

Isolating JavaScript with Filters, Rewriting, and Wrappers

Sergio Maffei¹, John C. Mitchell², and Ankur Taly²

¹ Imperial College London

² Stanford University

Abstract. We study methods that allow web sites to safely combine JavaScript from untrusted sources. If implemented properly, filters can prevent dangerous code from loading into the execution environment, while rewriting allows greater expressiveness by inserting run-time checks. Wrapping properties of the execution environment can prevent misuse without requiring changes to imported JavaScript. Using a formal semantics for the ECMA 262-3 standard language, we prove security properties of a subset of JavaScript, comparable in expressiveness to Facebook FBJS, obtained by combining three isolation mechanisms. The isolation guarantees of the three mechanisms are interdependent, with rewriting and wrapper functions relying on the absence of JavaScript constructs eliminated by language filters.

1 Introduction

Web sites such as OpenSocial [18] platforms, iGoogle [10], Facebook [7], and Yahoo!’s Application Platform [28] allow users of the site to build gadgets, which we will refer to as *applications*, that will be served to other users when they visit the site. In the general scenario represented by these sites, application developers would like to use an expressive implementation language like JavaScript, while the sites need to be sure that applications served to users do not present security threats. In the view of the hosting site and its visitors, the containing page (for example, an iGoogle page) is “trusted,” while applications included in it are not; untrusted applications could try to steal cookies, navigate the page or portions of it [3], replace password boxes with controls of their own, or mount other attacks [4]. While hosting sites can use browser *iframe* isolation, iframes require structured inter-frame communication mechanisms [3,4]. Just as OS inter-process isolation is useful in some situations, while others require language-based isolation between lightweight threads in the same address space, we expect that both iframes and language-based isolation will be useful in future Web applications. While some straightforward language-based checks make intuitive sense, JavaScript [6,8] provides many subtle ways for malicious code to subvert language-based isolation methods, as demonstrated here and in previous work [17]. We therefore believe it is important to develop precise definitions and techniques that support security proofs for mechanisms used critically in popular modern Web sites.

In this paper, we devise and analyze a combination of isolation mechanisms for a subset of ECMA 262-3 [11] JavaScript that is comparable in expressiveness to Facebook [7] FBJS [23]. Isolation from untrusted code in our subset of JavaScript is based on filtering out certain constructs (`eval`, `Function`, `constructor`), rewriting others (`this`, `e1[e2]`) to allow them to be used safely, and wrapping properties (*e.g.*, object and array prototype properties) of the execution environment to further limit the impact of untrusted code. Our analysis and security proofs build on a formal foundation for proving isolation properties of JavaScript programs [17], based on our operational semantics of the full ECMA-262 Standard language (3rd Edition) [11], available on the web [13] and described previously [14]. While we focus on one particular combination of filters, rewriting functions, and wrappers, our methods are applicable to variants of the specific subset we present. In particular, DOM functions such as `createElement` could be allowed, if suitable rewriting is used to insert checks on the string arguments to `eval` at run-time.

While Facebook FBJS uses filters, source-to-source rewriting, and wrappers, we have found several attacks on FBJS using our methods, presently and as reported in previous work [17]. These attacks allow a Facebook application to access arbitrary properties of the hosting page, violating the intent of FBJS. Each was addressed promptly by the Facebook team within hours of our reports to them. While the safe subset of JavaScript we present here is very close to current FBJS, we consider it a success that we were able to contribute to the security of Facebook through insights obtained by our semantic methods, and a success that in the end we are able to provide provable guarantees for a subset of JavaScript that is essentially similar to one used by external application developers for a hugely popular current site.

Related work on language-based methods for isolating the effects of potentially malicious web content include [21], which examines ways to inspect and cleanse dynamic HTML content, and [29], which modifies questionable JavaScript, for a more restricted fragment of JavaScript than we consider here. A short workshop paper [27] also gives an architecture for server-side code analysis and instrumentation, without exploring details or specific methods for constraining JavaScript. The Google Caja [4] project uses an approach based on transparent compilation of JavaScript code into a safe subset with libraries that emulate DOM objects. Additional related work on rewriting-based methods for controlling the execution of JavaScript include [19]. Foundational studies of limited subsets of JavaScript and dynamic languages in general are reported in [2,25,29,9,22,1,26]; see [14]. In previous work [17], we described problems with then-current FBJS and proposed a safe subset based on filtering alone. The present paper includes a new FBJS vulnerability related to rewriting and extends our previous analysis to rewriting and wrapper functions. This produces a far more expressive safe subset of JavaScript. The workshop paper [16] describes some intermediate results on rewriting without wrappers.

The rest of this paper is organized as follows. In Section 2, we describe the basic isolation problem, our threat model, and the isolation mechanisms we use.

In Section 3, we briefly review our previous work [14] on JavaScript operational semantics and discuss details of JavaScript that are needed to understand isolation problems and their solution. In Section 4, we motivate and define the specific filter, rewriting, and wrapper mechanism we use and state our main theorem about the isolation properties they provide. In Section 5, we compare our methods to those used in FBJS, with discussion of related work in Section 6. Concluding remarks are in Section 7.

2 The JavaScript Isolation Problem

The isolation problem we consider in this paper arises when a hosting page P_{host} includes content $P_1, \dots, P_k$ from untrusted origins that will execute in the same JavaScript environment as P_{host} . We assume that $P_1, \dots, P_k$ may try to maliciously manipulate properties of objects defined or used by P_{host} , and therefore consider $P_1, \dots, P_k$ under control of an attacker. The isolation mechanisms we provide are intended to be used by a site that has access to $P_1, \dots, P_k$ before they loaded in the browser execution environment. In practice, this may be achieved if the page and its constituents are aggregated at a site, or if there is some proxy in front of the browser that identifies and modifies trusted and untrusted JavaScript. While Facebook is a good example, with trusted content developed by Facebook containing untrusted user-defined applications, we develop general solutions that can be used in other scenarios that allow untrusted JavaScript to be identified and processed in advance of rendering and execution of content.

The basic defenses we provide involve changing the definitions of objects or properties in the hosting page P_{host} so that untrusted components $P_1, \dots, P_k$ run in a modified environment, filtering $P_1, \dots, P_k$ so that they must be expressed in a restricted subset of JavaScript, or rewriting $P_1, \dots, P_k$ to change their semantics in some way. While potentially dangerous constructs can be eliminated by filtering, allowing them to be rewritten may provide greater programming expressiveness. While generally there may be an arbitrary number of untrusted components, we will simplify notation and discuss the problem of a program P_{host} containing two untrusted subprograms P_1 and P_2 . We consider two untrusted subprograms instead of one because it is important to account for possible interaction between P_1 and P_2 .

Attacker model. An attacker may design malicious JavaScript code that runs in the context of a honest page. If the honest page contains two untrusted subprograms P_1 and P_2 from different origins, then these may both be under control of a single attacker, or one may be honest and the other provided by the attacker. In the event that P_1 is honest and P_2 malicious, for example, the attacker is considered successful if execution of P_2 accesses or modifies sensitive properties of either P_1 or the hosting page P_{host} .

Sensitive Properties and Challenges. In general, different hosting pages may have different security requirements, and application developers may wish to express security requirements in some way. However, expressing and enforcing

custom policies is beyond the scope of this paper. Instead, we focus on protecting a hosting page and any honest components in the following ways.

Restricting Access to Native Properties.. While memory safety is often the bottom line for language-based isolation mechanisms, JavaScript does not provide direct access to memory locations. The analogous bottom line for JavaScript isolation is preventing an attacker with control of one or more applications from accessing security-critical properties of native objects (in the context of web pages, this will also include DOM objects) used by the hosting page or by other applications. In JavaScript, there are three ways to directly access a property x of a generic object o : by $o.x$, by $o["x"]$, or by the identifier expression x if o is part of the current scope chain. Certain native objects such as `Array`, `Function`, and a few others can also be accessed indirectly, without naming a global variable. Although for certain purposes some of them may have to be made inaccessible, these objects themselves do not constitute sensitive resources *per se*. Therefore, we focus on direct access to native objects. In doing so, we assume that the hosting page has a list of security critical properties, which we call *blacklist* $\mathcal{B}$. Thus the first part of our isolation goal (formally stated in Section 4) is to prevent untrusted code from accessing any properties from the list $\mathcal{B}$. Although the isolation problem and the solution proposed in this paper are parametric on a blacklist, the way our solution is designed, it is completely straightforward to transform the solution to instead apply to a *whitelist* which is the set of all properties of native objects that *can* be exposed to untrusted code.

Isolating the Namespace of Untrusted Principals.. In our attacker model, a malicious application succeeds in attacking the system also if it can access properties defined by other honest applications. All untrusted application code is executed in the same global scope. Therefore, a secondary isolation goal is to separate out the set of global variables accessed by any two untrusted programs coming from different origins. In the solution we propose, we assume that each untrusted program P has an id pid_P associated with it which is unique for each origin, and we prefix all identifiers appearing in the program P with pid_P . This effectively separates the namespaces of two programs with different pids.

Enforcement Techniques. We analyze and prove the correctness of three techniques that are effective in protecting sensitive properties of honest code against an attacker that supplies code to be executed in the same JavaScript environment.

Filtering.. Untrusted code may be statically analyzed and rejected if it does not conform to certain criteria. In principle, filtering may range from simple syntactic checks to full-fledged static analysis, with obvious tradeoffs between efficiency and precision. Filtering takes place once, before untrusted code is loaded into the execution environment. Since filtering does not modify code, it does not affect the performance or the behavior of untrusted code that passes the filter.

Rewriting.. Selected constructs within untrusted code may be re-written. Typically, rewriting inserts run-time checks that prevent undesirable actions. While run-time checks impose a performance penalty, they are a valuable option for

constructs that are potentially dangerous but also useful when used appropriately in honest code. Rewritten code may execute differently from the original code, for example when a run-time security violation is detected.

Wrapping.. Sensitive resources of the trusted environment can be wrapped inside functions that use run-time checks to ensure that these resources are not used maliciously by untrusted code. Wrapper functions do not alter the untrusted code. When trusted code can access the wrapped resources directly, bypassing the wrappers, the run-time overhead or other down-sides of wrapping can be limited to untrusted code.

3 Design principles

In this Section we informally summarize the key features and insights that we gained while formalizing the operational semantics of JavaScript [13,14] based on the ECMA-262 standard [11].

We denote the ECMA-262 compliant subset of JavaScript by JS_{E2} . This paper deals with subsets of JS_{E2} . Our operational semantics consists of a set of rules written in a conventional meta-notation suitable for rigorous but (currently) manual proofs. Given the space constraints, we only describe informally the semantics of some of the unusual and interesting constructs which will help us in designing the isolation enforcement mechanisms in Section 4. Note that besides all terms derivable from the grammar (called *user* terms), our semantics introduces also certain internal terms, objects and properties useful to clearly express the evaluation semantics of user terms. None of these internal terms, objects and properties are visible in user code. Throughout the semantics, we use the symbol $@$ to distinguish user terms from internal terms.

Notations and Conventions. Our semantics is a small-step operational semantics ([20]). We represent objects as records of values *ov* indexed by strings *m* or internal identifiers $@x$. The record indexes are also called object properties. In JavaScript everything, including functions, is represented as an object. In our semantics the memory (or *heap* *H*) is a mapping from heap address (*l*) to objects. Object values (*ov*) are either pure values (*pv*) or function descriptions $\text{fun}(x, \dots)\{P\}$ or heap addresses. We refer to the union of the set of primitive values and heap addresses by *va*.

We use H_0 to denote the initial heap of JS_{E2} . It contains native objects for representing predefined functions, constructors and prototypes, and the global object $@Global$ that constitutes the initial scope, and is always the root of the scope chain. For example, the global object defines properties to store special values such as $\&NaN$ and $\&undefined$, functions such as *eval* and constructors to build generic objects, functions, numbers and arrays. In browsers, the global object is called *window*. We use l_g to denote the heap address of the global object.

The scope and prototype chains are two distinctive features of JavaScript. The stack is represented by a chain of objects whose properties represent the binding of local variables in the scope. Each scope object stores a pointer to its enclosing scope object in an internal $@Scope$ property. Representing the stack

as a chain of scope objects helps in dealing with the semantics of constructs that modify the scope chain, such as function calls and the `with` expression. JavaScript follows a prototype-based approach to inheritance. In our semantics, each object stores in an internal property `@Prototype` a pointer to its prototype object, and inherits its properties. At the root of the prototype tree there is `@Object.prototype`, that has a `null` prototype. There are also other native prototype objects such as `Function.prototype`, `Array.prototype` etc., which are present at the top of the prototype chains for function, array objects.

We represent a program state as a triple (H, l, t) where H denotes the heap mapping locations to objects, l denotes the heap address of the *current scope object* and t denotes the term being evaluated. Terms t can be expressions, statements and programs. We use the notation $\mathcal{H}(S)$, $\mathcal{S}(S)$ and $\mathcal{T}(S)$ to denote heap, scope and term component of the state respectively. The general form of an evaluation rule is $\frac{\langle \text{Premise} \rangle}{S_1 \rightarrow S_2}$, meaning that if a certain premise is true then the state S_1 evaluates to a state S_2 . A *reduction trace* τ is the (possibly infinite) maximal sequence of states $S_1, \dots, S_n, \dots$ such that $S_1 \rightarrow \dots \rightarrow S_n \rightarrow \dots$. Given a state S , we denote by $\tau(S)$ the (unique) trace originating from S and, if $\tau(S)$ is finite, we denote by $Final(S)$ the final state of $\tau(S)$.

Property access. We now describe the semantics of various constructs which involve accessing properties of objects. By “accessing a property” we refer to either reading or writing the contents of the property. The evaluation of certain constructs, such as `p in o`, involve checking if the object `o` has a property `p`. We do not consider those events instances of property access. Property accesses can be *explicit* or *implicit*.

Explicit property access. These take place when a term explicitly names the property that is being read.

Fact 1 *There are only three kinds of expressions in JS_{E2} which can be used for explicit property access: `x`, `e.x` and `e1[e2]`.*

We now discuss the semantics of the expressions `x`, `e.x` and `e1[e2]`. The semantics of the identifier expression `x` is based on the scope and prototype lookup mechanism. The evaluation involves successively looking at objects on the scope chain, starting from the current scope object until we find an object which has the property `x` (either in it or in one of its prototypes). Thus the expression `x` can potentially involve access to property “`x`” of one of the objects (or its prototype) present on the current scope chain. The semantics of the standard *dot* notation `e.x` results in accessing property “`x`” of the object obtained by evaluating the expression `e`. Finally, the semantics of `e1[e2]` involves accessing the property name corresponding to the string form of the value obtained by evaluating `e2`. Thus the property that is accessed is constructed dynamically by evaluating an expression. concretely, the evaluation of `e1[e2]` goes through the following steps (informally): first `e1` is evaluated to a value `va1`, then `e2` to `va2`, then if `va1` is not an object it is converted into an object `o`, and similarly if `va2` is not a string it is converted into a string `m`. Finally, property `m` of object `o` is accessed:

$$e1[e2] \longrightarrow va1[e2] \longrightarrow va1[va2] \longrightarrow o[va2] \longrightarrow o[m]$$

Each of these steps, which precede the actual access of property m in o , may raise an exception or have other side effects.

Implicit property access. These take place when the property accessed is *not* named explicitly by the term, but is accessed as part of an intermediate evaluation step in the semantics. For example, the `toString` property is accessed implicitly by evaluating the expression `"a" + o`, which involves resolving the identifier `o` and then type converting it to a string, by calling its `toString` property. There are many other expressions whose execution involves implicit property accesses to native properties, and the complete set is hard to characterize. Instead, we enumerate the set of all property names that can be implicitly accessed.

Fact 2 [14]. *The set of all property names $\mathcal{P}_{nat}$ that can be accessed implicitly by JS_{E2} constructs is $\{0, 1, 2, \dots\} \cup \{\text{toString, toNumber, valueOf, length, prototype, constructor, message, arguments, Object, Array, RegExp}\}$*

Dynamic code generation. For example, the native function `eval` takes a string as an argument, parses it as a program, and evaluates the resulting program returning its final value. According to the operational semantics, in JavaScript there are only two constructs which can dynamically generate new code.

Fact 3 *The only JS_{E2} constructs which involve dynamic code generation (from strings to Programs) are the native functions pointed to by the properties `eval` and `Function` of the global object.*

Accessing the global object. Since controlling access to global object is crucial in isolating untrusted from trusted code, we explore the set of constructs that can be used to access the global object.

As our semantics is formulated, the global object for the initial heap state is only accessible via the internal properties `@scope` and `@this`. These internal properties can only be accessed as a side effect of the execution of other instructions. An analysis of our semantics shows that the contents of the `@scope` property are never returned as the final result of any evaluation step, and the only construct whose evaluation involves access to the `@this` property is the expression `this`. Besides using `this`, the global object can be returned by calling in the global scope the functions `valueOf` of `Object.prototype`, and `concat`, `sort` or `reverse` of `Array.prototype`. For example, `var f=Object.prototype.valueOf; f()` evaluates to the global object.

Fact 4 *The only JS_{E2} constructs that can return a pointer to the global object are: the expression `this`, the native method `valueOf` of `Object.prototype` and native methods `concat`, `sort` and `reverse` of `Array.prototype`.*

4 Safe JavaScript subset

In this Section, we formally state the isolation problem introduced in Section 2, and propose a solution based on filtering, rewriting and wrapping techniques.

As mentioned in Section 2, we consider web pages which include untrusted content $P_1, \dots, P_k$ in the JavaScript environment of the host page. We associate to each untrusted user program P a unique identifier pid_P , which corresponds to the origin from which the program was loaded. Given a heap H , let $Acc(H, P)$ be the set of property names accessed when P is executed against H in the global scope, and let $Acc_l(H, P)$ ($l \in dom(H)$) be the set of properties of the object at address l , accessed when P is executed against the heap H in the global scope.

Isolation Problem *Given a blacklist $\mathcal{B}$ of property names and untrusted programs $P_1, \dots, P_k$ with program ids $pid_{P_1}, \dots, pid_{P_k}$, find a meaningful subset $J_{sub}(\mathcal{B}) \subseteq JS_{E2}$, an appropriate wellformed initial heap state H_0^{sub} and a function $Enf : pid * J_{sub}(\mathcal{B}) \rightarrow JS_{E2}$ such that: (Goal 1) For all user programs P in the subset $J_{sub}(\mathcal{B})$ with program ids pid_P , $Acc(H_0, Enf(pid_P, P)) \cap \mathcal{B} = \emptyset$. (Goal 2) For any two untrusted programs P_1 and P_2 in the subset J_{sub} with program ids pid_{P_1} and pid_{P_2} respectively*

$$Acc_{l_g}(H_0, Enf(pid_{P_1}, P_1)) \cap Acc_{l_g}(H_0, Enf(pid_{P_2}, P_2)) \subseteq \mathcal{P}_{nat} \cup \mathcal{P}_{noRen}.$$

Goal (2), as stated above, is the most precise property isolating different applications that we are able to support using the current proof techniques. In future work, we plan to generalize this property to enforce isolation when the execution of applications is interleaved, introducing proof techniques able to handle the combination of alternative safety properties for each application.

Isolating blacklisted properties. In order to achieve Goal 1, we need to control all possible ways in which object properties can be accessed. As discussed in Section 3, there are two kinds of property accesses: explicit and implicit access, and for isolating blacklisted properties we need to control both of them. The implicit accesses are in general very difficult to control because given a term t , it is undecidable to statically decide the precise list of property names that will be accessed implicitly. On the positive side, from Fact 2, we know that the set of property names that would be accessed implicitly would be contained in the set $\mathcal{P}_{nat}$. In this work, we therefore assume that none of the properties from the set $\mathcal{P}_{nat}$ are blacklisted or in other words all implicit property accesses are considered safe and are allowed. From Fact 1 we know that x , $e.mp$ and $e1[e2]$ are the only expressions which can be used for explicitly accessing user properties. Hence, in order to restrict access to blacklisted properties we have to restrict the behavior of these expressions. In this work we combine the filtering approach of [17] to restrict the behavior of expressions x and $e.x$ with a rewriting based approach to restrict the behavior of $e1[e2]$.

Restricting x and $e.x$. The expressions x and $e.x$ can access a blacklisted property if the identifier name " x " is contained in the blacklist. In order to restrict this behavior we conservatively disallow all such expressions where " x " is contained in the blacklist.

Filter 1 *Disallow all terms which contain an identifier from the blacklist $\mathcal{B}$.*

This restriction mechanism will fail if dynamically generated code can contain blacklisted identifiers. From Fact 3 we know that JS_{E2} includes two primitive

functions which can be used to generate code dynamically. One approach to fixing this problem is to restrict all ways of accessing such functions. In the initial heap, this can be achieved by disallowing the identifiers `eval`, `Function` and `constructor`. Although this may be a restriction for full-blown JavaScript applications that use `eval` to parse JSON code, a recent study [12] shows that a low percentage of widgets use constructs like `eval`. Thus, we propose the following filtering step.

Filter 2 *Disallow all terms containing any of the identifiers `eval`, `Function`, or `constructor`.*

An alternative to the above filtering step is to define safe wrappers for the functions `eval` and `Function`. Such wrappers need to use a JavaScript expression to parse, filter and rewrite the string passed as an argument to the original functions. Proving such a JavaScript expression correct would complicate severely our analysis, and we leave for future work.

Restricting $e1[e2]$. We restrict the behavior of $e1[e2]$ by rewriting it to a safe expression. The main idea is to insert a run-time check in each occurrence of $e1[e2]$ to make sure that $e2$ does not evaluate to a blacklisted property name. We transform every access to a blacklisted property of an object into an access to the property `"bad"` of the same object (assuming $\mathcal{B}$ does not contain `"bad"`). Although this transformation seems easy, it is complicated by subtle details of the semantics of the expression $e1[e2]$. In view of our operational semantics for $e1[e2]$ we propose the following rewriting step.

Rewrite 1 *Rewrite every occurrence of $e1[e2]$ in a term by $e1[IDX(e2)]$, where,*

$$\begin{aligned} IDX(e2) &= (\$=e2, \{toString: function() \{ return (\$=\$String(\$), CHECK.\$) \} \}) \\ CHECK.\$ &= (\$BL[\$] ? "bad": \\ &\quad (\$ == "constructor" ? "bad": \\ &\quad (\$ == "eval" ? "bad": \\ &\quad (\$ == "Function" ? "bad": \\ &\quad (\$[0] == "\$" ? "bad".\$)))))) \end{aligned}$$

where $\$String$ refers to the original `String` constructor, $\$BL$ is a (blacklisted) global variable containing an object with all blacklisted property names initialized to `true`, and $\$$ is a reserved variable name.

In order to initialize the variables $\$String$ and $\$BL$ to their appropriate values, we propose the following (trusted) initialization code, that must be executed in the global scope of the initial heap.

Initialization Code 1 (T_{idx}) *Let $\{p_1, \dots, p_n\}$ be the blacklist $\mathcal{B}$.*

$$\text{var } \$String = String; \text{ var } \$ = "" ; \text{ var } \$BL = \{p_1:true; \dots, p_n:true\}.$$

The IDX code defined in the rewrite rule work as follows: evaluates (once and for all) $e2$ to a value $va2$ that is saved in the variable $\$$. It then creates a new object with a specially crafted `toString` property, and returns the address of this object as the final value $l2$. These steps correspond to the internal execution trace

$e1[\text{IDX}(e2)] \longrightarrow \text{val}[\text{IDX}(e2)] \longrightarrow \text{val}[l2] \longrightarrow o[l2]$. According to the JavaScript semantics, the evaluation of $o[l2]$ involves converting the object at address $l2$ to a string by calling the `toString` method of $l2$ that will return the result of converting $\$$ to a (sanitized) string. The conversion to a string is faithfully implemented by the expression $\$String(\$)$, which calls the native `String` method on $\$$. The expression $\$CHECK_ \$$, uses nested conditional expressions to return the string saved in $\$$ only if it is not set a blacklisted property.

To protect this mechanism from tampering, we also need to ensure that the properties $\$, \$String$ and $\$BL$ cannot be accessed by untrusted code. Similar restrictions need to be imposed on other variables needed by similar enforced mechanisms. Therefore, we impose the restriction that untrusted code cannot use identifier names beginning with $\$$, thus separating the namespaces of trusted and untrusted code.

Filter 3 *Disallow all terms which involve an identifier name beginning with $\$$.*

Note that the condition $\$[0] == "\$"? "bad":\$$ in the $\$CHECK_ \$$ expression already imposes this restriction on dynamically generated property names.

Isolating one program from another. In order to achieve Goal 2, we need to make sure that for two programs P_1 and P_2 with ids $pid_{P_1} = pid_{P_2}$, it is the case that

$$Accl_g(H_0^{sub}, Enf(pid_{P_1}, P_1)) \cap Accl_g(H_0^{sub}, Enf(pid_{P_2}, P_2)) = \emptyset$$

where $Accl_g(H_0, Enf(P_1))$ refers to the set of global object properties (or global variables) that are accessed during the entire evaluation trace of program P . As discussed in the previous subsection, it is very difficult to control implicit property accesses. Therefore we assume that accessing the same properties from the set $\mathcal{P}_{nat}$ is safe for both programs and weaken our goal to the following

$$Accl_g(H_0^{sub}, Enf(pid_{P_1}, P_1)) \cap Accl_g(H_0^{sub}, Enf(pid_{P_2}, P_2)) \subseteq \mathcal{P}_{nat}.$$

On analyzing our semantics, we found that properties of the global object can be accessed in two ways: (i) If the program can get a pointer l_g to the global object, then it can access properties of the global object directly by using one of the two expressions $l_g.x$ or $l_g[x]$. We isolate the property names accessed using these expressions by conservatively disallowing explicit access to the global object by untrusted code. (ii) Since the global object is also the base scope object, variable names appearing in a program can resolve to the global object thereby resulting in access to the corresponding property. In other words, evaluation of the expression x can potentially involve accessing the property x of the global object. We isolate the set of property names accessed in this way by uniquely prefixing all identifiers appearing in a program by its id, thereby separating out the namespaces of two programs with different ids.

From Fact 4 we know that a pointer to the global object can potentially be obtained by using the expression `this` or calling method `valueOf` of `Object.prototype` or methods `sort`, `reverse`, `concat` of `Array.prototype`. In [17] we used the filtering approach

and conservatively disallowed `this` and the identifiers `valueOf`, `sort`, `reverse`, `concat` from the language. In this work, we use the rewriting technique for restricting the behavior of `this` and the wrapping technique for the native methods.

Rewriting `this`. The main idea is to rewrite every occurrence of `this` in the user code to the expression `NOGLOBALTHIS` which returns the result of evaluating `this`, if it is not the global object, and `null` otherwise.

Rewrite 2 *Rewrite every occurrence of `this` by `NOGLOBALTHIS`, where `NOGLOBALTHIS` = `(this== $\$g$ ?null;this)`. and $\$g$ is a blacklisted global variable, initialized with the address of the global object.*

In order to initialize correctly $\$g$ with the global object, we use the following initialization code that must be executed in the global scope.

Initialization Code 2 (T_{ng}) `var  $\$g$  = this;`

Note that, Filter 3 and Rewrite 1 already enforce that untrusted code cannot access the trusted variable name $\$g$.

Wrapping Native methods. As opposed to [17], in this work we take the less conservative approach of wrapping the native methods in order to ensure that the value returned by them is never the heap address of the global object. The following trusted initialization code demonstrates the wrapping for the method `valueOf`.

Initialization Code 3 ($T_{valueOf}$)
 `$\$OPvalueOf$ = Object.prototype.valueOf;
 $\$OPvalueOf.call$ = Function.prototype.call;
Object.prototype.valueOf =
function(){var $\$$ = $\$OPvalueOf.call(this)$; return ($\$==\g ?null:$)}`

The main idea is to redefine the method to a new function which calls the original `valueOf` method and returns the result only if it is not the global object. We store a pointer to the original `valueOf` and `call` methods and the global object using $\$$ -variable names. Since untrusted code is restricted from accessing $\$$ -properties (see Filter 3 and Rewrite 1), these are automatically isolated from untrusted code. Similarly we can define the appropriate initialization code for the methods `sort`, `concat`, `reverse` of `Array.prototype`. We denote these by T_{sort} , T_{concat} and $T_{reverse}$.

Restricting identifier names. In order to make sure that the identifier names appearing in a program P are distinct from the ones occurring in another program with a different pid, we essentially rewrite all identifiers x to `pid.x`.

Although this will completely separate the namespaces of any two programs with different pids, thereby achieving the isolation goal, blindly renaming all identifiers will drastically modify the semantics of the program including that of good programs. The most obvious example is the expression `toString()`, that

evaluates to "*[object.Window]*" in the un-renamed version, whereas it raises a reference error exception when it is evaluated as *a12345.toString()* in the renamed version. The main issue is that variable names are in fact properties of the scope object or of the prototypes of the scope objects. Since the native properties of the global object and prototype objects are not renamed, the corresponding variable names in the program should also not be renamed, in order to preserve this correspondence between them. By analyzing the semantics, we found the complete set of property names that should not be renamed as, denoted by $\mathcal{P}_{noRen}$, to be

$$\left\{ \begin{array}{l} \text{NaN, Infinity, undefined, eval, parseInt, parseFloat, isNaN,} \\ \text{isFinite, Object, Function, Array, String, Number, Boolean,} \\ \text{Date, RegExp, Error, RangeError, ReferenceError, TypeError,} \\ \text{SyntaxError, EvalError, constructor, toString, toLocaleString,} \\ \text{valueOf, hasOwnProperty, propertyIsEnumerable,} \\ \text{isPrototypeOf} \end{array} \right\}$$

Since we do not rename the variable whose names appear in $\mathcal{P}_{noRen}$, we can only enforce the weaker isolation

$$Acc_{l_g}(H_0^{sub}, Enf(pid_{P_1}, P_1)) \cap Acc_{l_g}(H_0^{sub}, Enf(pid_{P_2}, P_2)) \subseteq \mathcal{P}_{nat} \cup \mathcal{P}_{noRen}$$

and rely on the assumption that it is safe for two untrusted programs to access the same set of non-blacklisted native properties of the global object. In particular, *eval* and *Function* are always filtered out by Filter 2. Thus, we propose the following rewriting step.

Rewrite 3 *Given a program P , rewrite all identifiers $x \notin \mathcal{P}_{noRen}$, appearing in P to *pid_Px*.*

Defining $J_{sub}(\mathcal{B})$, H_0^{sub} and Enf . We now combine the filtering, rewriting and heap initialization steps mentioned in the previous section to define the subset $J_{sub}(\mathcal{B})$, the initial heap H_0^{sub} and the enforcement function *enf*, which together solve the isolation problem. By design, the steps proposed in the previous subsection are all compatible with each other and can be combined in a straightforward manner. Based on the filtering steps, we propose the following definition for the subset $J_{sub}(\mathcal{B})$.

Definition 1. *[$J_{sub}(\mathcal{B})$] Given a blacklist $\mathcal{B}$, the subset $J_{sub}(\mathcal{B})$ is defined as JS_{E2} MINUS: all terms containing identifiers from the set $\mathcal{B}$, all terms containing one or more of the identifiers $\{\text{eval}, \text{Function}, \text{constructor}\}$, all terms containing identifiers beginning with \$.*

Based on the rewriting steps, we define the function *Enf* as follows:

Definition 2. *[Enf] Given a program P we define, $Enf(pid_P, P)$ as program P with (i) Every occurrence of the expression *e1[e2]* is rewritten to *e1[IDX(e2)]*. (ii) Every occurrence of the expression *this* is rewritten to *NOGLOBAL(this)*. (iii) Every identifier x appearing in the program must be replaced with *pid_Px* if $x \notin \mathcal{P}_{noRen}$.*

Combining all the initialization steps we define the initialized heap H_0^{sub} as:

Definition 3. $[H_0^{sub}]$ Given the initial JS_{E2} heap H_0 , we define H_0^{sub} as the heap obtained after executing all the initialization codes in the global scope. Formally, $H_0^{sub} = \mathcal{H}(Final(H_0, l_g, T_{idx}; T_{ng}; T_{valueOf}; T_{sort}; T_{concat}; T_{reverse}))$.

Note that for correctness of our solution, it is very important to execute the trusted initialization code on the initial JS_{E2} heap H_0 (described in Section 3) and hence *before* any untrusted code is executed.

Theorem 1 (Isolation theorem). Given a blacklist $\mathcal{B}$, such that $\mathcal{B} \cap \mathcal{P}_{nat} = \emptyset$, and the subset $J_{sub}(\mathcal{B})$, function Enf and the heap H_0^{sub} as defined in Definitions 1, 2 and 3 respectively. (1) For all user programs P in the subset $J_{sub}(\mathcal{B})$ with program ids pid_P , $Acc(H_0, Enf(pid_P, P)) \cap \mathcal{B} = \emptyset$. (2) For all user programs P_1 and P_2 in the subset J_{sub} with program ids pid_{P_1} and pid_{P_2} respectively $Acc_{l_g}(H_0, Enf(pid_{P_1}, P_1)) \cap Acc_{l_g}(H_0, Enf(pid_{P_2}, P_2)) \subseteq \mathcal{P}_{nat} \cup \mathcal{P}_{noRen}$.

The proof of the above theorem is described in the online version [15].

5 Case Study: FBJS

We studied the isolation mechanisms of FBJS and Yahoo! ADsafe because of their importance to hundreds of millions of Web users, and their relative simplicity. As reported in [17], we initially studied isolation based on filtering alone, and made suggestions for improvement in FBJS and ADsafe that have been adopted in both systems. However, the provably safe JavaScript subset based on filtering of [17] is far too restrictive to be used as a satisfactory replacement for FBJS. In this paper, we therefore designed rewritings and wrapper functions to design a more expressive, provably safe subset of JavaScript. We believe that $J_{sub}(\mathcal{B})$ is now comparable to FBJS from the application developer viewpoint, has fewer semantic anomalies (as described below), and has the advantage of being provably safe.

Facebook. Facebook [7] is a well-known social networking Web site reporting 200 millions active users. Registered and authenticated users store private and public information on the Facebook website. Users can share information by sending messages, directly writing on a public portion of a user profile (called the wall), or interacting with Facebook applications. Facebook applications can be written by any user and can be deployed in various ways: as desktop applications, as external web pages displayed inside an `iframe` within a Facebook page, or as integrated components of a user profile.

Integrated Facebook applications are written in FBML [24], a variant of HTML designed to make it easy to write applications and also to restrict their possible behavior. A Facebook application is retrieved from the application publisher's server and embedded as a subtree of the Facebook page document. Since integrated Facebook applications are intended to interact with the rest of the user's profile, they are not isolated inside an `iframe`. As part of the Facebook

isolation mechanism, the scripts used by applications must be written in a subset of JavaScript called FBJS [23] that restricts them from accessing arbitrary parts of the DOM tree of the larger Facebook page. The source application code is checked to make sure it contains valid FBJS, rewriting is applied to limit the application’s behavior, and a specialized library is provided.

FBJS. While FBJS has the same syntax as JavaScript, a preprocessor consistently adds an application-specific prefix to all top-level identifiers in the code, isolating the effective namespace of an application from the namespace of other applications and of the rest of the Facebook page. For example, a statement `document.domain` may be rewritten to `a12345.document.domain`, where `a12345_` is the application-specific prefix. This renaming will prevent application code from directly accessing most of the host and native JavaScript objects, such as the `document` object, Facebook provides libraries that are accessible within the application namespace. For example, the libraries include the object `a12345.document`, which mediates interaction between the application code and the true `document` object. Additional steps are used to restrict the use of the `this` and `o[e]` in FBJS code. Occurrences of `this` are replaced with the expression `$FBJS.ref(this)`, which calls the function `$FBJS.ref` to check what object `this` refers to when it is used. If `this` refers to `window`, then `$FBJS.ref(this)` returns `null`. FBJS rewrites `o[e]` to `a12345.o[$FBJS.idx(e)]`, where `$FBJS.idx` enforces blacklisting on the string value of `e`. Other, indirect ways that malicious content might reach the `window` object involve accessing certain standard or browser-specific predefined object properties such as `__parent__` and `constructor`. Therefore, FBJS blacklists such properties and rewrites any explicit access to them in the code into an access to the useless property `__unknown__`. Finally, FBJS code runs in an environment where properties such as `valueOf`, which may access (indirectly) the `window` object, are redefined to something harmless.

Comparison. FBJS imposes essentially the same filtering restrictions as those we propose in Section 4, and the FBJS library appears to impose conditions similar to those we state in our wrapper conditions. However, there are some differences when it comes to renaming identifiers to place applications in separate namespaces and in the rewriting used to restrict `this` and `e[e]`.

The renaming issue is that the FBJS implementation renames properties in the set $\mathcal{P}_{noRen}$ of properties we suggest should not be renamed. For example, `toString()` is rewritten to `a12345.toString()`, with an application-specific prefix. While `toString()` normally evaluates to `"[object.Window]"`, the rewritten version throws a “reference error” exception when evaluated. As noted in [17], FBJS does not correctly support renaming because it does not prevent explicit manipulation of the scope; the subset we propose here does not completely prevent access to scope objects either (for greater expressiveness), but has fewer pathological cases, because we avoid renaming $\mathcal{P}_{noRen}$ properties. A minor point is that we show that a safe subset can contain `with`, which FBJS prohibits, although our safe subset removes or restricts constructs that appear in many `with` use-cases.

To discuss more substantive issues, we consider $FBJS_{09}^v$, the version of FBJS deployed on Facebook at the time of our analysis, in March 2009. This version

reflects repairs to the rewriting of `this` based on our earlier discovery of ways to redefine the run-time checking function [17]. The $FBJS_{09}^v$ `$FBJS.ref` function performs a check equivalent to `NOGLOBAL`, with some additional filtering to wrap DOM objects exposed to user code. Since `$FBJS` is effectively blacklisted in $FBJS_{09}^v$, we believe that `ref` prevents the `this` identifier from being evaluated to the `window` object; the check is semantically faithful to the requirements developed in Section 3.

On the other hand, the $FBJS_{09}^v$ `$FBJS.idx` function does not preserve the semantics of the property access, and as a result can be compromised in certain environments. More specifically, we report an attack we identified during the research reported here, a repair to prevent that attack, and a remaining problem. In the context of other filtering, `$FBJS.idx` is equivalent to

```
($=e2,($ instanceof Object||$blacklist[$])?"bad":$)
```

where `$blacklist` is the object `{caller:true,$:true,$blacklist:true}`. The main problem is that, in contrast to our definition of `IDX`, the expression `$blacklist[$]?"bad":$` converts `va` to a string two times. This is a problem if evaluation has a side effect. For example, the object

```
{toString:function(){this.toString=function(){return "caller"}; return "good"}}
```

can fool `FBJS` by first returning a good property `"good"`, and then returning the bad property `"caller"` on the second evaluation. To avoid this problem, $FBJS_{09}^v$ inserts the check `$ instanceof Object` that tries to detect if `$` contains an object. In general, however, this check is not sound – according to the JavaScript semantics, any object with a `null` prototype (such as `Object.prototype`) escapes this check. Moreover, in Firefox, Internet Explorer and Opera the `window` object also escapes the check. In $FBJS_{09}^v$, `Object.prototype` and `window` are not accessible by user code, so cannot be used to implement this attack.

We found that the scope objects described in Section 3 have a null prototype in Safari, and therefore we were able to mount attacks on `$FBJS.idx` that effectively let user application code escape the Facebook sandbox. Shortly after we notified Facebook of this problem, the `$FBJS.ref` function was been modified to include a check of current browser, and if it is Safari an additional check that `this` is not bound to an object able to escape the `instanceof` check described above. This solution is not completely satisfactory, for two reasons. First, some browsers may have other host objects that have a `null` prototype, and that can be accesses without using `this`. Such objects could still be used to subvert `$FBJS.idx`, which has not been changed. Second, `$FBJS.idx` prevents objects from being uses as arguments of member expressions. This restriction is unnecessary for the safety of blacklisting, as shown by our proof for `IDX`.

6 Other Language-Based Approaches to Isolation

In this Section, we describe a few other approaches to JavaScript isolation which have not been subjected to rigorous semantic analysis, and could therefore ben-

efit from the reasoning techniques presented in this paper. Due to space limitations, we do not discuss solutions based on idealized subsets of JavaScript with limited expressiveness, or that rely on browser modifications (for example [29]).

ADSafe. The Yahoo!ADsafe subset [5] is designed to allow advertising code to be placed directly on the host page, limiting interaction by a combination of static analysis and syntactic restrictions. The advertising code must satisfy very severe syntactical restrictions (including no `this`), and has access to an `ADSAFE` object, provided as a library, that mediates access to the DOM and other page services. Since we discovered that ADsafe was liable to prototype-poisoning attacks [14], the filtering process for ADsafe code has been complemented by a static analysis which gathers information about the objects that untrusted code may try to get access to. It is left to the page hosting the advertisement to make sure that those objects cannot be used to subvert the isolation mechanism. Our results show that some of the ADsafe restrictions are not strictly necessary, and the subset could be made more expressive.

Browsershield. Browsershield is a system that rewrites web pages in order to enforce run-time monitoring of the embedded scripts. The system takes an HTML page, adds a script tag to load a trusted library, rewrites embedded scripts so that they invoke a local rewriting function before being executed, and rewrites instructions to load remote scripts by making them load through a rewriting proxy. The run time monitoring is enforced by *policies* which are in effect functions that monitor the JavaScript execution. Common operations such as assignment suffer from a hundred-fold slowdown, and policies are arbitrary JavaScript functions for which there is no systematic way of guaranteeing correctness.

GateKeeper. Livshitz and Guarnieri [12] propose an approach to enforcing security and reliability policies in JavaScript based on static analysis based on two subsets. The first, JS_{Safe} , is obtained exclusively by filtering, and does not contain `with`, `eval`, `e[e]` or other dangerous constructs. The second subset, JS_{GK} reinstates `e[e]` after wrapping it in a run-time monitor. A static analysis approximates the call-graph and points-to relation of objects in these subsets. Unfortunately, the implementation of GateKeeper is not available for inspection, and the sparse details on the definition of JS_{Safe} and the run-time monitors in JS_{GK} are not sufficient for a formal comparison with our results.

Caja. The Google Caja [4] project is a substantial effort to provide a safe JavaScript subset. Caja uses a compilation process that takes untrusted JavaScript and produces code in Cajita, a well behaved capability-based safe subset of JavaScript. Our goal is to isolate certain variables in the heap, whereas Caja enforces a finer grained security policy, which allows untrusted code from different principals to interact safely, by leveraging the capability-based paradigm. The Caja enforcement mechanisms also include filtering and rewriting, but the additional expressive power is gained at the price of complexity and efficiency. The reasoning techniques introduced in this paper could be used to proof the correctness of such mechanisms, and possibly improve their implementations.

Lightweight Self-Protecting Javascript. Phung *et al.* [19] introduce a principled approach for enforcing safety properties on JavaScript native libraries. The enforcement mechanism involves wrapping each of the security critical native library methods and properties, before executing an untrusted script. Unfortunately, this approach is not sound for existing browsers. For example, by deleting certain properties of the global object, some native object are reinstated in the global environment, subverting the wrapping mechanism. Future versions of JavaScript may provide better support this implementation technique.

7 Conclusions

We systematically presented and analyzed a combination of isolation mechanisms for a subset of JavaScript that is comparable in expressiveness to Facebook FBJS [23]. Isolation from untrusted code in our subset of JavaScript is based on filtering out certain constructs (`eval`, `Function`, `constructor`), rewriting others (`this`, `e1[e2]`) to allow them to be used safely, and wrapping properties (*e.g.*, object and array prototype properties) of the execution environment to further limit the impact of untrusted code. Our analysis and security proofs build on a formal foundation for proving isolation properties of JavaScript programs [17], based on our operational semantics [14] of the full ECMA-262 Standard language (3rd Edition) [11]. While we focus on one particular combination of filters, rewriting functions, and wrappers, our methods are applicable to variants of the specific subset we present. For example, a DOM function such as `createElement` could be allowed, if suitable rewriting is used to insert checks on its string argument at run-time. In future work, we intend to examine Caja [4] and other systems, with the goal of providing provable security for practically useful language-based isolation mechanisms.

Acknowledgments. Mitchell and Taly acknowledge the support of the National Science Foundation. Maffei is supported by EPSRC grant EP/E044956/1.

References

1. Irem Aktug, Mads Dam, and Dilian Gurov. Provably correct runtime monitoring. In *FM 2008*, volume 5014 of *Lecture Notes in Computer Science*, pages 262–277. Springer, 2008.
2. C. Anderson, P. Giannini, and S. Drossopoulou. Towards type inference for JavaScript. In *Proc. of ECOOP’05*, page 429452, 2005.
3. A. Barth, C. Jackson, and J.C. Mitchell. Securing browser frame communication. In *17th USENIX Security Symposium*, 2008.
4. Google Caja Team. Google-Caja: A source-to-source translator for securing JavaScript-based web. <http://code.google.com/p/google-caja/>.
5. Douglas Crockford. ADsafe: Making JavaScript safe for advertising. <http://www.adsafe.org/>, 2008.
6. B. Eich. JavaScript at ten years. <http://www.mozilla.org/js/language/ICFP-Keynote.ppt>.

7. FaceBook. Web Site. <http://www.facebook.com/>.
8. D. Flanagan. *JavaScript: The Definitive Guide*. O'Reilly, 2006. <http://proquest.safaribooksonline.com/0596101996>.
9. P. Heidegger and P. Thiemann. Recency types for dynamically-typed, object-based languages. Foundations of Object-Oriented Languages (FOOL'09), 2009.
10. iGoogle. Web Site. <http://www.google.com/ig>.
11. ECMA International. ECMAScript language specification. standard ECMA-262, 3rd Edition. <http://www.ecma-international.org/publications/files/ECMA-ST/Ecma-262.pdf>, 1999.
12. B. Livshits and S. Guarnieri. Gatekeeper: Mostly static enforcement of security and reliability policies for JavaScript code. MSR-TR-2009-16, Feb. 2009.
13. S. Maffei, J. Mitchell, and A. Taly. Complete ECMA 262-3 operational semantics. <http://jssec.net/semantics/>.
14. S. Maffei, J.C. Mitchell, and A. Taly. An operational semantics for JavaScript. In *Proc. of APLAS'08*, volume 5356 of *LNCS*, pages 307–325, 2008. See also: Dep. of Computing, Imperial College London, Technical Report DTR08-13, 2008.
15. S. Maffei, J.C. Mitchell, and A. Taly. Isolating JavaScript with filters, rewriting, and wrappers. Dep. of Computing, Imperial College London, Technical Report DTR09-6, 2009.
16. S. Maffei, J.C. Mitchell, and A. Taly. Run-time enforcement of untrusted javascript subsets. In *Web 2.0 Security & Privacy (W2SP)*, 2009.
17. S. Maffei and A. Taly. Language-based isolation of untrusted Javascript. In *Proc. of CSF'09*, IEEE, 2009. See also: Dep. of Computing, Imperial College London, Technical Report DTR09-3, 2009.
18. OpenSocial. Web Site. <http://www.opensocial.org/>.
19. David Sands Phu H. Phung and Andrey Chudnov. Lightweight self protecting JavaScript. In *ASIACCS 2009*. ACM Press, 2009.
20. Gordon D. Plotkin. A structural approach to operational semantics. *J. Log. Algebr. Program.*, 60-61:17–139, 2004.
21. C. Reis, J. Dunagan, H. Wang, O. Dubrovsky, and S. Esmeir. BrowserShield: Vulnerability-driven filtering of Dynamic HTML. *ACM Transactions on the Web*, 1(3), 2007.
22. A. Sabelfeld and Aslan Askarov. Tight enforcement of flexible information-release policies for dynamic languages. Second International Workshop on Proof-Carrying Code'08, 2008.
23. The FaceBook Team. FBJS. <http://wiki.developers.facebook.com/index.php/FBJS>.
24. The FaceBook Team. FBML. <http://wiki.developers.facebook.com/index.php/FBML>.
25. P. Thiemann. Towards a type system for analyzing javascript programs. In *Proc. of ESOP'05*, volume 3444 of *LNCS*, page 408422, 2005.
26. P. Thiemann. A type safe DOM API. In *Proc. of DBPL*, pages 169–183, 2005.
27. K. Vikram and M. Steiner. Mashup component isolation via server-side analysis and instrumentation. In *Web 2.0 Security & Privacy (W2SP)*, 2008.
28. YahooApp. Web Site. <http://developer.yahoo.com/yap/>.
29. D. Yu, A. Chander, N. Islam, and I. Serikov. JavaScript instrumentation for browser security. In *Proc. of POPL'07*, pages 237–249, 2007.