates distinguished by η . Namely, when the program is observed by means of its (abstract) semantics the actual abstraction of predicates is not precisely η , but it is η affected in some way by $\llbracket \cdot \rrbracket^\rho$. By changing the point of view, we have that, in this case, the analysis cannot precisely interpret the abstract code, since η abstracts the code by distinguishing information that ρ cannot distinguish.

As an example, consider the sign domain above, when $\eta(x:=5) = \{x:=n \mid 1 \leq n \leq 5\}$ the equation does not hold since the concrete semantics of this set does not take *any* positive value for x . While, if $\eta(x:=5) = \{x:=n \mid n \in \mathbb{Z}^+ \cup \{0\}\}$, then Eq. 1 holds since its concrete semantics is precisely the set of non-negative values. It is worth noting that Eq. 1 is a forward completeness [16] of the code abstraction w.r.t. the semantic interpretation, meaning that the semantic abstraction does not add imprecision to the code one.

In order to investigate the relation existing between the code abstraction η and the semantic abstraction ρ , we observe that, whenever we have a semantic abstraction ρ , we have a natural code abstraction induced by ρ . Namely, by only observing (abstract) information about the computation, we cannot distinguish statements with the same (abstract) semantics, independently from what any possible code abstraction does. For instance, if we analyze parity of program variables, we are unable to distinguish $x:=2$ from $x:=4$, independently from how a potential code abstraction η is defined on $x:=2$. The first step consists in defining a code abstraction for expressions in terms of semantic one. Consider $\rho \in uco(\mathbb{V})$, we define $\widehat{\rho}(e)$ inductively on the expressions structure

$$\begin{aligned} \widehat{\rho}(a) &: \begin{cases} \widehat{\rho}(a_1 \text{ op } a_2) \stackrel{\text{def}}{=} \{a' \text{ op } a'' \mid a' \in \widehat{\rho}(a_1), a'' \in \widehat{\rho}(a_2)\} \stackrel{\text{def}}{=} \widehat{\rho}(a_1) \text{ op } \widehat{\rho}(a_2) \\ \widehat{\rho}(x) \stackrel{\text{def}}{=} x, \quad \widehat{\rho}(n) \stackrel{\text{def}}{=} \{m \mid m \in \rho(\{n\})\} \end{cases} \\ \widehat{\rho}(b) &: \begin{cases} \widehat{\rho}(b_1 \text{ bop } b_2) \stackrel{\text{def}}{=} \widehat{\rho}(b_1) \text{ bop } \widehat{\rho}(b_2), \quad \widehat{\rho}(\neg b) \stackrel{\text{def}}{=} \neg \widehat{\rho}(b) \\ \widehat{\rho}(x) \stackrel{\text{def}}{=} x, \quad \widehat{\rho}(\text{true}) \stackrel{\text{def}}{=} \{t \mid t \in \rho(\text{true})\}, \quad \widehat{\rho}(\text{false}) \stackrel{\text{def}}{=} \{t \mid t \in \rho(\text{false})\} \end{cases} \\ \widehat{\rho}(s) &: \begin{cases} \widehat{\rho}(\text{concat}(s_1, s_2)) \stackrel{\text{def}}{=} \text{concat}(\widehat{\rho}(s_1), \widehat{\rho}(s_2)), \\ \widehat{\rho}(\text{substr}(s, a_1, a_2)) \stackrel{\text{def}}{=} \text{substr}(\widehat{\rho}(s), \widehat{\rho}(a_1), \widehat{\rho}(a_2)) \\ \widehat{\rho}(x) \stackrel{\text{def}}{=} x, \quad \widehat{\rho}(\text{"}\sigma\text{"}) \stackrel{\text{def}}{=} \{\text{"}\delta\text{"} \mid \delta \in \rho(\sigma)\} \end{cases} \end{aligned}$$

At this point, we can characterize the CFG labels abstraction $\bar{\gamma}[\rho] : \wp(\Psi) \longrightarrow \wp(\Psi)$, as the additive lift of the function

$$\begin{aligned} \bar{\gamma}[\rho](x:=e) &\stackrel{\text{def}}{=} x := \hat{\rho}(e) \stackrel{\text{def}}{=} \{ x := e' \mid e' \in \hat{\rho}(e) \} \\ \bar{\gamma}[\rho](b) &\stackrel{\text{def}}{=} \hat{\rho}(b) \qquad \bar{\gamma}[\rho](\text{eval}(s)) \stackrel{\text{def}}{=} \text{eval}(\hat{\rho}(s)) \end{aligned}$$

where $\text{eval}(\hat{\rho}(s))$ is treated as the implicit representation of all the statements that it can execute, namely it represents the (potentially infinite) set $\{ c \mid \llbracket c \rrbracket_{\mathfrak{m}} \subseteq \llbracket (s)^{\rho} \cap \text{Imp} \rrbracket^{\rho}_{\mathfrak{m}} \}$.

The following result is immediate by construction.

Proposition 3.1 *Given $\rho \in \text{uco}(\mathbb{V})$, then $\bar{\gamma}[\rho] \in \text{uco}(\wp(\Psi))$ and it is additive.*

Finally, in order to show that this code abstraction can be used to force satisfiability of Eq. 1, we have first to characterize the meaning of interpreting an edge label abstracted by $\bar{\gamma}[\rho]$:

$$\begin{aligned} \llbracket x := \hat{\rho}(e) \rrbracket_{\mathfrak{m}} &= \sqcup \{ \llbracket x := e' \rrbracket_{\mathfrak{m}} \mid e' \in \hat{\rho}(e) \} & \llbracket \hat{\rho}(b) \rrbracket_{\mathfrak{m}} &= \sqcup \{ \llbracket b' \rrbracket_{\mathfrak{m}} \mid b' \in \hat{\rho}(b) \} \\ \llbracket \text{eval}(\hat{\rho}(s)) \rrbracket_{\mathfrak{m}} &= \sqcup \{ \llbracket c \rrbracket_{\mathfrak{m}} \mid \llbracket c \rrbracket_{\mathfrak{m}} \subseteq \llbracket (s)^{\rho} \cap \text{Imp} \rrbracket^{\rho}_{\mathfrak{m}} \} \end{aligned}$$

Then we have the following results

Lemma 3.2 *Given $\rho \in \text{uco}(\mathbb{V})$ additive, then $\forall e. \forall \mathfrak{m} \in \mathbb{M}^{\rho}. \llbracket \hat{\rho}(e) \rrbracket_{\mathfrak{m}} = \llbracket e \rrbracket^{\rho}_{\mathfrak{m}}$ (trivially implying $e' \in \hat{\rho}(e) \Leftrightarrow \forall \mathfrak{m} \in \mathbb{M}^{\rho}. \llbracket e' \rrbracket_{\mathfrak{m}} \subseteq \llbracket e \rrbracket^{\rho}_{\mathfrak{m}}$) and $\forall \Phi \in \wp(\Psi). \forall \mathfrak{m} \in \mathbb{M}^{\rho}. \llbracket \bar{\gamma}[\rho](\Phi) \rrbracket_{\mathfrak{m}} = \llbracket \Phi \rrbracket^{\rho}_{\mathfrak{m}}$.*

PROOF. Let us prove first the property for expressions by induction on the syntactic structure of e .

$e = n$: $\llbracket \hat{\rho}(e) \rrbracket_{\mathfrak{m}} = \llbracket \hat{\rho}(n) \rrbracket_{\mathfrak{m}} \stackrel{\text{def}}{=} \rho(n)$, while $\llbracket e \rrbracket^{\rho}_{\mathfrak{m}} = \llbracket n \rrbracket^{\rho}_{\mathfrak{m}} = \rho(n)$ (where $\llbracket n \rrbracket_{\mathfrak{m}} = n$);

$e = x$: $\llbracket \hat{\rho}(e) \rrbracket_{\mathfrak{m}} = \llbracket \hat{\rho}(x) \rrbracket_{\mathfrak{m}} \stackrel{\text{def}}{=} \llbracket x \rrbracket_{\mathfrak{m}} = \mathfrak{m}(x)$, while $\llbracket e \rrbracket^{\rho}_{\mathfrak{m}} = \llbracket x \rrbracket^{\rho}_{\mathfrak{m}} = \rho(\mathfrak{m}(x)) = \mathfrak{m}(x)$ (since $\mathfrak{m} \in \mathbb{M}^{\rho}$);

$e = e_1 \text{ op } e_2$: Suppose op any arithmetic or boolean operator.

$\llbracket \hat{\rho}(e) \rrbracket_{\mathfrak{m}} = \llbracket \hat{\rho}(e_1 \text{ op } e_2) \rrbracket_{\mathfrak{m}} \stackrel{\text{def}}{=} \llbracket \hat{\rho}(e_1) \text{ op } \hat{\rho}(e_2) \rrbracket_{\mathfrak{m}} = \llbracket \hat{\rho}(e_1) \rrbracket_{\mathfrak{m}} \text{ op } \llbracket \hat{\rho}(e_2) \rrbracket_{\mathfrak{m}} = \llbracket e_1 \rrbracket^{\rho}_{\mathfrak{m}} \text{ op } \llbracket e_2 \rrbracket^{\rho}_{\mathfrak{m}}$ by inductive hypothesis. But this is precisely $\llbracket e_1 \text{ op } e_2 \rrbracket^{\rho}_{\mathfrak{m}}$ since op is computed on the semantics as additive lift to sets.

Analogously, we can prove all the other cases.

Now, let us prove the fact for CFG single edge labels, again by induction on the syntactic structure. Note that, being ρ additive then also $\llbracket \cdot \rrbracket^{\rho}$ is additive, being also the concrete semantics additive on sets of statements.

$$\begin{aligned} \llbracket \bar{\gamma}[\rho](x:=e) \rrbracket_{\mathfrak{m}} &= \llbracket x := \hat{\rho}(e) \rrbracket_{\mathfrak{m}} \\ &= \sqcup \{ \llbracket x := e' \rrbracket_{\mathfrak{m}} \mid e' \in \hat{\rho}(e) \} \\ &= \sqcup \{ \mathfrak{m}[x / \llbracket e' \rrbracket_{\mathfrak{m}}] \mid e' \in \hat{\rho}(e) \} \\ &= \mathfrak{m}[x / \sqcup \{ \llbracket e' \rrbracket_{\mathfrak{m}} \mid e' \in \hat{\rho}(e) \}] \\ &= \mathfrak{m}[x / \sqcup \{ \llbracket e' \rrbracket_{\mathfrak{m}} \mid \llbracket e' \rrbracket_{\mathfrak{m}} \subseteq \llbracket e \rrbracket^{\rho}_{\mathfrak{m}} \}] \\ &= \mathfrak{m}[x / \llbracket e \rrbracket^{\rho}_{\mathfrak{m}}] = \llbracket x := e \rrbracket^{\rho}_{\mathfrak{m}} \end{aligned}$$

$$\begin{aligned} \llbracket \bar{\gamma}[\rho](b) \rrbracket_{\mathfrak{m}} &= \llbracket \hat{\rho}(b) \rrbracket_{\mathfrak{m}} \\ &= \sqcup \{ \llbracket b' \rrbracket_{\mathfrak{m}} \mid b' \in \hat{\rho}(b) \} \\ &= \sqcup \{ \mathfrak{m} \cap \sqcup \{ \mathfrak{m} \mid \llbracket b' \rrbracket_{\mathfrak{m}} = \text{true} \} \mid b' \in \hat{\rho}(b) \} \\ &= \mathfrak{m} \cap \sqcup \{ \mathfrak{m} \mid \llbracket b' \rrbracket_{\mathfrak{m}} = \text{true}, b' \in \hat{\rho}(b) \} \\ &= \mathfrak{m} \cap \sqcup \{ \mathfrak{m} \mid \llbracket b' \rrbracket_{\mathfrak{m}} = \text{true}, \llbracket b' \rrbracket_{\mathfrak{m}} \subseteq \llbracket b \rrbracket^{\rho}_{\mathfrak{m}} \} \\ &= \mathfrak{m} \cap \sqcup \{ \mathfrak{m} \mid \text{true} \in \llbracket b \rrbracket^{\rho}_{\mathfrak{m}} \} = \llbracket b \rrbracket^{\rho}_{\mathfrak{m}} \end{aligned}$$

$$\begin{aligned}
\llbracket \bar{\gamma}[\rho](\text{eval}(s)) \rrbracket_{\mathfrak{m}} &= \llbracket \text{eval}(\hat{\rho}(s)) \rrbracket_{\mathfrak{m}} \\
&= \bigsqcup \{ \llbracket c \rrbracket_{\mathfrak{m}} \mid \llbracket c \rrbracket_{\mathfrak{m}} \subseteq \llbracket (s)^\rho \mathbin{\frown} \text{Imp} \rrbracket_{\mathfrak{m}}^\rho \} \\
\text{By additivity of } \llbracket \cdot \rrbracket^\rho &= \llbracket (s)^\rho \mathbin{\frown} \text{Imp} \rrbracket_{\mathfrak{m}}^\rho = \llbracket \text{eval}(s) \rrbracket_{\mathfrak{m}}^\rho
\end{aligned}$$

Finally, for each set of labels Φ , we have that $\llbracket \bar{\gamma}[\rho](\Phi) \rrbracket_{\mathfrak{m}} = \bigsqcup_{\varphi \in \Phi} \llbracket \bar{\gamma}[\rho](\varphi) \rrbracket_{\mathfrak{m}} = \bigsqcup_{\varphi \in \Phi} \llbracket \varphi \rrbracket_{\mathfrak{m}}^\rho = \llbracket \Phi \rrbracket_{\mathfrak{m}}^\rho$, since all the involved functions are additive by definition or by construction. $\square$

Then we have that:

Theorem 3.3 *Let $\rho \in \text{uco}(\mathbb{V})$ additive, and $\eta \in \text{uco}(\wp(\Psi))$. Then $\bar{\eta}_\uparrow \stackrel{\text{def}}{=} \bar{\gamma}[\rho] \circ \eta$ satisfies Eq. 1.*

PROOF. It is worth noting that, we trivially have by abstraction that $\forall \varphi \in \Psi. \llbracket \eta_\uparrow(\varphi) \rrbracket \subseteq \llbracket \eta_\uparrow(\varphi) \rrbracket^\rho$. Let us prove the other implication: $\forall \varphi \in \Psi$

$$\begin{aligned}
\llbracket \eta_\uparrow(\varphi) \rrbracket &= \llbracket \bar{\gamma}[\rho] \circ \eta(\varphi) \rrbracket \\
&= \llbracket \bar{\gamma}[\rho] \circ \bar{\gamma}[\rho] \circ \eta(\varphi) \rrbracket && \text{[By properties of uco]} \\
&= \llbracket \bar{\gamma}[\rho] \circ \eta(\varphi) \rrbracket^\rho && \text{[By Lemma. 3.2]} \\
&= \llbracket \eta_\uparrow(\varphi) \rrbracket^\rho
\end{aligned}$$

$\square$

This result tells us that by taking a code abstraction more abstract than (or equal to) $\bar{\gamma}[\rho]$, we guarantee that the abstract interpretation ρ can be performed on the abstracted program (Eq. 1). We have so far proved that it is always possible to force Eq. 1, in order to make it possible to continue the analysis (observing ρ) also on the abstracted code. In the following we show how this framework can be integrated with the existing analysis of dynamic code [3] in order to improve its precision.

4 An improved dynamic code analysis

In this section we show how the constructive code abstraction characterization, provided in the previous section, can be used for representing the code approximation which soundly captures the potential code executed by a string-to-code statement. As we will show, without abstracting code, we cannot capture situations where the collecting semantics on strings generates sets of statements that cannot be represented by using the concrete syntax. Nevertheless, we must also observe that the analyzer cannot change dynamically with the generated code, hence the abstraction *must* be driven by the semantic property analyzed. This means that, without using the proposed framework, the analysis would surely be less precise in those situations where code abstraction becomes a necessity.

Let us summarize how we propose to exploit the framework:

- Consider a fixed semantic abstraction $\rho \in \text{uco}(\mathbb{V})$ and a corresponding static analyzer, designed in such a way that it can interpret also code abstracted by $\bar{\gamma}[\rho]$.
- Analyze the program, and when an `eval` is met, extract the language of its argument. If the language is infinite (under specific conditions that we will discuss) build the abstract CFG approximating it and extract the corresponding code abstraction η . In general, this code abstraction η is not more abstract than $\bar{\gamma}[\rho]$ (the code abstraction already embedded in the static analyzer, depending only on ρ);
- Build $\bar{\gamma}[\rho] \circ \eta$ in order to make also the generated code (approximated by the generated abstract CFG) analyzable by the static analysis for ρ .

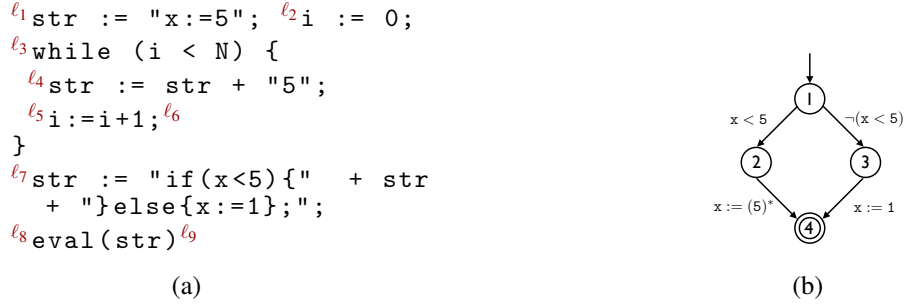


Figure 5: (a) Dynamically-generating code sample. (b) CFG associated to `str`.

Analyzing dynamic code. Let ρ be a static analysis performing in particular $\rho_s \in uco(\wp(\mathbb{S}))$ on strings, where $\mathbb{S} = \mathcal{K}^*$ denotes strings over a finite alphabet $\mathcal{K}$. Note that, our analyzer has to work on any (abstract) CFG that can be dynamically generated, hence it has to be designed with this purpose in mind. In particular, as we will show, we will generate only abstract CFGs with a code abstraction η complete w.r.t. ρ . This means, by construction, that η must be more abstract than $\bar{\Upsilon}[\rho]$, which means that each set of elements in η corresponds to a subset of the elements (abstract predicates) of $\bar{\Upsilon}[\rho]$. Hence, in order to guarantee to interpret predicates in any η complete, it is sufficient to design the analyzer soundly interpreting any abstract predicate in $\bar{\Upsilon}[\rho]$. For instance, $\bar{\Upsilon}[\text{Sign}]$ is the abstraction containing all the predicates, involving integers, of the form $x := S$, $x < S$, etc, with $S \in \text{Sign}$, e.g., an abstract predicate is $x := \mathbb{Z}^+$, and the analyzer for `Sign` should be able to interpret also such abstract predicates.

Let x be the input string parameter of an `eval` statement, we denote by $\mathcal{S}^{\rho_s}(x)$ the abstract value for x computed by the analysis on ρ_s . For example, suppose that the collection of values for the string x before the `eval` is $\{a:=0, a:=1\}$. By defining ρ_s as the k -bounded string set abstract domain [1], with $k = 2$, $\mathcal{S}^{\rho_s}(x) = \{a:=0, a:=1\}$, while by using the prefix abstract domain $\overline{\mathcal{PR}}$ [8], $\mathcal{S}^{\overline{\mathcal{PR}}}(x) = \{a:=s \mid s \in \mathbb{S}\}$. When the abstracted string and the abstraction is clear from the context, we simply denote this set by $\mathcal{S}$ and we assume (for the sake of simplicity) that any string in $\mathcal{S}$ is an executable language statement¹⁰. In the following, we abuse notation by denoting $\mathcal{S}$ also the automaton recognizing the language.

Consider for example, the program reported in Fig. 5a, a program building and manipulating the string `str` at run-time, which is, afterwards, interpreted as executable code, being the input parameter of the string-to-code statement `eval`. Since the value of `N` is unknown at compile-time, we cannot predict the precise number of iterations of the `while`-loop. In this case, a suitable string abstract analysis would approximate the value of `str`, before the `eval` execution, to an abstract value corresponding to an over-approximation of the possible values for `str`, which may be also, due to abstraction, an infinite set of strings, and therefore an infinite set of possible programs. For instance, in the example, if we abstract strings into the regular expression abstract domain [7] (or equivalently into the finite state automata abstract domain [2]), the value of `str` after the `while` loop will be the abstract value $x := 5(5)^*$; corresponding to an infinite set of programs, i.e., $x:=5$, $x:=55$, $x:=555$;... In this case, the common practice for analyzing `eval` is simply to give up with the analysis, for example by halting the analysis throwing an exception [18] or forbidding its usage [19].

Let $\rho_{\mathcal{CS}}$ be the abstract domain for all the possible values (integers, strings and booleans) [3]. Note

¹⁰Note that, this assumption corresponds to a decidable condition, hence it is possible to check it and to implement ad hoc solutions when it does not hold.

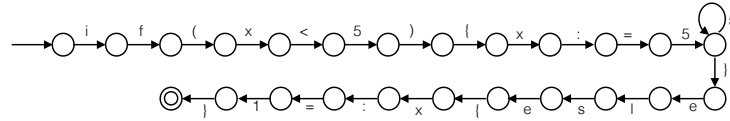


Figure 6: Finite state automaton corresponding to the abstract value of `str`.

that, $\bar{\Upsilon}[\rho_{\mathcal{E}, \mathcal{S}}]$ contains, for integers, predicates like the ones in the abstract CFG in Fig. 4.

The analysis $\rho_{\mathcal{E}, \mathcal{S}}$ at point ℓ_3 , due to widening¹¹ applied in the analysis of the while loop [2], abstracts the value of `str` in the infinite language $\{x:=s \mid s \in (5)^+\}$ (namely `x` is assigned to any value represented by a finite sequence of 5). Hence, at point ℓ_8 the analysis abstracts `str` to the strings set $\mathcal{S}_{\text{str}} = \{ \text{if}(x < 5) \{x:=s\} \text{else} \{x:=1\} \mid s \in (5)^+ \}$ meaning that, the true-branch of the string that may be transformed by `eval` may be either `x:=5`, or `x:=55`, or `x:=555`, ... The automaton corresponding to the abstract value of `str` is reported in Fig. 6, and it denotes an infinite language, i.e., an infinite set of possible statements. Unfortunately, this is a problem for the analysis provided in [3], where the language containing all the possible strings would be returned, losing any precision.

Generating the code: From automata to CFGs. At this point, we have the (potentially infinite) language of the `eval` argument (and hence an automaton $\mathcal{S}$), and the goal is to generate a CFG modeling an over-approximation of the executable code contained in the language of the automaton $\mathcal{S}$. The idea is to generate a CFG from a language of strings, i.e., from an automaton, by performing a parsing on the paths of the automaton. Indeed, we have defined and implemented an algorithm¹², reported in Alg. 1, performing an abstract parser on automata that, given an automaton $\mathcal{S}$, returns the CFG $\mathcal{P}$ that over-approximates, for each $s \in \mathcal{S}$ (executable), the concrete execution of `eval`.

The idea of Alg. 1 is to perform a depth-first search on the automaton and, when a language statement is recognized, to generate an edge in the CFG. This phase is handled by lines 3-13 of Alg. 1, building the set of nodes *Nodes* and the set of edges *Edges* of the resulting CFG $\mathcal{P}$. The set W contains the states of the finite state automaton for which we still have to generate edges in the CFG and it is initialized, at line 2, with the initial state q_0 . At this point, Alg. 1 looks for language statements readable from any path of the input automaton starting from a state q , taken from W , by means of the module `ReduceStmts` (line 5). In particular, `ReduceStmts` returns a set of triples (q', c, q'') , where each returned triple means that from $q' \in Q$ to $q'' \in Q$ a language statement c has been recognized. The set returned by `ReduceStmts` corresponds to the set of statements of $\mathcal{P}$ readable from the state q , hence they are added to *Edges*, substituting the reached states with the corresponding labels by means of the function `lab` (lines 7-8). At this point, we need to look for the statements that can be read from q'' , hence, q'' is added to W in order to be eventually processed at the next iterations of the while loop at lines 3-13. When there are no more states of $\mathcal{S}$ to be processed, namely when W is empty, the CFG $\mathcal{P} = \langle \text{Nodes}, \text{Edges} \rangle$ is returned (line 14), with entry label `lab`(q_0) and exit labels the ones associated with the states in F .

Problems arise when the automaton contains cycles (namely, when the automaton denotes an infinite language). In this case, Alg. 1 first transforms, at line 1, the input automaton, over the alphabet $\mathcal{K}$, in an automaton without cycles, over the alphabet $\mathcal{K} \cup \wp(\mathcal{K}^*)$, by means of the module `ReduceCycles`.

¹¹Widening is a fix-point accelerator used in infinite domains with infinite ascending chains, namely where the semantic fix-point computation may diverge. In this case we use a widening on automata defined in [10]

¹²In the following, we only discuss the main parts of the algorithm for space limitations.

Algorithm 1:

Input: $\mathcal{S} = (Q, \mathcal{K}, \delta, q_0, F)$
Output: CFG $\mathcal{P}$ over-approximating executable strings of $\mathcal{S}$

```

1  $\mathcal{S} = \text{ReduceCycles}(\mathcal{S})$ ;
2  $Nodes \leftarrow \emptyset$ ;  $Edges \leftarrow \emptyset$ ;  $W \leftarrow \{q_0\}$ ;  $visited \leftarrow \emptyset$ ;
3 while  $W \neq \emptyset$  do
4   select and remove  $q$  from  $W$ ;
5    $stmts \leftarrow \text{ReduceStmts}(\mathcal{S}, q)$ ;
6   foreach  $(q', c, q'') \in stmts$  do
7      $Nodes \leftarrow Nodes \cup \{\text{lab}(q'), \text{lab}(q'')\}$ ;
8      $Edges \leftarrow Edges \cup \{(\text{lab}(q'), c, \text{lab}(q''))\}$ ;
9      $visited \leftarrow visited \cup \{q'\}$ ;
10     $W \leftarrow W \cup \{q''\}$ ;
11     $W \leftarrow W \setminus visited$ ;
12  end
13 end
14 return  $\mathcal{P} = \langle Nodes, Edges \rangle$ ;

```

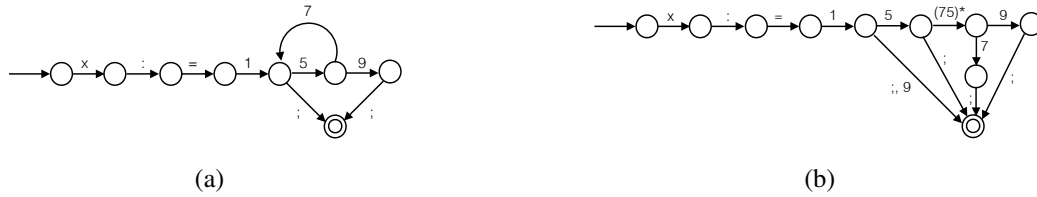


Figure 7: (a) Finite state automaton with cycle. (b) Result of ReduceCycles.

Given an input automaton $\mathcal{S}$, we retrieve the cycles of $\mathcal{S}$ using the well-known Tarjan's algorithm [27] for identifying cycles. Then, for each detected cycle of $\mathcal{S}$, we check whether the string read by the cycle is a whole statement r or not. In the first case, we substitute the cycle of the string r in the automaton, i.e., r^* , with the automaton reading the string corresponding to the statement `while(true){  $r$  }` over the alphabet $\mathcal{K}$. Otherwise, if the cycle does not read a whole statement, the idea is to collapse the cycle in a single transition, labeled with the regular expression corresponding to what is read in the cycle, i.e., denoting a set of string on $\mathcal{K}$ ($\wp(\mathcal{K}^*)$). Hence the resulting automaton is on the alphabet $\mathcal{K} \cup \wp(\mathcal{K}^*)$. In Fig. 7 we report an example of application of ReduceCycles algorithm. As example note that, by applying Alg. 1 to the automaton for $\mathcal{S}_{\text{str}}$ in Fig. 6, we generate the CFG $\mathcal{P}_{\text{str}}$, depicted in Fig. 5b. It is worth noting that the CFG obtained so far may contain abstract expressions on edges, hence edges may represent an infinite collection of statements. At this point, we need to approximate these edges for making it possible to analyze the CFG.

Making the code analyzable: Abstracting the CFG. Let us recall that we have to perform the analysis ρ also on the resulting code, in order to continue the static analysis. Hence, as observed before, we have to combine the code abstraction corresponding to the generated (abstract) CFG with the code abstraction induced by the semantic abstraction ρ , i.e., $\bar{\Upsilon}[\rho]$, which models, as code abstraction, the analysis. First of all, we have to formally characterize the abstraction η induced by the construction of the CFG given above, namely we characterize how the construction abstracts together different predicates. Let us build a code abstraction starting from the CFG $\mathcal{P} = \langle Nodes, Edges \rangle$ built in Alg. 1: In particular, let $\text{Merge} \stackrel{\text{def}}{=} \{ \{ \varphi \in \Psi \mid \langle \ell', \varphi, \ell'' \rangle \in Edges \} \mid \ell', \ell'' \in Nodes \} \subseteq \Psi$ be the set of collections of predicates

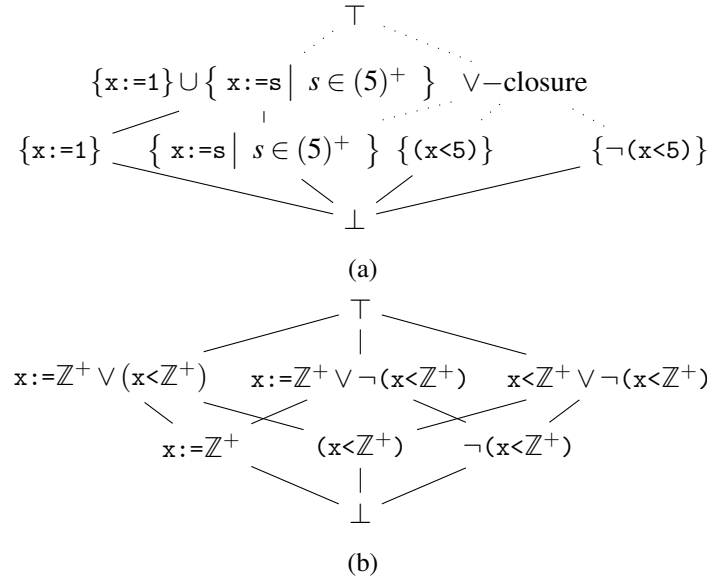


Figure 8: (a) Code abstraction $\eta^{\mathcal{P}_{\text{str}}}$ w.r.t. the CFG reported in Fig. 5b, (b) Code abstraction $\bar{\gamma}[\rho_{\ell, S}]^{\mathcal{P}_{\text{str}}}$

between any pair of states in the CFG, we define

$$\eta^{\mathcal{P}}(\wp(\Psi)) \stackrel{\text{def}}{=} \wp(\{ X \in \text{Merge} \mid \forall Y \in \text{Merge} \setminus \{X\}. X \cap Y = \emptyset \}) \in \text{uco}(\wp(\Psi)) \quad (2)$$

Note that, this abstraction, being characterized starting from the CFG is defined only in terms of a finite subset of Ψ , namely on the predicates in the given CFG, i.e., $\Psi^{\mathcal{P}} \stackrel{\text{def}}{=} \Psi \cap \{ \phi \mid \langle \ell', \phi, \ell'' \rangle \in \text{Edges} \}$.

In the example, $\Psi^{\mathcal{P}_{\text{str}}}(\wp(\Psi)) = \{ \{x:=s \mid s \in (5)^+\}, \{x:=1\}, \{(x<5)\}, \{\neg(x<5)\} \}$, hence we have that $\eta^{\mathcal{P}_{\text{str}}} = \wp(\Psi^{\mathcal{P}_{\text{str}}})$, being $\Psi^{\mathcal{P}_{\text{str}}}$ already a partition. In Fig. 8a this abstraction is partially depicted.

Finally, we need to satisfy Eq. 1 (completeness) between the code abstraction $\eta^{\mathcal{P}}$, built so far, and the static analysis, modeled as a semantic abstraction ρ , performing ρ_s (introduced above) on strings. Clearly we have no guarantee that $\eta^{\mathcal{P}}$ satisfies Eq. 1, hence, we have to (further) abstract the CFG in order to guarantee completeness w.r.t. the performed static analysis, namely in order to make it possible to perform the given static analysis on the code in the generated CFG. As observed in the previous section, in order to force completeness, we have to combine the desired abstraction $\eta^{\mathcal{P}}$ on predicates, with the code abstraction $\bar{\gamma}[\rho]$. Formally, in order to allow this operation, since $\eta^{\mathcal{P}}$ is defined on $\Psi^{\mathcal{P}}$, we have to restrict also $\bar{\gamma}[\rho]$ on $\Psi^{\mathcal{P}}$ (denoted $\bar{\gamma}[\rho]^{\mathcal{P}}$). This abstraction is obtained by intersecting the meaning of each one of its elements (i.e., its concretization) with the set of predicates in the CFG. In the running example, we have to compute $\bar{\gamma}[\rho_{\ell, S}]^{\mathcal{P}_{\text{str}}}$, which is the code abstraction induced by the Sign on the predicates in $\mathcal{P}_{\text{str}}$. For instance, all the predicates in $\{x:=s \mid s \in (5)^+\}$ and the predicate $x:=1$ cannot be distinguished when integers are abstracted by observing only their signs, hence the resulting abstraction is depicted in Fig. 8b, where the abstract predicate $x:=\mathbb{Z}^+$ corresponds, in the concrete, to the set of predicates $\{x:=s \mid s \in (5)^+\} \cup \{x:=1\}$, while $x<\mathbb{Z}^+$ and $\neg(x<\mathbb{Z}^+)$ correspond, respectively, to $\{(x<5)\}$ and to $\{\neg(x<5)\}$ (all the other elements corresponds to $\perp$).

Finally, we aim at building a code abstraction which can be interpreted by the initial abstract interpreter ρ , namely, that satisfies Eq. 1. By Th. 3.3 such an abstraction is $\bar{\eta}_{\uparrow}^{\mathcal{P}} = \bar{\gamma}[\rho]^{\mathcal{P}} \circ \eta^{\mathcal{P}}$.

Corollary 4.1 *Let $\rho \in \text{uco}(\mathbb{V})$ be additive. Then the code abstraction $\bar{\eta}_{\uparrow}^{\mathcal{P}} = \bar{\gamma}[\rho]^{\mathcal{P}} \circ \eta^{\mathcal{P}} \in \text{uco}(\Psi^{\mathcal{P}})$ is complete w.r.t. the semantic abstraction ρ , i.e., it satisfies Eq. 1.*

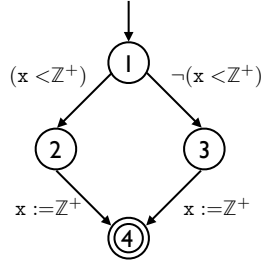


Figure 9: Abstract CFG generated by abstracting $\mathcal{P}_{\text{str}}$ by means of $\bar{\eta}_{\uparrow}^{\mathcal{P}_{\text{str}}}$

Hence, in our example, the code abstraction $\bar{\eta}_{\uparrow}^{\mathcal{P}_{\text{str}}} = \bar{\Upsilon}[\rho_{\mathcal{E}, \mathcal{S}}] \mathcal{P}_{\text{str}} \circ \eta^{\mathcal{P}_{\text{str}}}$ satisfies Eq. 1. In particular, we can observe that $\bar{\eta}_{\uparrow}^{\mathcal{P}_{\text{str}}} = \bar{\Upsilon}[\rho_{\mathcal{E}, \mathcal{S}}] \mathcal{P}_{\text{str}}$. Finally, we have to abstract the CFG $\mathcal{P}$, previously generated, by applying $\bar{\eta}_{\uparrow}^{\mathcal{P}}$ to each edge of the CFG. In our example, the so far resulting abstract CFG is reported in Fig. 9, where the abstract CFG generated by abstracting $\mathcal{P}_{\text{str}}$ by means of $\bar{\eta}_{\uparrow}^{\mathcal{P}_{\text{str}}}$ is depicted.

A taste of implementation. A static analyzer based on finite state automata is available at [2]. Moreover, we have implemented Alg. 1 in order to validate our approach¹³. The implementation of a static analysis of abstract CFGs is in an early stage development and it is left as future work. Nevertheless, it is able to parse executable automata and to abstract them into abstract CFGs, as we have previously described. In order to make these abstract CFGs effectively analyzable, we are currently extending the static analyzer, and the underlying abstract interpreter, to parse, and thus analyze, also abstract predicates.

5 Conclusion

We conclude highlighting the value, in the context of static analysis, of the framework presented in this paper. What we propose here is a precision improvement of [3], an analysis that attacks an extremely hard problem in static program analysis by abstract interpretation, since the standard static analysis assumption (i.e., the program code we want to analyze must be static) is broken when we have to deal with string-to-code statements. In [3], we have shown that even without this assumption, it is still possible for static analysis to semantically analyze dynamically mutating code in a meaningful and sound way. It has been the very first proof of concept for a sound static analysis for self-modifying code based on bounded reflection for a high-level script-like programming language. In this paper, we improve this approach by characterizing code transformations that do not lose precision w.r.t. a fixed abstract semantics/analysis of the code. The idea we develop consists in embedding the property to analyze in the code transformation in order to make the property analysis to work also on the transformed code (as it happens in dynamic code analysis). Hence, the main contribution is to make even more precise the first truly *dynamic static analyzer*, which has the feature to keep the analysis going on, even when code is dynamically built.

Clearly, the framework improved here is still at an early stage and surely there is much work to do, not only for the presented algorithm and the implementation, which has clearly to be further developed, but also for making the approach more precise and general. As far as the algorithm is concerned we have not explicitly provided soundness and completeness proofs or discussions. In particular, completeness

¹³Available at
<https://github.com/SPY-Lab/java-fsm-library/tree/abstract-parser>

holds under decidable hypotheses (the input automaton has to recognize only executable strings), here only briefly treated, and therefore these aspects need further formal development.

On the other hand, a direction for improving precision can be that of integrating the proposed static analysis in an hybrid solution, by using, for instance, taint analysis (or other dynamic analyses) for driving when to apply static analysis. Finally, we have considered only `eval` as string-to-code statement, while there are other ways, for dynamically executing code built out of strings, that should be investigated. However, we strongly believe that, the same approach used for `eval`, could be easily applied to any other string-to-code statement. Moreover, we believe that this framework could be instantiated in order to deal with other forms of code transformations, maybe by considering more general CFG abstractions.

From a more theoretical point of view, interesting future works consist in exploiting the proposed approach for analyzing code in order to investigate, on dynamic languages, several application contexts where static analysis by abstract interpretations has been exploited. First of all, we could trace (abstract) flows of information during execution [15, 22, 20, 21, 13, 17, 12] in order to tackle different security issues, such as the detection of (abstract) code injections [6, 5] or the formal characterization of dynamic code obfuscators and of their potency [11, 14]. Moreover, the ability to analyze malware code could be exploited for extracting code properties which could be used for analyzing code similarity [9], a technique useful for instance to identify or at least classify malicious code.

References

- [1] R. Amadini, G. Gange, F. Gauthier, A. Jordan, P. Schachte, H. Søndergaard, P. J. Stuckey & C. Zhang (2018): *Reference Abstract Domains and Applications to String Analysis*. *Fundam. Inform.* 158(4), pp. 297–326.
- [2] V. Arceri & I. Mastroeni (2018): *Static Program Analysis for String Manipulation Languages*. In: *Verification and Program Transformation 2018*, 299, EPTCS, pp. 19–33.
- [3] V. Arceri & I. Mastroeni (2021): *Analyzing Dynamic Code: A Sound Abstract Interpreter for evil eval*. *ACM Trans. Priv. Secur.* 24(2), pp. 10:1–10:38.
- [4] V. Arceri, I. Mastroeni & S. Xu (2020): *Static Analysis for ECMAScript String Manipulation Programs*. *Appl. Sci.* 10, p. 3525, doi:10.3390/app10103525.
- [5] M. Balliu & I. Mastroeni (2010): *A Weakest Precondition Approach to Robustness*. *Trans. Comput. Sci.* 10, pp. 261–297.
- [6] S. Buro & I. Mastroeni (2018): *Abstract Code Injection - A Semantic Approach Based on Abstract Non-Interference*. In Isil Dillig & Jens Palsberg, editors: *Verification, Model Checking, and Abstract Interpretation - 19th International Conference, VMCAI 2018, Los Angeles, CA, USA, January 7-9, 2018, Proceedings, Lecture Notes in Computer Science* 10747, Springer, pp. 116–137.
- [7] T. Choi, O. Lee, H. Kim & K. Doh (2006): *A Practical String Analyzer by the Widening Approach*. In: *Programming Languages and Systems, 4th Asian Symposium, APLAS 2006, Sydney, Australia, November 8-10, 2006, Proceedings*, pp. 374–388.
- [8] G. Costantini, P. Ferrara & A. Cortesi (2015): *A suite of abstract domains for static analysis of string values*. *Softw., Pract. Exper.* 45(2), pp. 245–287.
- [9] M. Dalla Preda, R. Giacobazzi, A. Lakhota & I. Mastroeni (2015): *Abstract Symbolic Automata: Mixed syntactic/semantic similarity analysis of executables*. In: *Proceedings of the 42nd ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages, POPL 2015*, ACM SIGPLAN Notices 50 (1), pp. 329–341.
- [10] V. D’Silva (2006): *Widening for Automata*. Diploma Thesis, Institut Fur Informatik, Universitat Zurich.

- [11] R. Giacobazzi, N. D. Jones & I. Mastroeni (2012): *Obfuscation by Partial Evaluation of Distorted Interpreters*. In O. Kiselyov & S. Thompson, editors: *Proc. of the ACM SIGPLAN Symp. on Partial Evaluation and Semantics-Based Program Manipulation (PEPM'12)*, ACM Press, pp. 63 – 72.
- [12] R. Giacobazzi & I. Mastroeni (2004): *Proving Abstract Non-Interference*. In A. Tarlecki J. Marcinkowski, editor: *Annual Conf. of the European Association for Computer Science Logic (CSL '04)*, 3210, Springer-Verlag, pp. 280–294.
- [13] R. Giacobazzi & I. Mastroeni (2010): *A Proof System for Abstract Non-interference*. *J. Log. Comput.* 20(2), pp. 449–479.
- [14] R. Giacobazzi & I. Mastroeni (2012): *Making Abstract Interpretation Incomplete: Modeling the Potency of Obfuscation*. In Antoine Miné & David Schmidt, editors: *Static Analysis - 19th International Symposium, SAS 2012, Deauville, France, September 11-13, 2012. Proceedings, Lecture Notes in Computer Science 7460*, Springer, pp. 129–145.
- [15] R. Giacobazzi & I. Mastroeni (2018): *Abstract Non-Interference: A Unifying Framework for Weakening Information-flow*. *ACM Trans. Priv. Secur.* 21(2), pp. 9:1–9:31.
- [16] R. Giacobazzi & E. Quintarelli (2001): *Incompleteness, Counterexamples, and Refinements in Abstract Model-Checking*. In: *Static Analysis, 8th International Symposium, SAS 2001, Paris, France, July 16-18, 2001, Proceedings*, pp. 356–373.
- [17] Roberto Giacobazzi & Isabella Mastroeni (2010): *Adjoining classified and unclassified information by abstract interpretation*. *J. Comput. Secur.* 18(5), pp. 751–797.
- [18] S. H. Jensen, P. A. Jonsson & A. Möller (2012): *Remedying the eval that men do*. In: *International Symposium on Software Testing and Analysis, ISSTA 2012, Minneapolis, MN, USA, July 15-20, 2012*, pp. 34–44.
- [19] V. Kashyap, K. Dewey, E. A. Kuefner, J. Wagner, K. Gibbons, J. Sarracino, B. Wiedermann & B. Hardekopf (2014): *JSAI: a static analysis platform for JavaScript*. In: *Proceedings of the 22nd ACM SIGSOFT International Symposium on Foundations of Software Engineering, (FSE-22), Hong Kong, China, November 16 - 22, 2014*, pp. 121–132.
- [20] I. Mastroeni (2013): *Abstract interpretation-based approaches to Security - A Survey on Abstract Non-Interference and its Challenging Applications*. In: *Semantics, Abstract Interpretation, and Reasoning about Programs: Essays Dedicated to David A. Schmidt on the Occasion of his Sixtieth Birthday, EPTCS 129*, pp. 41–65.
- [21] I. Mastroeni & D. Nikolic (2010): *Abstract Program Slicing: From Theory towards an Implementation*. In Jin Song Dong & Huibiao Zhu, editors: *Formal Methods and Software Engineering - 12th International Conference on Formal Engineering Methods, ICFEM 2010, Shanghai, China, November 17-19, 2010. Proceedings, Lecture Notes in Computer Science 6447*, Springer, pp. 452–467.
- [22] I. Mastroeni & D. Zanardini (2017): *Abstract Program Slicing: An Abstract Interpretation-Based Approach to Program Slicing*. *ACM Trans. Comput. Log.* 18(1), pp. 7:1–7:58.
- [23] N. Mavrogiannopoulos, N. Kisserli & B. Preneel (2011): *A taxonomy of self-modifying code for obfuscation*. *Computers & Security* 30(8), pp. 679–691. Available at <http://dx.doi.org/10.1016/j.cose.2011.08.007>.
- [24] Antoine Miné (2013): *Static analysis by abstract interpretation of concurrent programs. (Analyse statique par interprétation abstraite de programmes concurrents)*. Available at <https://tel.archives-ouvertes.fr/tel-00903447>.
- [25] Gregor Richards, Christian Hammer, Brian Burg & Jan Vitek (2011): *The Eval That Men Do - A Large-Scale Study of the Use of Eval in JavaScript Applications*. In: *ECOOP 2011 - Object-Oriented Programming - 25th European Conference, Lancaster, UK, July 25-29, 2011 Proceedings*, pp. 52–78.
- [26] H. Seidl, R. Wilhelm & S. Hack (2012): *Compiler Design - Analysis and Transformation*. Springer.
- [27] R. E. Tarjan (1972): *Depth-First Search and Linear Graph Algorithms*. *SIAM J. Comput.* 1(2), pp. 146–160, doi:10.1137/0201010.