

Morton Filters: Faster, Space-Efficient Cuckoo Filters via Biasing, Compression, and Decoupled Logical Sparsity

Alex D. Breslow
Advanced Micro Devices, Inc.
AMD Research
2485 Augustine Drive
Santa Clara, CA 95054
Alex.Breslow@amd.com

Nuwan S. Jayasena
Advanced Micro Devices, Inc.
AMD Research
2485 Augustine Drive
Santa Clara, CA 95054
Nuwan.Jayasena@amd.com

ABSTRACT

Approximate set membership data structures (ASMDSSs) are ubiquitous in computing. They trade a tunable, often small, error rate (ϵ) for large space savings. The canonical ASMDSS is the Bloom filter, which supports lookups and insertions but not deletions in its simplest form. Cuckoo filters (CFs), a recently proposed class of ASMDSSs, add deletion support and often use fewer bits per item for equal ϵ .

This work introduces the Morton filter (MF), a novel ASMDSS that introduces several key improvements to CFs. Like CFs, MFs support lookups, insertions, and deletions, but improve their respective throughputs by $1.3\times$ to $2.5\times$, $0.9\times$ to $15.5\times$, and $1.3\times$ to $1.6\times$. MFs achieve these improvements by (1) introducing a compressed format that permits a logically sparse filter to be stored compactly in memory, (2) leveraging succinct embedded metadata to prune unnecessary memory accesses, and (3) heavily biasing insertions to use a single hash function. With these optimizations, lookups, insertions, and deletions often only require accessing a single hardware cache line from the filter. These improvements are not at a loss in space efficiency, as MFs typically use comparable to slightly less space than CFs for the same ϵ .

PVLDB Reference Format:

Alex D. Breslow and Nuwan S. Jayasena. Morton Filters: Faster, Space-Efficient Cuckoo Filters via Biasing, Compression, and Decoupled Logical Sparsity. *PVLDB*, 11(9): 1041-1055, 2018. DOI: <https://doi.org/10.14778/3213880.3213884>

1. INTRODUCTION

As time has progressed, systems have added many more levels to the memory hierarchy. In today's enterprise servers, it is not uncommon to have three to four levels of hardware cache, a vast pool of DRAM, several SSDs, and a pool of disks. With each successive level of the hierarchy, latency and bandwidth typically increase by one or more orders of magnitude. To avoid accessing a slower medium unnecessarily, many applications make use of approximate set membership data structures (ASMDSSs). An ASMDSS like a set data

structure answers set membership queries (i.e., is an item e an element of the set S ?). However, unlike a set, which always reports with certitude whether e is in S , an ASMDSS is able to report false positives (i.e., falsely state that e is in S) with a worst-case expected error rate of $0 \leq \epsilon \leq 1$. An ASMDSS does not report false negatives [29]: if an ASMDSS reports e not in S , then it certainly is not. A core benefit of an ASMDSS is that its error rate ϵ is typically independent of the size of the data items that are encoded, so an ASMDSS can often reside one or two levels higher in the memory hierarchy than the slow medium to which it filters requests.

The most common ASMDSS is the Bloom filter [7], which in its simplest form supports insertions and a `LIKELY_CONTAINS` lookup primitive. Deletions and counting occurrences of an item are supported via a number of different Bloom filter variants, albeit with an increase in the storage cost (where a $2\times$ to $4\times$ increase is not uncommon [10, 30]). Bloom filters have been used in data storage systems such as Google's BigTable [17], distributed analytics platforms such as Apache Impala [41], bioinformatics applications such as the counting of k-mers during DNA sequencing [49], diverse networking applications [14], and more. One of the pitfalls of the Bloom filter is that its simplest version exhibits poor locality of reference, and more cache-friendly blocked variants are typically less space efficient [14].

Consequently, a number of other filters have been proposed, of which two of the most practical are the quotient filter [6, 59] and cuckoo filter [29]. Both the quotient filter and the cuckoo filter differ from the Bloom filter in that they store fingerprints, short hashes that each typically have a one-to-one correspondence with an item e that is encoded by the filter as belonging to a set S . Both cuckoo filters and quotient filters support deletions and when filled to high load factors (i.e., a 95% full filter) use less space than an equivalent Bloom filter when the desired false positive rate is less than about 1% to 3%, the usual case for a wide array of applications.

In this work, we focus on the cuckoo filter (CF) and present a novel variant, the *Morton filter* (MF).¹ Like a CF, MFs' storage is organized as a linear array of buckets, with each bucket containing a fixed number of slots that can each store a single fingerprint. Fingerprints are mapped to the table by emplacing the fingerprint in one of two *candidate buckets*, whose indices are independently determined by two hash functions (H_1 and H_2) that each operate on the key and output a different bucket index. Provided that one candidate bucket has spare capacity, the insertion trivially succeeds. Conflicts where both candidate buckets are full are

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Articles from this volume were invited to present their results at The 44th International Conference on Very Large Data Bases, August 2018, Rio de Janeiro, Brazil.

Proceedings of the VLDB Endowment, Vol. 11, No. 9

Copyright 2018 VLDB Endowment 2150-8097/18/05... \$ 10.00.

DOI: <https://doi.org/10.14778/3213880.3213884>

¹Named after a certain elephant's half-bird baby [23].

resolved via cuckoo hashing [58], a hashing technique that triggers a recursive chain of evictions.

Despite these similarities, MFs differ in several key ways. In contrast to CFs, MFs are able to more heavily bias insertions in favor of H_1 . This biasing makes subsequent retrieval of fingerprints require fewer hardware cache accesses because most of the time, the fingerprint is found in the first bucket. For *negative lookups* (queries to keys that were never inserted into the table), the filter employs an Overflow Tracking Array (OTA), a simple bit vector that tracks when fingerprints cannot be placed using H_1 . By checking the OTA, most negative lookups only require accessing a single bucket, even when the filter is heavily loaded. This biasing and tracking means that regardless of the type of lookup (positive, false positive, or negative), typically only one bucket needs to be accessed. When buckets are resident in a cache line, most often only 1 cache access is needed per probe of the filter and at most 2, a savings of close to 50%.

In addition to biasing, MFs decouple their logical representation from how their data are stored in virtual memory. They logically underload the filter and apply a simple compression algorithm that replaces storage of space-hungry empty slots with a series of *fullness counters* that track the load of each logical bucket. With fullness counters, reads and updates to an MF happen in-situ without explicit need for materialization. This zero-compression makes logically underloading the filter (1) space efficient because many mostly empty buckets can be packed into a single cache block and (2) high performance because accesses occur directly on the compressed representation and only on occupied slots.

With logically underloaded buckets, most insertions only require accessing a single cache line from the MF. For example, with an MF block of 60 slots, of which 59 are occupied, an insertion to any bucket is likely to directly succeed. However, with those same 60 slots arranged as 15 4-slot buckets in a CF, that same insertion operation only has a $1/15$ chance of direct success (i.e., where no accesses to additional buckets are necessary). Consequently, MFs much more efficiently utilize scarce cache and memory bandwidth and sustain high insertion throughput at much heavier loads than a CF (e.g., $3\times$ to $15\times$ higher for load factors exceeding 0.75, a filter that is more than 75% full).

Further, an MF typically performs many fewer fingerprint comparisons than a CF, with fewer than one fingerprint comparison per lookup not uncommon, even when the filter is heavily loaded. Instead, many lookups can be fully or partially resolved simply by examining one or two fullness counters and a bit in the OTA.

Due to this compression, sparsity, and biasing, MFs attain improved throughput, space usage, and flexibility. With fewer comparisons and a reduced number of cache accesses, MFs boost lookup, deletion, and insertion throughputs, respectively, by as much as $2.5\times$, $1.6\times$, and $15.5\times$ over a stock CF. Similarly, these traits permit using shorter fingerprints because false positives are integrally tied to the number of fingerprint comparisons. Consequently, the space overhead of the fullness counters and OTA can largely be hidden, and space per item can often be reduced by approximately 0.5 to 1.0 bits over a CF with the same ϵ .

Our contributions are as follows:

1. We present the design of and empirically evaluate the Morton filter, a novel ASMDS that uses compression, sparsity, and biasing to improve throughput without sacrificing on space efficiency or flexibility. MFs improve performance over CFs by making accesses to fewer cache lines during filter reads and updates.

2. We greatly ameliorate the insertion throughput performance collapse problem for fingerprint-based ASMDSS at high loads ($>100\times$ for a CF) by decoupling the logical representation of the filter from how it is stored.
3. We present a fast algorithm for computing reductions on fullness counters that is the key to the high performance and which can be applied in other contexts.
4. We present a hashing mechanism that reduces TLB misses, row buffer misses, and page faults and does away with requiring the total buckets be a power of 2.
5. We intend to release a C++ MF implementation to the wider research community, which we have tested on AMD and Intel X86 server processors on Linux.

2. CUCKOO FILTERS

In this section, we describe the cuckoo filter (CF) [29], the ancestor of the MF.

2.1 Baseline Design

CFs are hash sets that store *fingerprints*, where each fingerprint is computed by using a hash function H_F , which takes as input a key representing an item in the set and maps it to a fixed-width hash. The filter is structured as a 2D matrix, where rows correspond to fixed-width associative units known as *buckets* and cells within a row to *slots*, with each slot capable of storing a single fingerprint. Prior work typically uses 4-slot buckets [29].

To map each key's fingerprint to the filter and to largely resolve collisions, Fan et al. encode a key's membership in the set by storing its fingerprint in one of two *candidate buckets*. The key's two candidate buckets are independently computed using two hash functions H_1 and H_2 . H_1 takes the key as input and produces the index of one candidate bucket, and H_2 operates on the same key and produces the index of the other candidate bucket [29].

2.2 Insertions

On insertions, provided that at least one of the eight slots is empty across the two candidate buckets, the operation completes by storing the fingerprint in one of the empty locations. If no slot is free, *cuckoo hashing* [58] is employed. Cuckoo hashing picks a fingerprint within one of the two buckets, evicts that fingerprint and stores the new fingerprint in the newly vacated slot. The evicted fingerprint is then rehashed to its alternate bucket using its alternate hash function. To compute the alternate hash function simply by using the bucket index and fingerprint as inputs, they define a new hash function H' , which takes the fingerprint and its current bucket as input and returns the other candidate bucket. So, if the fingerprint is currently found in the first candidate bucket given by $H_1(key)$, H' yields the alternate candidate given by $H_2(key)$ and vice versa. Provided that the alternate candidate has a free slot, the evicted key is emplaced, and the operation succeeds. If no such slot exists, the initial evicted fingerprint takes its place, and the operation continues until a free slot is found.

Two example insertions are shown in Figure 1. The first is for key K_x , which succeeds because H_2 maps its fingerprint x to a Bucket 0, which has a vacant slot. The second is for key K_y , where it initially fails to find a vacant slot in either candidate bucket (6 and 4), and therefore uses cuckoo hashing to displace a chain of fingerprints beginning with 1011 in Bucket 6 and ending with 1101 in Bucket 2 moving to one of the free slots in Bucket 1. Note that in practice, the series of displacements may occur in reverse order in an optimized implementation to avoid storing displaced fingerprints at each step.

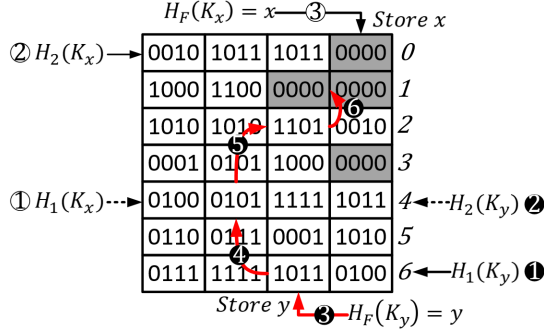


Figure 1: Insertion of two different keys K_x and K_y into the filter by storing their respective fingerprints x and y . Empty slots are shown in gray. K_x 's insertion only involves accessing its two candidate buckets (4 and 0) since 0 has a free slot, but K_y 's candidates (6 and 4) are both full, so a series of fingerprints are displaced each to their alternate candidate to make an empty slot for y in bucket 6. The updated filter is shown in Figure 2.

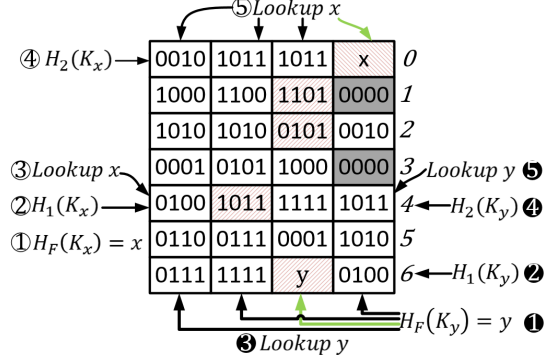


Figure 2: Lookup of two different keys K_x and K_y following the successful insertion of their respective fingerprints x and y in Figure 1. For K_y , steps 4 and 5 can optionally be skipped, since y is found in the first candidate bucket.

2.3 Lookups

On lookups, the algorithm computes H_1 and H_2 on the key to compute its candidate buckets, and H_F to determine its fingerprint. If the fingerprint appears in any of the eight slots across the two candidate buckets, the lookup returns *LIKELY_IN_SET*, else *NOT_IN_SET*. The certitude of *LIKELY_IN_SET* is subject to an expected error rate ϵ that is tunable by assigning an appropriate bit-width to each fingerprint. Increasing the width of each fingerprint by one bit roughly halves ϵ . It is worth noting that the actual incidence of false positives will interpolate between 0 (all queries are to items inserted in the filter) to roughly ϵ (none of the queried items were inserted in the filter) subject to whether lookups are to true members of the encoded set.

Figure 2 follows the insertion of keys K_x and K_y into the filter in Figure 1. Even though K_y triggered a series of displacements, since it is only allowed to insert y in one of its two candidate buckets, only Buckets 6 and 4 need to be searched. The same holds for any other queried key: at most two buckets need to be probed.

2.4 Modeling the Error Rate and Space Use

In this section, we present formulae for calculating the error rate of a cuckoo filter, its space usage per item, and show how the insights from this study can be leveraged to

Table 1: Glossary of Symbols

ϵ - false positive rate
S - slots per bucket
b - buckets searched per negative lookup
α - the load factor
f - fingerprint length in bits
I - bits per item in the filter

design an MF (see Section 3 for a high-level MF description). Table 1 provides a glossary of symbols.

A CF has an error rate ϵ which reports the expected ratio of false positives to total *LIKELY_CONTAINS* queries that are true negatives (i.e., no risk of a false positive on a true positive). To understand how ϵ is calculated given a filter, we first introduce several terms: S the slots per bucket, b the expected number of buckets that need to be searched per negative lookup, α the load factor, and f the bits per fingerprint. When comparing an f -bit fingerprint to a fingerprint stored in a bucket, the occurrence of *aliasing* (i.e., falsely matching on a fingerprint inserted by another key) is $1/2^f$ if all fingerprint values are equally probable. There are 2^f potential values and only one of those can alias. To compute the net probability of an alias, prior work by Fan et al. [29] observes that there are S slots per bucket, b is fixed at 2 (they always search both buckets), and therefore the rate of aliasing is $\epsilon = 1 - (1 - 1/2^f)^{bS}$, so the necessary f for a target ϵ is roughly $f = \log_2(bS/\epsilon)$, which for their parameters of $S = 4$ and $b = 2$ is $f = 3 + \log_2(1/\epsilon)$.

However, what this model discounts is the effect of α , that is, if one is careful and clearly marks empty slots (by reserving one of the 2^f states to encode an empty slot), then there is no way that empties can alias when performing a lookup. Marking empties changes the math slightly, to $\epsilon = 1 - (1 - 1/(2^f - 1))^{\alpha bS}$, which alters f to approximately

$$f = \log_2(\alpha bS/\epsilon) \quad (1)$$

for typical values of f (i.e., $f > 6$). For underloaded filters, it turns out that that extra α term is important because since $0 \leq \alpha \leq 1$, its logarithm is less than or equal to zero. For instance, filling the filter half full ($\alpha = 0.5$) means that α in the numerator decreases f 's required length by $\log_2(\alpha = 0.5) = 1$ bit for a target ϵ . Further, this effect is amplified in the savings in bits per item (shown in Equation 2). With the additional α and a fixed ϵ , $\alpha = 0.5$ would decrease the required bits per item by $\log_2(0.5)/0.5 = 2$ bits over Fan et al.'s more pessimistic model.

$$I = f/\alpha = \log_2(\alpha bS/\epsilon)/\alpha \quad (2)$$

However, because the α in the denominator of Equation 2 dwarfs the impact of the α in the numerator, prior work largely ignores this space savings opportunity. In particular, α needs to be close to 1 for a CF to be space-competitive with a Bloom filter [7], which largely negates the positive impact of the α in the numerator. To obtain a large value of α (e.g., > 0.95), there are several options for b and S , but practical constraints limit the viable options. For b , a high-performance implementation is limited to selecting 2. A value of $b = 1$ cannot yield a sufficiently high α even for large values of S . $b > 2$ is also undesirable because it results in additional memory traffic (e.g., $b = 3$ triggers an additional memory request per lookup). For S , larger values improve α but at the expense of a worsening error rate given a fixed f . With each additional fingerprint comparison, the likelihood of a false positive increases. In practice, $S = 4$ is the minimum value that permits an α that exceeds 0.95. Larger values of S could be used at the expense of increased bits per item. As we will see in the proceeding sections, our work gets around these limitations. For further analysis of

feasible parameters, we point the reader to Erlingsson et al.’s work on cuckoo hash tables (an analogous hash table rather than an approximate hash set) [26], which provides a concise table showing the trade-offs of b and S .

2.5 Bloom Filters and Relative Optimality

Given $b = 2$, $S = 4$, and $\alpha = 0.95$, we examine a CF’s relative space and performance optimality. CFs make very efficient use of space. Whereas a Bloom filter uses approximately $\log_2(1/\epsilon)/\ln(2) \approx 1.44\log_2(1/\epsilon)$ bits per item [14], these parameters for b , S , and α place the bits per item at about $3.08 + 1.05\log_2(1/\epsilon)$, clearly asymptotically better and not too far off from the information theoretic limit of $\log_2(1/\epsilon)$ (see Carter et al. [15]). Fan et al. show that the leading constant can be further improved via Bonomi et al.’s semi-sort optimization [11], which sorts fingerprints within a bucket by a fixed-width prefix and then replaces those prefixes with a code word. With 4-bit prefixes, four prefixes can be replaced with a 12-bit code word, a savings of one bit per fingerprint. That reduces the bits per item to $2.03 + 1.05\log_2(1/\epsilon)$, albeit with reduced performance from sorting, compression, and decompression (see Section 7).

On the performance front, an ideal filter only requires accessing a single cache line from the filter for each lookup, insertion, or deletion. For lookups, the CF is $2\times$ off of optimal since each invocation examines two buckets, which with high probability are in different cache lines. Deletions are better, since ASMDs only permit deletions to items in the filter (otherwise, false negatives are possible), the expected cost is typically 1 plus the fraction of items that are inserted using H_2 . Insertions are often the furthest from optimal. At a heavy load factor, it can take many cache accesses to insert a single item. In the proceeding sections, we will show how to largely get around these limitations and achieve lookups, insertions, and deletions that typically only access a single cache line from the filter. For a comparative analysis of MFs, which discusses the decoupling of the $\log_2(\alpha)$ term in the numerator from the α term in the denominator in Equation 2, see Section 5.

3. MORTON FILTERS

This section describes the MF and elaborates on the principal features that differentiate it from a cuckoo filter.

3.1 Optimizing for the Memory Hierarchy

The MF is a reengineered CF that is tuned to make more efficient use of cache and memory bandwidth. Today’s memory systems move data in coarse-grain units known as cache lines that are typically 64 to 128 bytes. On a load or store instruction, the entire cache line is fetched and brought up through the memory hierarchy to the L1 data cache. Subsequent accesses to words in the same cache line that occur in short sequence (known as *temporal locality* [5, 35]) are cheap because they likely *hit* in the high-bandwidth, low-latency L1 data cache. Typical ASMDs workloads are often cache- or memory-bound because they employ pseudorandom hash functions to set bits or fingerprints within the filter, which limits their data reuse and taxes the comparatively limited bandwidth of lower level caches (e.g., L3 or L4) and bandwidth to off-chip memory. In contrast to a CF, which optimizes for latency at the expense of performing two random memory accesses per lookup query, an MF probe performs a single cache access most of the time and at most two. In bandwidth-limited scenarios, these efficiency gains correspond to significant speedups (see Section 7). We point the interested reader to prior work by Ross [65], Polychroniou

and Ross [61], and Breslow et al. [13] for the related discussion of latency and bandwidth tradeoffs in hash tables.

3.2 Logical Interpretation

Like a cuckoo filter, the MF maintains a set of buckets and slots, fingerprints encoding keys are computed using H_F and mapped to one of two candidates using one of H_1 or H_2 . Collisions are resolved using a variant of cuckoo hashing (see Section 4.2 and Figure 7).

3.3 Compressed Structure: The Block Store

The MF stores its data in parameterizable units known as blocks. Blocks have a compressed storage format that stores both the fingerprints from a fixed number of buckets and accompanying metadata that permits recovering the MF’s logical interpretation while solely performing in-situ reads and updates to the block. Blocks are stored sequentially in a structure known as the *Block Store*. Block size within the Block Store is dictated by the physical block size of the storage medium for which the MF is optimized. If the filter resides entirely in cache and system memory, then block sizes that evenly divide a cache line are the natural choice (e.g., 256- or 512-bit blocks for 512-bit hardware cache lines). Similarly, block sizes that evenly divide an SSD block are logical choices when the filter is primarily SSD-resident.

3.4 Block Components

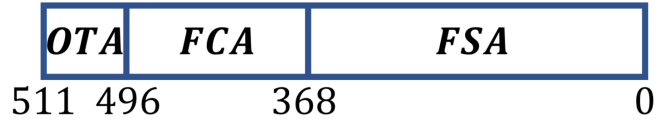


Figure 3: A sample block in an MF that is performance optimized for 512-bit cache lines. The block has a 46-slot FSA with 8-bit fingerprints, a 64-slot FCA with 2-bit fullness counters (64 3-slot buckets), and a 16-bit OTA with a single bit per slot.

Each MF block has three principal components (shown in Figure 3), which we detail below:

Fingerprint Storage Array (FSA) - The FSA is the array that stores the fingerprints from a block. Fingerprints from consecutive buckets within a block are stored one after another in compact, sequential order with no gaps. Empty slots within the FSA are entirely at the end of the buffer. An FSA typically has many fewer slots than the total *logical* slots across all buckets that it stores from the logical interpretation. For instance, in Figure 3, there are 46 slots for fingerprints in the FSA but a total of $64 * 3 = 192$ slots across the 64 buckets whose fingerprints it stores. Thus, the filter can be logically highly underloaded while allowing the FSA to be mostly full and accordingly conserve space.

Fullness Counter Array (FCA) - The FCA encodes the logical structure of the block by associating a *fullness counter* with each of its buckets that tracks how many slots are occupied by fingerprints. It enables in-situ reads and writes to the serialized buckets in the FSA without the need to materialize a full logical view of the associated block by summing the loads of the buckets prior to the bucket of interest to determine an offset in the FSA where reading should begin. Further, with an FCA, vacant slots in the logical interpretation no longer have to be stored in the FSA, and our implementation uses the FCA to completely skip comparisons to empty fingerprint slots.

Overflow Tracking Array (OTA) - The OTA in its simplest form is a bit vector that tracks overflows from the block by setting a bit every time a fingerprint overflows (see Section 3.8). By querying the OTA, queries determine whether

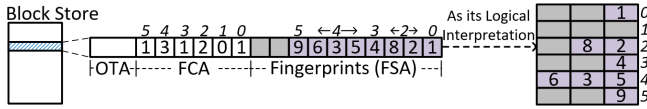


Figure 4: An MF’s Block Store and a sample block’s compressed format and logical interpretation, with corresponding buckets labeled 0 to 5. The FCA and FSA state dictates the logical interpretation of the block. Buckets and fingerprints are ordered right to left to be consistent with logical shift operations.

accessing a single bucket is sufficient for correctness or if both candidates need to be checked.

3.5 Accessing a Bucket

A sample block and its logical representation are shown in Figure 4. In the example, the least significant bit of the block is the farthest to the right. Fullness counters and fingerprints are labeled with the bucket to which they correspond. To show how in-situ lookups are possible, we consider the case where we want to examine the contents of bucket 4. Bucket 4 contains the fingerprints 5, 3, and 6. To determine where to read in the FSA, we can add the loads of buckets 0 to 3 to provide an offset. These are $1 + 0 + 2 + 1$ or 4, so the first fingerprint 5 appears at $FSA[4]$. We know to stop reading at $FSA[6]$ because $FCA[4] = 3$, and so we have already read all three fingerprints. See Section 4.1 for a more in-depth description of the algorithm with an accompanying figure.

Note that in the example, the logical interpretation of the block has 18 slots of which a mere 8 are occupied. In a CF, if the block were representative of the average load on the filter, then the bits per item would be $f/(8/18) \approx 2.25f$. However, in the MF, 8 of 10 FSA slots are occupied, so the block’s actual load is much higher, and the actual bits per item is $1.25f + FCA \text{ bits} + OTA \text{ bits}$, clearly asymptotically better. Thus, heavily logically underloaded MFs with densely filled FSAs conserve space while allowing inexpensive lookups and updates that are typically localized to a single block and thus a single cache line.

3.6 Primacy

In contrast to a CF, we differentiate between the two hash functions H_1 and H_2 . We call H_1 the *primary hash function* and for a given key say that a bucket is its *primary bucket* if its fingerprint would be stored there on an insertion using H_1 . We call H_2 the *secondary hash function* and for a given key say that a bucket is its *secondary bucket* if its fingerprint would be stored there on an insertion using H_2 . When inserting a fingerprint into an MF, we always first try to place it into its primary bucket and only fall back to the secondary function H_2 when that fails. By heavily biasing insertions in this way, most items in the MF can be found by examining a single bucket, and thus a single hardware cache line.

3.7 Filtering Requests to Secondary Buckets

For *negative lookups* (i.e., where the queried key never had its fingerprint inserted into the table), biasing still helps performance. The OTA tracks all overflows from a block. When H_2 gets used to map a key K ’s fingerprint to a bucket, we set a bit in the OTA corresponding to the block containing K ’s primary bucket. We select the bit’s index by computing a hash function H_{OTA} on K . On a subsequent query to a key K' that was never inserted into the filter but whose primary bucket is in the same block as K ’s, we compute H_{OTA} on K' . If we hash to an unset bit in the OTA, then the lookup only requires a single bucket access. A set bit requires accessing the other candidate bucket (likely two different cache

lines). Bits in the OTA that are previously set remain set on additional overflows that hash to the same bit.

3.8 Types of Overflows

At the time of filter initialization, all OTAs across all blocks begin zeroed out. When fingerprints are first inserted, they are inserted exclusively using H_1 and accordingly into their primary buckets. It is only after some time that one of two events will trigger the setting of a bit in the OTA. The first event is a *bucket overflow*. Bucket overflows occur when a key’s fingerprint maps to a bucket for which there is no spare logical capacity, that is, when its associated counter in the FCA has already hit its maximum value. The second event is a *block overflow*, which occurs when a key’s fingerprint is mapped to a bucket where its block has no spare FSA slots. In both cases, one fingerprint needs to be remapped from the block to make room for the new fingerprint. A bucket overflow requires the evicted fingerprint to come from the new fingerprint’s candidate bucket; however, when a block overflow occurs that is not also a bucket overflow, any fingerprint within the block’s FSA may be evicted. As it turns out, for most parameter values for the slots per bucket, the overwhelming vast majority of overflows are purely block overflows. With common fingerprint sizes of several to a few tens of bits, this affords tens of potential fingerprints to evict on any overflow and makes the filter particularly robust during insertions. See Section 4.2 for further detail.

3.9 Interplay Between Buckets and Blocks

The addition of the block abstraction is one of the defining features of the MF. By aggregating the loads across the many underloaded buckets that they store, blocks improve the space efficiency of the filter while permitting smaller, less heavily loaded buckets (e.g., a 3-slot bucket with fewer than 1 occupied slot). With small buckets that are mostly empty, most lookups require much fewer loads and comparisons and are thus cheaper. For example, an MF that employs the block parameters in Figure 3 requires fewer than 0.8 fingerprint comparisons per `LIKELY_CONTAINS` query even when 95% of the FSA slots are full, an improvement of more than 10× over a stock CF that checks 8 slots.

Further, small, underloaded buckets afford greater opportunity to batch work from multiple lookups and insertions of multiple items into a shared set of SIMD instructions [31] (see Section 4.1 for further discussion).

The large block size greatly benefits insertions. Because the logical interpretation of the filter is sparsely filled, bucket overflows are infrequent because most fullness counters never max out. As such, provided the block has at least one free slot, most insertions are likely to succeed on the first attempt. Thus, overwhelmingly most items are hashed with H_1 (>95% for the parameters in Figure 3 for FSA occupancies less than or equal to 0.95), so most insertions, deletions, and lookups only access a single cache line from the MF.

3.10 Hashing

The MF employs a different method for calculating H' and H_2 than a stock CF that reduces TLB misses and page faults. We show H_2 and H' below, where K is an arbitrary key, B is the buckets per block, β is the bucket index where K ’s fingerprint is placed, n is the total buckets, H is a hash function like MurmurHash [3], $map(x, n)$ maps a value x between 0 and $n - 1$ inclusive, and $H_F(K)$ is K ’s fingerprint:

$$H_1(K) = map(H(K), n)$$

$$H_2(K) = map(H_1(K) + (-1)^{H_1(K) \& 1} * offset(H_F(K)), n)$$

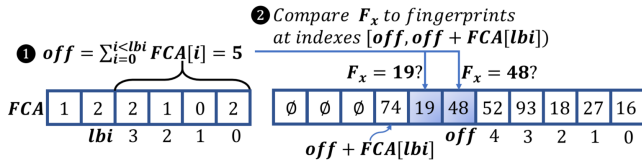


Figure 5: An example of checking for the presence of fingerprint F_x once the logical bucket index lbi within the block is known.

$$H'(\beta, H_F(K)) = \text{map}(\beta + (-1)^{\beta \& 1} * \text{offset}(H_F(K)), n)$$

$$\text{offset}(F_x) = [B + (F_x \% \text{OFF_RANGE})] \mid 1$$

This formulation logically partitions the filter into two halves: the even buckets and odd buckets. If a bucket is even, then we add an offset to the primary bucket. If it is odd, then we subtract that offset. Offsets are always odd (the 1 term) to enforce switching between partitions. By partitioning in this way, it makes it possible to compute $H'(K)$ and remap K 's fingerprint without knowing K itself. This property is also true of Fan et al.'s original scheme, which XORs $H_1(H_F(x))$ with the fingerprint's current bucket to determine its alternate candidate. However, their approach requires the number of buckets in the filter to be a power of two, which in the worst case increases memory use by almost 2 $\times$. Our approach does not make this assumption. Rather, we only mandate that the number of buckets be a multiple of two so that H' is invertible.

Our offset calculation has several nice properties. By adding B to the offset, it guarantees that for any key, its two candidate buckets fall in different blocks. This property ensures that rehashing a fingerprint always removes load from the originating block, which is crucial during block overflows. Further, the offset is at most $\pm(B + \text{OFF_RANGE})$. Thus, by tuning this value so that it is much less than the number of buckets in a physical memory page, we ensure that most pairs of candidate buckets fall within the same page of memory. This optimization improves both TLB and DRAM row buffer hit ratios, which are crucial for maximizing performance [8, 39, 83]. Further, we select an OFF_RANGE that is a power of two so that modulo operations can be done with a single logical AND. The $\text{map}(x, n)$ primitive is implemented two different ways that both get around performing an integer division. The first method [44] is given by $\text{map}(x, n) = (x * n) >> k$, where x is a k -bit integer that is uniformly random between 0 and $2^k - 1$. For the second method, because the offset is bounded, provided that it is smaller the total buckets in the MF, we do the following:

if $x \geq 0$ && $x \leq n - 1$, then $\text{map}(x, n) = x$
else if $x < 0$, then $\text{map}(x, n) = x + n$
else, then $\text{map}(x, n) = x - n$

4. ALGORITHMS

In this section, we describe the MF's core algorithms.

4.1 Lookups

This section describes how to determine the presence of a key K_x 's fingerprint F_x in an MF. A simplified algorithm is presented in Algorithm 1. We first both compute the primary hash function H_1 on K_x to determine the global bucket index for its primary bucket (call it $glbi1$) for F_x . Dividing $glbi1$ by the buckets per block B yields the block index. Computing $\text{mod}(glbi1, B)$ produces the block-local bucket index $lbi1$. From here, we check for the presence of x in its bucket using *table_read_and_compare*, which performs an in-situ check for F_x on K_x 's block. No materialization to a logical representation of the block is necessary.

Algorithm 1 Algorithm for LIKELY_CONTAINS function

```

1: function LIKELY_CONTAINS( $MF, K_x$ )
2:    $F_x = H_F(K_x)$ 
3:    $glbi1 = H_1(K_x)$ 
4:    $block1 = MF.BlockStore[glbi1/B]$ 
5:    $lbi1 = \text{mod}(glbi1, B)$ 
6:    $match = \text{table\_read\_and\_cmp}(block1, lbi1, F_x)$ 
7:   if ( $match$  or  $OTA\_bit\_is\_unset(block1, lbi1)$ ) then
8:     return match
9:   else
10:     $glbi2 = H_2(K_x)$ 
11:     $block2 = MF.BlockStore[glbi2/B]$ 
12:     $lbi2 = \text{mod}(glbi2, B)$ 
13:    return  $\text{table\_read\_and\_cmp}(block2, lbi2, F_x)$ 

```

Algorithm 2 Algorithm for INSERT function

```

1: function INSERT( $MF, K_x$ )
2:    $F_x = H_F(K_x)$ 
3:    $glbi1 = H_1(K_x)$ 
4:    $block1 = MF.BlockStore[glbi1/B]$ 
5:    $lbi1 = \text{mod}(glbi1, B)$ 
6:    $success = \text{table\_simple\_store}(block1, lbi1, F_x)$ 
7:   if ( $success$ ) then
8:     return success
9:   else
10:     $\text{set\_OTA}(block1, lbi1)$ 
11:     $glbi2 = H_2(K_x)$ 
12:     $block2 = MF.BlockStore[glbi2/B]$ 
13:     $lbi2 = \text{mod}(glbi2, B)$ 
14:     $success = \text{table\_simple\_store}(block2, lbi2, F_x)$ 
15:   if ( $success$ ) then
16:     return success
17:   return  $\text{res\_conflict}(MF, block1, block2, lbi1, lbi2, F_x)$ 

```

Figure 5 shows *table_read_and_compare* in action. ① We first compute K_x 's bucket's offset in fingerprints from the start of its primary block. In the example, K_x 's lbi is 4, so we sum the loads of all buckets that appear before bucket 4 (0 through 3 inclusive), which yields an offset (off) of 5 fingerprints. ② Since we use zero indexing, that indicates that bucket 4's first fingerprint appears at index 5. Since the $FCA[lbi] = 2$, that means that bucket 4 has 2 fingerprints. Therefore, since we begin reading at index 5, we stop reading prior to index $off + FCA[lbi] = 5 + 2 = 7$ (index 6). If any of the fingerprints in the queried range (i.e., 19 or 48) match F_x , then we return true, else false.

Having returned to Algorithm 1, we then check if a match was successful or if $lbi1$ maps to an unset bit in the OTA. If either hold, then we return the result. Otherwise, we probe the bucket in the second block and return the result.

To achieve high performance with this algorithm, we modify it slightly to perform multiple lookups in a flight. All lookups first compute prior to the *if* statement, we gather those lookups for which the statement evaluated to false, and then perform the *else* statement for the subset where accessing the secondary bucket is necessary. This batching improves performance by permitting improved SIMD vectorization of the code and by reducing the number of branch mispredictions [69]. Batching is akin to the vectorized query processing employed in columnar databases [9] and loop tiling [47, 57, 78] in high-performance computing.

4.2 Insertions

Algorithm 2 shows the high-level algorithm for insertions. We first attempt to insert into the first candidate bucket at $glbi1$ (lines 2 through 6). The function *table_simple_store*

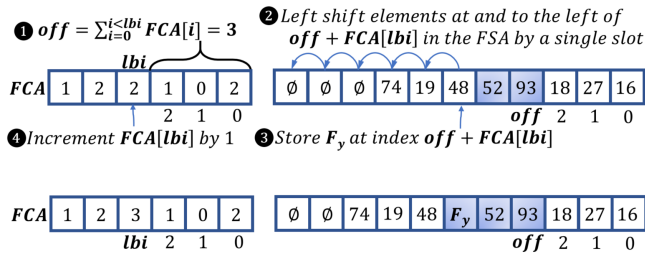


Figure 6: An example of inserting a fingerprint F_x into its logical bucket at index lbi within the block. The updated block is shown below. We leave out the OTA for clarity.

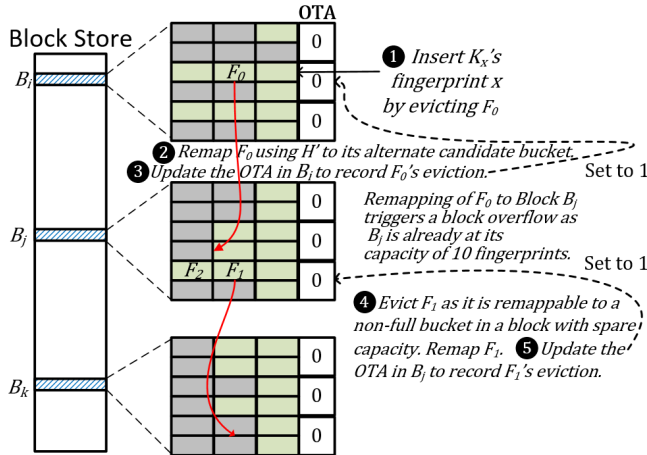


Figure 7: Insertion of a key K_x into an MF (visualized as the logical interpretation). This insertion is atypical as it requires two levels of evictions to resolve the conflict. In contrast, most insertions only need to update a single block and no cuckoo hashing is necessary, even when blocks are heavily loaded.

succeeds if both the candidate bucket and its block's FSA have spare capacity (i.e., no block nor bucket overflow). If *table_simple_store* fails, then the algorithm proceeds to the second candidate bucket (lines 10 through 14). Provided the block and candidate have sufficient capacity, the insertion succeeds. Otherwise, we proceed to the conflict resolution stage (line 17). In this stage, a series of cuckoo hashing displacements are made.

Figure 6 shows the block-local updates that occur during an insertion of a key K_y 's fingerprint F_y (the bulk of *table_simple_store*). ① We begin by computing the bucket offset within the FSA. In this case, K_y 's block-local bucket index lbi is 3, so we sum all fullness counters before index 3, which correspond to the loads in fingerprints of the 0th, 1st, and 2nd buckets in the block. ② Next, we shift all fingerprints to the left of the end of the bucket (at an offset of $off + FCA[lbi]$) to the left by one slot to vacate a slot for F_y . These operations are inexpensive because many elements are shifted via a single low-latency logical shift instruction, and because Block Store blocks are sized to evenly divide a cache line, only one cache line from the Block Store is accessed per block-level read or update. ③ F_y is then stored at $FSA[off + FCA[lbi]]$. ④ The final step is to increment the fullness counter at the bucket's logical index.

Figure 7 shows the core functionality of the function *res_conflict*, with the series of displacements that occur when inserting a sample key K_x 's fingerprint x into its primary bucket. ① In the example, K_x maps to a bucket that is full (a bucket overflow) within a block that has spare ca-

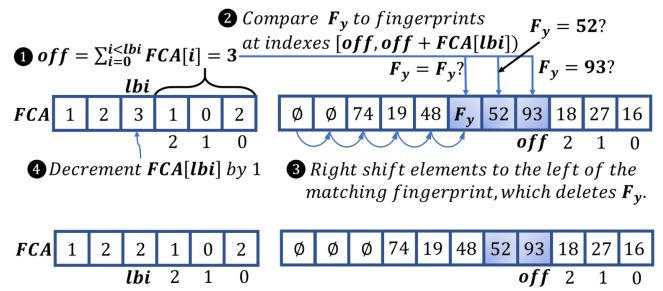


Figure 8: An example of checking for the presence of fingerprint F_y once the logical bucket index lbi within the block is known and deleting it on a match. The updated block is shown below. We leave out the OTA for clarity.

capacity. Cuckoo hashing evicts one of the fingerprints in K_x 's candidate (F_0). ② F_0 is remapped to its other candidate bucket, which is found in block B_j and ③ the OTA in block B_i is updated by setting the bit that H_{OTA} specifies to 1. In the example, blocks have capacity for 10 fingerprints, so B_j is already full even though F_0 's other candidate has spare capacity. In this case, B_j experiences a block overflow. In a block overflow without a bucket overflow, any of the existing fingerprints can be evicted to make space for F_0 . ④ In the example, F_1 is evicted from the bucket preceding F_0 's other candidate. F_1 remaps to its alternate candidate, a bucket in B_k , and because B_k is under capacity and F_1 's new bucket has a spare slot, the displacement completes. ⑤ The OTA in B_j is then updated by setting a bit to record F_1 's eviction.

We stress that these complex chains of displacements are infrequent in an MF, contrary to a CF, even at high load factors (e.g., 0.95). With the proper choice of parameters (see Section 5), over 95% of items are trivially inserted in their first or second buckets without triggering evictions.

4.3 Deletions

Deletions proceed similarly to lookups. Our example proceeds by deleting the fingerprint F_y that we inserted in Figure 6. We first compute H_1 on the key (call it K_y) to determine the primary bucket and $H_F(K_y)$ to calculate F_y . From there, we compute the block index and block-local bucket index. The next step is to search the key's primary bucket and delete its fingerprint provided there is a match. Figure 8 shows the block-local operations. ① We first sum the fullness counters from index 0 to $lbi - 1$ inclusive, which gives us the offset of the primary bucket's fingerprints within the FSA. ② We then perform a comparison between F_y and all the fingerprints in the primary bucket. ③ On a match, we right shift all fingerprints to the left of the matching fingerprint, which has the effect of deleting the fingerprint. If there are two or more matching fingerprints, we select one and delete it. ④ Finally, we update the FCA to reflect the primary bucket's new load by decrementing its fullness counter (at index lbi in the FCA) and return.

When the fingerprint is not found in the primary bucket, we calculate $H_2(K_x)$ to produce the secondary bucket's global logical index and proceed as before by computing the block ID and block-local bucket index. We then repeat the process in Figure 8 and delete a single fingerprint on a match. Note that contrary to lookups, we did not need to check the OTA before proceeding to the secondary bucket. Because ASMDs only permit deletions to items that have actually been stored in the filter (otherwise false negatives are possible), a failure to match in the primary bucket means the fingerprint must be in the alternate candidate and that the secondary deletion attempt will succeed [29].

Note, our implementation does not clear OTA bits. Repeat insertions and deletions will lead to a growing number of set bits. We combat this effect by biasing block overflows so that they overwhelmingly set the same bits in the OTA by biasing evictions on block overflows from lower order buckets. Given that typically only several percent of fingerprints overflow at load factors at or below 0.95 (less than 5% for the design in Figure 4), cotuning the OTA’s length and the MF’s load factor is sufficient for many applications.

For supporting many repeat deletions while sustaining near-maximal load factors (e.g., ≥ 0.95), one robust approach is to prepend each fingerprint with a bit that specifies whether the fingerprint is in its primary (or secondary) bucket and force all overflows from a block that map to the same OTA bit to remap to the same alternate block (but very likely different buckets). On deletions of a secondary item x , it is then possible to clear x ’s corresponding bit in its primary block’s OTA if no other fingerprints in its secondary block are simultaneously secondary, would map back to x ’s originating block, and would hash to the same bit that x would have set in the OTA when it overflowed. A probabilistic variant that saves space by forgoing tagging fingerprints at the expense of not being able to as aggressively clear OTA bits is also possible. We leave the implementation to future work.

4.4 Fast Reductions for Determining Bucket Boundaries in the FSA

One of the keys to a high-throughput MF is implementing fast reductions on the FCA. Specifically, determining the start of each bucket in the block requires summing all the counts of all fingerprints in buckets that precede it within its block’s FSA. A core challenge is that buckets at different indices require summing differing numbers of fullness counters in the FCA. A naive implementation that issues a variable number of instructions will lead to poor performance. Instead, a branchless algorithm with a fixed instruction sequence is necessary. At first, we considered doing a full exclusive scan (i.e., for every bucket computing the number of fingerprints that precede it). Efficient branchless parallel algorithms like Kogge-Stone [40] exist that require $O(\log(n))$ time and map well to the SIMD instruction sets of modern processors (e.g., SSE [64] and AVX variants [46]).

However, it turns out that a class of more efficient algorithms is possible that consists entirely of simple logic operations (e.g., NOT, AND, and OR), logical shifts, and the *population count* primitive (popcount for short). Popcount is a native instruction on almost every major processor and is a high-throughput, low-latency primitive [53]. Further, high-performance SIMD implementations of popcounts exist that use a combination of lookup tables and permute operations, so even if there is no native popcount instruction, performance-competitive workarounds like Mula et al.’s algorithm and AVX2 implementation are possible [53]. Popcount takes as input a multibyte integer and returns the number of bits that are set to 1. Prior work by González et al. [34] leverages the popcount primitive as part of a high-performance rank and select algorithm. Our algorithm generalizes these primitives to arrays of fixed-width counters.

Our approach is shown in Algorithm 3 which, given a bucket at bucket index lbi within the block, computes and returns the number of fingerprints that precede the bucket’s fingerprints in the block’s FSA.

We first perform a masked copy of the FCA where all counters that appear at lbi or greater are cleared (lines 2-3). This masking ensures that we only count fingerprints that precede the bucket at lbi . In our implementation, this op-

Algorithm 3 This algorithm takes as input a block’s fullness counter array FCA , a bucket index lbi within the block, the width of each counter in bits w , and then returns the index of the bucket’s initial element in the FSA.

```

1: function EXCLUSIVEREDVIAPOPCOUNT( $FCA, lbi, w$ )
2:    $fullnessCounterArrayMask = (1 \ll (w * lbi)) - 1$ 
3:    $mFCA = FCA \& fullnessCounterArrayMask$ 
4:    $pcMask = \text{getPopCountMask}(w)$ 
5:    $sum = 0$ 
6:   for ( $bitPos = 0$ ;  $bitPos < w$ ;  $bitPos++$ ) do
7:      $sum += \text{popCnt}(mFCA \& pcMask) \ll bitPos$ 
8:    $pcMask \ll= 1$ 
   return  $sum$ 

```

eration also masks out fingerprints that are packed into the same word as the fullness counter array. We next call `getPopCountMask` (line 4), which for w -bit fullness counters, returns a mask where the LSB and every w th bit thereafter are set to 1 and the remaining bits to 0. This mask when **anded** with the masked fullness counter array $mFCA$ selects out all bits that occur in the zeroth position of each of the counters and zeroes out the other digits. For instance, with 4-bit counters, four buckets per block, and hence a 16-bit fullness counter array, the mask $pcMask$ would be **0b0001000100010001** or equivalently **0x1111**. For any value of $bitPos$ between 0 and $w - 1$ inclusive, shifting the $pcMask$ to the left by $bitPos$ selects out the $bitPos$ th least significant digit of each w -bit counter (line 8).

The next phase incrementally counts each fingerprint that appears prior to the bucket digit-by-digit across all remaining counters in the masked copy. We initialize the sum to zero (line 5). During the $bitPos$ th pass of the algorithm (line 7), we count the $bitPos$ th least significant bit of each counter. That sum is then shifted left by the exponent of the power of two to which the digit corresponds ($bitPos$) before applying it to the growing partial sum sum , which we ultimately return once all passes are complete. Note that loop unrolling will eliminate the branch on line 6.

5. MODELING

In this section, we describe a set of models that are used to select the parameters for a filter given a set of constraints. Table 2 lists the parameters that we use in our models.

Table 2: MF Glossary of Symbols

- B - buckets per block
- C - the multiplicative slot compression ratio, where $C = 0.25$ corresponds to four slots in the logical interpretation for each physical slot in the FSAs
- S - logical slots per bucket
- α_L - logical load factor (e.g., $\alpha_L = 0.5$ for 4-slot buckets where on average two slots are full)
- α_C - block load factor (e.g., $\alpha_C = 0.8$ for blocks with 40-slot FSAs where on average 35 slots are occupied)
- O - number of bits in the OTA of each block
- m - expected number of items that overflow a block
- b - expected buckets accessed per negative lookup
- M - total fingerprints (net total occupied FSA slots)

5.1 False Positives and Storage Costs

The false positive rate ϵ for an MF is given by Equation 3.

$$\epsilon = 1 - (1 - 1/2^f)^{\alpha_L b S} \quad (3)$$

From Equation 3, we derive Equation 4 the formula for the fingerprint length f in bits for a target ϵ .

$$f \approx \log_2((\alpha_L b S)/\epsilon) = \log_2((\alpha_C b C S)/\epsilon) \quad (4)$$

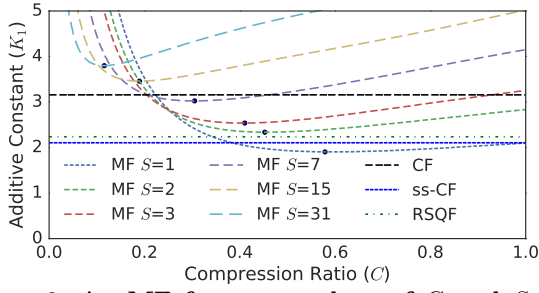


Figure 9: An MF for some values of C and S uses fewer bits per item than a CF, a CF with semi-sorting (ss-CF), and a rank-and-select quotient filter (RSQF). $\alpha_C = 0.95$ for the MFs and $\alpha = 0.95$ for the RSQF, ss-CF, and CF. For the MFs, we set the block size at 512 bits, and $O = 16$ and $f = 8$.

The bits per item is given by Equation 5, with the transformed expression obtained by substituting $\alpha_C C$ in for α_L .

$$\begin{aligned} I &= \text{OTA bits/item} + \text{FCA bits/item} + \text{FSA bits/item} \\ &= \frac{O}{\alpha_L BS} + \frac{\log_2(S+1)}{\alpha_L S} + \frac{Cf}{\alpha_L} \\ &= \frac{O/(BCS) + \log_2(S+1)/(CS) + \log_2((\alpha_C bCS)/\epsilon)}{\alpha_C} \end{aligned} \quad (5)$$

The OTA bits per item is O divided by the expected occupied fingerprint slots in an FSA ($\alpha_L BS$). The $\log_2(S+1)/(\alpha_L S)$ term counts the bits per FCA counter per item in the filter, and the $(Cf)/\alpha_L$ scales the fingerprint length f by the block load factor α_C as $(Cf)/\alpha_L = f/\alpha_C$. For the FSA bits per item term (f/α_C), the α_L in the numerator (i.e., $f = \log_2((\alpha_L bS)/\epsilon)$) works to reduce f by typically being a small value such that $\log_2(\alpha_L)$ shortens f by 1 to 3 bits while the α_C in f/α_C 's denominator can be tuned to be close to 1 if the MF's workload is known a priori. These gains are primarily from the FCA and FSA working in concert, which allows us to select α_L to be small (e.g., 0.2) and to shrink S from 4 to 3 or less, all while using comparable or slightly less space than a CF. Further, the OTA helps by reducing b from 2 to close to 1, enough to hide the OTA's space cost while also permitting some space savings.

If we fix all parameters except ϵ , I becomes $I = K_1 + K_2 \log_2(1/\epsilon)$, where K_1 and K_2 are constants. Figure 9 shows how K_1 varies with the compression ratio C and slots per bucket S for a fixed block load factor α_C of 0.95. We coplot the associated K_1 constants for a $S = 4$ $b = 2$ cuckoo filter (CF), a rank-and-select quotient filter, and a $S = 4$ $b = 2$ cuckoo filter with semi-sort compression (ss-CF) all at a load factor of 0.95. At this design point, $K_2 \log_2(1/\epsilon)$ is the same for all filters. The figure demonstrates several key points: (1) MFs use a comparable amount of storage to other filters, (2) there is a fair amount of flexibility when choosing the compression ratio without adversely affecting space usage when the slots per bucket is small, (3) optimizing for space as buckets scale in size requires reducing C , and (4) large buckets may be used with MFs with an additional storage overhead of 1 or 2 bits. This latter point contrasts with CFs, which use an extra bit for each power of two increase in S .

5.2 Lookup Costs

An important parameter when modeling the expected cost of each lookup in total buckets is m , the number of items expected to overflow an MF block. Equation 6 presents an approximation for m , which models both block and bucket overflows, as well as their intersection (i.e., an overflow that is both a block and bucket overflow). We ignore cascading evictions due to cuckoo hashing since they are not prevalent

(e.g., <1%) for typical parameter values. In our approximation, we model bucket and block overflows using models derived from the Poisson distribution. For $S \geq 2$, block overflows typically overwhelmingly dominate.

$$\begin{aligned} m &\approx \text{bucket overflows} + \text{block overflows} - \text{bucket and block overflows} \\ &\approx \alpha_L BS \left[\frac{\sum_{x=S+1}^M (x-S) Pr(\alpha_L S, x)}{\alpha_L S} + \frac{\sum_{x=CBS+1}^M (x-CBS) Pr(\alpha_L BS, x)}{\alpha_L BS} \right] \\ &\quad - \frac{1}{\alpha_L BS} \sum_{x=CBS+1}^M (x-CBS) Pr(\alpha_L BS, x) \sum_{y=S+1}^x \frac{(y-S) Pr(x/B, y)}{x/B} \end{aligned} \quad (6)$$

where $Pr(\lambda, x) = \frac{\lambda^x e^{-\lambda}}{x!}$

For lookups that are correctly characterized as negatives (excluding false positives), the cost of each such lookup is b . b is one plus the fraction of OTA bits that are set on average for each block. For instance, with a 12-bit OTA per block, if the mean bits that are set is 2, then b would be $1 + 2/12$ (i.e., about 1.167 buckets are expected to be accessed per lookup query that returns *false*).

The number of set bits is dependent on the mean fingerprints per block that overflow. These overflows occur both when a bucket does not have sufficient capacity (i.e., when its fullness counter maxes out) and when the block becomes full. Assuming m overflows per block, then Equation 7 gives the expected negative lookup cost in buckets.

$$\text{negative lookup cost} \approx b = 1 + 1 - (1 - \frac{O-1}{O})^m = 2 - (\frac{O-1}{O})^m \quad (7)$$

The final term is the expected fraction of OTA bits that are unset. It is also the probability that a single bit within the OTA is unset. We derive the model by using a balls-into-bins model (see Mitzenmacher and Upfal [52]) where each of the O bits in the OTA are bins and balls are the m overflow items. With m balls thrown uniformly randomly into O bins, the likelihood that an unset bit remains unset after one ball is thrown is $\frac{O-1}{O}$, and we exponentiate by m because the m balls are thrown independently of one another.

For positive lookups (excluding false positives), the lookup cost is shown in Equation 8 and is approximately one plus the expected fraction of items that overflow a block.

$$\text{positive lookup cost} \approx 1 + (1 - \frac{1}{2f})^{\alpha_L S} * \frac{m}{\alpha_L BS} \quad (8)$$

$\alpha_L BS$ is the mean occupied fingerprint slots in the FSA per block. The first term in the product corrects for an alias that occurs on the primary bucket when the item's actual fingerprint is in the secondary bucket. Equation 8 is also the expected cost of a deletion since well-formed deletions always succeed (otherwise, false negatives are possible).

The lookup cost for false positives is shown in Equation 9, which interpolates between 1.5 (the cost of a false positive with a completely full OTA) and 1.0 (the cost of a false positive with an empty OTA). For example, if zeros constitute three quarters of the OTA's bits, then we expect $1.5 - 0.5 * (0.75) = 1.125$ buckets to need to be checked.

$$\text{false positive lookup cost} \approx 1.5 - 0.5 * (1 - \frac{1}{O})^m \quad (9)$$

Given P , a ratio of true positives to total lookups, we can compute the expected lookup cost of Equation 10, the weighted average of the individual lookup costs.

$$\begin{aligned} \text{expected lookup cost} &= P * \text{positive lookup cost} + \\ &\quad (1-P)(1-\epsilon) * \text{negative lookup cost} + \\ &\quad (1-P)(\epsilon) * \text{false positive lookup cost} \end{aligned} \quad (10)$$

6. EXPERIMENTAL METHODOLOGY

We conduct our experiments on an AMD Ryzen™ Threadripper™ 1950X processor, which is composed of two 8-core dies for a total of 16 cores, each with 2-way simultaneous multithreading [74, 75]. Each core has a 32 KB L1

data cache, a 64 KB L1 instruction cache, a 512 KB L2, and there is a 32 MB L3 that is shared among all cores. We fix the CPU’s frequency at 3400 Mhz. Each 8-core die has two memory channels for a total of four. The machine has 128 GB of RAM that is clocked at 2133 MHz, and it runs Ubuntu 16.04.4 LTS (Linux 4.11.0).

We compare the MF’s throughput to three other filters and implementations from prior work. These are Fan et al.’s cuckoo filter (CF) [28,29], Fan et al.’s CF with semi-sorting (ss-CF) [11,28,29], and Pandey et al.’s rank-and-select quotient filter (RSQF) [59,60]. Fan et al. [29] already demonstrated the CF to be faster than Bloom filters [7], blocked Bloom filters [62], and d -left counting Bloom filters [10,11], so we do not evaluate these designs.

Unless stated otherwise, we run experiments on filters with $128 * 1024 * 1024$ slots. We configure the MF to use 8-bit fingerprints, and to have a 46-slot FSA, a 128-bit FCA (64×2 -bit counters), and a 16-bit OTA, for a total of 512 bits. It thus contains 64 buckets (each with three logical slots) and has a slot compression ratio (C) of about 0.24. The configuration is the same as the one in Figure 3 and at a load factor of $\alpha_C = 0.95$ produces an MF that uses $512 / (46 * 0.95) = 11.71$ bits per item. We thus compare our implementation to a CF that uses 12-bit fingerprints since Fan et al.’s code does not support 11-bit fingerprints [28,29]. Both filters have roughly equivalent error rates for similar load factors (α_C in the case of the MF).

We generate 64-bit random integers using C++’s standard library and benchmark the performance of each filter by running separate trials where we fill the filter up to a load factor that is a multiple of 0.05 and then proceed to insert, delete, or look up a number of fingerprints equal to 0.1% of the total slots in the filter. Both the generation of large filters that do not fit into cache and of uniformly random fingerprints are consistent with the evaluations of Fan et al. [29] and Pandey et al. [59]. We rebuild filters from scratch after each trial with a new set of random fingerprints to reduce noise. Results are the average of 5 trials. Filters are compiled using g++ (version 5.4.0-6) with the `-Ofast -march=native` flags, as they yield the best performance. We plot throughput in millions of operations per second (MOPS).

Fan et al.’s implementation packs four 12-bit fingerprints into every 64-bit word and pads the remaining 16 bits. Thus, their CF uses 8 bytes for every four 4 fingerprint slots [29]. The MF uses 64 bytes for every 46 fingerprint slots. Thus, the MF is about 186.7 MB in size whereas the CF is 256 MB (192 MB if it did not pad).

7. EVALUATION

In this section, we present our results, which show our MF implementation to sustain higher throughput than a CF for lookups, insertions, and deletions.

7.1 Lookup Throughput

Figure 10 presents the throughput of the MF when 100% of the lookups are true positives. In this configuration, at a load factor of 0.95, a mere 1.05 cache lines from the filter are accessed per query, a reduction of close to 50% over a CF. We observe that at low load factors, the MF is over $2\times$ faster than a CF and upwards of $1.6\times$ at high loads.

Figure 11 presents the throughput of the MF when 100% of the lookups are true negatives (a mix of negatives and false positives). The filter achieves a throughput that is $1.3\times$ to $2.5\times$ faster. At a load factor of 0.95, the 16-bit OTA has about 11% to 12% of its bits set to 1, so lookups require accessing about 1.11 cache lines (approximately twice as

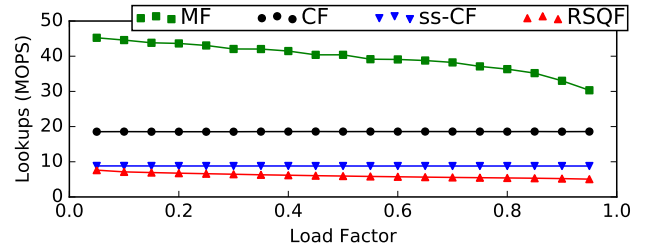


Figure 10: An MF is approximately $1.6\times$ to $2.4\times$ faster than a CF for queries that are true positives.

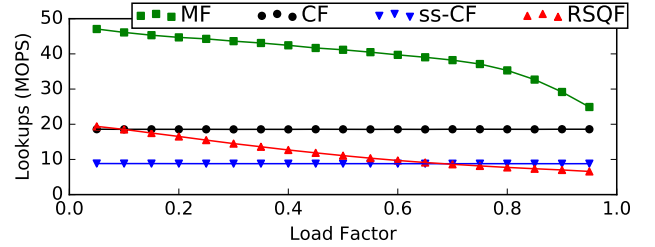


Figure 11: An MF is approximately $1.3\times$ to $2.5\times$ faster than a CF for queries that are true negatives.

many secondary lookups are required as the positive lookup case). This difference explains why positive lookups sustain higher throughput than negative lookups at heavy loads. To counter this drop in throughput and approximately match the performance of positive lookups, we could double the length of the OTA or reduce the load factor to 0.9.

On an idealized machine and implementation, performance would only drop by about 5% to 11% (i.e., in line with additional data movement). However, it is difficult to achieve that performance in practice due to microarchitectural limitations (e.g., branch mispredictions) and practical trade-offs in software engineering (e.g., trading off between writing code quickly that is cross-platform versus hand-tuning code for a specific processor). Since approximately 90% or more of the lookups never have to perform a secondary access, we focused our efforts on making that fast. An industrial implementation with hand tuned assembly or vector intrinsics is likely to achieve speedups at high load that are much closer to the reductions in pseudorandom cache accesses.

7.2 Insertion Throughput

Like lookups, MF insertion throughput realizes large improvements over CFs for most load factors. Figure 12 demonstrates that an MF is able to sustain high throughput for much longer than a CF. At high loads (e.g., a load factor of 0.75 or higher), the MF is approximately $3\times$ to $15\times$ faster than a comparable CF. This difference matches the simple intuition that buckets with empty slots are easier to insert into than those that are full. Imagine an MF with blocks with 48-slot FSAs and a sample block in which only a single FSA slot is free. Even in this extreme case, provided α_L is low, it is very likely that the block can receive another fingerprint (i.e., very few, if any, of the buckets are likely full). However, if we arrange those same 48 slots into 12 4-slot CF buckets, the probability that we map to a bucket with a free slot is $1/12$. Thus, whereas the MF very likely only needs to access one cache line from the filter’s storage, the CF is expected to access two or more for 11 out of every 12 insertions, as it needs to access at least one secondary bucket. The MF’s reduction in accesses is the key to improving throughput, as cache and memory bandwidth (the

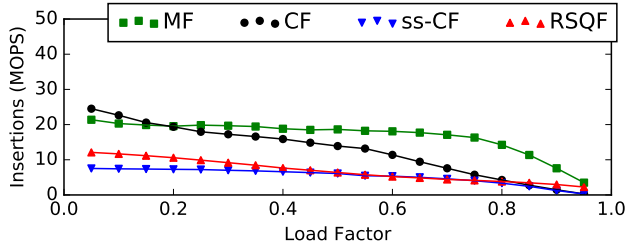


Figure 12: An MF is approximately $0.9\times$ to $15.5\times$ faster for insertions than a CF.

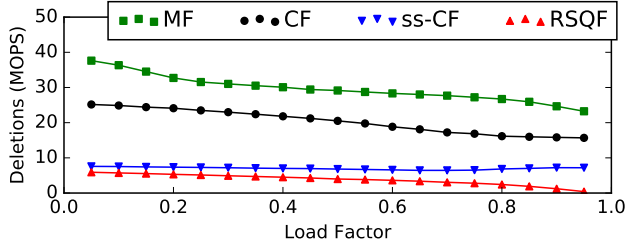


Figure 13: An MF is approximately $1.3\times$ to $1.6\times$ faster for deletions than a CF.

typical bottlenecks for an ASMDS) are much more efficiently used.

7.3 Deletion Throughput

MF deletion throughput is about $1.3\times$ to $1.6\times$ higher than a CF's (Figure 13). Like lookups and insertions, the improvement is driven by reducing cache and memory traffic per operation. With these parameters, over 95% of deletions never access more than one MF cache line (even for $\alpha_C=0.95$).

7.4 Throughput Impact of Optimizations

In Figure 14, we explore the vital role that batching and the branchless popcount-accelerated reduction optimizations play in yielding a high-performance MF implementation. For our baseline MF, we perform one operation at a time and use a naive accumulator loop for summing fullness counters to determine bucket starting offsets in the FSA. Adding batching improves throughput by roughly $2\times$ (except for insertions). Replacing the naive accumulator loop with Algorithm 3 yields another $2\times$ to $3\times$ improvement. Insertion results are the net throughput from filling an MF from empty to a block load factor α_C of 0.95. All other results are the throughput at $\alpha_C = 0.95$. The optimized MF fills to a load factor of 0.95 about $3\times$ faster than the CF.

7.5 Flexibility

Figure 15 plots lookup, insertion, and deletion throughput when varying fingerprint length. Despite the wide range in fingerprint lengths, throughput is relatively constant.

Figure 16 shows how throughput changes per operation as the slots per bucket is varied. As an MF logically underloads its buckets, the expected number of slots that require checks remains small, and throughput is only modestly affected.

7.6 Cross-Platform Performance Portability

In this section, we demonstrate the MF's strong performance portability across different microarchitectures by benchmarking on a Skylake-X server. Since the MF and CF were also the fastest on this platform, we leave out the RSQF and ss-CF, as their throughput as compared to the CF's and MF's did not appreciably change.

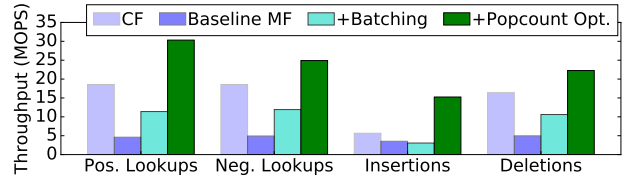


Figure 14: MF throughput surges with the cumulative addition of optimizations ($\alpha_C = 0.95$).

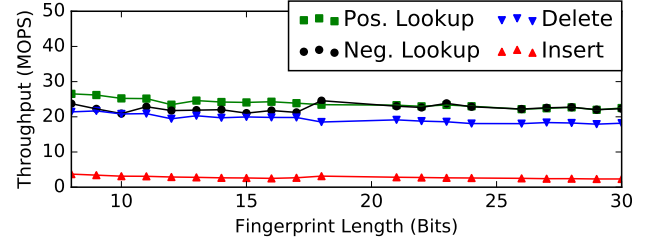
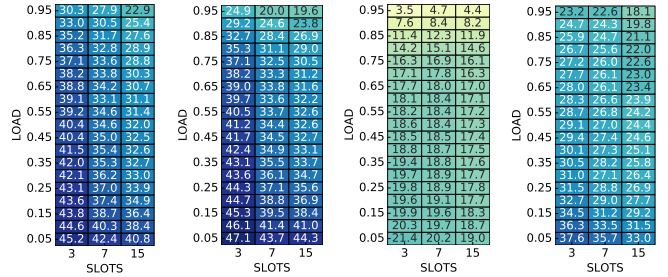
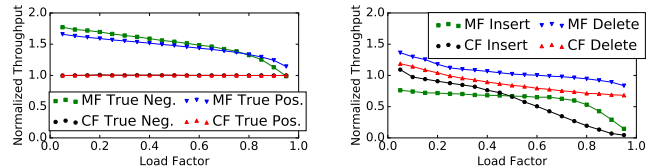


Figure 15: An MF sustains high throughput for a range of fingerprint lengths ($C \approx 0.24$, $S = 3$, $O = 16$, and $\alpha_C = 0.95$).



(a) Pos. (b) Neg. (c) Insert (d) Delete

Figure 16: MF throughput in MOPS as we covary the slots per bucket (S) and block load factor (α_C). High throughput is achieved even for large S .



(a) Lookup Throughput (b) Update Throughput

Figure 17: On a Skylake-X server, MF lookup throughput is on par with or nearly $1.8\times$ higher than a CF's. MF deletion throughput is about 20% higher than a CF's. MF insertion throughput is about $2/3$ to $3\times$ that of a CF. Results are normalized to a CF's lookup throughput on a Skylake-X CPU.

For Skylake-X, we perform no additional tuning of our MF implementation, yet despite that, it is still in many cases significantly faster than the CF; the MF's disciplined conservation of cache and memory bandwidth give it the edge. Given that the CF was developed and performance tuned on an Intel processor like the Skylake-X, it is logical that its performance relative to an MF improves. For lookups (Figure 17a), the MF attains throughput that matches or exceeds the CF even though it uses fewer bits per item for an equivalent ϵ . Deletion throughput is roughly $1.2\times$ higher (Figure 17b), and insertion throughput (Figure 17b) is higher for load factors at or exceeding 0.5 (up to $\approx 3\times$ better).

8. RELATED WORK

There continues to be significant interest in the design and application of ASMDs. Bloom first proposed his eponymous filter in 1970 [7]. Since then it has been used in many different contexts and has evolved into many different variants [1, 18, 21, 30, 66, 80]. Broder and Mitzenmacher provide a survey of some variants and their applications in computer networking [14]. Fan et al. propose the counting Bloom filter (CBF) for use in SummaryCache [30]. Cohen et al. develop the spectral Bloom filter [18], which like the CBF provides deletions and counting but is more resistant to skew.

ASMDs Use in Databases and Data Stores - In the database domain, Bloom filters have seen widespread use in accelerating hashed semijoins (e.g., the Bloomjoin [12, 48]). PostgreSQL [70] supports using a Bloom filter as an index and is visible as an SQL extension [68].

RocksDB employs Bloom filters and partitioned Bloom filters to speed querying of large data blocks [22]. OceanStore uses a multi-layer attenuated Bloom filter [42]. Similar hierarchies of Bloom filters are commonly found in Log-Structured Merge-trees [56, 67]. LSM-trees are employed in a number of data stores such as RocksDB [22], BigTable [17], LevelDB [20], HBase [33], and Cassandra [43].

Fingerprint-based ASMDs - A number of filters exist that use fingerprints in lieu of setting individual bits. Cuckoo filters [25, 29], d -left counting Bloom filters [10], quotient filters [6, 59], and TinySet [24] are some examples. TinySet truncates fingerprints to avert overflows that would reduce locality. As such, repeat deletions increase its error rate. d -left counting Bloom filters leverage the improved load distribution properties of d -left hashing [76] to improve space utilization approximately $2\times$ over a counting Bloom filter.

Compression for Sparse Matrices - The compression method used by the MF shares similarities with approaches for sparse matrices such as *compressed sparse rows* (CSR) [36, 73] and *compressed sparse columns* (CSC) [73] but requires less metadata because positional information within a row (such as in CSR) or within a column (such as in CSC) does not need to be encoded.

Sparse and Succinct Data Structures - Prior works present methods for storing sparse data structures [32, 38, 55, 63, 72]. Many use clever hashing or bit vectors with rank and select (see Jacobson [38] and Navarro [54]). Common applications are compressing sparse trees [38, 81] or tables [55, 72].

Compressed Bitmap Indices - Our compression approach differs from the typical compressed bitmap algorithms like BBC [2] and the WAH variants which primarily use run length encoding [16, 19, 37, 77, 79]. Since fingerprints are approximately uniformly random, the main compression opportunity is eliminating storing empty slots, which our simpler approach already does well.

Compression in Prior ASMDs - Other ASMDs have used compression in the past. Semi-sorting, a form of compression where a portion of each fingerprint is compressed has been proposed in the context of d -left counting Bloom filters [10, 11] and cuckoo filters [29]. Semi-sorting could be added to an MF to save additional space. We chose not to use semi-sorting because the gains in space come at significant cost to performance. Further, our implementation would have required additional complexity since we would have had to support encoding and decoding for differing numbers of occupied slots per bucket since empty slots are not explicitly stored. Mitzenmacher provides detailed analysis of design tradeoffs when compressing Bloom filters and discusses its applicability to web caching [50].

Cuckoo Hash Tables - A CF is highly related to cuckoo hash tables (CHTs) [58] and variants [13, 65, 71]. Rather than storing fingerprints, a CHT stores key-value pairs. Like CFs, CHTs typically have two candidate buckets with four or eight slots [13, 26, 27, 45, 65, 82]. A commonality of a baseline CF and CHT is that as the load on the table increases, insertion throughput decreases due to a rapid increase in the prevalence and mean length of cuckoo evictions. Prior work addresses this reduction in throughput in several ways. Li et al. employ a breadth-first search (BFS) to reduce the maximum chain length to one that is logarithmic in the maximum number of slots that are checked before declaring failure [45]. Their concurrent CHT outperforms MemC3's [27] that uses 2-choice hashing [4, 51] with the random knockout algorithm employed in Fan et al.'s CF [29]. A concurrent MF would likely similarly benefit from using BFS. Sun et al. [71] add metadata that explicitly tracks the graph-theoretic state of the table to prune the search space for cuckoo hashing. Our design avoids this complexity by using blocks that support storing tens of fingerprints. Horton tables [13] convert the final slot of buckets that overflow into a *remap entry array* (REA) that enables lookups that access close to a single bucket, provides many bucket candidates (e.g., 8), and keeps a worst-case lookup cost of two buckets. We implement the OTA as a bit vector rather than an REA for simplicity.

9. CONCLUSION

We have presented a high-throughput filter that supports improved throughput for lookups, insertions, and deletions without increasing memory usage. Perhaps most notable is that an MF's insertion throughput is about $3\times$ to $15\times$ higher than a comparable CF for load factors at or above 0.75. Further, lookup and deletion throughput are up to $2.5\times$ and $1.6\times$ faster, respectively. These properties are achieved while also using comparable or fewer bits per item than a CF for a target false positive rate. Key to these advances is the block abstraction and its compressed format, which allows for hiding the storage cost of additional metadata structures by logically underloading the filter and using smaller buckets. The OTA further decreases these costs by reducing aliasing that would require increasing the length of fingerprints by filtering out unnecessary accesses to secondary buckets. With the OTA and a reduction in bucket overflows due to packing many underloaded buckets into a single cache line, lookups most often only have to access a single bucket (one hardware cache line) even when the filter is heavily loaded. We look forward to applying the MF in a variety of contexts due to its memory friendliness. Further, the innovations of this work like the compression and performance optimizations can be applied to a broad range of other data structures such as hash tables (e.g., a cuckoo hash table), various fingerprint-based filters, and algorithms that employ reductions or scans on fixed-width narrow fields or counters.

10. ACKNOWLEDGMENTS

We thank the VLDB reviewers and our kind colleagues Shaizeen Aga, Joseph L. Greathouse, and Gabriel Loh for their time and superb feedback which substantially improved the paper's clarity and quality. We also thank John Kalamatianos and Jagadish Kotra for giving us access to the Skylake-X server and Karen Prairie for her copy edits. AMD is a trademark of Advanced Micro Devices, Inc. Other product names used in this publication are for identification purposes only and may be trademarks of their respective companies.

11. REFERENCES

- [1] P. S. Almeida, C. Baquero, N. M. Prego, and D. Hutchison. Scalable Bloom filters. *Information Processing Letters*, 101(6):255–261, 2007.
- [2] G. Antoshenkov. Byte-aligned bitmap compression. In *Data Compression Conference, 1995. DCC '95. Proceedings*, pages 476–, March 1995.
- [3] A. Appleby. MurmurHash. <https://sites.google.com/site/murmurhash>, 2008. Accessed: 2018-02-05.
- [4] Y. Azar, A. Z. Broder, A. R. Karlin, and E. Upfal. Balanced allocations. *SIAM Journal of Computing*, 29(1):180–200, 1999.
- [5] L. A. Belady. A study of replacement algorithms for a virtual-storage computer. *IBM Systems Journal*, 5(2):78–101, 1966.
- [6] M. A. Bender, M. Farach-Colton, R. Johnson, R. Kraner, B. C. Kuszmaul, D. Medjedovic, P. Montes, P. Shetty, R. P. Spillane, and E. Zadok. Don't thrash: how to cache your hash on flash. *PVLDB*, 5(11):1627–1637, 2012.
- [7] B. H. Bloom. Space/time trade-offs in hash coding with allowable errors. *Communications of the ACM*, 13(7):422–426, 1970.
- [8] P. A. Boncz, S. Manegold, and M. L. Kersten. Database architecture optimized for the new bottleneck: memory access. In *Proceedings of the 25th International Conference on Very Large Data Bases, VLDB '99*, pages 54–65, San Francisco, CA, USA, 1999. Morgan Kaufmann Publishers Inc.
- [9] P. A. Boncz, M. Zukowski, and N. Nes. MonetDB/X100: hyper-pipelining query execution. In *CIDR 2005, Second Biennial Conference on Innovative Data Systems Research, Asilomar, CA, USA, January 4-7, 2005, Online Proceedings*, pages 225–237, 2005.
- [10] F. Bonomi, M. Mitzenmacher, R. Panigrahy, S. Singh, and G. Varghese. An improved construction for counting Bloom filters. In *ESA*, volume 6, pages 684–695. Springer, 2006.
- [11] F. Bonomi, M. Mitzenmacher, R. Panigrahy, S. Singh, and G. Varghese. Bloom filters via d-left hashing and dynamic bit reassignment extended abstract. In *Forty-Fourth Annual Allerton Conference, Illinois, USA*, pages 877–883, 2006.
- [12] K. Bratberg. Hashing methods and relational algebra operations. In *Proceedings of the 10th International Conference on Very Large Data Bases, VLDB '84*, pages 323–333, San Francisco, CA, USA, 1984. Morgan Kaufmann Publishers Inc.
- [13] A. D. Breslow, D. P. Zhang, J. L. Greathouse, N. Jayasena, and D. M. Tullsen. Horton tables: fast hash tables for in-memory data-intensive computing. In A. Gulati and H. Weatherspoon, editors, *2016 USENIX Annual Technical Conference, USENIX ATC 2016, Denver, CO, USA, June 22-24, 2016.*, pages 281–294. USENIX Association, 2016.
- [14] A. Z. Broder and M. Mitzenmacher. Network applications of Bloom filters: a survey. *Internet Mathematics*, 1(4):485–509, 2003.
- [15] L. Carter, R. Floyd, J. Gill, G. Markowsky, and M. Wegman. Exact and approximate membership testers. In *Proceedings of the Tenth Annual ACM Symposium on Theory of Computing, STOC '78*, pages 59–65, New York, NY, USA, 1978. ACM.
- [16] S. Chambi, D. Lemire, O. Kaser, and R. Godin. Better bitmap performance with Roaring bitmaps. *Software: Practice and Experience*, 46(5):709–719, 2016.
- [17] F. Chang, J. Dean, S. Ghemawat, W. C. Hsieh, D. A. Wallach, M. Burrows, T. Chandra, A. Fikes, and R. E. Gruber. BigTable: a distributed storage system for structured data. *ACM Transactions on Computer Systems (TOCS)*, 26(2):4, 2008.
- [18] S. Cohen and Y. Matias. Spectral Bloom filters. In A. Y. Halevy, Z. G. Ives, and A. Doan, editors, *Proceedings of the 2003 ACM SIGMOD International Conference on Management of Data, San Diego, California, USA, June 9-12, 2003*, pages 241–252. ACM, 2003.
- [19] A. Colantonio and R. D. Pietro. Concise: compressed 'n' composable integer set. *Information Processing Letters*, 110(16):644–650, 2010.
- [20] J. Dean and S. Ghemawat. LevelDB: a fast persistent key-value store. <https://opensource.googleblog.com/2011/07/leveldb-fast-persistent-key-value-store.html>, July 27, 2011. Accessed: 2017-01-25.
- [21] F. Deng and D. Rafiei. Approximately detecting duplicates for streaming data using Stable Bloom filters. In *Proceedings of the ACM SIGMOD International Conference on Management of Data, Chicago, Illinois, USA, June 27-29, 2006*, pages 25–36, 2006.
- [22] S. Dong, M. Callaghan, L. Galanis, D. Borthakur, T. Savor, and M. Strum. Optimizing space amplification in RocksDB. In *CIDR 2017, 8th Biennial Conference on Innovative Data Systems Research, Chaminade, CA, USA, January 8-11, 2017, Online Proceedings*, 2017.
- [23] Dr. Seuss. *Horton Hatches the Egg*. Random House, 1940.
- [24] G. Einziger and R. Friedman. TinySet - an access efficient self adjusting Bloom filter construction. *IEEE/ACM Transactions on Networking*, 25(4):2295–2307, 2017.
- [25] D. Eppstein, M. T. Goodrich, M. Mitzenmacher, and M. R. Torres. 2-3 cuckoo filters for faster triangle listing and set intersection. In *Proceedings of the 36th ACM SIGMOD-SIGACT-SIGAI Symposium on Principles of Database Systems, PODS 2017, Chicago, IL, USA, May 14-19, 2017*, pages 247–260, 2017.
- [26] U. Erlingsson, M. Manasse, and F. McSherry. A cool and practical alternative to traditional hash tables. In *Proc. 7th Workshop on Distributed Data and Structures (WDAS06)*, 2006.
- [27] B. Fan, D. G. Andersen, and M. Kaminsky. MemC3: compact and concurrent memcache with dumber caching and smarter hashing. In N. Feamster and J. C. Mogul, editors, *Proceedings of the 10th USENIX Symposium on Networked Systems Design and Implementation, NSDI 2013, Lombard, IL, USA, April 2-5, 2013*, pages 371–384. USENIX Association, 2013.
- [28] B. Fan, D. G. Andersen, and M. Kaminsky. Cuckoo filter. <https://github.com/efficient/cuckoofilter>, 2017. Accessed: 2017-11-19.
- [29] B. Fan, D. G. Andersen, M. Kaminsky, and M. Mitzenmacher. Cuckoo filter: practically better than Bloom. In *Proceedings of the 10th ACM International Conference on Emerging Networking Experiments and Technologies, CoNEXT 2014*,

- Sydney, Australia, December 2-5, 2014, pages 75–88, 2014.
- [30] L. Fan, P. Cao, J. M. Almeida, and A. Z. Broder. Summary Cache: a scalable wide-area web cache sharing protocol. *IEEE/ACM Transactions on Networking*, 8(3):281–293, 2000.
 - [31] M. J. Flynn. Some computer organizations and their effectiveness. *IEEE Transactions on Computers*, C-21(9):948–960, Sept 1972.
 - [32] M. L. Fredman, J. Komlós, and E. Szemerédi. Storing a sparse table with $O(1)$ worst case access time. *Journal of the ACM*, 31(3):538–544, June 1984.
 - [33] L. George. *HBase: the definitive guide: random access to your planet-size data*. O’Reilly Media, Inc., 2011.
 - [34] R. González, S. Grabowski, V. Mäkinen, and G. Navarro. Practical implementation of rank and select queries. In *Poster Proc. Volume of 4th Workshop on Efficient and Experimental Algorithms (WEA)*, pages 27–38, 2005.
 - [35] J. R. Goodman. Using cache memory to reduce processor-memory traffic. In *Proceedings of the 10th Annual International Symposium on Computer Architecture*, ISCA ’83, pages 124–131, New York, NY, USA, 1983. ACM.
 - [36] J. L. Greathouse and M. Daga. Efficient sparse matrix-vector multiplication on GPUs using the CSR storage format. In T. Damkroger and J. J. Dongarra, editors, *International Conference for High Performance Computing, Networking, Storage and Analysis, SC 2014, New Orleans, LA, USA, November 16-21, 2014*, pages 769–780. IEEE Computer Society, 2014.
 - [37] G. Guzun, G. Canahuate, D. Chiu, and J. Sawin. A tunable compression framework for bitmap indices. In *IEEE 30th International Conference on Data Engineering, Chicago, ICDE 2014, IL, USA, March 31 - April 4, 2014*, pages 484–495, 2014.
 - [38] G. Jacobson. Space-efficient static trees and graphs. In *30th Annual Symposium on Foundations of Computer Science, Research Triangle Park, North Carolina, USA, 30 October - 1 November 1989*, pages 549–554, 1989.
 - [39] M. Kandemir, H. Zhao, X. Tang, and M. Karakoy. Memory row reuse distance and its role in optimizing application performance. In *Proceedings of the 2015 ACM SIGMETRICS International Conference on Measurement and Modeling of Computer Systems*, SIGMETRICS ’15, pages 137–149, New York, NY, USA, 2015. ACM.
 - [40] P. M. Kogge and H. S. Stone. A parallel algorithm for the efficient solution of a general class of recurrence equations. *IEEE Transactions on Computers*, 100(8):786–793, 1973.
 - [41] M. Kornacker, A. Behm, V. Bittorf, T. Bobrovitsky, C. Ching, A. Choi, J. Erickson, M. Grund, D. Hecht, M. Jacobs, I. Joshi, L. Kuff, D. Kumar, A. Leblang, N. Li, I. Pandis, H. Robinson, D. Rorke, S. Rus, J. Russell, D. Tsirogiannis, S. Wanderman-Milne, and M. Yoder. Impala: a modern, open-source SQL engine for Hadoop. In *CIDR 2015, Seventh Biennial Conference on Innovative Data Systems Research, Asilomar, CA, USA, January 4-7, 2015, Online Proceedings*, 2015.
 - [42] J. Kubiawicz, D. Bindel, Y. Chen, S. E. Czerwinski, P. R. Eaton, D. Geels, R. Gummadi, S. C. Rhea, H. Weatherspoon, W. Weimer, C. Wells, and B. Y. Zhao. OceanStore: an architecture for global-scale persistent storage. In L. Rudolph and A. Gupta, editors, *ASPLOS-IX Proceedings of the 9th International Conference on Architectural Support for Programming Languages and Operating Systems, Cambridge, MA, USA, November 12-15, 2000.*, pages 190–201. ACM Press, 2000.
 - [43] A. Lakshman and P. Malik. Cassandra: a decentralized structured storage system. *Operating Systems Review*, 44(2):35–40, 2010.
 - [44] D. Lemire. A fast alternative to the modulo reduction. <https://lemire.me/blog/2016/06/27/a-fast-alternative-to-the-modulo-reduction/>, June 27, 2016. Accessed: 2017-07-01.
 - [45] X. Li, D. G. Andersen, M. Kaminsky, and M. J. Freedman. Algorithmic improvements for fast concurrent cuckoo hashing. In *Ninth EuroSys Conference 2014, EuroSys 2014, Amsterdam, The Netherlands, April 13-16, 2014*, pages 27:1–27:14, 2014.
 - [46] C. Lomont. Introduction to Intel advanced vector extensions. *Intel White Paper*, pages 1–21, 2011.
 - [47] D. B. Loveman. Program improvement by source-to-source transformation. *Journal of the ACM*, 24(1):121–145, Jan. 1977.
 - [48] L. F. Mackert and G. M. Lohman. R* optimizer validation and performance evaluation for distributed queries. In *Proceedings of the 12th International Conference on Very Large Data Bases, VLDB ’86*, pages 149–159, San Francisco, CA, USA, 1986. Morgan Kaufmann Publishers Inc.
 - [49] P. Melsted and J. K. Pritchard. Efficient counting of k-mers in DNA sequences using a Bloom filter. *BMC Bioinformatics*, 12:333, 2011.
 - [50] M. Mitzenmacher. Compressed Bloom filters. In *Proceedings of the Twentieth Annual ACM Symposium on Principles of Distributed Computing, PODC 2001, Newport, Rhode Island, USA, August 26-29, 2001*, pages 144–150, 2001.
 - [51] M. Mitzenmacher. The power of two choices in randomized load balancing. *IEEE Transactions on Parallel and Distributed Systems*, 12(10):1094–1104, 2001.
 - [52] M. Mitzenmacher and E. Upfal. *Probability and Computing: randomization and Probabilistic Techniques in Algorithms and Data Analysis*. Cambridge University Press, 2017.
 - [53] W. Mula, N. Kurz, and D. Lemire. Faster population counts using AVX2 instructions. *The Computer Journal*, 61(1):111–120, 2018.
 - [54] G. Navarro. *Compact data structures: a practical approach*. Cambridge University Press, 2016.
 - [55] D. Okanohara and K. Sadakane. Practical entropy-compressed rank/select dictionary. In *Proceedings of the Meeting on Algorithm Engineering & Experiments*, pages 60–70. Society for Industrial and Applied Mathematics, 2007.
 - [56] P. E. O’Neil, E. Cheng, D. Gawlick, and E. J. O’Neil. The Log-Structured Merge-tree (LSM-tree). *Acta Informatica*, 33(4):351–385, 1996.

- [57] D. A. Padua and M. J. Wolfe. Advanced compiler optimizations for supercomputers. *Communications of the ACM*, 29(12):1184–1201, Dec. 1986.
- [58] R. Pagh and F. F. Rodler. Cuckoo hashing. *Journal of Algorithms*, 51(2):122–144, 2004.
- [59] P. Pandey, M. A. Bender, R. Johnson, and R. Patro. A general-purpose counting filter: making every bit count. In *Proceedings of the 2017 ACM International Conference on Management of Data, SIGMOD Conference 2017, Chicago, IL, USA, May 14-19, 2017*, pages 775–787, 2017.
- [60] P. Pandey and R. Johnson. A general-purpose counting filter: counting quotient filter. <https://github.com/splatlab/cqf>, 2017. Accessed: 2017-11-09.
- [61] O. Polychroniou, A. Raghavan, and K. A. Ross. Rethinking simd vectorization for in-memory databases. In *Proceedings of the 2015 ACM SIGMOD International Conference on Management of Data, SIGMOD '15*, pages 1493–1508, New York, NY, USA, 2015. ACM.
- [62] F. Putze, P. Sanders, and J. Singler. Cache-, hash-, and space-efficient Bloom filters. *ACM Journal of Experimental Algorithmics*, 14, 2009.
- [63] R. Raman, V. Raman, and S. S. Rao. Succinct indexable dictionaries with applications to encoding k-ary trees and multisets. In *Proceedings of the Thirteenth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA '02*, pages 233–242, Philadelphia, PA, USA, 2002. Society for Industrial and Applied Mathematics.
- [64] S. K. Raman, V. Pentkovski, and J. Keshava. Implementing streaming SIMD extensions on the Pentium III Processor. *IEEE Micro*, 20(4):47–57, July 2000.
- [65] K. A. Ross. Efficient hash probes on modern processors. In R. Chirkova, A. Dogac, M. T. Özsu, and T. K. Sellis, editors, *Proceedings of the 23rd International Conference on Data Engineering, ICDE 2007, The Marmara Hotel, Istanbul, Turkey, April 15-20, 2007*, pages 1297–1301. IEEE Computer Society, 2007.
- [66] O. Rottenstreich, Y. Kanizo, and I. Keslassy. The variable-increment counting Bloom filter. *IEEE/ACM Transactions on Networking*, 22(4):1092–1105, 2014.
- [67] R. Sears and R. Ramakrishnan. bLSM: a general purpose Log Structured Merge tree. In K. S. Candan, Y. Chen, R. T. Snodgrass, L. Gravano, and A. Fuxman, editors, *Proceedings of the ACM SIGMOD International Conference on Management of Data, SIGMOD 2012, Scottsdale, AZ, USA, May 20-24, 2012*, pages 217–228. ACM, 2012.
- [68] T. Sigaev, A. Korotkov, and O. Bartunov. PostgreSQL 10 documentation: F.5. bloom. <https://www.postgresql.org/docs/10/static/bloom.html>, 2017. Accessed: 2018-01-25.
- [69] J. E. Smith. A study of branch prediction strategies. In *Proceedings of the 8th Annual Symposium on Computer Architecture, ISCA '81*, pages 135–148, Los Alamitos, CA, USA, 1981. IEEE Computer Society Press.
- [70] M. Stonebraker, L. A. Rowe, and M. Hirohama. The implementation of POSTGRES. *IEEE Transactions on Knowledge and Data Engineering*, 2(1):125–142, 1990.
- [71] Y. Sun, Y. Hua, S. Jiang, Q. Li, S. Cao, and P. Zuo. SmartCuckoo: a fast and cost-efficient hashing index scheme for cloud storage systems. In *2017 USENIX Annual Technical Conference, USENIX ATC 2017, Santa Clara, CA, USA, July 12-14, 2017.*, pages 553–565. USENIX Association, 2017.
- [72] R. E. Tarjan and A. C. Yao. Storing a sparse table. *Communications of the ACM*, 22(11):606–611, 1979.
- [73] W. F. Tinney and J. W. Walker. Direct solutions of sparse network equations by optimally ordered triangular factorization. *Proceedings of the IEEE*, 55(11):1801–1809, 1967.
- [74] D. M. Tullsen, S. J. Eggers, J. S. Emer, H. M. Levy, J. L. Lo, and R. L. Stamm. Exploiting choice: instruction fetch and issue on an implementable simultaneous multithreading processor. In *Proceedings of the 23rd Annual International Symposium on Computer Architecture, Philadelphia, PA, USA, May 22-24, 1996*, pages 191–202, 1996.
- [75] D. M. Tullsen, S. J. Eggers, and H. M. Levy. Simultaneous multithreading: maximizing on-chip parallelism. In *Proceedings of the 22nd Annual International Symposium on Computer Architecture, ISCA '95, Santa Margherita Ligure, Italy, June 22-24, 1995*, pages 392–403, 1995.
- [76] B. Vöcking. How asymmetry helps load balancing. In *40th Annual Symposium on Foundations of Computer Science, FOCS '99, 17-18 October, 1999, New York, NY, USA*, pages 131–141, 1999.
- [77] J. Wang, C. Lin, Y. Papakonstantinou, and S. Swanson. An experimental study of bitmap compression vs. inverted list compression. In *Proceedings of the 2017 ACM International Conference on Management of Data, SIGMOD '17*, pages 993–1008, New York, NY, USA, 2017. ACM.
- [78] M. Wolfe. More iteration space tiling. In *Proceedings of the 1989 ACM/IEEE Conference on Supercomputing, Supercomputing '89*, pages 655–664, New York, NY, USA, 1989. ACM.
- [79] K. Wu, E. J. Otoo, and A. Shoshani. Optimizing bitmap indices with efficient compression. *ACM Transactions on Database Systems*, 31(1):1–38, 2006.
- [80] M. Yoon. Aging Bloom filter with two active buffers for dynamic sets. *IEEE Transactions on Knowledge and Data Engineering*, 22(1):134–138, 2010.
- [81] H. Zhang, H. Lim, V. Leis, D. G. Andersen, M. Kaminsky, K. Keeton, and A. Pavlo. SuRF: practical range query filtering with fast succinct tries. In *Proceedings of the 2018 ACM SIGMOD International Conference on Management of Data*, 2018.
- [82] K. Zhang, K. Wang, Y. Yuan, L. Guo, R. Lee, and X. Zhang. Mega-KV: a case for GPUs to maximize the throughput of in-memory key-value stores. *PVLDB*, 8(11):1226–1237, 2015.
- [83] Z. Zhang, Z. Zhu, and X. Zhang. A permutation-based page interleaving scheme to reduce row-buffer conflicts and exploit data locality. In *Proceedings 33rd Annual IEEE/ACM International Symposium on Microarchitecture. MICRO-33 2000*, pages 32–41, 2000.