



CacheFX: A Framework for Evaluating Cache Security

Daniel Genkin
genkin@gatech.edu
Georgia Institute of Technology
Atlanta, USA

William Kosasih
william.kosasih@adelaide.edu.au
University of Adelaide
Adelaide, Australia

Fangfei Liu
fangfei.liu@intel.com
Intel Labs
Hillsboro, USA

Anna Trikalinou*
atrikalinou@microsoft.com
Microsoft
Seattle, USA

Thomas Unterluggauer
thomas.unterluggauer@intel.com
Intel Labs
Villach, Austria

Yuval Yarom[†]
yuval.yarom@rub.de
Ruhr University Bochum
Bochum, Germany

ABSTRACT

Over the last two decades, the danger of sharing resources between programs has been repeatedly highlighted. Multiple *side-channel attacks*, which seek to exploit shared components for leaking information, have been devised, mostly targeting shared caching components. In response, the research community has proposed multiple cache designs that aim at curbing the source of side channels.

With multiple competing designs, there is a need for assessing the level of security against side-channel attacks that each design offers. Several metrics have been suggested for performing such evaluations. However, these tend to be limited both in terms of the potential adversaries they consider and in the applicability of the metric to real-world attacks, as opposed to attack techniques. Moreover, all existing metrics implicitly assume that a single metric can encompass the nuances of side-channel security.

In this work we propose CacheFX, a flexible framework for assessing and evaluating the resilience of cache designs to side-channel attacks. CacheFX allows the evaluator to implement various cache designs, victims, and attackers, as well as to exercise them for assessing the leakage of information via the cache.

To demonstrate the power of CacheFX, we implement multiple cache designs and replacement algorithms, and devise three evaluation metrics that measure different aspects of the caches: (1) the entropy induced by a memory access; (2) the complexity of building an eviction set; (3) protection against cryptographic attacks. Our experiments highlight that different security metrics give different insights to designs, making a comprehensive analysis mandatory. For instance, while eviction-set building was fastest for randomized skewed caches, these caches featured lower eviction entropy and higher practical attack complexity. Our experiments show that all non-partitioned designs allow for effective cryptographic attacks. However, in state-of-the-art secure caches, eviction-based

attacks are more difficult to mount than occupancy-based attacks, highlighting the need to consider the latter in cache design.

CCS CONCEPTS

• **Hardware** → **Simulation and emulation**; • **Security and privacy** → **Side-channel analysis and countermeasures**.

KEYWORDS

secure caches, side-channel attacks, security metrics

ACM Reference Format:

Daniel Genkin, William Kosasih, Fangfei Liu, Anna Trikalinou, Thomas Unterluggauer, and Yuval Yarom. 2023. CacheFX: A Framework for Evaluating Cache Security. In *ACM ASIA Conference on Computer and Communications Security (ASIA CCS '23)*, July 10–14, 2023, Melbourne, VIC, Australia. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3579856.3595794>

1 INTRODUCTION

Memory caches, which store recently accessed memory contents, became a standard feature of mainstream computer processors. While instrumental to the needs of contemporary computing, sharing caches between multiple untrusted programs can lead to undesired information leaks, in that the contents of caches necessarily depend on past computations and by their nature are intended to enhance the speed of future computations [22].

By monitoring the timing of memory operations, an attacker can infer the state of the cache and learn information about the behavior of the victim. Such *side channels* can result from any of the caches in the processor [1, 27, 37, 43, 50, 53], and using such side channels, a malicious actor may seek to infer sensitive information such as cryptographic keys [2, 6, 24, 48, 50, 60, 83], user keystrokes and their timing [32, 59, 61], address space information [20, 28, 35], and others [4, 49, 65, 79]. The shared use of caches has also been shown to enable efficient *covert channels* [8, 44, 46], where a malicious Trojan colludes with an attacker to bypass the system's security policy.

The two main types of cache attacks are contention-based attacks [1, 27, 37, 43, 50, 53, 81], which seek to exploit the limited storage space in the cache, and reload-based attacks [28, 32, 33, 83], which seek to exploit the attacker's ability to evict memory it shares with the victim from the cache. Because reload-based attacks rely on shared memory, preventing memory from being shared across security domains can be an effective countermeasure.

*Work done while affiliated with Intel Labs.

[†]Work partially done while affiliated with the University of Adelaide.



This work is licensed under a [Creative Commons Attribution International 4.0 License](https://creativecommons.org/licenses/by/4.0/).

ASIA CCS '23, July 10–14, 2023, Melbourne, VIC, Australia

© 2023 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-0098-9/23/07.

<https://doi.org/10.1145/3579856.3595794>

Many cache designs have been suggested to address contention-based cache attacks: partitioned caches aim to prevent contention [18, 76] while randomized caches [42, 57, 58, 67, 77] aim to introduce noise and prevent the attacker from analyzing the side-channel signal. Randomized caches often try to prevent the attacker from mapping addresses to predictable cache line indexes, a step that is considered essential for the attack. Finally, some proposals try to prevent cross-core attacks by tweaking the inclusion properties of shared cache levels in modern processors [29, 40, 80, 82].

With multiple proposals for protecting against contention, processor vendors need some method of assessing their security. Several approaches for evaluating secure caches have been suggested [8, 10, 11, 14, 15, 19, 21, 23, 25, 34, 41, 75, 86, 87]. For instance, [19, 25, 41] use formal methods and model checking to determine cache leakage, [34, 75] model cache attacks to obtain attack success probabilities, and [14, 15] use a three-step attack model to exhaustively test for vulnerable attack patterns and apply it to various Arm devices [13]. However, all of those suffer from some limitations as they only work with simple cache models, focus on theoretical analysis, cannot be automated, or do not cover the full range of cache attacks. In addition, to strengthen confidence in the security of cache designs, the evaluation of multiple metrics is mandatory. Thus, in this work, we investigate the following question.

How can we evaluate the security that cache designs offer against contention-based cache attacks?

1.1 Our Contribution.

To address this question, this paper presents CacheFX, a framework for evaluating the security of caches. CacheFX provides an interface for emulating the operation of cache designs with different victims and attackers, and measuring the leakage for each combination of attacker, victim, and cache design. We demonstrate the flexibility of CacheFX by implementing nine cache designs, and evaluating them using three metrics.

Cache Designs. The cache designs we implement include traditional fully-associative and set-associative caches, PLCache [76], Newcache [42], PhantomCache [67], ScatterCache [77], way-partitioned caches [18], and the two variants of CEASER [57, 58]. For caches that do not stipulate a replacement policy, we support four replacement algorithms: random replacement, least recently used (LRU), and two variants of pseudo-LRU.

Evaluation Metrics. We further design and implement three evaluation metrics. These not only add to the existing portfolio of metrics proposed in prior works, allowing cache designers more options for cache evaluation, but also demonstrate the flexibility of CacheFX and its ability to measure a variety of metrics.

- **Relative Eviction Entropy.** The *Relative Eviction Entropy* (REE) is a new metric we propose to measure the information leakage from a single victim access via the cache side channel. We then use an unrealistic attacker that can set the cache to a known configuration and accurately observe the cache state. We combine this attacker with a victim that accesses a single cache line only and then calculate the amount of information that the attacker receives from observing which cache line the victim evicted.
- **Measuring Eviction-Set Creation.** Most cache-based side channel attacks require the attacker to find a minimal set of

addresses such that accessing them results in evicting a specific victim line from the cache. Our second metric measures the difficulty for an attacker to find such sets. Here, we implemented three eviction-set building strategies: Single Holdout Method (SHM) [58], Group Elimination Method (GEM) [58] and Prime+Prune+Probe (PPP) [55].

- **Cryptographic Attack.** Cryptographic attacks evaluate the protection that the cache provides for cryptographic code. Our victims perform cryptographic functions using implementations known to be vulnerable. We let the victims encrypt data with one of two keys, and task the attacker with distinguishing between the keys. We evaluate both traditional attacks that aim to exploit eviction sets, and occupancy-based attacks [63, 65]. In both cases, we use the number of encryptions the attacker needs to observe to distinguish between the keys as the security metric.

Results. Using CacheFX we can easily compare a large number of cache design and evaluate them with multiple metrics. The framework allows for easy implementation of new strategies and designs and for comparison across the field. We now show some of the new insights about caches and cache attacks, which CacheFX highlights.

Multiple Metrics. Our experiments show that different metrics highlight different aspects of the caches. In particular, we find that building eviction sets is faster in skewed caches such as ScatterCache and CEASER-S, than other randomized caches, such as fully associative caches or PhantomCache. Faster eviction-set construction reduces the effort required for mounting attacks. At the same time, our experiments show that, with the right parameters, skewed caches are not less secure when it comes to cryptographic attacks.

Evaluating Cryptographic Attacks. We find that the security against cryptographic attackers depends not only on the design, but also on other parameters, such as the replacement policy and the cache associativity. We also show that *all* non-partitioned caches are vulnerable to both eviction-set and occupancy attacks.

Comparing Attacks. Evaluating the metrics with CacheFX enables us to compare attack strategies across the various cache designs and parameters. For example, we find that most non-partitioned secure cache designs offer protection against eviction-set attacks. However, cache-occupancy attacks are left unconsidered and for highly secure designs occupancy attacks are no less effective than eviction-set attacks. Our evaluation thus demonstrates that partitioning is preferable from a side-channel perspective and the resistance to eviction-set attacks of non-partitioned solutions may be tuned to match the respective complexity of cache-occupancy attacks to balance overall cache side-channel resistance and cost.

2 BACKGROUND

2.1 Cache Attacks

Modern processors use an array of *caches* to speed up accesses to memory by exploiting program locality. Caches are architecturally transparent—whether a specific piece of data is cached does not affect the architectural behavior of a program. It does, however, affect the performance of programs, thus monitoring program performance can reveal information about the state of the cache.

Tsunoo et al. [69] were the first to demonstrate that this information can be used to recover secret cryptographic keys. Early attacks focused on the L1 data cache [50, 53, 68], but attacks targeting other

caches soon emerged [1, 27, 33, 37, 43, 81, 83]. Cache attacks target symmetric cryptography [7, 24, 33, 37, 48, 50, 68, 69], public-key schemes [1, 9, 27, 43, 53, 60, 83], Post-quantum cryptography [30, 54], and non cryptographic software [4, 20, 28, 32, 35, 49, 61, 65, 79]. There are two main groups of cache attacks:

Reload-Based Attacks monitor accesses to a shared memory address [32, 33, 83]. The attack first evicts the data from the cache either via a dedicated instruction [33, 83] or by forcing contention on the cache set containing the data [32]. The attacker then waits a bit and measures the time to access the previously evicted data. If while waiting the victim accesses the data, the data will be cached and the attacker's access be fast. As not sharing memory across domains can be an effective mitigation, this attack type is not the main target of this work.

Contention-Based Attacks, which seek to exploit the limited storage in the cache, and in particular in each of the cache sets [1, 27, 37, 43, 50, 53, 81], are the main focus of this work. The most common contention-based attack technique is Prime+Probe, where the attacker first primes the cache by filling some or all of the cache sets with their data and, after letting the victim some time to execute, measures the time to access the cached data. A slow access indicates that the data is no longer cached, suggesting eviction from a cache set due to victim activity. Variants of the attack avoid using timing information by relying on performance counters [4, 70] or transaction aborts [16] for contention detection. Another contention-based attack is Evict+Time [27, 38, 50, 72], where the attacker evicts data from the cache before measuring the execution time of a victim. The victim's execution time will be longer if the evicted data is used by the victim, revealing information on cache sets that the victim uses.

Other Types of Cache Attacks. Some cache attacks do not fit into either of the two groups. Such attacks seek to exploit implementation aspects of the cache, such as port contention [85], cache flushing time [31], replacement policies [56, 73, 78], cache inspection operations [5], or variations in power consumption based on caching information [51]. Due to their specific requirements, such attacks are outside the scope of this work.

Eviction-Set Construction. For many of the attacks mentioned above, the attacker must be able to repeatedly evict specific contents from the cache. Typically, attackers achieve this by constructing an *eviction set*, which consists of memory locations that all map to the same cache set as the data to evict. When mapping information for the cache is available, constructing an eviction set tends to be straightforward. However, when the mapping function is undocumented or when the information it uses for indexing the cache is not available to the attacker, further techniques are required to recover the missing information. Past research shows how to reverse-engineer undocumented mapping functions [27, 36, 45, 47, 84], and how to build eviction sets without physical address information [43, 74].

2.2 Secure Caches

Several proposed cache designs aim to mitigate contention-based attacks. Their mitigation strategies are either based on partitioning [18, 76] or randomization [42, 57, 58, 77].

Partitioned Caches Way-partitioned caches [18] enforce a strong partitioning between security domains by letting each security domain use a different subset of the cache ways. Hence, domains not sharing cache ways will not see any interference. Alternatively, Partition-Locked (PL) [76] caches share the whole cache among all security domains, but offer to pin cache lines in the cache. These pinned lines cannot be evicted by other security domains, preventing contention-based attacks. However, aggressive pinning can starve other domains and severely degrade their performance.

CEASER. The CEASER cache [57] is based on an ordinary set-associative cache and uses encryption to randomize the mapping of addresses to cache sets. As a result, attackers need to first profile the victim's accesses of interest to find a suitable eviction set before they can perform contention-based attacks. To limit the attacker's time for finding such eviction set, CEASER regularly changes the encryption key. However, in this work we only measure information leakage in each key epoch (i.e., during the time period the cache uses the same key) and do not model re-keying. This allows assessing the security of pure cache-set randomization as it is needed to appropriately tune the re-keying interval for long-term security.

CEASER-S and ScatterCache. With improving eviction set building techniques [58, 74], CEASER required higher key refresh rates to maintain security, resulting in increased overheads. CEASER-S [58] and ScatterCache [77] thus use a skewed cache [62] to improve cache randomization. These skewed caches split the cache into partitions along its ways and use a different key to encrypt the address to index into each partition. The partition count can vary between 1 (CEASER) and the number of ways (ScatterCache) and allows to control the degree of randomization. As before, we omit re-keying to assess the security of pure cache randomization with skewing.

PhantomCache. PhantomCache [67] builds upon set-associative caches and maps each address to multiple sets via multiple hash functions, i.e., it looks in multiple sets for a cache hit. On a cache miss, PhantomCache randomly selects one of the sets the address maps to and inserts the cache line into the chosen set. The number of cache sets looked up in parallel determines the degree of randomization and the cost of lookup. As before, we evaluate PhantomCache without re-keying.

NewCache. Rather than randomizing the cache mapping, NewCache [42] is a more efficient implementation of a fully associative cache. NewCache allows every cache line to be stored in any of the physical lines of the cache. Compared to a standard associative design, it optimizes power using a two-step look-up procedure: For a cache that can hold 2^n physical cache lines, NewCache first looks up $n + k$ index bits of the cache line address in a 2^n -element Content-Addressable Memory (CAM), which has a 1:1 mapping to the actual cache lines. Only if these $n + k$ bits match, this *index hit* is secondly followed by checking the tag for the respective entry. If the *tag hits*, the cache line is found and returned. If there is a *tag miss* for the same security domain in the second step, the tag and cache line are simply replaced. If there is a *index miss*, any of the 2^n cache lines in the cache is randomly replaced. While for large k NewCache resembles a traditional fully associative cache, a smaller k significantly reduces power and implementation cost.

Secure Cache Hierarchies. For cross-core attacks, attackers must be able to evict data not only from the caches they use, but also from caches at the victim's core. Most cross-core attacks rely on

an inclusive LLC, which ensures that the contents of the shared LLC is a superset of the contents of all private caches. Inclusiveness guarantees that when data is evicted from the LLC, it is also evicted from all core-private caches. If the LLC is not inclusive, evictions from it do not necessarily translate to evictions from core-private caches, and may hamper cache attacks. Yan et al. [81] use a similar property of cache directories in non-inclusive caches. Several cache designs and features [29, 52, 82] that prevent cross-core eviction are outside the scope for this work.

3 PROBLEM DESCRIPTION

With the abundance of secure cache designs, there is a clear need for systematically evaluating the security of caches to ensure that emerging cache architectures deliver the promised protection. Tackling this task, previous works [8, 10–12, 14, 15, 17, 34, 75, 86, 87] have suggested several metrics. However, all of these tend to suffer from some limitations to their practicality. For example, measuring the amount of information that can be transferred via a cache side channel [8] or the correlation between a specific victim’s activity and attacker observation [10] may not translate easily to cryptographic attack scenarios. Possibly the most common limitation of these approaches is the attempt to provide a single metric that somehow represents the security of the cache.

A General Cache Evaluation Framework. Instead of focusing on a single metric, this work proposes CacheFX, a framework for evaluating the security of cache designs. The main design aim of CacheFX is flexibility: CacheFX is extensible and allows evaluating various combinations of victims, cache designs, and attack strategies. As a proof of CacheFX’s generality, this paper implements and evaluates three security metrics on nine different cache designs.

A Leakage Upper Bound. While we try to evaluate in realistic scenarios, CacheFX aims to provide an upper bound on the amount of leakage an attacker can obtain from a cache design. We thus assume an attacker who has significant control over the victim and is tightly synchronized with the victim’s execution. We further assume that the attacker has access to victim’s memory layout and code and thus knows the position of “interesting” data in the cache (e.g., cache lines containing secrets). This allows the attacker to craft inputs to the victim that may cause specific cache footprints. Unless required by the cache design, we assume a noise-free environment without any system activity besides the attacker and victim.

Using such strong assumptions allows CacheFX to properly evaluate the security offered by the cache design, as opposed to being misled by security guarantees stemming from other components’ noise. We note that previous works have demonstrated that attackers can find interesting cache lines [43] and overcome noise [8].

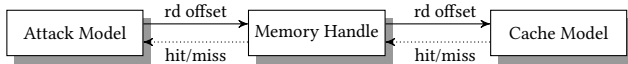


Figure 1: CacheFX design overview.

4 CACHEFX DESIGN

As mentioned above, CacheFX is designed to provide an easily extendable framework for evaluating (1) the security of emerging

cache designs and (2) the applicability and complexity of new attack strategies to both deployed and emerging caches. To facilitate these goals, CacheFX is split into three major components as depicted in Figure 1. First, the *attack model* provides a set of interfaces and their implementations to model different attack and security evaluation strategies. Attack models use a *memory handle* to request reads, writes, and cache line invalidations to the memory system by specifying a certain offset into a memory region that is associated with the memory handle. The memory handle translates the requests to cache line addresses and queries the *cache model* correspondingly. The cache model returns whether the request hit or missed in the cache via the memory handle to the attack model, which then proceeds with the attack accordingly. Finally, the cache model provides a generic cache interface allowing for multiple different cache implementations. In the following, we give additional details about each of these components. The CacheFX code and a more comprehensive documentation is available at <https://github.com/0xADE1A1DE/CacheFX>.

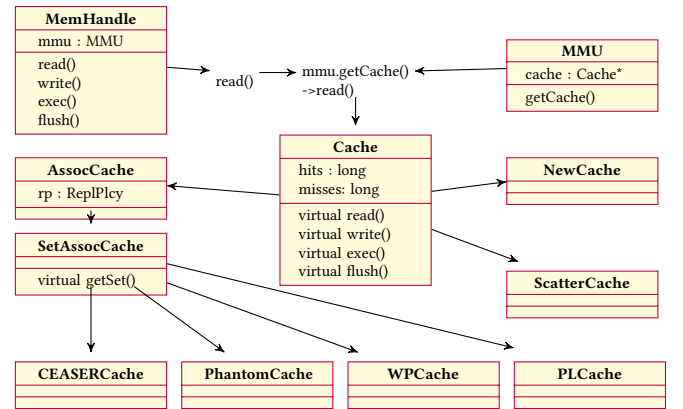


Figure 2: CacheFX overall architecture

4.1 Cache Model

CacheFX’s cache model offers a generic cache interface that the memory handle and the attacker model can use to issue read, write, and invalidation requests to the cache under test. For each of these requests, the cache responds with whether the request hit or missed. This indication removes the need to distinguish between hits and misses using (potentially noisy) timing measurements, providing an upper bound on the amount of leakage available to the attacker and consequently lower bounding the attack’s complexity.

Supported Cache Designs. The cache model currently provides multiple implementations of security-oriented cache designs: fully associative cache, set-associative cache, way-partitioned cache, partition-locked cache, CEASER and CEASER-S [58, 74], ScatterCache [77], NewCache [42], and PhantomCache [67]. These cache implementations are parameterized by the number of sets, ways, replacement policy, and cache-specific parameters. Unless a cache design mandates a specific replacement policy, all the implementations support LRU, Bit-PLRU, Tree-PLRU, and random replacement.

The Cache Interface. The internal interface of the cache model is defined by the *Cache* virtual class which acts as an interface with

mainly four functions, namely: *read*, *write*, *exec*, and *flush* all of which embody actual requests that are sent to the cache hardware. Being an interface, these functions' specific implementations are delegated to the derived classes of the *Cache* virtual class, such as, *NewCache*, *ScatterCache*, *SetAssocCache*, as shown in Figure 2. Each inherited class devises distinct mechanisms for dealing with the four requests which are carefully modeled based on the actual cache hardware functionality. This design promotes effective encapsulation and abstraction along with ease of implementation when adding new models as CacheFX's simulation logic interfaces only with these top level functions. To support generic set-associative classes, CacheFX supports the *SetAssocCache* class, where each set is implemented as an *AssocCaches*. This allows for easy creation of the different set-based caches and reuses the code of the associative cache class, e.g., to support multiple replacement algorithms.

The mechanism for selecting a set is implemented by the function *getSet()*. For the *SetAssocCache*, this is a simple modulus operation of the cache line and the number of cache sets. For *CEASERCache* and *PhantomCache* the set selection mechanism is based on a hashing algorithm. For both of these cases, the *getSet()* function is simply overloaded while the underlying implementation of *SetAssocCache* remains unchanged. Similarly, *WPCache* separates its context into two partitions, one for sensitive, and another for general data. This model simply consists of two *SetAssocCache* instances that are chosen based on the security context of a data access. Finally, *PLCache*'s design is based on the *SetAssocCache*, but with a minor change to the replacement algorithm to facilitate pinning of specific lines.

Statistics Generation. The abstract cache model automatically tracks the number of cache hits and misses for each security domain. In addition, the cache model can return the evicted address, if a cache access causes an eviction. While attackers usually do not have direct access to such information, providing the address allows us to apply novel and efficient techniques, such as the *Relative Eviction Entropy* (REE) in Section 5.1, for analyzing cache security.

4.2 Attack Model

CacheFX's attack model implements the actual adversarial strategy and evaluates the cache design under test. Currently, CacheFX supports three security evaluation strategies:

The Attacker. CacheFX allows to model synchronized pairs of victims and attackers, aiming to evaluate the security of cache designs with respect to realistic attacks, such as cache attacks against cryptographic block ciphers. CacheFX supports two types of attackers, *EvictionAttacker*, and *OccupancyAttacker*, both are subclasses of a generic *Attacker* class that manages the attack and collects the success statistics.

Information Leakage Assessment. CacheFX supports entropy-based security metrics that quantify information leakage during cache attacks (e.g., mutual information analysis). Most noteworthy, CacheFX implements a novel technique for evaluating information leakage in cache designs via the REE, by efficiently analyzing the statistical properties of a cache's cross-domain eviction behavior.

Eviction-Set Profiling. CacheFX provides an environment that allows for the evaluation of strategies to construct eviction sets for different cache designs.

Experiment Randomization and Automation. CacheFX allows to conduct each of these experiments multiple times with randomized address ranges to automatically obtain statistical data like maximum, minimum, etc. CacheFX hereby collects data such as cache statistics and attack success rates.

4.3 Victim Model

The purpose of this model is to simulate the behavior of victim applications within CacheFX's simulation. CacheFX implements a number of victim models including:

- **SingleAccessVictim** is the simplest victim model that repeatedly accesses a single address in the memory.
- **AESVictim** simulates the behavior of AES encryption.
- **SquareMultVictim** imitates the square-and-multiply routine used in popular algorithms such as RSA.

All models are carefully crafted to resemble their actual attack characteristics. Take for example the *AESVictim* model whose code is taken directly from the original AES implementation, but is adapted to call into CacheFX's API on each T-table access. In other words, all memory operations of the victim are redirected to CacheFX for further simulation. A similar approach is taken for *SquareMultVictim* where the cache line containing the multiplication code is executed conditionally based on the exponent. In this context, a call to CacheFX's cache line read function is invoked at the beginning of the multiplication basic block to notify the simulator of the instruction cache read. With this methodology, we ensure high precision of CacheFX's simulated model characteristic in comparison to the actual implementation.

Note that CacheFX victims currently focus on cryptographic code as its properties are well understood and are well suited to analyzes the properties of the underlying cache design. However, CacheFX is generic enough to similarly model other leaking code, e.g., (de-)compression algorithms, data en-/decoders [66] or neural networks [79].

4.4 The Attack Controller Function

As its name suggests, the purpose of this function is to moderate interactions between *Attack Model* and *Victim Model* (both of which subsequently interact with *Cache Model*).

Listing 1: Attack Controller main loop

```
// Eviction or Occupancy Attacker
Attacker* a = createAttacker(attackerModel);
// Single, AES, or SquareMult Victim
Victim* v = createVictim(victimModel);

while (controller_run) {
    a->prime(); v->cipher(); a->probe();
}
```

Listing 1 outlines the main workings of the attack controller function. At the outset, pointers to both the *Attacker* and *Victim* classes are instantiated to their desired model. At the heart of this function is a loop that interleaves the execution of the attacker and the victim i.e. the prime, cipher, and probe methods. Note that the controller is agnostic of the specific implementations of these victim and attack functions and that the choice is left up to polymorphism.

The controller consolidates all major simulation components and can be thought of as the main driver of CacheFX.

5 EVALUATION

Recognizing that no single metric is sufficient for measuring the resilience of caches to side-channel attacks, we evaluate emerging cache designs w.r.t. multiple metrics using our framework. First, the *Relative Eviction Entropy* (REE) metric measures the amount of information (in bits) that an attacker can deduce following a single memory access performed by the victim. Our second metric measures the complexity of creating eviction sets in randomized caches. Our third metric measures the complexity of performing cache attacks on cryptographic implementations. It evaluates both traditional attacks that seek to exploit eviction sets and cache-occupancy attacks [8, 65], which do not require eviction sets.

We now discuss each metric in detail and compare different designs according to each of the measurement metrics.

5.1 Relative Eviction Entropy

In this section we introduce our *Relative Eviction Entropy* (REE) technique for effectively measuring the amount of information available to an attacker following a single memory access performed by the victim. We begin by observing that traditional mutual information analysis [8, 86] achieves such estimation for general side channels by computing a 2-dimensional joint probability distribution, which describes the likelihood of each victim activity (side channel input) to be mapped to an effect observable by an attacker (side channel output). For the case of caches, this means that for any address i accessed by the victim, and for all cached addresses a , we need to compute $p_e(a, i)$ which is the probability that a is evicted from the cache assuming that the victim accesses address i . We note that mutual information techniques typically measure average leakage across accesses, and thus do not capture the worst-case leakage.

Avoiding Quadratic Overheads. To avoid the quadratic overhead of computing the 2-dimensional joint probability distribution, we start by observing that natural cache designs typically do not have different eviction behavior between cache line addresses, and instead use the same replacement policy constantly across all cache lines. In addition to simplifying cache designs, this property implies that all cache line addresses exhibit the same leakage behavior. Leveraging this fact, we can thus fix an arbitrary address i to be accessed by the victim, and simply sample $p_e(a, i)$ for all other addresses a . This avoids iterating over all possible values of i and thus makes the evaluation time of our metric linear in the size of the victim and attacker address spaces. When the value of i is fixed and clear from the context, we will simply omit i from the notation.

Quantifying Information Leakage. To capture the amount of leakage available to the attacker (in bits), we start from the intuition that fully associative caches with a random replacement policy leak the least amount of information among all cache designs that share cache lines between security domains, i.e., without consideration of partitioned caches. We argue that this assumption is reasonable, since fully associative caches with uniformly-random replacement only leak whether an address a was evicted or not, and do not reveal any information about which address i accessed by the victim caused the eviction of a . To evaluate leakage of new cache designs, we thus measure the REE as the statistical distance (in bits) of the eviction behavior of the tested cache design from the eviction behavior of ideal fully associative caches with random replacement.

Computing Relative Eviction Entropy. Our strategy for computing a cache design’s REE is as follows. First, we allocate a chunk of memory in the adversary’s address space, typically a small multiple of the cache size. We denote the set of cache line addresses within that adversary’s memory as $a \in [0, \dots, N-1]$. Second, for a single victim access to some fixed address i , we estimate the eviction probability $p_e(a)$ for each cache line $a \in [0, \dots, N-1]$ in the adversary’s memory, using our implementation of the cache design under test. The distribution $p_e(a)$ will reflect the cache’s placement policy: e.g., if the single victim access can evict every adversary address a , as in a fully associative cache with random replacement, $p_e(a)$ will be uniform among all adversary addresses a . If the single victim access can only evict adversary addresses a mapping to the same cache set in a set-associative cache, $p_e(a)$ will be uniform among those addresses a mapping to the same set as the victim address i and zero otherwise. The reference eviction distribution of a fully associative cache with random replacement is set to $p_u(a) = 1/N$ for all addresses a , reflecting that every adversary address is equally likely to get evicted. Finally, we compute the REE as the statistical distance in bits between the eviction probability distributions $p_e(a)$ and $p_u(a)$ using the Kullback-Leibler (KL) divergence to measure,

$$D_{KL}(p_e||p_u) = \sum_{a \in [0, \dots, N-1]} p_e(a) \log_2 \frac{p_e(a)}{p_u(a)}. \quad (1)$$

Note that the KL divergence does not fulfill the requirements of a metric and is asymmetric. Nevertheless, $D_{KL}(p_e||p_u)$ describes the relative entropy of $p_e(a)$ with respect to $p_u(a)$ and is a measure of the information lost if $p_u(a)$ was used to approximate $p_e(a)$. Mapped to cache side channels, the KL divergence thus nicely characterizes the leakage of a cache design with an eviction probability distribution $p_e(a)$ relative to the distribution $p_u(a)$ in a fully associative cache design.

Sampling $p_e(a)$. As $p_e(a)$ is generally unknown, we sample $p_e(a)$ and use the plug-in estimator [89] for the KL divergence to estimate the REE: we simply count the number of evictions for the attacker’s cache lines when the victim repeatedly accesses a fixed, randomly chosen address. More specifically, we first fill the cache by randomly accessing cache line addresses from the memory chunk corresponding to the attacker’s security domain. To keep track of self-evictions and hence the attacker’s lines that are actually cached, we utilize our cache model’s capability to return which cache line is evicted with each access, as described in Section 4.1. We note that this is an over-approximation of the attacker’s capabilities, as on real systems this translates to an attacker who can perfectly monitor cache evictions and accurately determine address collisions in the cache. Once the cache is entirely filled with the attacker’s data, we access a fixed secret address from the victim’s security domain, forcing an eviction of one of the attacker’s addresses. We then increment the eviction counter for the attacker address that is being reported as evicted from the cache. We repeat this sampling step multiple times and finally divide the per-address eviction counts by the total number of observed evictions, thereby obtaining $p_e(a)$. The repeated sampling procedure reduces the error of the sampled eviction probabilities proportional to $\sqrt{r}$, where r is the number of samples collected.

Definition

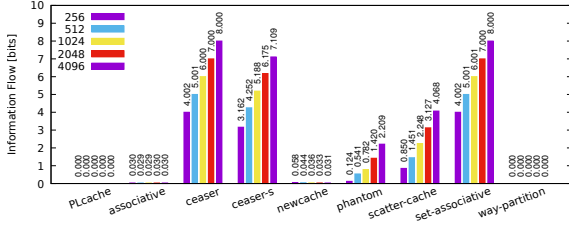


Figure 3: REE across cache designs with random replacement. All but NewCache and the fully associative cache use 16 ways.

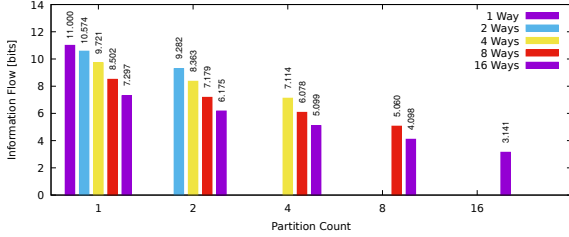


Figure 4: REE for CEASER-S with 2048 lines depending on ways and partitions.

Relative Eviction Entropy is evaluated as follows:

- (1) Select a victim address v and initialize a pool of memory $\mathcal{P}$.
- (2) Sample $p_e(a) \forall a \in \mathcal{P}$, which is the probability that v evicts address a from the cache.
- (3) Compute the REE via Equation 1, where $p_u(a)$ denotes a uniform distribution over all addresses $a \in \mathcal{P}$.

Evaluation Results. Figure 3 depicts the information leakage in the analyzed cache designs for various cache sizes and using random replacement. While the partitioned cache designs exhibit zero leakage, the leakages for CEASER and set-associative caches is the number of sets, i.e., $\log_2(\#sets)$ bits, thereby confirming the validity of our results. Next, we attribute the slightly above-zero leakages in NewCache and the fully associative cache to statistical noise. Note that CEASER-S and ScatterCache (with 2 and 16 partitions, respectively) show considerably lower leakage than standard set-associative caches. Moreover, as PhantomCache is looking up 8 sets, i.e., 128 lines, in parallel, PhantomCache stands out with significantly lower leakage per access than other designs, but also hurts chip area and power consumption.

Figure 4 analyzes the leakage in skewed caches like CEASER-S depending on way and partition count. Figure 4 clearly shows that increasing the number of ways and partitions effectively reduces leakage, with the difference between the best and worst configuration being 8 bits per access.

Supporting More General Cache Designs. We note that our Relative Eviction Entropy method can be computed in linear time, allowing us to evaluate different cache designs within minutes. However, we do assume some properties of the replacement policy of the cache being tested, namely that every line in the considered cache design exhibits the same leakage behavior, which in turn is independent from the specific address accessed by the victim. We rely

on this assumption in our procedure for sampling $p_e(a)$, evaluating the eviction distribution using only a single fixed address accessed by the victim. We argue that this assumption is natural and holds for most cache designs, including all the caches considered in this paper, as typical replacement policies do not differentiate between cache line addresses. While a single access does not reflect practical attack scenarios, it gives strong insight into the theoretical leakage caused by the caches’s structural mapping of addresses to cache lines. However, while the REE is a highly efficient tool to approximate leakage, we recognize that its underlying assumptions may be limiting its use in various corner cases, e.g., when replacement decisions are based on the actual address.

To better understand the practical exploitability of leakage determined via the REE, we conduct application-specific tests using cryptographic routines later in Section 5.3. However, note that the REE metric can be easily adapted to other cases as well, by simply testing multiple victim addresses and reporting the range of the occurring leakage as a function of victim’s address.

5.2 Eviction-Set Creation

To perform contention-based cache attacks, attackers first construct suitable eviction sets, i.e., minimal sets of addresses in their own address space that collide with the victim’s accesses of interest. Due to its perceived importance, multiple cache designs aim at randomizing the cache to prevent efficient eviction-set creation and thus contention-based attacks.

Definition

Eviction-set creation is evaluated as follows:

- (1) Select an eviction-set construction algorithm $\mathcal{A}$, an address a , a target eviction set size T , and number of repetitions R
- (2) Select a pool $\mathcal{P}$ of candidate addresses.
- (3) Run $\mathcal{A}$ on pool $\mathcal{P}$ to find an eviction set $\mathcal{E} \subset \mathcal{P}$ for address a until the target size $|\mathcal{E}| = T$ is reached and repeat R times.
- (4) Evaluate minimum/maximum/median/average for R samples of key metrics, e.g., number of attacker accesses and final set size.

Constructing Eviction Sets on Randomized Caches. Previous works proposed a range of methods for finding eviction sets in randomized caches. Taking a top-down approach, the Single Holdout Method (SHM) and the Group Elimination Method (GEM) [58, 74] both start from a large set of attacker addresses that evicts a certain victim address and then shrink this conflict set to a minimal eviction set by trying to remove (groups of) addresses while continuously verifying that the cache conflict remains. Taking a bottom-up approach, the Prime+Prune+Probe (PPP) method [55] pre-fills the cache with a set of candidate addresses, and subsequently triggers the victim access of interest. PPP then tests for cache misses in its candidate set, thereby locating conflicting addresses. Note that all of these approaches allow for optimizations specific to the cache replacement strategy in use.

Evaluating Difficulty of Eviction Set Construction. As protecting against eviction set construction is a major design goal for randomized caches, CacheFX allows to evaluate the effectiveness of SHM, GEM, and PPP on a candidate cache design. In particular, CacheFX quantifies the number of memory accesses required by an attacker, the number of conflicting addresses found, and the

success rate of using the found addresses for evicting the victim address. These figures eventually allow to configure cache re-keying intervals, e.g., for CEASER and CEASER-S. To set a level playing field and support an equal comparison across cache designs, we use the same implementations of eviction-set construction techniques for all evaluated cache designs. We intentionally avoided cache-specific optimizations, opting for comparable results rather than for optimal strategies. Specifically, all of our implementations iterate until they find (or shrink a conflict set to) the minimum number of addresses required for an eviction set, or until a predefined maximum iteration count is reached. The latter is a necessity to perform bulk testing as some algorithms do not terminate for every cache design. We leave the question of identifying optimal strategies and evaluating them to future work.

Measurement Setup. To measure the success rate, we set up a clean cache environment 1000 times and count the number of successful evictions of the cached victim address given the found eviction set. We extracted the cache hit/miss statistics to evaluate the number of attacker accesses needed for eviction-set creation. We determine the number of true conflicts in the eviction set by testing every found address for a collision with the victim address in the cache. While this is not directly possible on real systems, CacheFX provides this feature to assess how well each algorithm works for every cache design.

In our experiments, we used random replacement, 2048 lines and 16 ways where applicable, i.e., except for NewCache and the fully associative cache, which only have one set. We operated CEASER-S with 2 partitions, NewCache with $k = 2$, and PhantomCache with 8 parallel set lookups. We set up the algorithms to look for as many addresses as there are cache ways. For PhantomCache, however, we require 8x the number of ways, because it can place lines in 8 different sets.

Evaluating the Number of Memory Accesses for Eviction Set Construction. Figure 5 shows the number of memory operations done by SHM, GEM, and PPP for different cache designs. As L1 cache accesses take about five CPU cycles, these results give an indication about the execution time of each technique when used against a specific cache design.

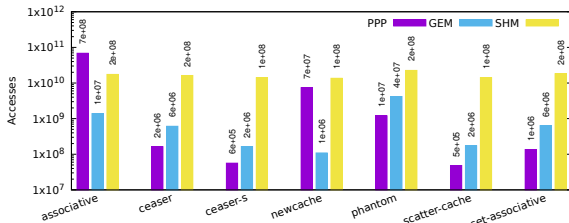


Figure 5: Number of memory accesses required by eviction-set building techniques for different 2048-line caches.

As the figure shows, the number of memory accesses for SHM is the highest, and in the same order of magnitude for all designs. In contrast, the complexity of PPP scales with the eviction set size, e.g., PPP is two orders of magnitude faster for ScatterCache than for NewCache. PPP also tends to be more efficient for skewed caches, as it is 3x faster for CEASER-S than for CEASER. The performance of GEM is mostly in between PPP and SHM, but tends to be faster than

PPP in the case of large eviction sets (e.g., for NewCache). Figure 6 gives further performance figures for CEASER-S and shows the linear increase in complexity with the number of cache lines.

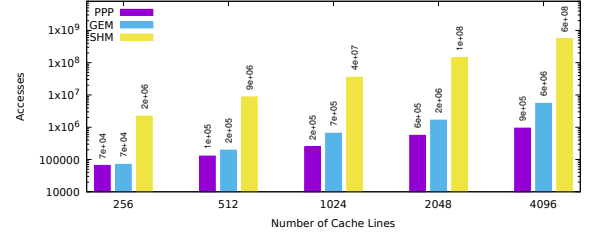


Figure 6: Number of memory accesses required by eviction-set building techniques for CEASER-S depending on cache size.

Evaluating Eviction Coverage. Different eviction set construction techniques can also produce eviction sets of different quality. Figure 7 thus shows the number of addresses in the found eviction sets that truly conflict with the victim address: PPP works best for all of the tested cache designs, producing eviction sets where all of its addresses truly conflict with the victim address. In contrast, SHM and GEM are less reliable, producing eviction sets where many of the addresses do not conflict with the victim address. The main reason for this is that SHM and GEM are highly susceptible to noise, which stems from both random replacement and cache skewing.

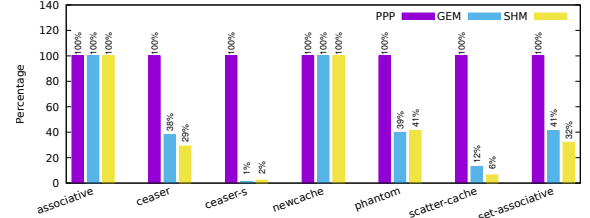


Figure 7: Percentage of addresses in the constructed eviction sets that conflict with the victim's address, using different eviction-set construction techniques and 2048-line caches.

To verify this, Figure 8 shows the constructed eviction sets' sizes for SHM, GEM and PPP. Except for NewCache and fully associative caches, both SHM and GEM stop shrinking the conflict set before it becomes minimal, which results in eviction sets where many addresses do not conflict with the victim address. This effect is particularly strong for the skewed cache designs CEASER-S and ScatterCache. Moreover, SHM and GEM also fail on PhantomCache, where both algorithms terminate with 10x as many addresses as needed. Finally, for NewCache and fully associative caches every address is equally suitable for an eviction set, which automatically results in 100% of the addresses conflicting with the victim address.

Evaluating Eviction Success Rate. We also evaluate the constructed eviction sets for their ability to effectively evict the victim address of interest. As Figure 9 shows, the eviction sets found by all three eviction set construction techniques perform equally well for CEASER, NewCache, set- and fully associative caches. For CEASER-S and ScatterCache, PPP yields better eviction success rates than

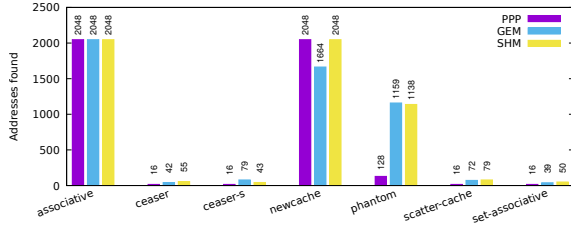


Figure 8: Eviction set sizes found by eviction-set building techniques for different 2048-line caches.

SHM and GEM, because PPP is generally more accurate (cf. Figure 7). For PhantomCache, however, GEM and SHM yielded better eviction rates as the found eviction set makes up roughly 50% of the cache. As skewed caches exhibit a significantly smaller probability of successful eviction (e.g., 2-4% for ScatterCache), eviction sets might be chosen larger to obtain high eviction probabilities and and Prime+Probe observability.

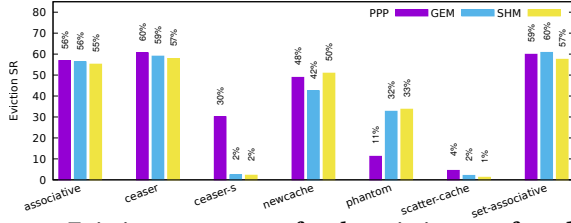


Figure 9: Eviction success rate for the eviction sets found for different 2048-line caches.

Obtaining a Specific Eviction Probability. To learn how many addresses would be needed to yield a certain eviction probability α , we start with an empty eviction set and successively add conflicting addresses until the eviction probability reaches α . Figure 10 presents the results of this routine for $\alpha = 90\%$, across different caches and replacement policies. It shows that LRU and Tree-PLRU allow for smaller eviction sets than Bit-PLRU and random replacement. In addition, skewing significantly increases the number of conflicting addresses needed, e.g., ScatterCache requires 10x more addresses than CEASER with equal sets and ways.

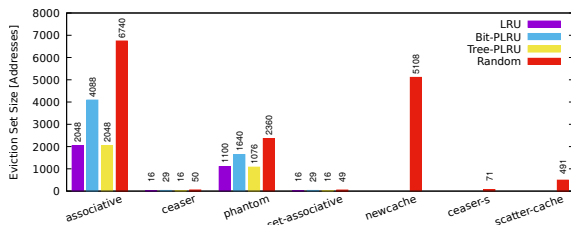


Figure 10: Eviction set size for 2048-line caches and 90% eviction probability.

5.3 Eviction-Set Attack

This section focuses on measuring the security offered by cache designs when performing attacks on cryptographic implementations.

To that aim, we simulate victims that use a cryptographic algorithm while the attacker tries to learn enough information to distinguish between two keys used by the victim. We use two cryptographic algorithms, each representing a different type of cache attack.

The AES Victim. Our AES victim is based on code from OpenSSL, which uses a set of tables, called T-tables, implemented as arrays. The attack focuses on the first four accesses made to the first T-table during the encryption. The two keys are selected such that, when encrypting some *vulnerable* plaintexts with the first key, all of these accesses fall in the first cache line of the T-table. Conversely, when encrypting vulnerable plaintexts with the second key, each of the four accesses falls in a different cache line. Finally, to further facilitate the attack, we allow the attacker to choose as many random vulnerable plaintexts as required for the attack.

In a more realistic scenario, the attacker can guess the characteristics of the vulnerable plaintexts. Specifically, the attacker can fix the first byte of the plaintext, and test every combination of plaintext values for the other three bytes that affect access to the first T-table. For each such combination, the attacker then performs the attack. If any of the combinations show statistical difference between the keys, the attack succeeds. With T-tables that span 16 cache lines, there are 16 possible values for each of these bytes. Thus, allowing to select vulnerable plaintexts represents a constant factor of $16^3 = 4096$ improvement in attack complexity.

Modular Exponentiation Victim Our second victim implements modular exponentiation, a core operation in multiple public-key schemes, e.g., RSA. Our modular exponentiation victim gets a 2048-bit base b , a 2048-bit modulus m , and a 32-bit exponent e . The victim then uses the square-and-multiply algorithm [26] to calculate $b^e \bmod m$. The square-and-multiply algorithm maintains an accumulator a that is initialized to 1. For each bit of the exponent e , the algorithm squares a , and if the bit is set the algorithm also multiplies a by b , reducing a modulo m as necessary. Thus, the multiplication code is only executed when the exponent bit is 1, and the effect on the cache is that when the bit is 1, more cache lines are accessed.

The keys are selected so that the value of a bit at a specific index (7 in our tests) of the exponent is 0 in the first key and 1 in the second. The other bits of each exponent are randomly chosen. We simulate an attacker that runs concurrently with the victim. The attacker can manipulate the cache whenever the victim finishes processing an exponent bit to distinguish between the number of cache lines accessed depending on the exponent bit.

Attacker Setup In the attack setup phase, the attacker is provided with an eviction set that evicts a monitored victim cache line with a probability 90%. We construct this eviction set by successively adding conflicting addresses to an initially empty set as outlined in Section 5.2. See there for an analysis of eviction-set construction.

Attacker Procedure. The attack proceeds as a sequence of rounds. In each round, the attacker asks the victim to encrypt a plaintext with the two selected keys, randomizing the order of using the keys in each round to avoid cache effects that depend on the order of the use of keys. Before each encryption, the attacker accesses the eviction set three times to prime the cache. After each encryption, the attacker accesses the eviction set, counting the number of cache misses during these accesses. Finally, the attacker calculates the average number of cache misses for each key, and stops when

achieving a 95% confidence that the averages differ, or when hitting a predefined number of rounds. (10^3 for the modular exponentiation and 10^5 for AES.) To overcome the case where the eviction set and the victim all fit in the cache, the attacker accesses some arbitrary memory when no cache evictions are observed during a round.

Selecting Cache Designs for Evaluation. We perform the attack on a sample of the cache designs considered in this paper. First, we do not test partitioned caches, because these do not leak information as there is no resource contention between the attacker and the victim. Secondly, to ensure that results are comparable, we limit our experiments to a cache size of 256 lines. Where applicable, we vary the associativity, testing all powers of two between 1 and 16. For each configuration, we run the attack 1,000 times and report the median of the number of encryptions required for distinguishing the keys. We use the median rather than the mean because in some cases the distribution has a long tail, skewing the mean towards a small number of cases where many encryptions are required.

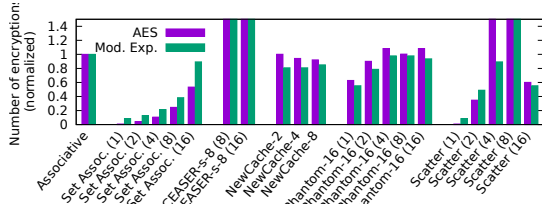


Figure 11: Eviction-set attack: Number of encryptions required to break AES and modular exponentiation with random replacement. Designs which show behavior similar to set associative caches, have been omitted from the figure. (Normalized to a random-replacement associative cache.)

Observing Key Leakage. Figure 11 shows the median number of encryptions required for the attacks when using random replacement. We normalize the results to the number of encryptions required for the fully associative cache. (10,590 and 94 for AES and modular exponentiation, respectively.) For brevity, we omit the results of CEASER, CEASER-S with one partition, and PhantomCache with one set lookup, all of which do not seem to offer any advantage over a set-associative cache with the same associativity.

The figure shows that all NewCache variants and PhantomCache with 16 set lookups are mostly equivalent to the fully associative cache. CEASER-S with 8 partitions provides a stronger protection: the majority of the AES attacks on CEASER-S with 8 ways and of the modular exponentiation attacks with 16 ways were not successful.

The results with ScatterCache are a mixed bag. When the associativity is four or eight, the design provides a good protection, equivalent or surpassing the fully associative cache. (In particular, the AES attack fails in most cases on an 8-way cache.) However, the protection is lower for the other cases.

5.4 Cache-Occupancy Attack

We now turn our attention to an emerging cache attack strategy that ignores spatial information and instead only utilizes the victim’s overall cache usage [44, 63, 65]. To measure resistance against so called *cache-occupancy attacks*, we use the same cryptographic victims as in Section 5.3. The attacker is still tasked with distinguishing between two keys, but instead of using an eviction set

targeting a specific cache line, the attacker uses a cache-size buffer and counts the number of cache misses when scanning the buffer. (A different sized buffer may also work [64], but requires further investigation.) Most other aspects of the attack are the same as in our eviction-set attack. We do not, however, handle failed eviction because using a cache-size buffer ensures contention on the cache.

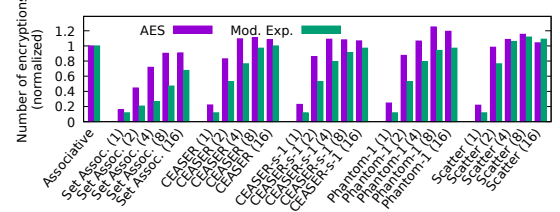


Figure 12: Occupancy attack: Number of encryptions required to break AES and modular exponentiation with random replacement. Designs which show behavior similar to fully associative caches, have been omitted from the figure. (Normalized to a random-replacement associative cache.)

Observing Key Leakage. Figure 12 shows the median number of encryptions required for the cache occupancy attacks when using a random replacement strategy. As in Figure 11, we normalize the results to the number of encryptions required for the fully associative cache. Similar to the eviction-set attack, NewCache, CEASER-S with 8 partitions, and PhantomCache with 16 set lookup achieve a protection similar to that of fully associative cache. (We omitted these three from the figure for brevity.) Most other configurations achieve a protection level which is significantly better than set-associative caches, in particular for the attack on AES.

Due to normalization, Figure ?? do not show that occupancy attacks on the fully associative cache require significantly less encryptions than eviction-set attacks. (5664 and 68 for AES and modular exponentiation, compared to 10590 and 94.) The cause is that the eviction set algorithm targets 90% eviction rate, which for fully associative caches leads to eviction sets that are larger than the cache-sized buffer used in the occupancy attack and thus more self evictions.

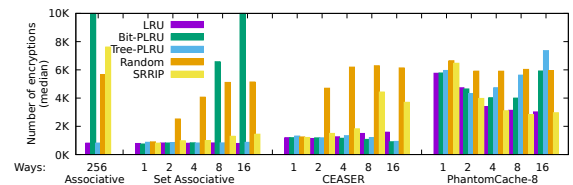


Figure 13: Median number of encryptions required to break AES. Fully associative and 16-way set associative caches are not fully represented, requiring 16,984 and 22,116 encryptions for Bit-PLRU, respectively.

Comparing Different Replacement Algorithms. Figure 13 shows the effect of changing the replacement policy on the attack complexity. As the figure demonstrates, in most cases, caches with a random replacement policy offer significantly better protection than those with deterministic replacement.

For deterministic replacement policies, we observe that CEASER only provides marginal benefit over set-associative caches, whereas PhantomCache provides a significantly better protection than other cache designs. We believe that the reason is that PhantomCache is inherently non-deterministic, hence, even with deterministic replacement algorithms, PhantomCache can reduce the correlation between the victim's access and the attacker's observation.

Attacks on deterministic cache designs that use bit-based pseudo-LRU replacement exhibit an anomaly that increases the number of encryptions required for statistical confidence. The cause is that the algorithm experiences some rare cases where a single cache miss causes cascading evictions of the eviction set. These rare cases increase the variance of the number of evictions observed, and with it the number of samples required. Modifying the attack to ignore outliers will eliminate these rare cases and significantly improve the attack. Hence, the results do not indicate that bit-based pseudo-LRU is more secure than random replacement.

6 CONCLUSION

This work fills a gap in the practice of evaluating cache designs for security. It presents CacheFX, a flexible framework that supports multiple metrics. We experiment with three different, albeit related, metrics, to evaluate and compare multiple secure cache designs.

We observe that all of the non-partitioned caches leak information and note that the leak is sufficient to implement cryptographic attacks. However, partitioned caches are likely not practical for many use cases. Moreover, we show that a single metric may fail to capture all of the intricacies. Thus, the main conclusion of this work is that there is no “best” cache design. Instead, we believe that caches need to be designed for the anticipated use cases.

The flexibility of CacheFX allows us to also compare attack strategies against existing caches. In particular, we show that the Prime+Prune+Probe approach for eviction set construction achieves more precise results than the Single Holdout and the Group Elimination Methods. Moreover, we show that for caches with low randomization, constructing an eviction set is a good strategy for cryptographic attacks. However, in highly random designs the cache-occupancy attack presents a more efficient strategy. Hence we recommend that secure cache designers consider the attack.

ACKNOWLEDGMENTS

This research was supported by the Air Force Office of Scientific Research (AFOSR) under award number FA9550-20-1-0425; an ARC Discovery Early Career Researcher Award DE200101577; an ARC Discovery Project number DP210102670; CSIRO's Data61; the Blavatnik ICRC at Tel-Aviv University; the Defense Advanced Research Projects Agency (DARPA) under contract HR00112390029, the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under Germany's Excellence Strategy - EXC 2092 CASA - 390781972; the National Science Foundation under grant CNS-1954712. the Phoenix HPC service at the University of Adelaide; and gifts by Qualcomm and Intel.

The views expressed in this paper are those of the authors and do not reflect the official policy or position of the U.S. Government or any of its agencies.

The strictly-enforced page limit has contributed not one, not two, but three whole appendices to this paper.

REFERENCES

- [1] Onur Acıçmez, Çetin Kaya Koç, and Jean-Pierre Seifert. 2007. Predicting Secret Keys Via Branch Prediction. In *CT-RSA*. 225–242.
- [2] Alejandro Cabrera Aldaya, Cesar Pereida García, Luis Manuel Alvarez Tapia, and Billy Bob Brumley. 2019. Cache-Timing Attacks on RSA Key Generation. *TCHES* 2019, 4 (2019), 213–242.
- [3] Thomas Bourgeat, Jules Drean, Yuheng Yang, Lillian Tsai, Joel Emer, and Mengjia Yan. 2020. CaSA: End-to-end Quantitative Security Analysis of Randomly Mapped Caches. In *MICRO 2020*. 1110–1123.
- [4] Ferdinand Brasser, Urs Müller, Alexandra Dmitrienko, Kari Kostinen, Srdjan Capkun, and Ahmad-Reza Sadeghi. 2017. Software Grand Exposure: SGX Cache Attacks Are Practical. In *WOOT*.
- [5] Billy Bob Brumley. 2015. Cache Storage Attacks. In *CT-RSA*. 22–34.
- [6] Billy Bob Brumley and Risto M. Hakala. 2009. Cache-Timing Template Attacks. In *ASIACRYPT*. 667–684.
- [7] Chithanok Chuengsatiansup, Daniel Genkin, Yuval Yarom, and Zhiyuan Zhang. 2022. Side-Channeling the Kalyna Key Expansion. In *CT-RSA*. 272–296.
- [8] David Cock, Qian Ge, Toby C. Murray, and Gernot Heiser. 2014. The Last Mile: An Empirical Study of Timing Channels on seL4. In *CCS*. 570–581.
- [9] Fergus Dall, Gabrielle De Micheli, Thomas Eisenbarth, Daniel Genkin, Nadia Heninger, Ahmad Moghimi, and Yuval Yarom. 2018. CacheQuote: Efficiently Recovering Long-term Secrets of SGX EPID via Cache Attacks. *TCHES* 2018, 2 (2018), 171–191.
- [10] John Demme, Robert Martin, Adam Waksman, and Simha Sethumadhavan. 2012. Side-channel vulnerability factor: A metric for measuring information leakage. In *ISCA*. 106–117.
- [11] John Demme, Robert Martin, Adam Waksman, and Simha Sethumadhavan. 2013. A Quantitative, Experimental Approach to Measuring Processor Side-Channel Security. *IEEE Micro* 33, 3 (2013), 68–77.
- [12] John Demme and Simha Sethumadhavan. 2014. Side-Channel Vulnerability Metrics: SVF vs. CSV. In *Workshop on Duplicating, Deconstructing and Debunking*.
- [13] Shuwen Deng, Nikolay Matyunin, Wenjie Xiong, Stefan Katzenbeisser, and Jakub Szefer. 2021. Evaluation of Cache Attacks on Arm Processors and Secure Caches. arxiv 2106.14054.
- [14] Shuwen Deng, Wenjie Xiong, and Jakub Szefer. 2019. Analysis of Secure Caches Using a Three-Step Model for Timing-Based Attacks. *J. Hardware and Systems Security* 3, 4 (2019), 397–425.
- [15] Shuwen Deng, Wenjie Xiong, and Jakub Szefer. 2020. A Benchmark Suite for Evaluating Caches' Vulnerability to Timing Attacks. In *ASPLOS*. ACM, 683–697.
- [16] Craig Disselkoen, David Kohlbrenner, Leo Porter, and Dean M. Tullsen. 2017. Prime+Abort: A Timer-Free High-Precision L3 Cache Attack using Intel TSX. In *USENIX Sec*. 51–67.
- [17] Leonid Domnitsker, Nael B. Abu-Ghazaleh, and Dmitry Ponomarev. 2010. A Predictive Model for Cache-Based Side Channels in Multicore and Multithreaded Microprocessors. In *MMM-ACNS*. 70–85.
- [18] Leonid Domnitsker, Aamer Jaleel, Jason Loew, Nael B. Abu-Ghazaleh, and Dmitry Ponomarev. 2012. Non-monopolizable caches: Low-complexity mitigation of cache side channel attacks. *TACO* 8, 4 (2012), 35:1–35:21.
- [19] Goran Doychev, Dominik Feld, Boris Köpf, Laurent Mauborgne, and Jan Reineke. 2013. CacheAudit: A Tool for the Static Analysis of Cache Side Channels. In *USENIX Sec*. 431–446.
- [20] Dmitry Evtushkin, Dmitry V. Ponomarev, and Nael B. Abu-Ghazaleh. 2016. Jump over ASLR: Attacking branch predictors to bypass ASLR. In *MICRO*. 40:1–40:13.
- [21] Qian Ge, Yuval Yarom, Tom Chothia, and Gernot Heiser. 2019. Time Protection: The Missing OS Abstraction. In *EuroSys*. 1:1–1:17.
- [22] Qian Ge, Yuval Yarom, David Cock, and Gernot Heiser. 2018. A survey of microarchitectural timing attacks and countermeasures on contemporary hardware. *JCEC* 8, 1 (2018), 1–27.
- [23] Qian Ge, Yuval Yarom, and Gernot Heiser. 2018. No Security Without Time Protection: We Need a New Hardware-Software Contract. In *APSys*. 1:1–1:9.
- [24] Daniel Genkin, Romain Poussier, Rui Qi Sim, Yuval Yarom, and Yuanjing Zhao. 2020. Cache vs. Key-Dependency: Side Channeling an Implementation of Pilsung. *TCHES* 2020, 1 (2020), 231–255.
- [25] Tara Ghasempouri, Jaan Raik, Kolin Paul, Cezar Reinbrecht, Said Hamdioui, and Mottaqiallah Taouil. 2020. A Security Verification Template to Assess Cache Architecture Vulnerabilities. In *DDECS*. 1–6.
- [26] Daniel M. Gordon. 1998. A Survey of Fast Exponentiation Methods. *Journal of Algorithms* 27, 1 (1998), 129–146.
- [27] Ben Gras, Kaveh Razavi, Herbert Bos, and Cristiano Giuffrida. 2018. Translation Leak-aside Buffer: Defeating Cache Side-channel Protections with TLB Attacks. In *USENIX Sec*. 955–972.
- [28] Ben Gras, Kaveh Razavi, Erik Bosman, Herbert Bos, and Cristiano Giuffrida. 2017. ASLR on the Line: Practical Cache Attacks on the MMU. In *NDSS*.

- [29] Marc Green, Leandro Rodrigues Lima, Andreas Zankl, Gorka Irazoqui, Johann Heyszl, and Thomas Eisenbarth. 2017. AutoLock: Why Cache Attacks on ARM Are Harder Than You Think. In *USENIX Sec.* 1075–1091.
- [30] Leon Groot Bruinderink, Andreas Hülsing, Tanja Lange, and Yuval Yarom. 2016. Flush, Gauss, and Reload - A Cache Attack on the BLISS Lattice-Based Signature Scheme. In *CHES*. 323–345.
- [31] Daniel Gruss, Clémentine Maurice, Klaus Wagner, and Stefan Mangard. 2016. Flush+Flush: A Fast and Stealthy Cache Attack. In *DIMVA*. 279–299.
- [32] Daniel Gruss, Raphael Spreitzer, and Stefan Mangard. 2015. Cache Template Attacks: Automating Attacks on Inclusive Last-Level Caches. In *USENIX Sec.* 897–912.
- [33] David Gullasch, Endre Bangerter, and Stephan Krenn. 2011. Cache Games - Bringing Access-Based Cache Attacks on AES to Practice. In *IEEE SP*. 490–505.
- [34] Zecheng He and Ruby B. Lee. 2017. How secure is your cache against side-channel attacks?. In *MICRO*. 341–353.
- [35] Ralf Hund, Carsten Willems, and Thorsten Holz. 2013. Practical Timing Side Channel Attacks Against Kernel Space ASLR. In *NDSS*.
- [36] Mehmet Sinan Inci, Berk Gülmehzoglul, Gorka Irazoqui, Thomas Eisenbarth, and Berk Sunar. 2016. Cache Attacks Enable Bulk Key Recovery on the Cloud. In *CHES*. 368–388.
- [37] Gorka Irazoqui Apecechea, Thomas Eisenbarth, and Berk Sunar. 2015. S\$A: A Shared Cache Attack That Works across Cores and Defies VM Sandboxing - and Its Application to AES. In *IEEE SP*. 591–604.
- [38] Himanshi Jain, D. Anthony Balaraju, and Chester Rebeiro. 2019. Spy Cartel: Parallelizing Evict+Time-Based Cache Attacks on Last-Level Caches. *J. Hardw. Syst. Secur.* 3, 2 (2019), 147–163.
- [39] Daniel Katzman, William Kosasih, Chitchanok Chuengsatiansup, Eyal Ronen, and Yuval Yarom. 2023. The Gates of Time: Improving Cache Attacks with Transient Execution. In *USENIX Security*.
- [40] Mehmet Kayaalp, Khaled N. Khasawneh, Hodjat Asghari Esfeden, Jesse Elwell, Nael B. Abu-Ghazaleh, Dmitry V. Ponomarev, and Aamer Jaleel. 2017. RIC: Relaxed Inclusion Caches for Mitigating LLC Side-Channel Attacks. In *DAC*. 7:1–7:6.
- [41] Boris Köpf, Laurent Mauborgne, and Martin Ochoa. 2012. Automatic Quantification of Cache Side-Channels. In *Computer Aided Verification*.
- [42] Fangfei Liu, Hao Wu, Kenneth Mai, and Ruby B. Lee. 2016. Newcache: Secure Cache Architecture Thwarting Cache Side-Channel Attacks. *IEEE Micro* 36, 5 (2016), 8–16.
- [43] Fangfei Liu, Yuval Yarom, Qian Ge, Gernot Heiser, and Ruby B. Lee. 2015. Last-Level Cache Side-Channel Attacks are Practical. In *IEEE SP*. 605–622.
- [44] Clémentine Maurice, Christoph Neumann, Olivier Heen, and Aurélien Francillon. 2015. C5: Cross-Cores Cache Covert Channel. In *DIMVA*. 46–64.
- [45] Clémentine Maurice, Nicolas Le Scouarnec, Christoph Neumann, Olivier Heen, and Aurélien Francillon. 2015. Reverse Engineering Intel Last-Level Cache Complex Addressing Using Performance Counters. In *RAID*. 48–65.
- [46] Clémentine Maurice, Manuel Weber, Michael Schwarz, Lukas Giner, Daniel Gruss, Carlo Alberto Boano, Stefan Mangard, and Kay Römer. 2017. Hello from the Other Side: SSH over Robust Cache Covert Channels in the Cloud. In *NDSS*.
- [47] John D. McCalpin. 2021. *Mapping Addresses to L3/CHA Slices in Intel Processors*. Technical Report TR-2021-03. UT Austin.
- [48] Ahmad Moghimi, Gorka Irazoqui, and Thomas Eisenbarth. 2017. CacheZoom: How SGX Amplifies the Power of Cache Attacks. In *CHES*. 69–90.
- [49] Yossef Oren, Vasileios P. Kemerlis, Simha Sethumadhavan, and Angelos D. Keromytis. 2015. The Spy in the Sandbox: Practical Cache Attacks in JavaScript and their Implications. In *ACM CCS*. 1406–1418.
- [50] Dag Arne Osvik, Adi Shamir, and Eran Tromer. 2006. Cache Attacks and Countermeasures: The Case of AES. In *CT-RSA*. 1–20.
- [51] Dan Page. 2002. Theoretical Use of Cache Memory as a Cryptanalytic Side-Channel. ePrint 2002/169.
- [52] Biswabandan Panda. 2019. Fooling the Sense of Cross-Core Last-Level Cache Eviction Based Attacker by Prefetching Common Sense. In *PACT*. 138–150.
- [53] Colin Percival. 2005. Cache Missing for Fun and Profit. In *BSDCan 2005*. <http://css.csail.mit.edu/6.858/2014/readings/ht-cache.pdf>
- [54] Peter Pessl, Leon Groot Bruinderink, and Yuval Yarom. 2017. To BLISS-B or not to be: Attacking strongSwan's Implementation of Post-Quantum Signatures. In *ACM CCS*. 1843–1855.
- [55] Antoon Purnal, Lukas Giner, Daniel Gruss, and Ingrid Verbauwhede. 2021. Systematic Analysis of Randomization-based Protected Cache Architectures. In *IEEE SP*.
- [56] Antoon Purnal, Furkan Turan, and Ingrid Verbauwhede. 2021. Prime+Scope: Overcoming the Observer Effect for High-Precision Cache Contention Attacks. In *CCS*. 2906–2920.
- [57] Moinuddin K. Qureshi. 2018. CEASER: Mitigating Conflict-Based Cache Attacks via Encrypted-Address and Remapping. In *MICRO*. 775–787.
- [58] Moinuddin K. Qureshi. 2019. New attacks and defense for encrypted-address cache. In *ISCA*. 360–371.
- [59] Thomas Ristenpart, Eran Tromer, Hovav Shacham, and Stefan Savage. 2009. Hey, you, get off of my cloud: exploring information leakage in third-party compute clouds. In *ACM CCS*. 199–212.
- [60] Eyal Ronen, Robert Gillham, Daniel Genkin, Adi Shamir, David Wong, and Yuval Yarom. 2019. The 9 Lives of Bleichenbacher's CAT: New Cache ATtacks on TLS Implementations. In *IEEE SP*. 435–452.
- [61] Martin Schwarzl, Erik Kraft, and Daniel Gruss. 2022. Layered Binary Templating: Efficient Detection of Compiler- and Linker-introduced Leakage. arxiv 2208.02093.
- [62] André Seznec. 1993. A Case for Two-Way Skewed-Associative Caches. In *ISCA*. 169–178.
- [63] Anatoly Shusterman, Ayush Agarwal, Sioli O'Connell, Daniel Genkin, Yossi Oren, and Yuval Yarom. 2021. Prime+Probe 1, JavaScript 0: Overcoming Browser-based Side-Channel Defenses. In *USENIX Sec.* 2863–2880.
- [64] Anatoly Shusterman, Zohar Avraham, Eliezer Croitoru, Yarden Haskal, Lachlan Kang, Dvir Levi, Yosef Meltser, Prateek Mittal, Yossi Oren, and Yuval Yarom. 2021. Website Fingerprinting Through the Cache Occupancy Channel and its Real World Practicality. *TDSC* 18, 5 (2021), 2042–2060.
- [65] Anatoly Shusterman, Lachlan Kang, Yarden Haskal, Yosef Meltser, Prateek Mittal, Yossi Oren, and Yuval Yarom. 2019. Robust Website Fingerprinting Through the Cache Occupancy Channel. In *USENIX Sec.* 639–656.
- [66] Florian Sieck, Sebastian Berndt, Jan Wichelmann, and Thomas Eisenbarth. 2021. Util.: Lookup: Exploiting Key Decoding in Cryptographic Libraries. In *CCS. ACM*, 2456–2473.
- [67] Qinhan Tan, Zhihua Zeng, Kai Bu, and Kui Ren. 2020. PhantomCache: Obfuscating Cache Conflicts with Localized Randomization. In *NDSS*.
- [68] Yukiyasu Tsunoo, Teruo Saito, Tomoyasu Suzuki, Maki Shigeri, and Hiroshi Miyauchi. 2003. Cryptanalysis of DES Implemented on Computers with Cache. In *CHES*. 62–76.
- [69] Yukiyasu Tsunoo, Etsuko Tsujihara, Kazuhiko Minematsu, and Hiroshi Miyauchi. 2002. Cryptanalysis of Block Ciphers Implemented on Computers with Cache. In *ISITIA*.
- [70] Leif Uhsadel, Andy Georges, and Ingrid Verbauwhede. 2008. Exploiting Hardware Performance Counters. In *FDTC*. 59–67.
- [71] Stephan van Schaik, Marina Minkin, Andrew Kwong, Daniel Genkin, and Yuval Yarom. 2021. CacheOut: Leaking Data on Intel CPUs via Cache Evictions. In *IEEE SP*. 339–354.
- [72] Stephan van Schaik, Kaveh Razavi, Ben Gras, Herbert Bos, and Cristiano Giuffrida. 2017. RevAnC: A Framework for Reverse Engineering Hardware Page Table Caches. In *EUROSEC*. 3:1–3:6.
- [73] Pepe Vila, Andreas Abel, Marco Guarnieri, Boris Köpf, and Jan Reineke. 2020. Flushgeist: Cache Leaks from Beyond the Flush. arxiv 2005.13853.
- [74] Pepe Vila, Boris Köpf, and José F. Morales. 2019. Theory and Practice of Finding Eviction Sets. In *IEEE SP*. 39–54.
- [75] Limin Wang, Ziyuan Zhu, Zhanpeng Wang, and Dan Meng. 2019. Colored Petri Net Based Cache Side Channel Vulnerability Evaluation. *IEEE Access* 7 (2019), 169825–169843.
- [76] Zhenghong Wang and Ruby B. Lee. 2007. New cache designs for thwarting software cache-based side channel attacks. In *ISCA*. 494–505.
- [77] Mario Werner, Thomas Unterluggauer, Lukas Giner, Michael Schwarz, Daniel Gruss, and Stefan Mangard. 2019. ScatterCache: Thwarting Cache Attacks via Cache Set Randomization. In *USENIX Sec.* 675–692.
- [78] Wenjie Xiong and Jakub Szefer. 2020. Leaking Information Through Cache LRU States. In *HPCA*. 139–152.
- [79] Mengjia Yan, Christopher W. Fletcher, and Josep Torrellas. 2020. Cache Telepathy: Leveraging Shared Resource Attacks to Learn DNN Architectures. In *USENIX Sec.*
- [80] Mengjia Yan, Bhargava Gopireddy, Thomas Shull, and Josep Torrellas. 2017. Secure Hierarchy-Aware Cache Replacement Policy (SHARP): Defending Against Cache-Based Side Channel Attacks. In *ISCA*. 347–360.
- [81] Mengjia Yan, Read Sprabery, Bhargava Gopireddy, Christopher W. Fletcher, Roy H. Campbell, and Josep Torrellas. 2019. Attack Directories, Not Caches: Side Channel Attacks in a Non-Inclusive World. In *IEEE SP*. 888–904.
- [82] Mengjia Yan, Jen-Yang Wen, Christopher W. Fletcher, and Josep Torrellas. 2019. SecDir: a secure directory to defeat directory side-channel attacks. In *ISCA*. 332–345.
- [83] Yuval Yarom and Katrina Falkner. 2014. Flush+Reload: A High Resolution, Low Noise, L3 Cache Side-Channel Attack. In *USENIX Sec.* 719–732.
- [84] Yuval Yarom, Qian Ge, Fangfei Liu, Ruby B. Lee, and Gernot Heiser. 2015. Mapping the Intel Last-Level Cache.
- [85] Yuval Yarom, Daniel Genkin, and Nadia Heninger. 2017. CacheBleed: a Timing Attack on OpenSSL Constant-Time RSA. *JCEN* 7, 2 (2017), 99–112.
- [86] Tianwei Zhang and Ruby B. Lee. 2014. New models of cache architectures characterizing information leakage from cache side channels. In *ACSAC*, Charles N. Payne Jr., Adam Hahn, Kevin R. B. Butler, and Micah Sherr (Eds.), 96–105.
- [87] Tianwei Zhang, Fangfei Liu, Si Chen, and Ruby B. Lee. 2013. Side channel vulnerability metrics: the promise and the pitfalls. In *HASP@ISCA*. 2.
- [88] Tianwei Zhang, Yinqian Zhang, and Ruby B. Lee. 2018. Analyzing Cache Side Channels Using Deep Neural Networks. In *ACSAC*. 174–186.

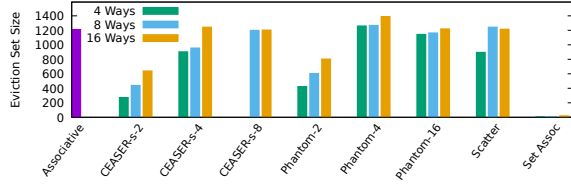


Figure 14: Optimal eviction set sizes for 1024-lines caches

[89] Zhiyi Zhang and Michael Grabchak. 2014. Nonparametric Estimation of Kullback-Leibler Divergence. *Neural Comput.* 26, 11 (2014), 2570–2593.

A OPTIMAL EVICTION-SET SIZE

In Section 5.2 we evaluate eviction sets based on the probability of evicting a victim cache line from the cache. However, as discussed in Section 5.4, larger eviction sets can result in lower attack efficiency, apparently due to self evictions. Specifically, increasing the size of the eviction set increases the probability of cache conflicts between elements of the eviction set. These self evictions introduce measurement noise that increases the variance in the measurements and consequently the number of samples the attacker needs to observe to distinguish the keys. As a secondary effect, larger eviction sets require more memory accesses for both the prime and the probe steps of the attack, reducing attack efficiency.

As a final example of the flexibility of CacheFX, we now use it to find the eviction-set size that allows for the most efficient attack. Specifically, we experiment with various cache designs, all with size 1024, our AES victim, and our eviction-set attacker. We vary the eviction-set size between 1 and 2048, and measure the median number of encryptions required for distinguishing the keys. Figure 14 reports the eviction-set size that allows the attack with the minimal median. As we can see, a lower associativity allows for smaller eviction-set sizes. However, when the associativity grows to 16, in most cache designs the best eviction-set size is similar to that of a fully associative cache, indicating that occupancy-based attacks are as effective as eviction-set attacks.

B THREATS TO VALIDITY AND LIMITATIONS

At the moment, CacheFX does not support the evaluation of cache hierarchies. Consequently, designs that rely on the hierarchy for defense are outside the scope for this work. Moreover, evaluations using CacheFX currently assume a noise-free scenario, which provides a conservative security estimate as the absence of noise is the best case for attackers. However, practical cache attacks also face systematic and random noise stemming from other system activity. To assess the impact of noise and cache hierarchies to security, we aim to add such models to CacheFX in the future.

For our cryptographic attack evaluations, CacheFX models a strong, synchronized attacker and an artificial victim that computes (and leaks) upon the attacker’s request. As for noise, this is a very strong attack model that allows to obtain a lower bound for security. While a simple model like CacheFX cannot capture all complexities involved in real-world attacks, we consider modeling more realistic scenarios as future work.

Another aspect of secure caches is their performance. CacheFX currently does not support the evaluation of cache performance, but

its mechanisms to collect data about memory accesses, cache hits and misses allows easy future extension of CacheFX to measure, e.g., the cache hit rate, based on memory traces of relevant workloads.

CacheFX simplifies cache models to efficiently analyze the security of caches against contention-based attacks. As a result, CacheFX does not lend itself to model cache-based attacks that relate to speculative execution [39] or other microarchitectural structures, e.g., cache ports [85] and fill buffers [71]. Note that CacheFX allows to model Flush+Reload attacks, but as Flush+Reload is well understood we do not expect new insights from doing so.

C RELATED WORKS

We now review past works that evaluate the security of caches against side channel attacks.

Formal Cache Model and Theoretical Analysis. This line of research [19, 41] tries to formally model the state change of the cache and extend the program execution semantics to include cache state changes by leveraging prior work on formal analysis of cache miss rates. Eventually they can estimate the number of reachable cache states and give an upper bound on the leakage in terms of channel capacity, for a given program under analysis. Similarly, [25] models caches and cache attacks as automata to verify cache security using model checking. Due to the restrictions of formal methods, these works are limited to simple cache models (e.g. set-associative cache with LRU replacement) and can only give a very loose upper bound of leakage. Hence, they are not suitable for comparing the security of various complex secure cache designs. In contrast, CacheFX empirically evaluates a number of metrics to quantify side-channel leakage in software cache models and evaluates the exploitability of cache leakage for programs such as cryptographic algorithms.

Metrics for Empirical Quantification of Information Leakage.

Another line of research introduces metrics to empirically evaluate the security of cache designs and implementations, such as by using mutual information and min-leakage [8], by using a linear correlation coefficient between oracle traces and the attacker’s observations [10–12, 87], by measuring the accuracy of deep learning models trained to learn the relationship between victim accesses and the attacker’s cache observations [88], or by modeling and statistically analyzing cache side channels using communication theory [3]. CacheFX as well tries to empirically characterize the leakage of cache designs. However, as we point out, a single metric is insufficient to entirely capture cache security. Moreover, none of these works looks at cache occupancy channels or tries to assess security by using well-studied cryptographic targets.

Modeling of Cache Side Channel Attacks. Some works tried to model caches and cache attacks such as to detect and quantify cache leakage. For instance, Zhang and Lee [86] model the cache as a finite state machine to identify interference and determine the mutual information. He and Lee [34] model cache attacks as a Probabilistic Information Flow Graph (PIFG) to derive for each cache and attack an overall probability of success. Wang et al. [75] derive a risk score from modeling attacks using Petri nets and calculating the success probabilities of concrete attacks. Deng et al. [13–15] model cache attacks as a series of three consecutive read/invalidation steps, identify vulnerable three-step patterns using a simulator,

and use the model for evaluating the security of the caches in multiple Arm devices. In addition, their work introduces a Cache Timing Vulnerability Score (CTVS) from running vulnerable patterns on real machines. While these prior works greatly improve the understanding of cache attacks, many are based on simple cache models.

CacheFX thus takes another step forward and automatically evaluates arbitrary software models of cache designs w.r.t. to a number of different metrics and attack complexity to provide a comprehensive security report.