

# Cuckoo Hashing

Rasmus Pagh<sup>\*</sup>

*BRICS<sup>†</sup>, Department of Computer Science, Aarhus University  
Ny Munkegade Bldg. 540, DK-8000 Århus C, Denmark.  
E-mail: pagh@daimi.au.dk*

and

Flemming Friche Rodler<sup>‡</sup>

*ON-AIR A/S, Digtervejen 9, 9200 Aalborg SV, Denmark.  
E-mail: ffr@onair-dk.com*

We present a simple dictionary with worst case constant lookup time, equaling the theoretical performance of the classic dynamic perfect hashing scheme of Dietzfelbinger et al. (*Dynamic perfect hashing: Upper and lower bounds. SIAM J. Comput.*, 23(4):738–761, 1994). The space usage is similar to that of binary search trees, i.e., three words per key on average.

Besides being conceptually much simpler than previous dynamic dictionaries with worst case constant lookup time, our data structure is interesting in that it does not use perfect hashing, but rather a variant of open addressing where keys can be moved back in their probe sequences.

An implementation inspired by our algorithm, but using weaker hash functions, is found to be quite practical. It is competitive with the best known dictionaries having an average case (but no nontrivial worst case) guarantee.

*Key Words:* data structures, dictionaries, information retrieval, searching, hashing, experiments

<sup>\*</sup> Partially supported by the Future and Emerging Technologies programme of the EU under contract number IST-1999-14186 (ALCOM-FT). This work was initiated while visiting Stanford University.

<sup>†</sup> Basic Research in Computer Science (www.brics.dk), funded by the Danish National Research Foundation.

<sup>‡</sup> This work was done while at Aarhus University.

## 1. INTRODUCTION

The *dictionary* data structure is ubiquitous in computer science. A dictionary is used for maintaining a set  $S$  under insertion and deletion of elements (referred to as *keys*) from a universe  $U$ . Membership queries (“ $x \in S?$ ”) provide access to the data. In case of a positive answer the dictionary also provides a piece of *satellite data* that was associated with  $x$  when it was inserted.

A large theory, partially surveyed in Section 2, is devoted to dictionaries. It is common to study the case where keys are bit strings in  $U = \{0, 1\}^w$  and  $w$  is the word length of the computer (for theoretical purposes modeled as a RAM). Section 3 discusses this restriction. It is usually, though not always, clear how to return associated information once membership has been determined. E.g., in all methods discussed in this paper, the associated information of  $x \in S$  can be stored together with  $x$  in a hash table. Therefore we disregard the time and space used to handle associated information and concentrate on the problem of maintaining  $S$ . In the following we let  $n$  denote  $|S|$ .

The most efficient dictionaries, in theory and in practice, are based on hashing techniques. The main performance parameters are of course lookup time, update time, and space. In theory there is no trade-off between these: One can simultaneously achieve constant lookup time, expected amortized constant update time, and space within a constant factor of the information theoretical minimum of  $B = \log \binom{|U|}{n}$  bits [6]. In practice, however, the various constant factors are crucial for many applications. In particular, lookup time is a critical parameter. It is well known that one can achieve performance arbitrarily close to optimal if a sufficiently sparse hash table is used. Therefore the challenge is to combine speed with a reasonable space usage. In particular, we only consider schemes using  $O(n)$  words of space.

The contribution of this paper is a new, simple hashing scheme called CUCKOO HASHING. A description and analysis of the scheme is given in Section 4, showing that it possesses the same theoretical properties as the dynamic dictionary of Dietzfelbinger et al. [10]. That is, it has worst case constant lookup time and amortized expected constant time for updates. A special feature of the lookup procedure is that (disregarding accesses to a small hash function description) there are just two memory accesses, which are *independent* and can be done in parallel if this is supported by the hardware. Our scheme works for space similar to that of binary search trees, i.e., three words per key in  $S$  on average.

Using weaker hash functions than those required for our analysis, CUCKOO HASHING is very simple to implement. Section 5 describes such an implementation, and reports on extensive experiments and comparisons with the most commonly used methods, having no nontrivial worst case guarantee

on lookup time. It seems that an experiment comparing the most commonly used methods on a modern multi-level memory architecture has not previously been described in the literature. Our experiments show CUCKOO HASHING to be quite competitive, especially when the dictionary is small enough to fit in cache. We thus believe it to be attractive in practice, when a worst case guarantee on lookups is desired.

## 2. PREVIOUS WORK ON LINEAR SPACE DICTIONARIES

Hashing, first described in public literature by Dumey [12], emerged in the 1950s as a space efficient heuristic for fast retrieval of information in sparse tables. Knuth surveys the most important classical hashing methods in [18, Section 6.4]. The most prominent, and the basis for our experiments in Section 5, are CHAINED HASHING (with separate chaining), LINEAR PROBING and DOUBLE HASHING. Judging from leading textbooks on algorithms, Knuth’s selection of algorithms is in agreement with current practice for implementation of general purpose dictionaries. In particular, the excellent cache usage of LINEAR PROBING makes it a prime choice on modern architectures. A more recent scheme called TWO-WAY CHAINING [2] will also be investigated. All schemes are briefly described in Section 5.

### 2.1. Analysis of early hashing schemes

Early theoretical analysis of hashing schemes was done under the assumption that hash function values are uniformly random and independent. Precise analyses of the average and expected worst case behaviors of the abovementioned schemes have been made, see for example [14, 18]. We mention just a few facts, disregarding asymptotically vanishing terms. Note that some figures depend on implementation details – the below hold for the implementations described in Section 5.

We first consider the expected number of memory probes needed by the two “open addressing” schemes to insert a key in a hash table where an  $\alpha$  fraction of the table,  $0 < \alpha < 1$ , is occupied by keys. For LINEAR PROBING the expected number of probes during insertion is  $\frac{1}{2}(1 + \frac{1}{(1-\alpha)^2})$ . This coincides with the expected number of probes for unsuccessful lookups, and with the number of probes needed for looking up the key if there are no subsequent deletions. A deletion rearranges keys to the configuration that would occur if the deleted key had never been inserted. In DOUBLE HASHING the expected cost of an insertion is  $\frac{1}{1-\alpha}$ . As keys are never moved, this coincides with the number of probes needed for looking up the key and for deleting the key. If a key has not been inserted in the hash table since the last rehash, the expected cost of looking it up (unsuccessfully) is  $\frac{1}{1-\beta}$ , where  $\beta$  is the fraction of keys and “deleted” markers in the hash table. If

the key still has a “deleted” marker in the table, the expected cost of the unsuccessful lookup is one probe more.

For CHAINED HASHING with hash table size  $n/\alpha$ , the expected length of the list traversed during an unsuccessful lookup is  $\alpha$ . This means that the expected number of probes needed to insert a new key is  $1 + \alpha$ , which will also be the number of probes needed to subsequently look up the key (note that probes to pointers are not counted). A deletion results in the data structure that would occur if the key had never been inserted.

In terms of number of *probes*, the above implies that, for any given  $\alpha$ , CHAINED HASHING is better than DOUBLE HASHING, which is again better than LINEAR PROBING. It should be noted, however, that the space used by CHAINED HASHING is larger than that in the open addressing schemes for the same  $\alpha$ . The difference depends on the relative sizes of keys and pointers.

The *longest* probe sequence in LINEAR PROBING is of expected length  $\Omega(\log n)$ . For DOUBLE HASHING the longest successful probe sequence is expected to be of length  $\Omega(\log n)$ , and there is in general no sublinear bound on the length of unsuccessful searches. The expected maximum chain length in CHAINED HASHING is  $\Theta(\log n / \log \log n)$ .

Though the above results seem to agree with practice, the randomness assumptions used for the analyses are questionable in applications. Carter and Wegman [7] succeeded in removing such assumptions from the analysis of CHAINED HASHING, introducing the concept of *universal* hash function families. When implemented with a random function from Carter and Wegman’s universal family, chained hashing has constant expected time per dictionary operation (plus an amortized expected constant cost for resizing the table). For a certain efficient hash function family of Siegel [32], LINEAR PROBING and DOUBLE HASHING provably satisfy the above performance bounds [30, 31]. Siegel’s hash functions, summarized in Theorem 3.1, are also used in CUCKOO HASHING.

## 2.2. Key rearrangement schemes

A number of (open addressing) hashing schemes have been proposed that share a key feature with the scheme described in this paper, namely that keys are moved around during insertions [3, 15, 19, 20, 28]. The main focus in these schemes is to reduce the average number of probes needed for finding a key in a (nearly) full table to a constant, rather than the  $O(\log n)$  average exhibited by standard open addressing. This is done by occasionally moving keys forward in their probe sequences.

In our algorithm we rearrange keys in order to reduce the *worst case* number of probes to a constant. A necessary condition for this is reuse of hash function values, i.e., that keys are moved back in their probe sequence. Backward moves were not used in any previous rearrangement scheme,

presumably due to the difficulty that moving keys back does not give a fresh, “random” placement. The thing that allows us to obtain worst case efficient lookups is that we do not deal with full hash tables, but rather hash tables that are less than half full. It was shown in [24] that in this case there exists, with high probability, an arrangement that allows any key to be found in two hash table probes. We show how to efficiently maintain such an arrangement under updates of the key set.

Arrangements of keys with optimal worst case retrieval cost were in fact already considered by Rivest in [28], where a polynomial time algorithm for finding such an arrangement was given. Also, it was shown that if one updates the key set, the expected number of keys that need to be moved to achieve a new optimal arrangement is constant. (The analysis requires that the hash table is sufficiently sparse, and assumes the hash function to be truly random.) This suggests a dictionary that solves a small assignment problem after each insertion and deletion. It follows from [24] and this paper, that Rivest’s dictionary achieved worst case constant lookup time and expected amortized constant update time, 8 years before an algorithm with the same performance and randomness assumption was published by Aho and Lee [1]. Further, we show that Siegel’s hash functions suffice for the analysis. Last but not least, the algorithm we use for rearranging keys is much simpler and more efficient than that suggested by Rivest.

Another key rearrangement scheme with similarities to CUCKOO HASHING is LAST-COME-FIRST-SERVED HASHING [27], which has low variance on search time as its key feature. It uses the same greedy strategy for moving keys as is used in this paper, but there is no reuse of hash function values.

### 2.3. Hashing schemes with worst case lookup guarantee

TWO-WAY CHAINING is an alternative to CHAINED HASHING that offers  $O(\log \log n)$  maximal lookup time with high probability (assuming truly random hash functions). The implementation that we consider represents the lists by fixed size arrays of size  $O(\log \log n)$  (if a longer chain is needed, a rehash is performed). To achieve linear space usage, one must then use a hash table of size  $O(n/\log \log n)$ , implying that the *average* chain length is  $\Omega(\log \log n)$ .

Another scheme with this worst case guarantee is *Multilevel Adaptive Hashing* [5]. However, lookups can be performed in  $O(1)$  worst case time if  $O(\log \log n)$  hash function evaluations, memory probes and comparisons are possible in parallel. This is similar to the scheme described in this paper, though we use only *two* hash function evaluations, memory probes and comparisons.

A dictionary with worst case *constant* lookup time was first obtained by Fredman, Komlós and Szemerédi [13], though it was *static*, i.e., did not support updates. It was later augmented with insertions and deletions

in amortized expected constant time by Dietzfelbinger et al. [10]. Dietzfelbinger and Meyer auf der Heide [11] improved the update performance by exhibiting a dictionary in which operations are done in constant time with high probability, namely at least  $1 - n^{-c}$ , where  $c$  is any constant of our choice. A simpler dictionary with the same properties was later developed [8]. When  $n = |U|^{1-o(1)}$  a space usage of  $O(n)$  words is not within a constant factor of the information theoretical minimum. The dictionary of Brodnik and Munro [6] offers the same performance as [10], using  $O(B)$  bits in all cases.

### 3. PRELIMINARIES

We assume that keys from  $U$  fit exactly in a single machine word, that is,  $U = \{0, 1\}^w$ . A special value  $\perp \in U$  is reserved to signal an empty cell in hash tables. For DOUBLE HASHING an additional special value is used to indicate a deleted key.

Our algorithm uses hash functions from a *universal* family.

**DEFINITION 3.1.** A family  $\{h_i\}_{i \in I}$ ,  $h_i : U \rightarrow R$ , is  $(c, k)$ -*universal* if, for any  $k$  distinct elements  $x_1, \dots, x_k \in U$ , any  $y_1, \dots, y_k \in R$ , and uniformly random  $i \in I$ ,  $\Pr[h_i(x_1) = y_1, \dots, h_i(x_k) = y_k] \leq c/|R|^k$ .

As an example, the family of all functions is  $(1, |U|)$ -universal. However, for implementation purposes one needs families with much more succinct memory representations. A standard construction of a  $(2, k)$ -universal family for range  $R = \{0, \dots, r-1\}$  and prime  $p > \max(2^w, r)$  is

$$\{x \mapsto ((\sum_{l=0}^{k-1} a_l x^l) \bmod p) \bmod r \mid 0 \leq a_0, a_1, \dots, a_{k-1} < p\} . \quad (1)$$

If  $U$  is not too large compared to  $k$ , there exists a space-efficient  $(2, k)$ -universal family due to Siegel [32] that has *constant* evaluation time (however, the constant is not a small one).

**THEOREM 3.1 (Siegel).** *There is a constant  $c$  such that for,  $k = 2^{\Omega(w)}$ , there exists a  $(2, k)$ -universal family that uses space and initialization time  $O(k^c)$ , and which can be evaluated in constant time.*

Our restriction that keys are single words is not a serious one. Longer keys can be mapped to keys of  $O(1)$  words by applying a random function from a  $(O(1), 2)$ -universal family. There is such a family whose functions can be evaluated in time linear in the number of input words [7]. It works by evaluating a function from a  $(O(1), 2)$ -universal family on each word, computing the bitwise exclusive or of the function values. (See [34] for an efficient implementation). Such a function with range  $\{0, 1\}^{2^{\log(n)+c}}$  will,

with probability  $1 - O(2^{-c})$ , be injective on  $S$ . In fact, with constant probability the function is injective on a given *sequence* of  $\Omega(2^{c/2}n)$  consecutive sets in a dictionary of initial size  $n$  (see [10]). When a collision between two elements of  $S$  occurs, everything is rehashed. If a rehash can be done in expected  $O(n)$  time, the amortized expected cost of this is  $O(2^{-c/2})$  per insertion. In this way we can effectively reduce the universe size to  $O(n^2)$ , though the full keys still need to be stored to decide membership. When  $c = O(\log n)$  the reduced keys are of length  $O(\log n)$ . For any  $\epsilon > 0$ , Theorem 3.1 then provides a family of constant time evaluable  $(2, n^{\Omega(1)})$ -universal hash functions, whose functions can be stored using space  $O(n^\epsilon)$ .

#### 4. CUCKOO HASHING

CUCKOO HASHING is a dynamization of a static dictionary described in [24]. The dictionary uses two hash tables,  $T_1$  and  $T_2$ , each of length  $r$ , and two hash functions  $h_1, h_2 : U \rightarrow \{0, \dots, r-1\}$ . Every key  $x \in S$  is stored in cell  $h_1(x)$  of  $T_1$  or  $h_2(x)$  of  $T_2$ , but never in both. Our lookup function is

```

function lookup( $x$ )
    return  $T_1[h_1(x)] = x \vee T_2[h_2(x)] = x$ 
end

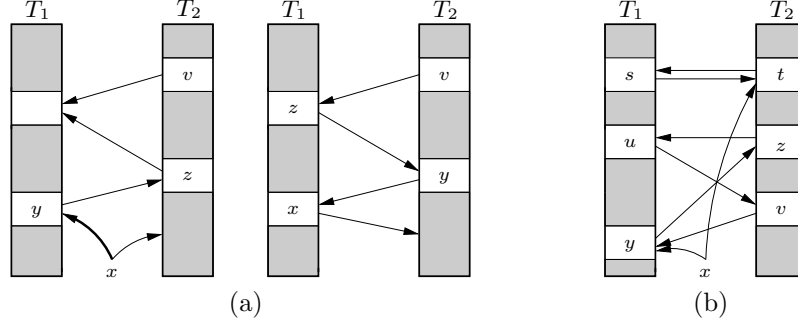
```

Two table accesses for lookup is in fact optimal among all dictionaries using linear space, except for special cases, see [24].

**Remark:** The idea of storing keys in one out of two places given by hash functions previously appeared in [16] in the context of PRAM simulation, and in [2] for TWO-WAY CHAINING.

It is shown in [24] that if  $r \geq (1 + \epsilon)n$  for some constant  $\epsilon > 0$  (i.e., the tables are to be a bit less than half full), and  $h_1, h_2$  are picked uniformly at random from an  $(O(1), O(\log n))$ -universal family, the probability that there is no way of arranging the keys of  $S$  according to  $h_1$  and  $h_2$  is  $O(1/n)$ . By the discussion in Section 3 we may assume without loss of generality that there is such a family, with constant evaluation time and negligible space usage. A suitable arrangement of the keys was shown in [24] to be computable in expected linear time, by a reduction to 2-SAT.

We now consider a simple dynamization of the above. Deletion is of course simple to perform in constant time, not counting the possible cost of shrinking the tables if they are becoming too sparse. As for insertion, it turns out that the “cuckoo approach”, kicking other keys away until every key has its own “nest”, works very well. Specifically, if  $x$  is to be inserted we first see if cell  $h_1(x)$  of  $T_1$  is occupied. If not, we are done. Otherwise we set  $T_1[h_1(x)] \leftarrow x$  anyway, thus making the previous occupant “nestless”.



**FIG. 1.** Examples of CUCKOO HASHING insertion. Arrows show possibilities for moving keys. (a) Key  $x$  is successfully inserted by moving keys  $y$  and  $z$  from one table to the other. (b) Key  $x$  cannot be accommodated and a rehash is necessary.

This key is then inserted in  $T_2$  in the same way, and so forth iteratively, see Figure 1(a). It may happen that this process loops, as shown in Figure 1(b). Therefore the number of iterations is bounded by a value “MaxLoop” to be specified in Section 4.1. If this number of iterations is reached, everything is rehashed with new hash functions, and we try once again to accommodate the nestless key.

Using the notation  $x \leftrightarrow y$  to express that the values of variables  $x$  and  $y$  are swapped, the following code summarizes the insertion procedure.

```

procedure insert( $x$ )
  if lookup( $x$ ) then return
  loop MaxLoop times
    if  $T_1[h_1(x)] = \perp$  then {  $T_1[h_1(x)] \leftarrow x$ ; return }
     $x \leftrightarrow T_1[h_1(x)]$ 
    if  $T_2[h_2(x)] = \perp$  then {  $T_2[h_2(x)] \leftarrow x$ ; return }
     $x \leftrightarrow T_2[h_2(x)]$ 
  end loop
  rehash(); insert( $x$ )
end

```

The above procedure assumes that each table remains larger than  $(1 + \epsilon)n$  cells. When no such bound is known, a test must be done to find out when a rehash to larger tables is needed. Resizing of tables can be done in amortized expected constant time per update by the usual doubling/halving technique (see, e.g., [10]). The hash functions used will be  $(O(1), \text{MaxLoop})$ -universal, which means that they will act almost like

truly random functions on any set of keys processed during the insertion loop.

The lookup call preceding the insertion loop ensures robustness if the key to be inserted is already in the dictionary. A slightly faster implementation can be obtained if this is known not to occur.

Note that the insertion procedure is biased towards inserting keys in  $T_1$ . As will be seen in Section 5 this leads to faster successful lookups, due to more keys being found in  $T_1$ . This effect is even more pronounced if one uses an *asymmetric* scheme where  $T_1$  is larger than  $T_2$ . In both cases, the insertion time is only slightly worse than that of a completely symmetric implementation. Another variant is to use a single table for both hash functions, but this requires keeping track of the hash function according to which each key is placed. In the following we consider just the symmetric two table scheme.

#### 4.1. Analysis

Our analysis of the insertion procedure has three main parts:

1. We first exhibit some useful characteristics of the behavior of the insertion procedure.
2. We then derive a bound on the probability that the insertion procedure uses at least  $t$  iterations.
3. Finally we argue that the procedure uses expected amortized constant time.

#### Behavior of the Insertion Procedure

The simplest behavior of the insertion procedure occurs when it does not visit any hash table cell more than once. In this case it simply runs through a sequence  $x_1, x_2, \dots$ , of nestless keys with no repetitions, moving each key from one table to the other.

If, at some point, the insertion procedure returns to a previously visited cell, the behavior is more complicated, as shown in Figure 2. The key  $x_i$  in the first previously visited cell will become nestless for the second time (occurring at positions  $i$  and  $j > i$  in the sequence) and be put back in its original cell. Subsequently all keys  $x_{i-1}, \dots, x_1$  will be moved back where they were at the start of the insertion (assuming that the maximum number of iterations is not reached). In particular,  $x_1$  will end up nestless again, and the procedure will try placing it in the second table. At some point after this there appears a nestless key  $x_l$  that is either moved to a vacant cell or a previously visited cell (again assuming that the maximum number of iterations is not reached). In the former case the procedure terminates. In the latter case a rehash must be performed, since we have a “closed

loop” of  $l - i + 1$  keys hashing to only  $l - i$  cells. This means that the loop will run for the maximum number of iterations, followed by a rehash.

**LEMMA 4.1.** *Suppose that the insertion procedure does not enter a closed loop. Then for any prefix  $x_1, x_2, \dots, x_p$  of the sequence of nestless keys, there must be a subsequence of at least  $p/3$  consecutive keys without repetitions, starting with an occurrence of the key  $x_1$ , i.e., the key being inserted.*

*Proof.* In the case where the insertion procedure never returns to a previously visited cell, the prefix itself is a sequence of  $p$  distinct nestless keys starting with  $x_1$ . Otherwise, the sequence of nestless keys is as shown in Figure 2. If  $p < i + j$ , the first  $j - 1 \geq \frac{i+j-1}{2} \geq p/2$  nestless keys form the desired sequence. For  $p \geq i + j$ , one of the sequences  $x_1, \dots, x_{j-1}$  and  $x_{j+i-1}, \dots, x_p$  must have length at least  $p/3$ . ■

### Probability Bounds

We now consider the probability that the insertion loop runs for at least  $t$  iterations. For  $t > \text{MaxLoop}$  the probability is of course 0. Otherwise, by the above analysis, iteration number  $t$  is performed in two (not mutually exclusive) situations:

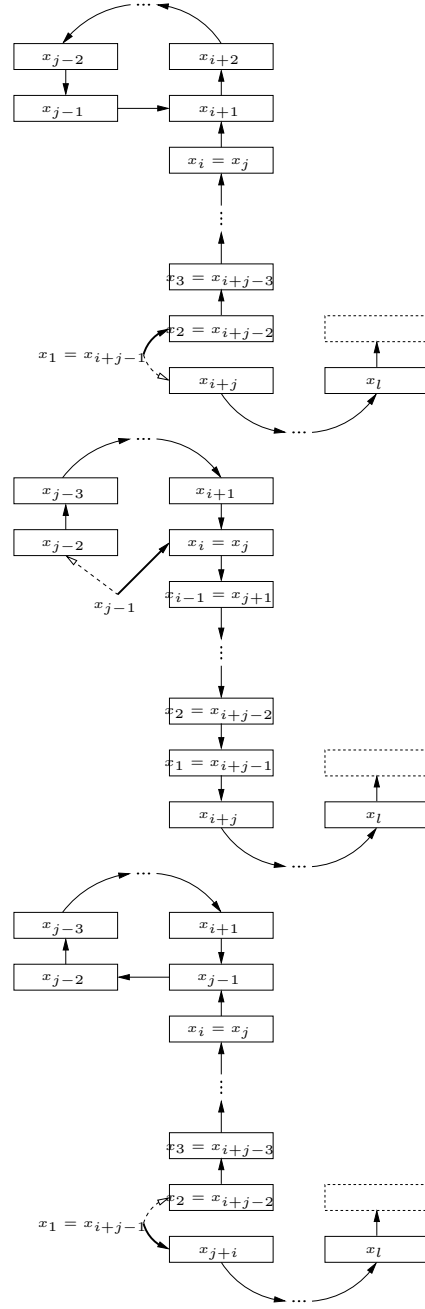
1. The insertion procedure has entered a “closed loop”, i.e.,  $x_l$  in Figure 2 was moved to a previously visited cell, for  $l \leq 2t$ .
2. The insertion procedure has processed a sequence of at least  $(2t - 1)/3$  consecutive nestless keys starting with the newly inserted key.

In the first situation let  $v \leq l$  denote the number of distinct nestless keys. The number of ways in which the closed loop can be formed is less than  $v^2 r^{v-1} n^{v-1}$  ( $v^2$  possible values for  $i$  and  $j$ ,  $r^{v-1}$  possible choices of cells, and  $n^{v-1}$  possible choices of keys other than  $x_1$ ). Since  $v \leq \text{MaxLoop}$ , the hash functions are  $(c, v)$ -universal. This means that each possibility occurs with probability at most  $c^2 r^{-2v}$ . Summing over all possible values of  $v$ , and using  $r/n > 1 + \epsilon$ , we get that the probability of situation 1 is at most:

$$\sum_{v=3}^l v^2 r^{v-1} n^{v-1} c^2 r^{-2v} \leq \frac{c^2}{rn} \sum_{v=3}^{\infty} v^2 (n/r)^v < \frac{13 c^2 / \epsilon}{n^2} = O(1/n^2) .$$

The above derivation follows a suggestion of Sanders and Vöcking [29], and improves the  $O(1/n)$  bound in the conference version of this paper [25].

In the second situation there is a sequence of distinct nestless keys  $b_1, \dots, b_v$ ,  $v \geq (2t - 1)/3$ , such that  $b_1$  is the key to be inserted, and



**FIG. 2.** Stages of an insertion of key  $x_1$ , involving the movement of keys  $x_1, \dots, x_l$ . Boxes correspond to cells in either of the two tables, and arcs show possibilities for moving keys. A bold arc shows where the nestless key is to be inserted.

such that for either  $(\beta_1, \beta_2) = (1, 2)$  or  $(\beta_1, \beta_2) = (2, 1)$ :

$$h_{\beta_1}(b_1) = h_{\beta_1}(b_2), h_{\beta_2}(b_2) = h_{\beta_2}(b_3), h_{\beta_1}(b_3) = h_{\beta_1}(b_4), \dots \quad (2)$$

Given  $b_1$  there are at most  $n^{v-1}$  possible sequences of  $v$  distinct keys. For any such sequence and any of the two choices of  $(\beta_1, \beta_2)$ , the probability that the  $b-1$  equations in (2) hold is bounded by  $c r^{-(v-1)}$ , since the hash functions were chosen from a  $(c, \text{MaxLoop})$ -universal family. Hence the probability that there is *any* sequence of length  $v$  satisfying (2), and thus the probability of situation 2, is bounded by

$$2c^2 (n/r)^{v-1} \leq 2c^2 (1 + \epsilon)^{-(2t-1)/3+1} . \quad (3)$$

### Concluding the Analysis

From now on we restrict attention to  $\text{MaxLoop} = O(n)$ . From (3) it follows that the expected number of iterations in the insertion loop is bounded by

$$\begin{aligned} 1 + \sum_{t=2}^{\text{MaxLoop}} 2c^2 (1 + \epsilon)^{-(2t-1)/3+1} + O(1/n^2) \\ \leq 1 + O\left(\frac{\text{MaxLoop}}{n^2}\right) + 2c^2 \sum_{t=0}^{\infty} ((1 + \epsilon)^{-2/3})^t \\ = O\left(1 + \frac{1}{1 - (1 + \epsilon)^{-2/3}}\right) \\ = O(1 + 1/\epsilon) . \end{aligned} \quad (4)$$

Finally, we consider the cost of rehashing, which occurs if the insertion loop runs for  $t = \text{MaxLoop}$  iterations. By the previous section, the probability that this happens because of entering a closed loop is  $O(1/n^2)$ . Setting  $\text{MaxLoop} = \lceil 3 \log_{1+\epsilon} n \rceil$ , the probability of rehashing without entering a closed loop is, by (3), at most

$$2c^2 (1 + \epsilon)^{-(2 \text{MaxLoop} - 1)/3+1} = O(1/n^2)$$

Altogether, the probability that any given insertion causes a rehash is  $O(1/n^2)$ . In particular, the  $n$  insertions performed during a rehash all succeed (i.e., cause no further rehash) with probability  $1 - O(1/n)$ . The expected time used per insertion is  $O(1)$ , so the total expected time for trying to insert all keys is  $O(n)$ . As the probability of having to start over with new hash functions is bounded away from 1, the total expected time for a rehash is  $O(n)$ . Thus, for any insertion the expected time used for rehashing is  $O(1/n)$ .

Summing up, we have shown that the expected time for insertion is bounded by a constant. The small probability of rehashing in fact implies that also the *variance* of the insertion time is constant.

## 5. EXPERIMENTS

To examine the practicality of CUCKOO HASHING we experimentally compare it to three well known hashing methods, as described in [18, Section 6.4]: CHAINED HASHING (with separate chaining), LINEAR PROBING and DOUBLE HASHING. We also consider TWO-WAY CHAINING [2].

The first three methods all attempt to store a key  $x$  at position  $h(x)$  in a hash table. They differ in the way collisions are resolved, i.e., in what happens when two or more keys hash to the same location.

CHAINED HASHING. A chained list is used to store all keys hashing to a given location.

LINEAR PROBING. A key is stored in the next empty table entry. Lookup of key  $x$  is done by scanning the table beginning at  $h(x)$  and ending when either  $x$  or an empty table entry is found. When deleting, some keys may have to be moved back in order to fill the hole in the lookup sequence, see [18, Algorithm R] for details.

DOUBLE HASHING. Insertion and lookup are similar to LINEAR PROBING, but instead of searching for the next position one step at a time, a second hash function value is used to determine the step size. Deletions are handled by putting a “deleted” marker in the cell of the deleted key. Lookups skip over deleted cells, while insertions overwrite them.

The fourth method, TWO-WAY CHAINING, can be described as two instances of CHAINED HASHING. A key is inserted in one of the two hash tables, namely the one where it hashes to the shortest chain. A cache-friendly implementation, as recently suggested in [4], is to simply make each chained list a short, fixed size array. If a longer list is needed, a rehash must be performed.

### 5.1. Previous Experimental Results.

Although the dictionaries with worst case constant lookup time surveyed in Sect. 2 leave little to improve from a theoretical point of view, large constant factors and complicated implementation hinder their direct practical use. For example, in the “dynamic perfect hashing” scheme of [10] the upper bound on space is  $35n$  words. The authors of [10] refer to a more practical variant due to Wenzel that uses space comparable to that of binary search trees.

According to [17] the implementation of this variant in the LEDA library [23], described in [35], has average insertion time larger than that

of AVL trees for  $n \leq 2^{17}$ , and more than four times slower than insertions in chained hashing<sup>1</sup>. The experimental results listed in [23, Table 5.2] show a gap of more than a factor of 6 between the update performance of chained hashing and dynamic perfect hashing, and a factor of more than 2 for lookups<sup>2</sup>.

Silverstein [33] reports that the space upper bound of the dynamic perfect hashing scheme of [10] is quite pessimistic compared to what can be observed when run on a subset of the DIMACS dictionary tests [22]. He goes on to explore ways of improving space as well as time, improving both the observed time and space by a factor of roughly three. Still, the improved scheme needs 2 to 3 times more space than an implementation of linear probing to achieve similar time per operation. Silverstein also considers versions of the data structures with packed representations of the hash tables. In this setting the dynamic perfect hashing scheme was more than 50% slower than linear probing, using roughly the same amount of space.

It seems that recent experimental work on “classical” dictionaries (that do not have worst case constant lookup time) is quite limited. In [17] it is reported that chained hashing is superior to an implementation of dynamic perfect hashing in terms of both memory usage and speed.

## 5.2. Data Structure Design and Implementation

We consider positive 32 bit signed integer keys and use 0 as  $\perp$ . The data structures are *robust* in that they correctly handle attempts to insert an element already in the set, and attempts to delete an element not in the set. During rehashes this is known not to occur and slightly faster versions of the insertion procedure are used.

Our focus is on achieving high performance dictionary operations with a reasonable space usage. By the *load factor* of a dictionary we will understand the size of the set relative to the memory used<sup>3</sup>. As seen in [18, Figure 44] the speed of LINEAR PROBING and DOUBLE HASHING degrades rapidly for load factors above  $1/2$ . On the other hand, none of the schemes improve much for load factors below  $1/4$ . As CUCKOO HASHING only works when the size of each table is larger than the size of the set, we can only perform a comparison for load factors less than  $1/2$ . To allow for doubling and halving of the table size, we allow the load factor to vary between  $1/5$  and  $1/2$ , focusing especially on the “typical” load factor of  $1/3$ . For CUCKOO HASHING and TWO-WAY CHAINING there is a chance that an insertion may fail, causing a “forced rehash”. If the load factor is larger

---

<sup>1</sup>On a Linux PC with an Intel® Pentium® 120 MHz processor.

<sup>2</sup>On a 300 MHz SUN ULTRA SPARC.

<sup>3</sup>For CHAINED HASHING, the notion of load factor traditionally disregards the space used for chained lists, but we desire equal load factors to imply equal memory usage.

than a certain threshold, somewhat arbitrarily set to  $5/12$ , we use the opportunity to double the table size. By our experiments this only slightly decreases the average load factor.

Apart from CHAINED HASHING, the schemes considered have in common the fact that they have only been analyzed under randomness assumptions that are currently impractical to realize. However, experience shows that rather simple and efficient hash function families yield performance close to that predicted under stronger randomness assumptions. We use a function family from [9] with range  $\{0,1\}^q$  for positive integer  $q$ . For every odd  $a$ ,  $0 < a < 2^w$ , the family contains the function  $h_a(x) = (ax \bmod 2^w) \div 2^{w-q}$ . Note that evaluation can be done very efficiently by a 32 bit multiplication and a shift. However, this choice of hash function restricts us to consider hash tables whose sizes are powers of two. A random function from the family (chosen using C's `rand` function) appears to work fine with all schemes except CUCKOO HASHING. For CUCKOO HASHING we experimented with various hash functions and found that CUCKOO HASHING was rather sensitive to the choice of hash function. It turned out that the exclusive or of three independently chosen functions from the family of [9] was fast and worked well. We have no good explanation for this phenomenon. For all schemes, various alternative hash families were tried, with a decrease in performance.

All methods have been implemented in C. We have striven to obtain the fastest possible implementation of each scheme. Specific choices made and details differing from the references are:

CHAINED HASHING. C's `malloc` and `free` functions were found to be a performance bottleneck, so a simple “freelist” memory allocation scheme is used. Half of the allocated memory is used for the hash table, and half for list elements. If the data structure runs out of free list elements, its size is doubled. We store the first element of each linked list directly in the hash table. This often saves one cache miss. It also slightly improves memory utilization, in the expected sense. This is because every non-empty chained list is one element shorter and because we expect more than half of the hash table cells to contain a linked list for the load factors considered here.

DOUBLE HASHING. To prevent the tables from clogging up with deleted cells, resulting in poor performance for unsuccessful lookups, all keys are rehashed when  $2/3$  of the hash table is occupied by keys and “deleted” markers. The fraction  $2/3$  was found to give a good tradeoff between the time for insertion and unsuccessful lookups.

LINEAR PROBING. Our first implementation, like that in [33], employed deletion markers. However, we found that using the deletion method described in [18, Algorithm R] was considerably faster, as far fewer rehashes were needed.

**TWO-WAY CHAINING.** We allow four keys in each bucket. This is enough to keep the probability of a forced rehash low for hundreds of thousands of keys, by the results in [4]. For larger collections of keys one should allow more keys in each bucket, resulting in general performance degradation.

**CUCKOO HASHING.** The architecture on which we experimented could not parallelize the two memory accesses in lookups. Therefore we only evaluate the second hash function after the first memory lookup has shown unsuccessful.

Some experiments were done with variants of CUCKOO HASHING. In particular, we considered **ASYMMETRIC CUCKOO**, in which the first table is twice the size of the second one. This results in more keys residing in the first table, thus giving a slightly better average performance for successful lookups. For example, after a long sequence of alternate insertions and deletions at load factor  $1/3$ , we found that about 76% of the elements resided in the first table of **ASYMMETRIC CUCKOO**, as opposed to 63% for **CUCKOO HASHING**. There is no significant slowdown for other operations. We will describe the results for **ASYMMETRIC CUCKOO** when they differ significantly from those of **CUCKOO HASHING**.

### 5.3. Setup

Our experiments were performed on a PC running Linux (kernel version 2.2) with an 800 MHz Intel® Pentium® III processor, and 256 MB of memory (PC100 RAM). The processor has a 16 KB level 1 data cache and a 256 KB level 2 “advanced transfer” cache. As will be seen, our results nicely fit a simple model parameterized by the cost of a cache miss and the expected number of probes to “random” locations. They are thus believed to have significance for other hardware configurations. An advantage of using the Pentium® processor for timing experiments is its `rdtsc` instruction which can be used to measure time in clock cycles. This gives access to very precise data on the behavior of algorithms. In our case it also supplies a way of discarding measurements significantly disturbed by interrupts from hardware devices or the process scheduler, as these show up as a small group of timings significantly separated from all other timings. Programs were compiled using the `gcc` compiler version 2.95.2, using optimization flags `-O9 -DCPU=586 -march=i586 -fomit-frame-pointer -finline-functions -fforce-mem -funroll-loops -fno-rtti`. As mentioned earlier, we use a global clock cycle counter to time operations. If the number of clock cycles spent exceeds 5000, and there was no rehash, we conclude that the call was interrupted, and disregard the result (it was empirically observed that no operation ever took between 2000 and 5000 clock cycles). If a rehash is made, we have no way of filtering away time spent in interrupts. However, all tests were made on a machine with no irrelevant user processes, so disturbances should be minimal. On our machine it took

32 clock cycles to call the `rdtsc` instruction. These clock cycles have been subtracted from the results.

#### 5.4. Results

##### Dictionaries of Stable Size

Our first test was designed to model the situation in which the size of the dictionary is not changing too much. It considers a sequence of mixed operations generated at random. We constructed the test operation sequences from a collection of high quality random bits publicly available on the Internet [21]. The sequences start by insertion of  $n$  distinct random keys, followed by  $3n$  times four operations: A random unsuccessful lookup, a random successful lookup, a random deletion, and a random insertion. We timed the operations in the “equilibrium”, where the number of elements is stable. For load factor  $1/3$  our results appear in Figure 3, which shows an average over 10 runs. As LINEAR PROBING was consistently faster than DOUBLE HASHING, we chose it as the sole open addressing scheme in the plots. Time for forced rehashes was added to the insertion time. The results had a large variance, over the 10 runs, for sets of size  $2^{12}$  to  $2^{16}$ . Outside this range the extreme values deviated from the average by less than about 7%. The large variance sets in when the data structure starts to fill the level 2 cache. We believe it is due to other processes evicting parts of the data structure from cache.

As can be seen, the time for lookups is almost identical for all schemes as long as the entire data structure fits in level 2 cache, i.e., for  $n < 2^{16}/3$ . After this the average number of random memory accesses (with the probability of a cache miss approaching 1) shows up. This makes linear probing an average case winner, with CUCKOO HASHING and TWO-WAY CHAINING following about 40 clock cycles behind. For insertion the number of random memory accesses again dominates the picture for large sets, while the higher number of in-cache accesses and more computation makes CUCKOO HASHING, and in particular TWO-WAY chaining, relatively slow for small sets. The cost of forced rehashes sets in for TWO-WAY CHAINING for sets of more than a million elements, at which point better results may have been obtained by a larger bucket size. For deletion CHAINED HASHING lags behind for large sets due to random memory accesses when freeing list elements, while the simplicity of CUCKOO HASHING makes it the fastest scheme. We suspect that the slight rise in time for the largest sets in the test is due to saturation of the bus, as the machine runs out of memory and begins swapping.

At this point we should mention that the good cache utilization of LINEAR PROBING and TWO-WAY CHAINING depends on the cache lines being considerably larger than keys (and any associated information placed to-

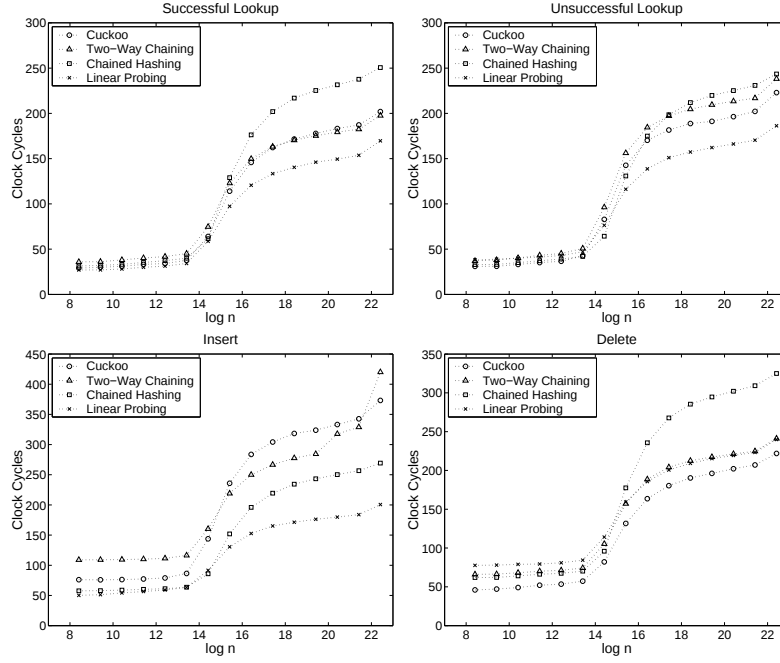
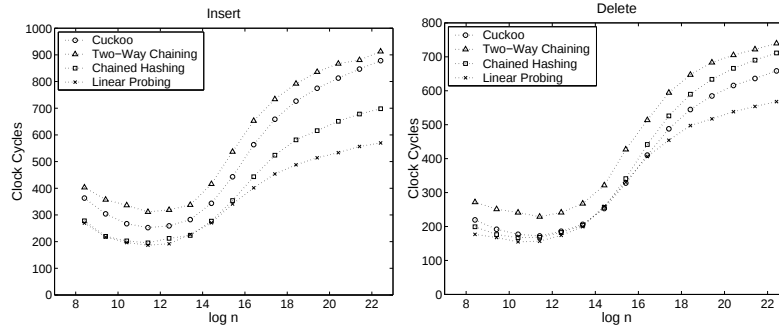


FIG. 3. The average time per operation in equilibrium for load factor 1/3.

gether with keys). If this is not the case, it causes the number of cache misses to rise significantly. The other schemes discussed here do not deteriorate in this way.

### Growing and Shrinking Dictionaries

The second test concerns the cost of insertions in growing dictionaries and deletions in shrinking dictionaries. This will be different from the above due to the cost of rehashes. Together with Figure 3 this should give a fairly complete picture of the performance of the data structures under general sequences of operations. The first operation sequence inserts  $n$  distinct random keys, while the second one deletes them. The plot is shown in Figure 4. For small sets the time per operation seems unstable, and dominated by memory allocation overhead (if minimum table size  $2^{10}$  is used, the curves become monotone). For sets of more than  $2^{12}$  elements the largest deviation from the averages over 10 runs was about 6%. Disregarding the constant minimum amount of memory used by any dictionary, the average load factor during insertions was within 2% of 1/3 for all schemes except CHAINED HASHING whose average load factor was about 0.31. Dur-



**FIG. 4.** The average time per insertion/deletion in a growing/shrinking dictionary for average load factor  $\approx 1/3$ .

ing deletions all schemes had average load factor 0.28. Again the fastest method is LINEAR PROBING, followed by CHAINED HASHING and CUCKOO HASHING. This is largely due to the cost of rehashes.

### DIMACS Tests

Access to data in a dictionary is rarely random in practice. In particular, the cache is more helpful than in the above random tests, for example due to repeated lookups of the same key, and deletion of short-lived keys. As a rule of thumb, the time for such operations will be similar to the time when all of the data structure is in cache. To perform actual tests of the dictionaries on more realistic data, we chose a representative subset of the dictionary tests of the 5th DIMACS implementation challenge [22]. The tests involving string keys were preprocessed by hashing strings to 32 bit integers, as described in Section 3. This preserves, with high probability, the access pattern to keys. For each test we recorded the average time per operation, not including the time used for preprocessing. The minimum and maximum of six runs can be found in Tables 5 and 6, which also lists the average load factor. Linear probing is again the fastest, but mostly just 20-30% faster than the CUCKOO schemes.

### The Number of Cache Misses During Insertion

We have seen that the number of random memory accesses (i.e., cache misses) is critical to the performance of hashing schemes. Whereas there is a very precise understanding of the probe behavior of the classic schemes (under suitable randomness assumptions), the analysis of the expected time for insertions in Section 4.1 is rather crude, establishing just a constant upper bound. One reason that our calculation does not give a very tight bound is that we use a pessimistic estimate on the number of key moves

|          | Joyce   |       | Eddington |       |
|----------|---------|-------|-----------|-------|
| LINEAR   | 42 - 45 | (.35) | 26 - 27   | (.40) |
| DOUBLE   | 48 - 53 | (.35) | 32 - 35   | (.40) |
| CHAINED  | 49 - 52 | (.31) | 36 - 38   | (.28) |
| A.CUCKOO | 47 - 50 | (.33) | 37 - 39   | (.32) |
| CUCKOO   | 57 - 63 | (.35) | 41 - 45   | (.40) |
| TWO-WAY  | 82 - 84 | (.34) | 51 - 53   | (.40) |

**FIG. 5.** Average clock cycles per operation and load factors for two DIMACS string tests.

|          | 3.11-Q-1  |       | Smalltalk-2 |       | 3.2-Y-1   |       |
|----------|-----------|-------|-------------|-------|-----------|-------|
| LINEAR   | 99 - 103  | (.30) | 68 - 72     | (.29) | 85 - 88   | (.32) |
| DOUBLE   | 116 - 142 | (.30) | 77 - 79     | (.29) | 98 - 102  | (.32) |
| CHAINED  | 113 - 121 | (.30) | 78 - 82     | (.29) | 90 - 93   | (.31) |
| A.CUCKOO | 166 - 168 | (.29) | 87 - 95     | (.29) | 95 - 96   | (.32) |
| CUCKOO   | 139 - 143 | (.30) | 90 - 96     | (.29) | 104 - 108 | (.32) |
| TWO-WAY  | 159 - 199 | (.30) | 111 - 113   | (.29) | 133 - 138 | (.32) |

**FIG. 6.** Average clock cycles per operation and load factors for three DIMACS integer tests.

needed to accommodate a new element in the dictionary. Often a free cell will be found even though it could have been occupied by another key in the dictionary. We also pessimistically assume that a large fraction of key moves will be spent backtracking from an unsuccessful attempt to place the new key in the first table.

Figure 7 shows experimentally determined values for the average number of probes during insertion for various schemes and load factors below  $1/2$ . We disregard reads and writes to locations known to be in cache, and the cost of rehashes. Measurements were made in “equilibrium” after  $10^5$  insertions and deletions, using tables of size  $2^{15}$  and truly random hash function values. It is believed that this curve is independent of the table size (up to vanishing terms). The curve for LINEAR PROBING does not appear, as the number of non-cached memory accesses depends on cache architecture (length of the cache line), but it is typically very close to 1. The curve for CUCKOO HASHING seems to be  $2 + 1/(4 + 8\alpha) \approx 2 + 1/(4\epsilon)$ . This is in good correspondence with (4) of the analysis in Section 4.1. As noted in Section 4, the insertion algorithm of CUCKOO HASHING is biased towards inserting keys in  $T_1$ . If we instead of starting the insertion in  $T_1$  choose the start table at random, the number of cache misses decreases slightly for insertion. This is because the number of free cells in  $T_1$  increases as the load balance becomes even. However, this also means a slight increase in lookup time. Also note that since insertion checks if the element is already inserted, CUCKOO HASHING uses at least two cache misses. The initial lookup can be exploited to get a small improvement in insertion performance, by inserting right away when *either* cell  $T_1[h_1(x)]$  or  $T_2[h_2(x)]$  is vacant. It should be remarked that the highest possible load factor for TWO-WAY CHAINING is  $O(1/\log \log n)$ .

Since lookup is very similar to insertion in CHAINED HASHING, one could think that the number of cache misses would be equal for the two operations. However, in our implementation, obtaining a free cell from the freelist may result in an extra cache miss. This is the reason why the curve for CHAINED HASHING in the figure differs from a similar plot in Knuth [18, Figure 44].

### 5.9. Model

In this section we look at a simple model of the time it takes to perform a dictionary operation, and note that our results can be explained in terms of this model. On a modern computer, memory speed is often the bottleneck. Since the operations of the investigated hashing methods mainly perform reads and writes to memory, we will assume that cache misses constitute the dominant part of the time needed to execute a dictionary operation.

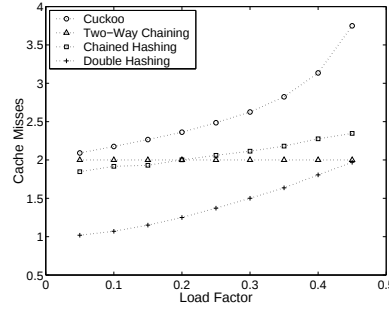


FIG. 7. The average number of random memory accesses for insertion.

This leads to the following model of the time per operation.

$$\text{Time} = O + N \cdot R \cdot (1 - C/T) , \quad (5)$$

where the parameters of the model are described by

- $O$  – Constant overhead of the operation.
- $R$  – Average number of memory accesses.
- $C$  – Cache size.
- $T$  – Size of the hash tables.
- $N$  – Cost of a non-cache read.

The term  $R \cdot (1 - C/T)$  is the expected number of cache misses for the operations with  $(1 - C/T)$  being the probability that a random probe into the tables results in a cache miss. Note that the model is not valid when the table size  $T$  is smaller than the cache size  $C$ . The size  $C$  of the cache and the size  $T$  of the dictionary are well known. From Figure 7 we can, for the various hashing schemes and for a load factor of  $1/3$ , read the average number  $R$  of memory accesses needed for inserting an element. Note that several accesses to consecutive elements in the hash table are counted as one random access, since the other accesses are then in cache. The overhead of an operation,  $O$ , and the cost of a cache miss,  $N$ , are unknown factors that we will estimate.

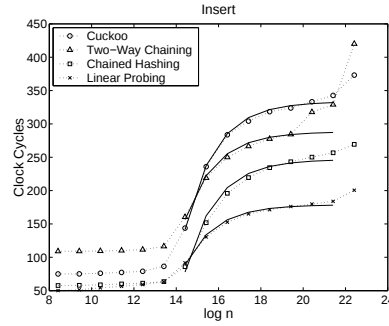
Performing experiments, reading and writing to and from memory, we observed that the time for a read or a write to a location known not to be in cache could vary dramatically depending on the state of the cache. For example, when a cache line is to be used for a new read, the time used is considerably higher if the old contents of the cache line has been written to, since the old contents must then first be moved to memory. For this reason we expect parameter  $N$  to depend slightly on both the particular dictionary

methods and the combination of dictionary operations. This means that  $R$  and  $T$  are the only parameters not dependent on the methods used.

| Method  | $N$ | $O$ |
|---------|-----|-----|
| CUCKOO  | 71  | 142 |
| TWO-WAY | 66  | 157 |
| CHAINED | 79  | 78  |
| LINEAR  | 88  | 89  |
| Average | 76  | -   |

**FIG. 8.** Estimated parameters according to the model for insertion.

Using the timings from Figure 3 and the average number of cache misses for insert observed in Figure 7, we estimated  $N$  and  $O$  for the four hashing schemes. As mentioned, we believe the slight rise in time for the largest sets in the tests of Figure 3 to be caused by other non-cache related factors. So since the model is only valid for  $T \geq 2^{16}$ , the two parameters were estimated for timings with  $2^{16} \leq T \leq 2^{23}$ . The results are shown in Table 8. As can be seen from the table, the cost of a cache miss varies slightly from method to method. The largest deviation from the average is about 15%.



**FIG. 9.** Model versus observed data.

To investigate the accuracy of our model we plotted in Figure 9 the estimated curves for insertion together with the observed curves used for estimating the parameters. As can be seen, the simple model explains the observed values quite nicely. The situation for the other operations is similar.

Having said this, we must admit that the values of  $N$  and  $O$  estimated for the schemes cannot be accounted for. In particular, it is clear that the true behavior of the schemes is more complicated than suggested by the model.

## 6. CONCLUSION

We have presented a new dictionary with worst case constant lookup time. It is very simple to implement, and has average case performance comparable to the best previous dictionaries. Earlier schemes with worst case constant lookup time were more complicated to implement and had worse average case performance. Several challenges remain. First of all an explicit practical hash function family that is provably good for the scheme has yet to be found. For example, future advances in explicit expander graph construction could make Siegel's hash functions practical. Secondly, we lack a precise understanding of why the scheme exhibits low constant factors. In particular, the curve of Figure 7 needs to be explained. Another point to investigate is whether using more tables yields practical dictionaries. Experiments in [26] suggest that space utilization could be improved to more than 80%, but it remains to be seen how this would affect insertion performance.

**Acknowledgements.** The authors would like to thank Andrei Broder, Peter Sanders and Berthold Vöcking for useful comments and discussions on CUCKOO HASHING.

## REFERENCES

1. Alfred V. Aho and David Lee. Storing a dynamic sparse table. In *Proceedings of the 23rd Annual Symposium on Foundations of Computer Science (FOCS '82)*, pages 55–60, Los Alamitos, CA, 1986. IEEE Comput. Soc. Press.
2. Yossi Azar, Andrei Z. Broder, Anna R. Karlin, and Eli Upfal. Balanced allocations. *SIAM J. Comput.*, 29(1):180–200 (electronic), 1999.
3. Richard P. Brent. Reducing the retrieval time of scatter storage techniques. *Communications of the ACM*, 16(2):105–109, February 1973. Modification of open addressing with double hashing to reduce the average number of probes for a successful search.
4. Andrei Broder and Michael Mitzenmacher. Using multiple hash functions to improve IP lookups. *Proceedings of INFOCOM 2001*, 2001.
5. Andrei Z. Broder and Anna R. Karlin. Multilevel adaptive hashing. In *Proceedings of the 1st Annual ACM-SIAM Symposium on Discrete Algorithms (SODA '90)*, pages 43–53. ACM Press, New York, 2000.
6. Andrej Brodnik and J. Ian Munro. Membership in constant time and almost-minimum space. *SIAM J. Comput.*, 28(5):1627–1640 (electronic), 1999.
7. J. Lawrence Carter and Mark N. Wegman. Universal classes of hash functions. *J. Comput. System Sci.*, 18(2):143–154, 1979.

8. Martin Dietzfelbinger, Joseph Gil, Yossi Matias, and Nicholas Pippenger. Polynomial hash functions are reliable (extended abstract). In *Proceedings of the 19th International Colloquium on Automata, Languages and Programming (ICALP '92)*, volume 623 of *Lecture Notes in Computer Science*, pages 235–246. Springer-Verlag, Berlin, 1992.
9. Martin Dietzfelbinger, Torben Hagerup, Jyrki Katajainen, and Martti Penttonen. A reliable randomized algorithm for the closest-pair problem. *Journal of Algorithms*, 25(1):19–51, 1997. doi:10.1006/jagm.1997.0873.
10. Martin Dietzfelbinger, Anna Karlin, Kurt Mehlhorn, Friedhelm Meyer auf der Heide, Hans Rohnert, and Robert E. Tarjan. Dynamic perfect hashing: Upper and lower bounds. *SIAM J. Comput.*, 23(4):738–761, 1994.
11. Martin Dietzfelbinger and Friedhelm Meyer auf der Heide. A new universal class of hash functions and dynamic hashing in real time. In *Proceedings of the 17th International Colloquium on Automata, Languages and Programming (ICALP '90)*, volume 443 of *Lecture Notes in Computer Science*, pages 6–19. Springer-Verlag, Berlin, 1990.
12. Arnold I. Dumey. Indexing for rapid random access memory systems. *Computers and Automation*, 5(12):6–9, 1956.
13. Michael L. Fredman, János Komlós, and Endre Szemerédi. Storing a sparse table with  $O(1)$  worst case access time. *J. Assoc. Comput. Mach.*, 31(3):538–544, 1984.
14. Gaston Gonnet. *Handbook of Algorithms and Data Structures*. Addison-Wesley Publishing Co., London, 1984.
15. Gaston H. Gonnet and J. Ian Munro. Efficient ordering of hash tables. *SIAM J. Comput.*, 8(3):463–478, 1979.
16. Richard M. Karp, Michael Luby, and Friedhelm Meyer auf der Heide. Efficient PRAM simulation on a distributed memory machine. *Algorithmica*, 16(4-5):517–542, 1996.
17. Jyrki Katajainen and Michael Lykke. Experiments with universal hashing. Technical Report DIKU Report 96/8, University of Copenhagen, 1996.
18. Donald E. Knuth. *Sorting and Searching*, volume 3 of *The Art of Computer Programming*. Addison-Wesley Publishing Co., Reading, Mass., second edition, 1998.
19. J. A. T. Maddison. Fast lookup in hash tables with direct rehashing. *The Computer Journal*, 23(2):188–189, May 1980.
20. E. G. Mallach. Scatter storage techniques: A uniform viewpoint and a method for reducing retrieval times. *The Computer Journal*, 20(2):137–140, May 1977.
21. George Marsaglia. The Marsaglia random number CDROM including the diehard battery of tests of randomness. <http://stat.fsu.edu/pub/diehard/>.
22. Catherine C. McGeoch. The fifth DIMACS challenge dictionaries. <http://cs.amherst.edu/~ccm/challenge5/dicto/>.
23. Kurt Mehlhorn and Stefan Näher. *LEDA. A platform for combinatorial and geometric computing*. Cambridge University Press, Cambridge, 1999.
24. Rasmus Pagh. On the Cell Probe Complexity of Membership and Perfect Hashing. In *Proceedings of the 33rd Annual ACM Symposium on Theory of Computing (STOC '01)*. ACM Press, New York, 2001.
25. Rasmus Pagh and Flemming Friche Rodler. Cuckoo hashing. In *Proceedings of the 9th European Symposium on Algorithms (ESA '01)*, volume 2161 of *Lecture Notes in Computer Science*, pages 121–133, Berlin, 2001. Springer-Verlag.

26. Rasmus Pagh and Flemming Friche Rodler. Lossy dictionaries. In *Proceedings of the 9th European Symposium on Algorithms (ESA '01)*, volume 2161 of *Lecture Notes in Computer Science*, pages 300–311, Berlin, 2001. Springer-Verlag.
27. Patricio V. Poblete and J. Ian Munro. Last-come-first-served hashing. *J. Algorithms*, 10(2):228–248, 1989.
28. Ronald L. Rivest. Optimal arrangement of keys in a hash table. *J. Assoc. Comput. Mach.*, 25(2):200–209, 1978.
29. Peter Sanders and Berthold Vöcking, 2001. Personal communication.
30. Jeanette P. Schmidt and Alan Siegel. On aspects of universality and performance for closed hashing (extended abstract). In *Proceedings of the 21st Annual ACM Symposium on Theory of Computing (STOC '89)*, pages 355–366. ACM Press, New York, 1989.
31. Jeanette P. Schmidt and Alan Siegel. The analysis of closed hashing under limited randomness (extended abstract). In *Proceedings of the 22nd Annual ACM Symposium on Theory of Computing (STOC '90)*, pages 224–234, New York, 1990. ACM Press.
32. Alan Siegel. On universal classes of fast high performance hash functions, their time-space tradeoff, and their applications. In *Proceedings of the 30th Annual Symposium on Foundations of Computer Science (FOCS '89)*, pages 20–25. IEEE Comput. Soc. Press, Los Alamitos, CA, 1989.
33. Craig Silverstein. A practical perfect hashing algorithm. Manuscript, 1998.
34. Mikkel Thorup. Even strongly universal hashing is pretty fast. In *Proceedings of the 11th Annual ACM-SIAM Symposium on Discrete Algorithms (SODA '00)*, pages 496–497. ACM Press, New York, 2000.
35. M. Wenzel. Wörterbücher für ein beschränktes Universum. Diplomarbeit, Fachbereich Informatik, Universität des Saarlandes, 1992.