

Mining Security-Sensitive Operations in Legacy Code using Concept Analysis

Vinod Ganapathy
University of Wisconsin
Madison, WI-53706
vg@cs.wisc.edu

David King and Trent Jaeger
Pennsylvania State University
University Park, PA-16802
{dhking,tjaeger}@cse.psu.edu

Somesh Jha
University of Wisconsin
Madison, WI-53706
jha@cs.wisc.edu

University of Wisconsin, Computer Sciences Department Technical Report 1580, September 2006

Abstract

We present an approach based on concept analysis to retrofit legacy servers with mechanisms for authorization policy enforcement. Our approach is based upon the observation that security-sensitive operations are characterized by idiomatic resource manipulations, called fingerprints. We statically mine fingerprints using concept analysis and then use them to identify security-sensitive operations and locate where they are performed by the server. Case studies with three real-world servers show that our approach is affordable and effective. We were able to identify security-sensitive operations for each of these servers with a few hours of manual effort and modest domain knowledge.

1 Introduction

Software systems must protect shared resources that they manage from unauthorized access. This is achieved by formulating and enforcing an appropriate authorization policy (also called access control policy). The policy specifies the set of *security-sensitive operations* that a user can perform on a resource. For example, a popular policy on UNIX-like systems allows only the root to perform the security-sensitive operations Read and Write on the `/etc/passwd` file (the resource). Operating systems have historically had mechanisms such as reference monitors [4] to enforce authorization policies. It is also important for user-space servers, such as middleware, web-, proxy and window-management servers, to implement such mechanisms because they manage shared resources on behalf of their clients. Unfortunately, economic and practical considerations force developers to choose functionality and performance over security. As a result, several legacy servers often completely lack policy enforcement mechanisms. For example, the X11 server [31] can simultaneously manage

Server	LOC	Number of fingerprints	Avg. size of fingerprints
ext2	4,476	18	3.67
X server/dix	30,096	115	3.76
PennMUSH	94,014	38	1.42

Figure 1: Highlights of our results. Concept analysis distills servers of several-thousand lines into a manageable number of candidate fingerprints for manual examination.

multiple X client windows, but was not built with mechanisms to isolate one X client from another, leading to several published attacks (see [20] for example attacks).

This paper investigates techniques for retrofitting legacy servers with authorization policy enforcement mechanisms. The main questions to be addressed when retrofitting a server are *what are the security-sensitive operations to be mediated?* and *where does the server perform these operations?* In current practice, these questions are answered manually. A team of software engineers inspects the code of the server to determine locations where authorization checks must be placed, and the set of security-sensitive operations that must be checked at each location. Not surprisingly, this painstaking process is time-consuming and error prone [18, 32]. For example, it took almost two years each for the Linux Security Modules (LSM) project [30], where additional authorization checks were added to the Linux kernel to enable enforcement of mandatory access control policies, and the X11/SELinux [20] project, where authorization checks were added to the X11 server. Similar recent efforts have also been time-consuming [13, 17]. In short, there are no automated techniques to aid the process of securing legacy servers for authorization.

We build on prior work [15] and develop a novel approach using concept analysis [29] to drastically reduce the

manual effort involved in retrofitting legacy servers. Key to our approach is the observation that security-sensitive operations performed by a server are associated with idiomatic ways in which resources are manipulated by the server. Such idioms, which we call *fingerprints*, are code-level descriptions of the security-sensitive operations that they represent. Each fingerprint is expressed as a combination of several abstract syntax trees (ASTs), called *code-patterns*. We use static program analysis in combination with concept analysis to automatically mine candidate fingerprints. These are then examined and refined manually by a domain expert. After refinement, we statically match each fingerprint against the code of the server to determine locations where the corresponding security-sensitive operation is performed. We then weave hooks to a reference monitor at all these locations to authorize that security-sensitive operation. This ensures that all security-sensitive operations performed by the server are mediated by authorization policy lookups.

Our results demonstrate the effectiveness of our approach. We conducted case studies on three real-world systems of significant complexity: the ext2 file system, a subset of X server (its dispatch loop), and PennMUSH, an online game server [2]. Figure 1 highlights our results, and shows the size of each of these servers, the number of candidate fingerprints automatically mined by concept analysis, and the average size of each candidate fingerprint (the number of code-patterns it contains). An analyst’s manual effort is reduced to refining candidate fingerprints and determining whether each refined fingerprint is indeed security-sensitive or not. As our results show, in each case we were able to reduce the analysis of several thousand lines of code to the analysis of under 115 candidate fingerprints with fewer than 4 code-patterns each (on average). For example, we were able to reduce the analysis of PennMUSH, a server with 94,014 lines of C code, to the analysis of 38 candidate fingerprints, with an average of 1.42 code-patterns each. It took just a few hours of manual effort and modest domain knowledge to find security-sensitive operations in each of our case studies. Without our approach, the entire code-base must be examined to find such security-sensitive operations.

The approach presented in this paper overcomes two important limitations of our prior work [15]. While we introduced fingerprints in that work, our approach for finding fingerprints (i) required a high-level description of security-sensitive operations, and (ii) used dynamic program analysis to find fingerprints. Both (i) and (ii) prevented our approach from easily being applied to a wide variety of servers. In particular, while a high-level description of security-sensitive operations was available for the case study that we considered (the X server), this may not be the case with other servers, as indeed was the case with

PennMUSH. A dynamic approach to fingerprint-finding meant that we could not guarantee that all fingerprints were found. This paper directly addresses both these shortcomings. Concept analysis automatically mines candidate fingerprints without the need for an *a priori* description of security-sensitive operations. Static program analysis ensures near-complete code-coverage, and identifies all possible ways in which a shared resource is accessed by the server, thus ensuring that all fingerprints have been mined.

In summary, our main technical contributions are:

- A fully static approach to retrofit policy enforcement into legacy servers. The key observation here is that security-sensitive operations performed by a server are associated with idiomatic resource manipulations, called fingerprints.
- A novel algorithm using concept analysis to automatically mine fingerprints of security-sensitive operations. To our knowledge, this is the first application of concept analysis to mine security properties of software.
- Case studies on three real-world servers of significant complexity. Our case studies demonstrate that our approach is efficient and effective. Our analysis completed in just over 310 seconds even for the largest of our benchmarks and produced manageable concept lattices. In each case, we were able to inspect the lattice and identify security-sensitive operations with a few hours of manual effort and modest domain knowledge.

Note that our approach to retrofit legacy servers follows the aspect-oriented paradigm. In particular, each fingerprint denotes a region of code before which a reference monitor hook must be placed, and thus helps identify join-points [5, 19]. The reference monitor query that executes as a result of the hook call is the body of the advice at that join point. Fingerprint-mining is thus aspect-mining to find join points relevant to security.

2 Approach overview

We give a high-level overview of our approach, depicted in Figure 2. Using a running example, we show how a software engineer would use our approach to mine fingerprints of security-sensitive operations and place hooks. We have currently implemented our analysis to work with C programs, but the underlying principles apply to servers written in other languages as well.

2.1 Running example

We use a subset of ext2, a Linux file system, and one of the case studies in Section 5 as our running example. In partic-

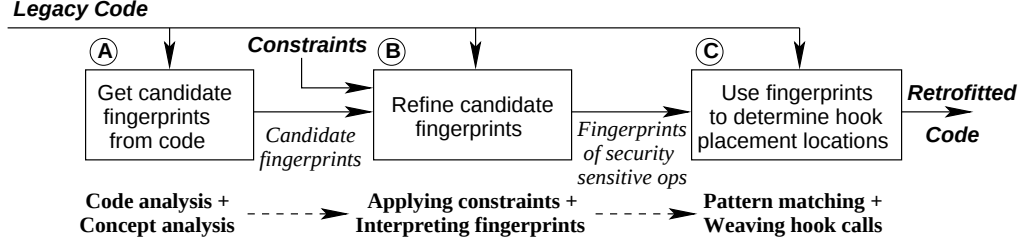


Figure 2: Steps to retrofit policy enforcement to legacy code, and the techniques used in each step.

ular, ext2 is responsible for laying out and interpreting disk blocks as belonging to specific files or directories. It uses several internal data structures to represent metadata information using which it retrieves files and directories from raw disk blocks.

File systems on Linux are “pluggable”, and must thus export a standard API to the kernel. A system call that manipulates files or directories ultimately resolves to one or more calls to this API. The relevant file system functions then serve this request. Thus a file system is a server that manages files and directories. For ext2, we considered 10 API functions related to manipulation of directories (e.g., `ext2_rmdir`, `ext2_mkdir` and `ext2_readdir`). We show how our approach can identify security-sensitive operations that ext2 performs on directories.

2.2 Step A: From source code to candidate fingerprints

In the first step, we employ static source code analysis and identify different ways in which ext2 accesses shared resources in response to client requests. This analysis is based upon two assumptions.

First, we assume that it suffices to examine accesses to internal data structures that ext2 uses to represent files and directories. These data structures are specified by a domain expert, and for ext2 they are variables of type `inode`, `ext2_dirent`, `ext2_dir_entry_2` and `address_space`, each of which is a C struct. Second, we assume that a client accesses server resources only via the server’s API. With ext2, this is indeed the case, and as mentioned earlier ext2 exports a well-defined API to the kernel. The inputs to our static analyzer are thus the source code of ext2, and two files, specifying, respectively, the types of critical data structures to be tracked, and a set of API functions. The static analyzer identifies how these tracked data structures are manipulated by the ext2 API.

To do so, the analyzer walks the source code of ext2 and distills each statement into a (possibly empty) set of

Code-pattern	:=	Call AST Read AST Write Value to AST
Value	:=	constant AST $\perp$ (unknown)
AST	:=	(type-name->)* field

Figure 3: Grammar for code-patterns.

code-patterns, the grammar for which is shown in Figure 3. A code-pattern is either a *Read*, *Write* or a *Call* and is expressed in terms of abstract syntax trees. For example, the C statement `de->file_type = 0`, where `de` is a variable of type `ext2_dirent` is distilled to *Write 0 To ext2_dirent->file_type*. Note in particular that this transformation ignores specific variable names and focuses instead on types of variables. As a result, we identify generic resource manipulations but not the specific instance of the resource (e.g., the instance `de`) that they happen on. Statements that do not manipulate tracked data structures are ignored. *Call* code-patterns correspond to calls via unresolved function pointers. For each function `ext2_api` in the ext2 API the static analyzer then aggregates code-patterns of all statements potentially reachable via a call to `ext2_api`. Thus, at the end of this step each ext2 API function `ext2_api` is associated with a set of code-patterns *CodePats*(`ext2_api`). Intuitively, *CodePats*(`ext2_api`) denotes all possible ways in which `ext2_api` can potentially manipulate tracked resources.

The next step is to identify idiomatic resource manipulations by the ext2 API. The goal here is to find sets of code-patterns that always appear together. That is, if one code-pattern from a set of code-patterns appears in an execution of ext2, then all the other code-patterns from that set appear in that execution as well. Note that we can have sets $\{pat\}$ with singleton code-patterns as well, denoting that no other code-pattern always appears together with $\{pat\}$. Each set of such code-patterns denotes an idiomatic way in which a resource is manipulated by ext2, and potentially indicates a security-sensitive operation. We call each such set a *fingerprint*.

We identify candidate fingerprints using concept analysis [29], a well-known hierarchical clustering technique. At a high-level (details are presented in Section 3), concept analysis identifies candidate fingerprints, as well as the API functions whose code-pattern sets contain these candidate fingerprints. We use the term *candidate fingerprints* because as described in Step B, imprecisions introduced in the program analysis step means that each candidate fingerprint may contain multiple fingerprints.

For example, concept analysis inferred that the set of six code-patterns shown in Figure 4 is a candidate fingerprint, and that it appears in *CodePats*(ext2.rename), *CodePats*(ext2.rmdir) and *CodePats*(ext2.unlink).

```
(1) Read address_space->host
(2) Read ext2_dir_entry_2->rec_len
(3) Write 0 To ext2_dir_entry_2->inode
(4) Read inode->i_mtime
(5) Read inode->u->ext2_inode_info->i_dir_start_lookup
(6) Write 1 To inode->u->ext2_inode_info->i_dir_start_lookup
```

Figure 4: One of the candidate fingerprints that concept analysis identifies for ext2.

For ext2, we identified 18 such candidate fingerprints, each denoting a unique way in which ext2 manipulates files and directories. While concept analysis is asymptotically inefficient—its complexity is exponential in $\max_i(|\text{CodePats}(\text{ext2_api}_i)|)$ —our experiments showed that it is efficient in practice. In particular, our analysis completed in about 2 seconds for ext2, and in just over 310 seconds even for the largest of our case studies.

2.3 Step B: Refining candidate fingerprints

In the second step, a domain expert (i) refines candidate fingerprints obtained from Step A and (ii) post refinement, determines, for each fingerprint, whether it embodies a security-sensitive operation that must be mediated by an authorization policy lookup.

Refinement of candidate fingerprints is necessary for two reasons. The first reason is because code analysis employed in Step A is imprecise. As a result, multiple fingerprints may be combined into a single candidate fingerprint. There are two ways in which precision is lost:

1. Code analysis is *flow-insensitive*. A candidate fingerprint may contain a pair of code-patterns pat_1 , pat_2 that do not always appear together in all executions of the server.
2. We ignore specific instances of resources that are manipulated and focus instead on their types. Thus, a candidate fingerprint may contain manipulations of multiple, possibly unrelated, resources.

We employ *precision constraints* to identify such cases and enable refinement of each candidate fingerprint, separating the code-patterns that it contains into several fingerprints. Intuitively, a precision constraint is a rule that determines the set of code-patterns that can be grouped together in a fingerprint. The second reason why refinement is necessary is because a domain expert may deem that a set of code-patterns is irrelevant for the authorization policies to be enforced for the server, or may wish to separate (or group together) a pair of code-patterns in a fingerprint of a security-sensitive operation. Such *domain-specific constraints* further refine candidate fingerprints.

For example, consider the candidate fingerprint shown in Figure 4. Using the output of our static analysis tool, we were able to determine that the code-patterns (1)-(4) appear together in each successful invocation of the ext2 function *ext2.delete_entry* and that the code-patterns (5) and (6) appear together in each successful invocation of the function *ext2.find_entry*. Each of the three API functions, *ext2.rename*, *ext2.rmdir* and *ext2.unlink*, that contain this candidate fingerprint call both these functions. Both *ext2.rmdir* and *ext2.unlink* call these functions on the *same* resource instance, namely the directory being removed (or unlinked). However, as Figure 5 shows, while *ext2.rename* calls both these functions on the instances *old_dir* and *old_dentry*,¹ it calls *ext2.find_entry* only on the instances *new_dir* and *new_dentry* when a certain predicate *new_inode* is satisfied.

```
1. int ext2_rename (inode *old_dir, dentry *old_dentry,
2.                 inode *new_dir, dentry *new_dentry) {
3.     /* declarations of old_page, new_page, old_de and new_de */
4.     new_inode = new_dentry->d_inode; ...
5.     old_de = ext2_find_entry(old_dir, old_dentry, &old_page);
6.     if (new_inode) { ...
7.         new_de = ext2_find_entry(new_dir, new_dentry, &new_page);
8.     } else { ...
9.         /* no call to ext2_find_entry */
10.    }; ...
11.    ext2_delete_entry(old_de, old_page); ...
12. }
```

Figure 5: Example showing the need for precision constraints.

Because *ext2.rename* performs the resource manipulations corresponding to code-patterns (5) and (6) on additional resource instances as compared to the code-patterns (1)-(4), code-patterns (1)-(4) and (5)-(6) represent potentially different security-sensitive operations. Imposing the constraint that code-patterns on different resource instances must be part of separate fingerprints, the candidate fingerprint shown in Figure 4 is split into two fingerprints, as shown in Figure 6. Additional examples of the use of precision constraints appear in Section 4. Note that such con-

¹The variable *old_de*, which *ext2.delete_entry* is invoked with on line 11 is derived from *old_dir* and *old_dentry*.

straints can potentially be avoided with sophisticated program analyses, that we plan to explore in future work. However, in our case studies we found that more than 50% of the candidate fingerprints did not require refinement. Thus our current approach provides a good tradeoff between precision of results and simplicity of the code analysis algorithm.

Fingerprint (1)	
(1)	<i>Read</i> address.space->host
(2)	<i>Read</i> ext2_dir.entry_2->rec_len
(3)	<i>Write</i> 0 To ext2_dir.entry_2->inode
(4)	<i>Read</i> inode->i_mtime
Fingerprint (2)	
(5)	<i>Read</i> inode->u->ext2_inode_info->i_dir.start_lookup
(6)	<i>Write</i> 1 To inode->u->ext2_inode_info->i_dir.start_lookup

Figure 6: Fingerprints obtained after refinement with precision constraints.

Domain-specific constraints encode rules that are formulated by a domain-expert. In particular, whether the resource manipulation embodied by a fingerprint is security-sensitive depends on the set of policies that must be enforced on clients. For example, it may only be necessary to protect the integrity of directories, and not their confidentiality. In this case, fingerprints that embody a write operation on directories are security-sensitive, while fingerprints that embody a read operation are not. Fingerprints expose all possible operations on resources, and let an administrator decide whether an operation is security-sensitive or not. For example, an analyst may decide that Fingerprint (2) in Figure 6, which corresponds to a directory lookup, is not interesting for a specific set of policies to be enforced.

After refinement, the domain expert assigns semantics to each fingerprint, associating it with a security-sensitive operation. For example, Fingerprint (1) in Figure 6 embodies the directory removal operation, while Fingerprint (2) embodies the lookup operation. The LSM project [30] has identified a comprehensive set of security-sensitive operations for Linux by considering a wide range of policies to be enforced, including security-sensitive operations on the file system. It turns out that Fingerprint (1) embodies the LSM operation `Dir.Remove.Name`, while Fingerprint (2) embodies the LSM operation `Dir.Search`. Thus, at the end of the second step, we have a set of fingerprints, each of which is associated with a security-sensitive operation.

2.4 Step C: From fingerprints to hooks

The final step is to place reference monitor hooks and implement the appropriate policy lookups for each hook (*i.e.*, the advice at each join point). Because we borrow from prior work [15] for this step, we only discuss it briefly here.

Each fingerprint is a set of code-patterns that can be matched against the server’s source code. Each code-fragment that matches a fingerprint is deemed as performing the security-sensitive operation associated with that fingerprint. In [15], we had presented an approach to place hooks at the granularity of function calls, *i.e.*, for each fingerprint that matched the set of code-patterns in a function, we would place a reference monitor hook to guard calls to this function with the appropriate security-sensitive operation. For example, using the fingerprints from our running example, we would place a hook guarding the call to `ext2_find.entry` on line (5) of Figure 5 to check that the LSM operation `Dir.Search` is authorized as follows:

```
if (check.policy(current process, old.dir, Dir.Search))
{ ext2_find.entry(old.dir, old.dentry, &old.page); }
else { Notify current process of failed authorization check }
```

A similar hook will also be placed for the call on line (7). The call to `ext2_delete.entry` on line (11) will be protected with a hook that checks that the client is authorized to perform the LSM operation `Dir.Remove.Name` on the directory being removed. Several optimizations are possible to this basic hook placement technique, *e.g.*, placing hooks so as to minimize the number of reference monitor queries executed at runtime. We leave such optimizations for future work.

Note that fingerprints are useful even when hook placements have been decided in advance. For example, if a team of software engineers decides to place just one hook guarding calls to `ext2_rename` (as was done in LSM), then fingerprints determine the security-sensitive operations that must be authorized by that hook. In this case, the hook must authorize `Dir.Remove.Name` on the old directory (instance `old.dir`) and `Dir.Search` on both the old (`old.dir`) and new directories (`new.dir`). Indeed, these security-sensitive operations are authorized in the implementation of the hook in the LSM implementation of security-enhanced Linux (SELinux) [22].²

3 Extracting candidate fingerprints from code

This section discusses Step A in detail. We discuss the use of static analysis to identify all resource manipulations potentially performed by each API function, and concept analysis to find candidate fingerprints.

²The SELinux LSM hook authorizes additional security-sensitive operations, corresponding to fingerprints that match code-fragments that were omitted from Figure 5.

3.1 Static analysis

Algorithm 1 describes the static code analysis that we have implemented as a plugin to the CIL toolkit [23]. Lines (1)-(5) employ a simple flow-insensitive analysis to extract for each function a set of code-patterns describing how the function manipulates tracked data structures. While this step sacrifices precision, it simplifies the rest of the analysis by making the output amenable to concept analysis. As described earlier, we recover some of the precision lost in this step by applying precision constraints. While we intend to explore in future work how a flow-sensitive program analysis can interact with concept analysis, we have found that our current implementation offers a reasonable tradeoff between simplicity of analysis and precision of the results obtained. Lines (6)-(9) compute $CodePats(api_i)$, the set of resource manipulations performed by api_i , for each API function api_i of the server by finding functions in the call-graph reachable from api_i . We resolve calls through function pointers using a simple pointer analysis: each function pointer can resolve to any function whose address is taken and whose type signature matches that of the function pointer. This analysis is conservative in the absence of type-casts, but may miss potential targets in the presence of type-casts.

```

Algorithm : EXTRACT_CODE-PATTERNS(Server, API, RSC)
Input : (i) Server: source code of server, (ii) API={api1, ..., apin}:
        set of API functions of Server, and (iii) RSC: data types of
        sensitive resources.
Output : CodePats(api1), ..., CodePats(apin), for
        api1, ..., apin ∈ API.

1 foreach (function f in Server) do
2   Summary(f) := ∅;
3   foreach (statement s ∈ f that affects a data structure of type ∈ RSC) do
4     CP := Breakdown of s into code-patterns (see Figure 3);
5     Summary(f) := Summary(f) ∪ CP;

6 foreach (apii ∈ API) do
7   CodePats(apii) := ∅;
8   foreach (function f reachable from apii) do
9     CodePats(apii) := CodePats(apii) ∪ Summary(f);

10 return CodePats(api1), ..., CodePats(apin);

```

Algorithm 1: Static analysis algorithm to extract resource manipulations.

Recall that $CodePats(api_i)$ is the set of resource manipulations that a client can perform by invoking API function api_i . However, we would like to identify idiomatic resource manipulations. Each such idiom is a set of code-patterns $FP = \{pat_1, \dots, pat_m\}$ satisfying the following property: if one of the code-patterns $pat_i \in FP$ appears in any valid execution trace of the server, then *all* the patterns in FP appear in that trace. Each such idiom is called a *fingerprint* and denotes a potential security-sensitive operation performed on the resource. Note that the above property implies that each fingerprint FP is such that $FP \subseteq CodePats(api_i)$

or $FP \cap CodePats(api_i) = \emptyset$, for each API function api_i . As described below, we use concept analysis to identify a set of *candidate fingerprints*. Each candidate fingerprint may possibly contain multiple fingerprints, and must be refined to yield the actual fingerprints.

3.2 Background on concept analysis

Concept analysis is a well-known hierarchical clustering technique that has found use in software engineering [3, 8, 9, 21, 24, 25, 26, 27, 28]. We give a brief overview of concept analysis and describe how we adapt it to find candidate fingerprints.

The inputs to concept analysis are (i) a set of *instances* I , (ii) a set of *features* F , and (iii) a binary relation $R : I \rightarrow F$ that associates instances with features. It produces a *concept lattice* as output. Intuitively, each node in the concept lattice pairs a set of instances X with a set of features Y , such that Y is the largest set of features in common to *all of the instances* in X . Formally, each node is a pair $\langle X, Y \rangle$, where $X \in I$ and $Y \in F$, such that $\alpha(X) = Y$ and $\gamma(Y) = X$, where $\alpha(X) = \{f \in F \mid \forall x \in X (x, f) \in R\}$, and $\gamma(Y) = \{i \in I \mid \forall y \in Y (i, y) \in R\}$. A node $\langle X, Y \rangle$ appears as an ancestor of a node $\langle P, Q \rangle$ in the concept lattice if $P \subset X$. In fact, this ordering also implies $Y \subset Q$. This is because a smaller set of instances will share a larger set of features in common. Thus, the root node shows the set of features common to all instances in I , while the leaf node shows the set of instances that share all features in F .

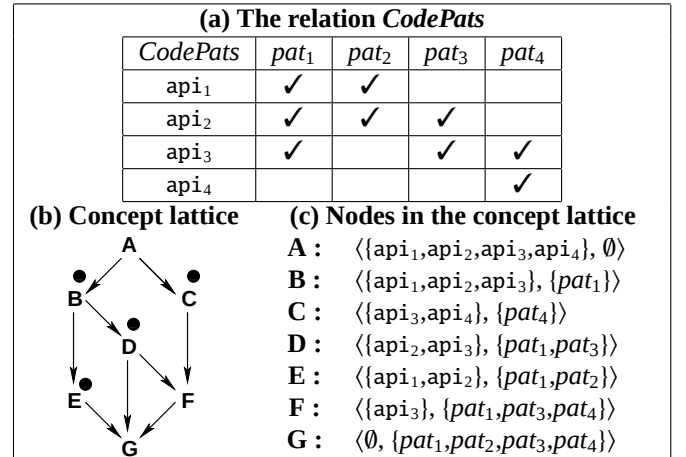


Figure 7: Concept analysis example.

Figure 7 shows an example of a concept lattice, as applied to our problem. Each API function api_1 , api_2 , api_3 and api_4 is considered an instance, and each code-pattern pat_1 , pat_2 , pat_3 , pat_4 is considered a feature. They are related by *CodePats*, which is obtained from static analysis, depicted in Figure 7(a) as a table. Each node $\langle X, Y \rangle$ is such that *all*

the code-patterns in Y appears in each $CodePats(api_i)$ for $api_i \in X$. This lattice shows, for example, that (i) there are no code-patterns in common to all API functions (node A in the lattice), (ii) Both pat_1 and pat_3 appear in both $CodePats(api_2)$ and $CodePats(api_3)$, and these are the only such API functions (node D), and that (iii) No API functions have all code-patterns (node G).

3.3 Using concept analysis

We compute candidate fingerprints using Algorithm 2. It first invokes concept analysis (line 1) on the set of API functions and the set of code-patterns to obtain a concept lattice as shown in Figure 7. It then finds candidate fingerprints, in lines 2-7, by finding nodes in the lattice where new code-patterns are introduced. Each such node is marked, and the set of new code-patterns introduced in that node is considered as a candidate fingerprint.

For the example in Figure 7, the nodes B, C, D, and E are marked because these nodes introduce the code-patterns pat_1 , pat_4 , pat_3 and pat_2 —i.e., any node containing one of these patterns *must* have the corresponding node as an ancestor. Each of these code-patterns is classified as a candidate fingerprint.

Intuitively, Algorithm 2 works because each fingerprint FP satisfies $FP \subseteq CodePats(api_i)$ or $FP \cap CodePats(api_i) = \emptyset$, for each API function api_i . Concept analysis ensures that the node of the concept lattice in which a new code-pattern $pat_i \in FP$ is introduced will introduce *all* of the code-patterns in FP . Line 7 identifies and marks nodes where a new code-pattern pat is introduced into the lattice. Because of the property above, all the code-patterns that appear in the same fingerprint as pat appear in that node. Note however, that code-patterns from other fingerprints may also be introduced in the same node. Thus, Algorithm 2 only computes candidate fingerprints: each candidate fingerprint may contain multiple fingerprints that must be obtained via refinement (in Step B).

Algorithm	: FIND_CANDIDATE_FINGERPRINTS($CodePats, API$)
Input	: (i) $CodePats$: The relation obtained from Algorithm 1, and (ii) $API = \{api_1, \dots, api_n\}$, set of API functions of the server.
Output	: $CFP_1, \dots, CFP_k$, a set of candidate fingerprints.
1	Run concept analysis with the set of instances $I = API$, the set of features $F = \bigcup_{i \in [1..n]} CodePats(api_i)$, and the relation $R = CodePats$;
2	$count := 1$;
3	foreach (node $\langle X, Y \rangle$ in the concept lattice) do
4	Let $\{\langle X_j, Y_j \rangle\}$ be the set of parents of $\langle X, Y \rangle$ in the concept lattice;
5	$Diff := Y - \bigcup_j Y_j$;
6	if ($Diff \neq \emptyset$) then
7	$CFP_{count} := Diff$; $count := count + 1$; Mark the node $\langle X, Y \rangle$;
8	return $CFP_1, \dots, CFP_{count}$ /* Note: k is the value of $count$ in this line. */;

Algorithm 2: Finding candidate fingerprints.

Candidate fingerprints $CFP_1, \dots, CFP_k$ identified by Algorithm 2 can also be shown to satisfy the following:

1. Each code-pattern appears in exactly one CFP_i .
2. Each $CodePats(api_i)$ can be constructed by combining candidate fingerprints using set union.

In fact, the properties above suggest an upper bound for k , the number of candidate fingerprints. In particular, $k \leq |\bigcup_{i \in [1..n]} CodePats(api_i)|$. Note that while the concept lattice can be exponentially large in the number of API functions (because asymptotically, it is a lattice on the power set of API functions), this upper bound places a restriction on the number of nodes that will be marked in line 7 of Algorithm 2. This is key, because these nodes introduce candidate fingerprints, and as discussed in Section 2, they must be manually examined for refinement in Step B.

Several algorithms have been proposed in the literature to compute concept lattices. We chose to implement the incremental algorithm by Godin *et al.* [16] because it has been shown to work well in practice (Ammons *et al.*[3]). While this algorithm is asymptotically exponential—its complexity is $O(2^{2p}|I|)$, where p is an upper bound on the number of features of any instance in I —the algorithm scaled well in our case studies.

4 Refining fingerprints with constraints

As described in Section 2.2, candidate fingerprints obtained from concept analysis are imprecise for two reasons. First, because of flow-insensitivity, a pair of code-patterns pat_1 and pat_2 that are not part of the same fingerprint may appear in the same candidate fingerprint. Second, the resource manipulations in a candidate fingerprint may be associated with multiple, possibly unrelated resource instances. Thus, candidate fingerprints must be refined using precision constraints. Domain-specific constraints can additionally be applied to refine constraints with domain-specific requirements.

This section presents a unified framework to express constraints and refine candidate fingerprints (Step B of our approach). Both precision constraints and domain-specific constraints can be expressed in this framework.

Constraint	:=	$Separate(PatSet, PatSet) \mid Ignore(PatSet)$ $\mid Combine(PatSet, PatSet)$
PatSet	:=	Set of code-patterns (as defined in Figure 3)

Figure 8: Grammar for constraints.

As Figure 8 shows, each constraint is either a $Separate(X, Y)$, an $Ignore(X)$ or a $Combine(X, Y)$, where X and Y are sets of code-patterns. $Separate(X, Y)$ refines candidate fingerprints by separating code-pattern sets X and Y into separate fingerprints. $Ignore(X)$ refines candidate fin-

gerprints by discarding the code-pattern set X from candidate fingerprints. $Combine(X, Y)$, for which we have only felt occasional need, combines code-pattern sets X and Y in two candidate fingerprints into a single fingerprint, thus “coarsening” the results of concept analysis. For example, the constraint $Separate(\{1,2,3,4\}, \{5,6\})$ refines the candidate fingerprint in Figure 4 to yield the fingerprints in Figure 6. We now discuss precision and domain-specific constraints in this framework.

Precision constraints are $Separate(X, Y)$ constraints and as discussed in Section 2, they serve two goals. The first goal is to refine candidate fingerprints based upon resource instances manipulated. $Separate(\{1,2,3,4\}, \{5,6\})$, the use of which was illustrated earlier, serves this goal. Formally, each set of code-patterns can be associated with one or more resource instances that it manipulates. We use a constraint $Separate(X, Y)$ to separate code-pattern sets X and Y that manipulate different sets of resource instances. For example, consider the code-patterns (1)-(4) in Figure 4, that appear in the function `ext2.delete_entry`, and the code-patterns (5) and (6), that appear in the function `ext2.find_entry`. Because of the way these functions are invoked in `ext2.rename` (see Figure 5), code-patterns (5) and (6) are associated with the resource instances `old_dir`, `old_dentry`, `new_dir` and `new_dentry`, while code-patterns (1)-(4) are associated with resource instances `old_dir` and `old_dentry`. Because the code-patterns (5) and (6) are applied to additional resource instances, they are separated out using the constraint above. We currently manually identify resource instances associated with a set of code-patterns. However, this can potentially be automated using a program analysis that is sensitive to resource instances manipulated.

The second goal of precision constraints is to identify and remove imprecision introduced because of flow-insensitive program analysis. In particular, a pair of code-patterns pat_1 and pat_2 may appear together in a candidate fingerprint, but may not appear together in all executions of the server. In such cases, a $Separate(pat_1, pat_2)$ constraint separates these code-patterns into different fingerprints. For example, one of the candidate fingerprints that we obtained in the analysis of `ext2` is shown below; it appeared in $CodePats(ext2.ioctl)$.

- | | |
|-----|--|
| (1) | <code>Write ⊥ To inode->i.flags</code> |
| (2) | <code>Write ⊥ To inode->i.generation</code> |

However, `ext2.ioctl` either performs the resource manipulation corresponding to code-pattern (1) or (2), but not both, in each execution, based upon the value of a flag that it is invoked with. Thus, a constraint $Separate(\{1\}, \{2\})$ is used to refine the candidate fingerprint above.

Note that precision constraints are not necessary if more precise program analysis is employed. Algorithm 1 currently lacks flow-sensitivity and data-flow information that

can potentially avoid the imprecisions reported above. However, in each of our case studies we needed precision constraints for fewer than 50% of the candidate fingerprints mined—9/18 for `ext2`, 24/115 for X server, and 4/38 for PennMUSH. Thus, we believe that our current approach strikes a good balance between simplicity and precision of candidate fingerprints.

Domain-specific constraints encode domain knowledge to further refine fingerprints. A domain specific constraint that we have found useful is $Ignore(Pat)$, using which we can eliminate certain code-patterns that we deem irrelevant for security from the set of fingerprints. For example, in the X server, which is an event-based server, each request from an X client is converted into a one or more events that are processed by the server. It may only be necessary to enforce an authorization policy governing the set of events that an X client can request on a resource. In such cases, all code-patterns except those related to event-processing can be filtered out from fingerprints using $Ignore$ constraints.

The use of $Combine$ constraints is relatively infrequent, and may be used if the fingerprints mined by concept analysis are at too fine a granularity. For example, in PennMUSH, we found that 30 of the 38 candidate fingerprints contained only one code-pattern. An administrator may wish to write authorization policies at a higher-level of granularity—where the fingerprint of each security-sensitive operation contains multiple code-patterns. $Combine$ constraints can be used to group together code-patterns to get such fingerprints.

5 Case studies

We conducted case studies on three complex systems, each of which has been in development for at least 10 years. We used (i) the `ext2` file system from Linux kernel distribution 2.4.21, (ii) a subset of the X server (X11R6.8), and (iii) PennMUSH, an online game server (v1.8.1p9). Our goal was to identify security-sensitive operations in each of them. Our results lead us to conclude that our approach is:

1. *Effective.* We were able to mine fingerprints of security-sensitive operations in each of our case studies. For `ext2`, we could correlate many of the security-sensitive operations that we found with the security-sensitive operations identified in the LSM project [30], while for the X server, we found all the fingerprints (and more) that we did in a prior study [15].
2. *Efficient.* Our analysis completed in just over 310 seconds (on a machine with a 1GHz AMD Athlon processor, with 1GB RAM) even for PennMUSH, our largest benchmark. Concept lattices produced for manual examination were succinct (available online [1]) and identified

Benchmark	LOC	Analysis time (secs)	Concept lattice		Num. of cand. fings.	Avg. size of cand. fings.	Refinement needed for
			# Nodes	# Edges			
ext2	4,476	2.1	21	32	18	3.67	9 (50%)
X server/dix	30,096	28.43	319	944	115	3.76	24 (20.87%)
PennMUSH	94,014	318.91	127	301	38	1.42	4 (10.53%)

Figure 9: Results for each of our case studies. Concept lattices are also available online [1].

under 115 candidate fingerprints, with fewer than 4 code-patterns on average. In each case, refining candidate fingerprints, semantically interpreting and associating each fingerprint with a security-sensitive operation took a few hours of manual effort with modest domain knowledge.

Figure 9 presents statistics on the size of the concept lattice and number of candidate fingerprints found. As these results show, concept analysis is effective at distilling several thousand lines of code into concept lattices of manageable size. The number and size of candidate fingerprints indicates the amount of manual effort involved. There were under 115 candidate fingerprints of average size under 4 across all our benchmarks. Note that because our approach is static, these fingerprints cover all the security-sensitive operations on a critical resource. Without our approach, an analyst will have to examine the entire code-base to find such security-sensitive operations. These results show that our approach drastically reduces the amount of manual effort required to mine security-sensitive operations. We now discuss notable points from each case study.

5.1 The ext2 file system

As discussed in Section 2, we focused on how directories are manipulated by the ext2 file system. Concept analysis produced 18 candidate fingerprints, of which 9 had to be refined with precision constraints (examples of which appear in Section 2.3 and Section 4).

The LSM project has identified a set of 11 operations on directories. These operations are used to write SELinux policies governing how processes can manipulate directories. We were able to identify at least one fingerprint for each of these LSM operations from the fingerprints that we mined. Section 2 presented two such fingerprints (for `Dir.Remove.Name` and `Dir.Search`). Two more examples are discussed below.

1. The fingerprint `{Write 0 To inode->i.blocks, Write 1 To inode->u->ext2.inode_info@i.new_inode, Write 4096 To inode->i.blksize}` appears in `CodePats(ext2.create)`, `CodePats(ext2.mkdir)`, `CodePats(ext2.mknod)` and `CodePats(ext2.symlink)`. The code-patterns in this fingerprint were all extracted from the function called `ext2.new_inode` and embody creation

and initialization of a new inode. We believe this is the fingerprint of the LSM operation `File.Create` (note that some LSM directory operations have the `File_` prefix).

2. The fingerprint `{Write 0 To inode->i.size}` appears in `CodePats(ext2.rmdir)`. This code-pattern embodies a key step in directory removal, and is a fingerprint for the LSM operation `Dir.Rmdir`.

5.2 The X11 server

The X server is a popular window-management server. X clients can connect to the X server, which manages resources such as windows and fonts on behalf of these X clients. The X server has historically lacked mechanisms to isolate X clients from each other, and has been the subject of several attacks [20]. Such attacks can be prevented with an authorization policy enforcement, that determines the set of security-sensitive operations that an X client can perform on a resource. While there have been efforts to secure the X server, these efforts have all been manual [7, 10, 20], often taking several years.

We focused on a subset of the X server, its main dispatch loop (called `dix`) that contains code to accept client requests and translate them to lower layers of the server. We focused on this subset because it contains the bulk of code that processes client windows, represented by the `Window` data structure, the resource on which we wanted to identify security-sensitive operations. In addition to `Window`, we also included the `xEvent` data structure, because the X server uses it extensively to process client requests. The API that we used contains 274 functions that the X server exposes to clients.

Concept analysis produced 115 candidate fingerprints with 3.76 code-patterns, on average. We had previously identified fingerprints for 15 security-sensitive operations on the `Window` resource [15]. However, as discussed in Section 1, that work used dynamic analysis, and could potentially miss fingerprints. Further, that work could automate fingerprint-finding only up to the granularity of function calls—these were then manually refined to the granularity of code-patterns. Concept analysis not only identified the fingerprints from prior work at the granularity of code-patterns, but did so automatically. Two examples fingerprints are discussed below.

1. The fingerprint {*Write 20 To xEvent->u->type, Write 1 To xEvent->u->mapRequest->window*}, contained in *CodePats* of 5 API functions, embodies an X client request to map a window on the screen, and potentially represents a security-sensitive operation.
2. The fingerprint {*Write 0 To Window->mapped, Write 18 To xEvent->u->type*, contained in *CodePats* of 7 API functions embodies unmapping a visible X client window from the screen, also a potential security-sensitive operation.

5.3 The PennMUSH server

PennMUSH is an open-source online game server. Clients connecting to a PennMUSH server assume the role of a virtual character, as in other popular massively-multiplayer online roleplaying games. For this work, it suffices to think of PennMUSH as a collaborative database of objects that clients can modify. Objects are shared resources, and an authorization policy must govern the set of security-sensitive operations that a client can perform on each object.

Clients interact with PennMUSH by entering commands to a text server, which activates one or more of 603 internal functions, that we used as the API of PennMUSH. Most of these API functions modify a database of objects. Thus, we tracked how the PennMUSH API manipulates resources of type `struct object`. Concept analysis produced 38 candidate fingerprints. Most of them had only one or two code-patterns, so we only had to refine 4 of these candidate fingerprints using precision constraints to obtain a total of 43 fingerprints, two examples of which we discuss below.

1. The fingerprint *Write 1 To object->name* potentially modifies an object name, and was contained in *CodePats* of 16 API functions, representing creation, destruction and modification of objects. Clearly, unauthorized clients must be disallowed from changing the name of an object, indicating that this is a fingerprint of a security-sensitive operation.
2. The fingerprint {*Write 8 To object->type, Write 0 To object->modification.time, Write 1118743 To object->warnings*} appears in *CodePats*(`cmd.pcreate`) and *CodePats*(`fun.pcreate`), both of which are API functions associated with creation of a “character” object. Here, the number 1118743 represents a flag that signifies that a character should be warned about problems with the objects that they own, and the number 8 written to the field `type` indicates that the newly created object is a character. These code-patterns represent necessary steps in character creation in PennMUSH, and thus indicate that this is fingerprint of a security-sensitive operation.

In PennMUSH, the object data structure has just 18 fields, while the API contains 603 functions. Each security-sensitive operation is performed at the granularity of accesses to just one or two of the fields of object. This explains the smaller number and size of candidate fingerprints extracted by concept analysis (as compared to X server).

While the security-sensitive operations that we extracted for PennMUSH can definitely form the basis for writing policies, site-specific policies may be created by combining several security-sensitive operations. For example, an administrator might decide that reading an object’s name is as security-sensitive as determining the kind of object. He can then use the domain-specific constraint *Combine*(*Read object->name, Read object->type*) to combine these code-patterns together into a single fingerprint that embodies this security-sensitive operation.

6 Related work

This paper overcomes two important shortcomings that we had identified in prior work [15]. The need for an *a priori* description of security-sensitive operations hindered the application of the techniques developed there to a wide variety of servers. Further, a dynamic trace-based approach to fingerprint-finding meant that large portions of the legacy server may go unexplored, thus resulting in an incomplete set of fingerprints.

As discussed in Section 1, our approach follows the aspect-oriented paradigm. Several other tools, such as PoET/PSLang [11], Naccio [12], Polymer [6] and our own prior work on Tahoe [14] also follow an aspect-oriented approach to enforce authorization policies on legacy code. In all these tools, a security analyst provides a description of locations to be protected (join points) as well as the policy check at each location (advice). These tools then weave calls to a reference monitor at each of these locations. However, when legacy servers manage their own resources, identifying locations where policy checks must be weaved becomes a challenge. The techniques developed in this paper can benefit the above tools by reducing the manual effort involved in identifying locations for reference monitoring, as well as the advice to be integrated at these locations.

Concept analysis has previously been used in software engineering, including aspect mining (Ceccato *et al.* present a survey of such techniques [8]) and software modularization. For example, concept analysis has been used on identifier names to find methods and classes that implement similar functionality [27]. Dynamic analysis in conjunction with concept analysis has been used to find methods that implement a particular feature [9, 26]. The idea here is to run an instrumented version of the program under dif-

ferent use-cases and label the traces with these use cases. Each trace contains information about the methods executed. Traces are then clustered using concept analysis to find crosscutting concerns, and thus identify aspects. Concept analysis has also found use to identify modular structure in legacy programs [21, 24, 25, 28]. The modular structure so identified can be used to refactor legacy software (e.g., convert non-object-oriented programs into object-oriented ones [24]). Another recent use of concept analysis is in the context of debugging mined specifications [3]. Automatically mined temporal specifications may often be buggy, and the problem here is for an analyst to classify each mined specification as correct or buggy. “Similar” traces can be clustered using concept analysis, so the analyst can decide en-masse whether an entire cluster is buggy.

7 Summary and future work

We presented an approach to reduce the manual effort involved in mining security-sensitive operations in legacy servers. Our approach uses concept analysis to mine fingerprints, which are code-level descriptions of security-sensitive behavior. Our experiments with three complex real-world servers show that our approach is efficient and effective at finding security-sensitive operations.

In the future, we plan to enhance the static analysis phase with flow-sensitivity and data-flow information. This will potentially enable extraction of temporal fingerprints, i.e., those that also identify the order in which fields of a data structure representing a resource are manipulated.

References

- [1] Concept lattices for our case studies. Available online at: <http://www.cs.wisc.edu/~vg/papers/tr1580>.
- [2] The PennMUSH multi-user dungeon online game server. <http://pennmush.org>.
- [3] AMMONS, G., MANDELIN, D., BODÍK, R., AND LARUS, J. R. Debugging temporal specifications with concept analysis. In *Proc. PLDI* (June 2003).
- [4] ANDERSON, J. P. Computer security technology planning study, volume II. Tech. Rep. ESD-TR-73-51, Deputy for Command and Management Systems, L. G. Hanscom Field, Bedford, MA, October 1972.
- [5] Aspect-oriented software development glossary. <http://aosd.net/wiki/index.php?title=Glossary>.
- [6] BAUER, L., LIGATTI, J., AND WALKER, D. Composing security policies with Polymer. In *Proc. PLDI* (June 2005).
- [7] BERGER, J., PICCIOTTO, J., WOODWARD, J., AND CUMMINGS, P. Compartmented mode workstation: Prototype highlights. *IEEE Transactions on Software Engineering* 16, 6 (June 1990).
- [8] CECCATO, M., MARIN, M., MENS, K., MOONEN, L., TONELLA, P., AND TOURWE, T. A quantitative comparison of three aspect mining techniques. In *Proc. 13th Intl. Wkshp. on Program Comprehension* (May 2005).
- [9] EISENBARTH, T., KOSCHKE, R., AND SIMON, D. Locating features in source code. *IEEE TSE* 29, 3 (March 2003).
- [10] EPSTEIN, J., MCHUGH, J., ORMAN, H., PASCALE, R., SQUIRES, A.-M., DANNER, B., MARTIN, C., BRANSTAD, M., BENSON, G., AND ROTHNIE, D. A high assurance window system prototype. *Journal of Computer Security* 2, 2-3 (1993).
- [11] ERLINGSSON, U. *The Inlined Reference Monitor Approach to Security Policy Enforcement*. PhD thesis, Cornell University, January 2004.
- [12] EVANS, D., AND TWYMAN, A. Flexible policy-directed code safety. In *Proc. 1999 IEEE S&P* (May 1999).
- [13] FLETCHER, B. Case study: Open source and commercial applications in a Java-based SELinux cross-domain solution. In *Second Annual Security-enhanced Linux Symp.* (March 2006).
- [14] GANAPATHY, V., JAEGER, T., AND JHA, S. Automatic placement of authorization hooks in the Linux security modules framework. In *Proc. ACM CCS* (November 2005).
- [15] GANAPATHY, V., JAEGER, T., AND JHA, S. Retrofitting legacy code for authorization policy enforcement. In *Proc. 2006 IEEE S&P* (May 2006).
- [16] GODIN, R., MISSAOUI, R., AND ALAOUI, H. Incremental concept formation algorithms based on Galois (concept) lattices. *Computational Intelligence* 11, 2 (1995), 246–267.
- [17] HOCKING, M., MACMILLAN, K., AND SHANKAR, D. Case study: Enhancing IBM Websphere with SELinux. In *Second Annual Security-enhanced Linux Symp.* (March 2006).

- [18] JAEGER, T., EDWARDS, A., AND ZHANG, X. Consistency analysis of authorization hook placement in the Linux security modules framework. *ACM TISSEC* 7, 2 (May 2004), 175–205.
- [19] KICZALES, G., LAMPING, J., MENDHEKAR, A., MAEDA, C., LOPES, C., LOINGTIER, J.-M., AND IRWIN, J. Aspect-oriented programming. In *Proc. ECOOP* (June 1997).
- [20] KILPATRICK, D., SALAMON, W., AND VANCE, C. Securing the X Window system with SELinux. Tech. Rep. 03-006, NAI Labs, March 2003.
- [21] LINDIG, C., AND SNELTING, G. Assessing modular structure of legacy code based on mathematical concept analysis. In *Proc. ICSE* (May 1997).
- [22] LOSCOCO, P., AND SMALLEY, S. Integrating flexible support for security policies into the Linux operating system. In *FREENIX track: USENIX Annual Technical Conf.* (June 2001).
- [23] NECULA, G. C., MCPeAK, S., RAHUL, S. P., AND WEIMER, W. CIL: Intermediate language and tools for analysis and transformation of C programs. In *International Conf. on Compiler Construction* (April 2002). Available at: <http://manju.cs.berkeley.edu/cil>.
- [24] SIFF, M. *Techniques for Software Renovation*. PhD thesis, University of Wisconsin-Madison, 1998.
- [25] SNELTING, G., AND TIP, F. Reengineering class hierarchies using concept analysis. In *Proc. ACM SIGSOFT FSE* (November 1998).
- [26] TONELLA, P., AND CECCATO, M. Aspect mining through the formal concept analysis of execution traces. In *Proc. 11th Conf. on Reverse Engineering* (November 2004).
- [27] TOURWE, T., AND MENS, K. Mining aspectual views using formal concept analysis. In *Proc. 4th Intl. Wkshp. on Source code Analysis and Manipulation* (September 2004).
- [28] VAN DUERSEN, A., AND KUIPERS, T. Identifying objects using cluster and concept analysis. In *Proc. ICSE* (May 1999).
- [29] WILLE, R. Restructuring lattice theory: An approach based on hierarchies of concepts. *Ordered Sets* (1982), 445–470.
- [30] WRIGHT, C., COWAN, C., SMALLEY, S., MORRIS, J., AND KROAH-HARTMAN, G. Linux security modules: General security support for the Linux kernel. In *USENIX Security Symp.* (August 2002).
- [31] The X11 Server, version X11R6.8 (X.Org Foundation).
- [32] ZHANG, X., EDWARDS, A., AND JAEGER, T. Using CQUAL for static analysis of authorization hook placement. In *USENIX Security Symp.* (August 2002).