

DACHash: A Dynamic, Cache-Aware and Concurrent Hash Table on GPUs

Hao Zhou, David Troendle, Byunghyun Jang

Computer and Information Science

The University of Mississippi

University, MS USA

hzhou3@go.olemiss.edu, {david, bjang}@cs.olemiss.edu

Abstract—GPU acceleration of hash tables in high-volume transaction applications such as computational geometry and bio-informatics are emerging. Recently, several hash table designs have been proposed on GPUs, but our analysis shows that they still do not adequately factor in several important aspects of a GPU’s execution environment, leaving large room for further optimization.

To that end, we present a dynamic, cache-aware, concurrent hash table named DACHash. It is specifically designed to improve memory efficiency and reduce thread divergence on GPUs. We propose several novel techniques including a GPU-friendly data structure & sizing, a reorder algorithm, and dynamic thread-data mapping schemes that make the operations of hash table more amendable to GPU architecture. Testing DACHash on an NVIDIA GTX 3090 achieves a peak performance of *8.65 billion queries/second* in static searching and *5.54 billion operations/second* in concurrent operation execution. It outperforms the state-of-the-art SlabHash by 41.53% and 19.92% respectively. We also verify that our proposed technique improves L2 cache bandwidth and L2 cache hit rate by $9.18\times$ and $2.68\times$ respectively.

Keywords-hash table; GPGPU; concurrent data structure

I. INTRODUCTION

GPUs have become the platform of choice for many compute and data intensive applications in various fields. Traditionally CPU centric data structures are finding GPU solution. Hash table algorithms offering fast data access in near constant time are important for the fields of computational geometry and bio-informatics, but not well researched. Designing a high-performance hash table on massively multi-threaded GPUs is a challenging task. Tens of thousands of active threads attempting simultaneous hash table access can cause severe performance degradation unless carefully designed. Traditional lock-based implementations suffer from high thread contention [1], leaving non-blocking methods a better choice for the GPU environment [2], [3], [4]. Nonetheless, any approach must accommodate and address the fact that GPUs are very sensitive to memory access patterns and thread divergence [5].

In this paper, we present a hash table specifically designed and optimized for a GPU architecture. We propose several novel techniques to address two major sources of GPU inefficiency - memory access patterns and thread divergence.

First, we introduce a GPU-friendly chaining structure to support hash collisions. This enables mutability via dynamic memory management for new data to be stored or old data to be deleted, while avoiding the need for repeated rebuilds from scratch. We optimize the chaining structure into a GPU-friendly linked-list of *super nodes*, where each super node is a small array of key-value pairs. Our cache-aware super node sizing improves memory access patterns, which is an important design consideration for a GPU’s SIMT (Single Instruction Multiple Threads) execution model.

Second, we improve the efficiency of dynamic memory management by pre-allocating a large memory pool and using a concurrent stack to manage memory buffer allocation and deallocation dynamically. This helps reduce the overhead of searching candidates to delete, and the cost of memory allocation and deallocation on GPUs.

Third, we reorder input data elements based on their hash values to improve cache performance. Rather than using expensive traditional sorting, our proposed reorder algorithm efficiently groups operations on the fly, increasing the likelihood of data reuse and coalesced memory transactions. To our knowledge, this is the first attempt to study and improve the locality of hash table data structures on GPUs.

Lastly, we design a novel dynamic mapping scheme that can switch between two different thread-data mapping schemes depending on the shape of hash table: A *one-to-one mapping scheme* maps each thread to a key so that threads process their keys individually; and a *many-to-one mapping scheme* maps each thread to a key, but an entire warp (32 threads) cooperatively processes 32 keys sequentially. Our proposed dynamic mapping scheme automatically switches between these two mapping schemes to achieve better performance.

Our experiments show that on a latest NVIDIA GPU, GTX 3090, our proposed DACHash achieves a static searching throughput and concurrent operations throughput of *8.65 billion queries/second* and *5.54 billion operations/second* respectively. It outperforms the state-of-the-art SlabHash [6] (*7.55 billion queries/second* and *4.41 billion operations/second*) with all overheads included. On average, DACHash is 41.53% and 19.92% faster than SlabHash under these two categories. We also profile and verify the

cache performance of DACHash using the NVIDIA Visual Profiler. It shows our proposed technique improves L2 cache bandwidth and hit rate by $9.18\times$ and $2.68\times$, demonstrating that the improved cache performance can yield a significant overall performance boost.

II. RELATED WORKS

Several hash table designs and implementations have recently been reported for GPUs in the literature.

Alcantara et al. [7] built a hash table on GPUs, which performs parallel insertions and retrievals. Their work is based on Cuckoo Hashing [8] and relies on atomic operations during multi-threads table construction. The authors use a set of hash functions to find a key in multiple candidate locations for insertion as Cuckoo Hashing does. Evicted keys need to be inserted into another location until no more evicted keys exist. A careful design of a set of hash functions is required since hash functions determine the frequency of rebuilding from scratch. The order of hash functions also matters.

Garcia et al. [9] presented a parallel hashing method where their hashing could reach high load factor but with a low rebuilding failure rate. The authors designed a coherent hash function to leverage coherence in memory and further increase locality in memory. In addition, coherent hashing also makes groups of threads execute consistent paths.

Khorasani et al. [10] proposed a hashing method called Stadium Hashing (Stash) and Stash with collaborative lanes (clStash). Stash Hashing has two basic structures: a table for keeping all keys and values, and a compact auxiliary structure called a ticket-board to maintain a ticket (consists of the availability bit and the info bits) for every bucket in the table. The availability bit determines if the bucket is occupied and the info bits store information of the key. This design reduces unnecessary accesses to the actual table content according to the availability bit and the info bits, which speeds up retrievals. By solving collisions via double-hashing (primary and secondary hash functions), Stash allows concurrent execution of mixed insertions and retrievals. The secondary hash function generates a step size that could hurt the memory performance on GPUs. clStash improves warp execution efficiency by redistributing tasks to early-finished threads in a warp.

SlabHash [6] proposes further improvements to the efficiency of warp execution and memory coalescing. The authors proposed a warp-cooperative work-sharing (WCWS) strategy, where all threads in a warp process one operation at a time by utilizing warp-synchronous programming and warp-wide communications. This design presents less thread divergence when compared to other hash tables. The authors also take advantage of array and linked-list structures to further serve their WCWS strategy. The SlabHash designs slabs which are arrays with key-value pairs stored. Each slab has the size of 128 bytes that matches the size of a cache line

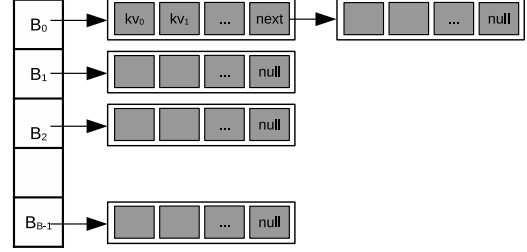


Figure 1: Basic structure of DACHash. Bucket 0 (B_0) has two super nodes and other buckets have one super node. Each super node has a small array of key-value pairs as well as a next pointer.

on GPUs. The SlabHash also designs a specialized memory pool to implement dynamic allocation.

Gao et al. [11] adopted a structure similar to SlabHash. In addition, the paper discusses the throughput of the WCWS strategy. They show that when the number of elements stored in the table is large, it could achieve higher throughput. Otherwise, the throughput is relatively low. Gao et al. solves the problem by proposing an adaptive model. In addition, the authors present a reader-writer lock based synchronization and bucket-level synchronization to ensure atomicity of hash operations (individual hash operations or groups of hash operations).

WarpCore [12] proposed a fast hash table on GPUs. They propose a memory-compact bucket list to support flexible multi-value storage, a hashing scheme to improve global memory access patterns by leveraging CUDA cooperative groups, and an efficient techniques to support multi-GPUs.

III. BASIC DESIGN AND IMPLEMENTATION

In this section, we introduce the basic design and implementation of DACHash, including base data structure, organization, supported operations, and memory management.

A. Base data structure and organization

Each DACHash bucket is designed as a linked-list of small arrays consisting of key-value pairs as shown in Fig. 1. The array offers contiguous, linear memory access patterns, while a linked-list chain offers easy, concurrent modification.

In our design, each node in the linked-list chain holds multiple interleaved key-value pairs. We call these nodes *super nodes*. The first super node connected to a bucket head is pre-allocated. Subsequent super nodes are dynamically allocated or deallocated at run time as needed.

The base data structure and organization of DACHash offer several optimization opportunities. First, the combined array/linked-list structure enables a natural way to support collisions. Second, a chaining technique allows dynamic allocation, instead of needing to rebuild the hash table from scratch. Third, array structures achieve GPU-friendly memory access patterns compared to a linked-list's scattered

memory accesses. Fourth, super nodes offer flexible thread-data mapping scheme options, e.g., *one-to-one* or *many-to-one* mapping schemes.

B. Operations supported

We implement five basic hash table operations on unique keys. CUDA atomic functions ensure correctness. Note that, although not implemented, duplicate keys can be accommodated without significant design changes. The supported operations are:

Search is responsible for finding a key in the hash table and returning its value. If no key is found, it returns null. The operation starts by hashing a key to a bucket. Searching begins at the bucket's first super node. If no key matches, it continues traversing the bucket's super nodes until a matching key is either found or it reaches the end of the bucket list.

Insert adds a key-value pair to the hash table. Since keys must be unique, we must first ensure the key exists in the hash table. If it does, the operation acts as *update*, replacing the old value with a new value. If it does not exist, it is inserted into the hash table. A new super node may be dynamically allocated if needed. When inserting a new key-value pair into the table, it first looks for an empty slot in the first super node of the bucket. If an empty slot exists, an *atomicCAS()* (a CUDA atomic function) ensures a correct insertion. If the first super node is full, the thread traverses the super node list until it finds an empty slot. If it reaches the last node in the bucket, the thread dynamically allocates a new super node and connects the new node after the bucket's last super node using an *atomicCAS()*. Although multiple threads may try to connect their super nodes after the last super node simultaneously, only one thread will succeed. The failing threads deallocate their nodes and retry until they succeed. Once successful, the threads redo their *insert* operation using the bucket's new last super node.

Update finds the key to update its value. If the key is found, it replaces its value with the new value using an *atomicExch()*. Otherwise, the new pair is inserted.

Delete is similar to the *search* operation, but returns no value. It starts its traversal at the first super node in the bucket the key maps to. If found, it marks the key as logically *deleted*. If not, it continues traversing super nodes looking for the matching key until it reaches the last super node in the bucket. This operation does not deallocate empty super node. Deallocation is done by the *clean* operation (see below).

Clean compacts the bucket's super node linked-list, ensuring only the last super node has any empty slots. We implement the *clean* operation as a separate kernel, so no other operations interfere when cleaning the hash table. Deallocated super nodes are pushed back to our memory stack for later use. The *clean* operation is only required when the memory stack is empty.

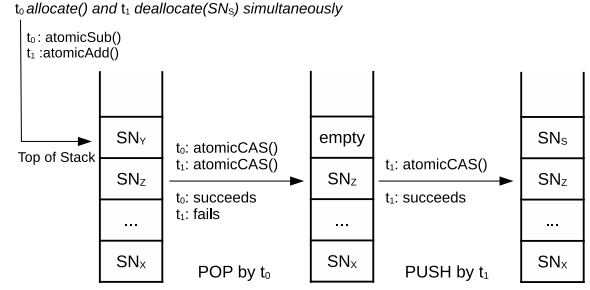


Figure 2: DACHash uses a concurrent stack to support dynamic memory allocation. All pre-allocated super nodes are pushed into the stack at the beginning. Two threads t_0 and t_1 compete for the same stack *top* index in this example.

C. Memory stack

The *insert*, *update*, *delete* and *clean* operations may require dynamic super node allocation or deallocation. To support this, we pre-allocate a large number of super nodes and place them on a concurrent stack. A pop allocates a super node and a push deallocates a super node concurrently as shown in Fig. 2. This is a simple, fast, GPU friendly alternative to a CPU-side *malloc()* or *free()*.

IV. REORDER

Input keys are hashed to different buckets. When they are mapped to threads, a warp suffers from poor locality because the super nodes within and across buckets are likely scattered in memory. Memory requests from threads in a warp are highly likely to reside in different cache lines (uncoalesced) rather than a single line (coalesced). Such poor spatial locality causes multiple memory transactions, which in turn, significantly increases memory traffic. Slab-Hash [6] proposed a work-sharing strategy within warps to increase coalesced memory accesses. However, it still suffers repeated linked-list traversals, which is another source of poor locality. Suppose a warp processes a list of query keys $\{0, 3, 7, 8\}$ and the size of warp is 4, as visualized in Fig. 3. The warp processes *key* 0, finds the key in *bucket* 0, and loads the first super node. With this access, other keys stored in that super node are loaded together but not used. The same issue occurs for *key* 3 and *key* 7. In the mean time, the super node loaded for *key* 0 may have been evicted from the cache. When processing *key* 8, the warp starts over from the beginning, and loads the first super node of *bucket* 0, then it loads the second super node that contains *key* 8. In this case, the first super node is loaded twice. Even if the warp issues coalesced memory requests, it may not end up being an efficient use of data as it can still cause extra global memory transactions and waste potential use of other keys in the loaded super nodes.

We observed that such inefficiencies are caused by the randomly arranged input data list. Reordering can improve

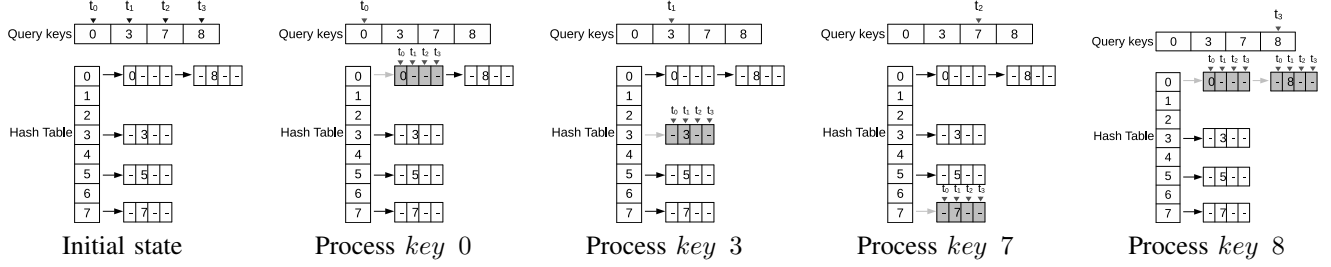


Figure 3: Example of poor data reuse. A warp of 4 threads processes keys 0, 3, 7, 8. The shaded nodes indicate the visiting super nodes by the warp.

the data locality and cache performance (i.e., hit rate and bandwidth). Sorting is the traditional way of reorganizing data. However, not only is a sort expensive on GPUs, but also a strictly ordered sort is not required for our purpose. We need only group the keys with the same hash value together. To this end, we propose a *reorder* algorithm. In our reorder algorithm, we take advantage of a pre-allocated memory buffer on GPUs to partition data according to their hash values. Keys with the same hash value compute their indexes in this pre-allocated memory buffer so that keys with the same hash value are physically close and adjacent to each other in memory. By doing so, when each thread claims its key, the adjacent threads are more likely mapped to the same bucket. In this way, when traversing bucket's super nodes, threads in a warp probably request the same or nearby super nodes, resulting in fewer memory transactions. Compared to scattered memory requests, the total number of global memory transactions can decrease significantly. Since the input keys are partitioned based on their hash values, when warps process these keys, the same or nearby super nodes are likely loaded from the cache directly, so that the cache performance increases as well.

Algorithm 1 details our proposed reorder algorithm. It utilizes a pre-allocated GPU memory buffer named *ReorderSpace* with the same or larger size than *keysList* but with an additional dimension. Note that the size of *ReorderSpace* is $B * M$, where B is the number of buckets and M is calculated by $\lceil N/B \rceil$ (line 1-4) where N is the total number of keys in *keysList*, and the size of *Record* is B in order to keep track of the latest index of buckets. The row index is defined by bucket (hash) value, and an *atomicAdd()* is used to find the corresponding index of that key in that row (line 8). By doing so, keys with the same hash value will be mapped to the same row in *ReorderSpace*. However, for our input keys, we cannot ensure that keys have a perfect uniform distribution across buckets. In this case, we have to arrange keys with the indexes that are greater than M to a nearby location (line 9-11). It checks whether the next row has an empty spot to add. If so, it updates *index*. If not, check the row after the next row until an empty spot is found. Lastly, we add keys to *ReorderSpace* (line 12). Since the size of

Algorithm 1: Pseudocode for the proposed reorder algorithm.

input: ReorderSpace
 keysList
 Record

```

1  $N = \text{lengthOf}(\text{keysList});$ 
2  $B = \text{totalNumberOfBucket}();$ 
3 //  $B = \text{totalNumberOfBucket}() /$   

     $\text{lengthOfCombinedBuckets};$ 
4  $M = \lceil N/B \rceil;$ 
5  $\text{key} = \text{keysList}[\text{threadId}];$ 
6  $\text{bucket} = \text{hash}(\text{key});$ 
7 //  $\text{bucket} = \text{hash}(\text{key}) / \text{lengthOfCombinedBuckets};$ 
8  $\text{index} = \text{atomicAdd}(\text{Record}[\text{bucket}], 1);$ 
9 while  $\text{index} \geq M$  do
10 |    $\text{index} = \text{atomicAdd}(\text{Record}[(++\text{bucket}) \bmod B],$   

    1);
11 end
12  $\text{ReorderSpace}[\text{bucket}][\text{index}] = \text{key};$ 
```

ReorderSpace is equal to or greater than the size of *keysList*, it guarantees all keys will have a spot to be added.

However, the reorder algorithm delays operations until the reorder process completes. This causes a synchronization problem. To this end, we utilize host-side device-synchronization via CUDA API. In order to speed up the reorder algorithm, we also improve the algorithm. Instead of grouping keys with the same hash value, we group keys with nearby hash values. Originally, B is equal to the total number of buckets in our hash table. By decreasing B (line 3), we can group the keys with nearby hash values into one row. For instance, if we decrease B to $B/4$, every four hash values will be grouped together, e.g., 0, 1, 2 and 3 (4 values). In *keysList*, the probability of adjacent data elements that can be mapped to the same row of *ReorderSpace* increases, so the throughput of reorder is further improved.

V. DYNAMIC MAPPING SCHEMES

There are two common thread-data mapping schemes used in hash table design: 1) a *one-to-one mapping* scheme, which

maps each thread to a single key so that threads can process their 32 keys individually. Note that threads in a warp could possibly exit at different times due to the different operations they may have. The threads tend to have scattered memory access patterns when reading from or writing to GPU memory; and 2) a *many-to-one mapping* scheme, which maps a warp to 32 keys but 32 threads in a warp work and communicate with each other to cooperatively process a single key at a time. Threads in a warp will converge at the same time because all 32 threads in the warp will execute the same low-level instructions and the memory access pattern is likely to be coalesced. In general, each scheme shows different performance characteristics. The *one-to-one mapping* scheme outperforms when thread divergence is not an issue (e.g., threads in a warp process same type of operations), while the *many-to-one mapping* scheme is a better choice when execution efficiency matters (e.g., threads within a warp process different types of operations).

In order to better utilize two different schemes, we propose our dynamic mapping schemes where DACHash selects one scheme over the other based on, so-called *expected length*. The expected length ϵ , is defined as the average number of super nodes per buckets. Therefore the selection of the proposed dynamic mapping schemes D , is expressed as the following.

$$D = \begin{cases} O & \text{if } \epsilon < \tau \\ M & \text{otherwise} \end{cases}$$

where O is *one-to-one mapping* scheme, M is *many-to-one mapping* scheme, and τ is a threshold.

Algorithm 2: Pseudocode for *one-to-one mapping* scheme.

```

1 key = keysList[threadId];
2 value = valuesList[threadId];
3 bucket = hash(key);
4 superNode = getNextSuperNode(bucket);
5 do
6   for  $i \leftarrow 0$  to  $NODE\_SIZE-1$  do
7     if  $superNode[i].key == key$  then
8       SEARCH();
9       or DELETE();
10      or UPDATE();
11    end
12  end
13  superNode = getNextSuperNode(superNode);
14 while  $superNode \neq nullptr$ ;
15 if  $UPDATE() \text{ failed \&\& } isUPDATE$  then
16   INSERT();
17 end

```

Algorithm 3: Pseudocode for *many-to-one mapping* scheme.

```

1 key = keysList[threadId];
2 value = valuesList[threadId];
3 bucket = hash(key);
4 operation = operationsList[threadId];
5 superNode = getNextSuperNode(bucket);
6 tile =
   tiled_partition( $CG\_SIZE$ )(this_thread_block());
   laneId = tile.thread_rank();
7 for  $lane \leftarrow 0$  to  $CG\_SIZE-1$  do
8   // a tile processes operations from 0 to
   CG_SIZE-1;
9   share data in the  $lane^{th}$  of tile by  $tile.shfl()$ ;
10  do
11    tile reads superNode and finds target key by
    tile.ballot();
12    if found then
13      SEARCH();
14      or DELETE();
15      or UPDATE();
16    else
17      share the next superNode by  $tile.shfl()$ ;
18    end
19  while  $superNode \neq nullptr$ ;
20  if  $UPDATE() \text{ failed \&\& } isUPDATE$  then
21    INSERT();
22  end
23 end

```

Our implementations of these mapping schemes are shown in Algorithm 2 and 3 respectively. In the *one-to-one mapping* scheme, groups of threads may have different operations to perform and their finishing times are likely to be different, resulting in low execution efficiency. The many-to-one mapping schemes partitions thread blocks into tiles and the size of tile is specified by CG_SIZE which can be $\{1, 2, 4, 8, 16, 32\}$. It enables threads communication by using CUDA primitives such as $shfl()$ and $ballot()$, and lastly tiles would perform different hash operations, i.e., *search*, *delete*, *update*, and *insert*. In the *many-to-one mapping* scheme, tiles are more likely to execute the same low-level instructions so that thread divergence is minimized, and the memory access pattern is also improved.

VI. PERFORMANCE EVALUATION

We evaluated our proposed DACHash on an NVIDIA GTX 3090, which is based on the *AMPERE* microarchitecture [13] with compute capability of 8.6, 10496 CUDA cores, and 24 GB off-chip global memory (Our experiments on older NVIDIA GPU architectures showed very similar performance behavior therefore those are not included in

this paper). We compiled our code with the CUDA 11.2 compiler. We divide our performance evaluation into three parts. First, we show the performance impact of various parameter values for the super node size, the length of combined buckets in our proposed reorder algorithm, and the threshold τ for dynamic mapping schemes. Second, we demonstrate the effectiveness of our reorder algorithm. Lastly, we compare our results using the state-of-the-art SlabHash [6] as a baseline for different categories, such as build rate, static search throughput, and concurrent operation throughput. In our experiments, we assume all keys are unique. All overheads such as data copy and reordering are included in our performance metrics. Experiments on other GPU architectures with OpenCL implementation remain as future work.

A. Impact of Parameters

1) *Super node size*: We measured the impact of the number of elements per super node¹ on performance while keeping the total input keys size fixed. We built a hash table with randomly generated keys and different super node sizes. Then we created the shuffled query list with the same keys existing in the hash table to perform *search* operations. We kept our reorder algorithm enabled for this experiment. Fig. 4 shows the results of super node sizes of 4*4, 8*4, 16*4, and 32*4 bytes.

Our experiment results show that super node size significantly affects the performance of our hash table. When super node size is small, the traversal speed is faster as each thread in the operation may need to traverse the entire bucket to find its target. When a bucket holds more than one super node, the performance impact of super node size decreases relatively due to dynamic memory allocation overhead and the bucket traversal time. In the following experiments, we choose 16*4 bytes to be our super node size for the following reasons. First, the way that we optimize the reorder algorithm (e.g., combining keys with nearby hash values) leaves smaller super node more reasonable. Suppose that the size of the super node is 32*4 bytes (e.g., the size of a cache line), when some threads in a warp load a super node, other threads are not likely to be mapped to the same super node, which results in more global loads. With a smaller super node (e.g., 16*4 bytes), the reorder algorithm decreases the repeated global loads because threads in a warp may load their super nodes directly from the cache. Second, smaller super nodes (e.g., 4*4 or 8*4 bytes) may trigger more dynamic allocations as the number of total key-value pairs stored in the hash table increases, which decreases the performance.

¹Note that the choices of super node sizes are guided by CUDA cooperative groups [14] where cooperative groups only allow finer granularity at level {1, 2, 4, 8, 16, 32}. Expected length indicates how many super nodes exist per bucket.

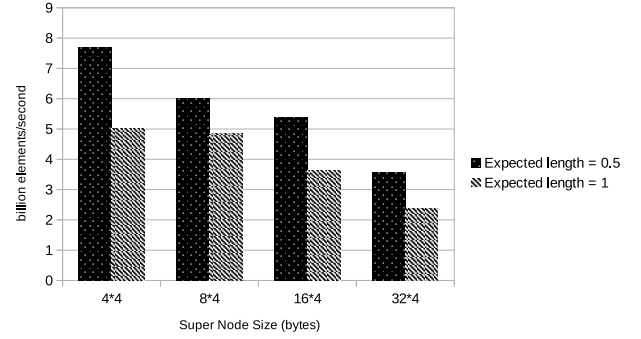


Figure 4: The performance impact of super node size. Expected lengths 0.5 and 1 indicate that there are 1 and 2 super nodes per bucket respectively. Note that when each bucket owns more than one node, dynamic memory allocation is necessary.

2) *Length of combined buckets in the reorder algorithm*: The objective of our reorder algorithm is to combine nearby buckets and partition data elements based on those combined buckets (line 7 in Algorithm 1). We measured the efficiency of our reorder algorithm on various lengths of combined buckets using *search* operations. Fig. 5 shows the experimental results. When the reorder algorithm is disabled (i.e., length of combined buckets is 0), we form a baseline for our *search* operation where the keys in the query list are randomly organized. Note that the total time consists of two components - the search and the reorder time. As we increase the number of combined buckets, the total time decreases because our reorder algorithm makes searching more cache-friendly. However, from the figure one can see that when too few or too many buckets are combined, the efficiency of our reorder algorithm diminishes. We observed two reasons for that. When there are too few combined buckets, the reorder algorithm suffers from poor cache performance as analyzed in Section IV. The input query list for reorder resides in contiguous memory but the data elements are randomly located. So, every time a cooperative group or a warp reads memory (e.g., 128 bytes), it achieves coalesced memory access. But when writing to different buckets (line 12 in Algorithm 1), the threads in the same warp or group are more likely to write to different cache lines, which results in an inefficient memory access pattern. When there are too many combined buckets, contention becomes an issue even though writing to memory could be efficient. This causes many threads to modify the same memory unit (line 8 and line 10 in Algorithm 1). In Fig. 5, when the number of combined buckets is roughly between 32 and 256, our hash table shows best performance with the total 2^{19} buckets. In general, we find the ratio between total buckets and the number of combined buckets, $1/256$ (i.e., combining 256 buckets for the different number of total buckets), is a good

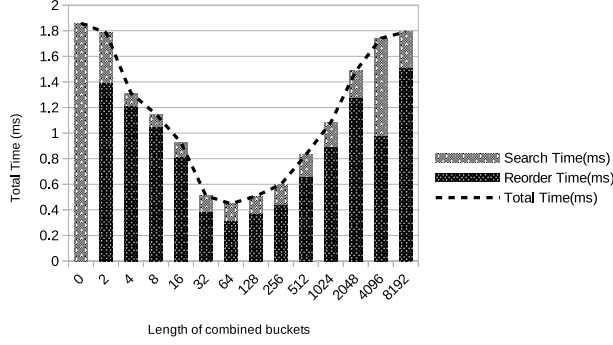


Figure 5: The performance impact of the length of combined buckets. The total number of buckets is 2^{19} .

choice.

3) *Threshold for dynamic mapping schemes*: We conducted this experiment to determine when we should switch between the *one-to-one mapping* and the *many-to-one mapping* schemes. Note that the *many-to-one mapping* scheme is designed to solve the thread divergence issue. So in order to better understand the difference of two schemes on the thread divergence issue, we conducted the experiment on two different settings. In the uniform operation setting (e.g., all searches), Fig. 6a shows that the *one-to-one mapping* and *many-to-one mapping* schemes have no performance crossover point, which indicates that *one-to-one mapping* scheme consistently outperforms when all threads perform the same operation. In the mixed operation setting (e.g., searches and updates) of Fig. 6b, however, we see a performance crossover point roughly at an expected length of 0.6. We believe that before the crossover point, the *one-to-one mapping* scheme wins due to its higher throughput, and after the crossover point, the *many-to-one mapping* scheme presents less thread divergence by having threads in warps work together to finish an operation at a time so that the threads possibly converge at the same time. Based on our experiment, we chose an expected length of 0.6 as the value of the threshold τ for our dynamic mapping schemes. Note that the dynamic mapping schemes go in effect only when a setting of the mixed operations is required.

B. Contribution of Reorder Algorithm

We performed experiments to show the effectiveness of our proposed reorder algorithm by enabling and disabling it, and measuring changes in cache hit rate and cache bandwidth using the performance counters provided in the NVIDIA Visual Profiler [14]². When disabled, the keys in the input query list are randomly arranged. As shown in Fig. 7, reorder enabled always outperforms reorder disabled

²Note that the data is retrieved from NVIDIA GTX 1070 since either the visual profiler or the *nvp* doesn't support GTX 3090 yet with a compute capability of 8.6 that is higher than that of 7.0.

across different expected lengths. The average speedup is around $4.33\times$.

1) *Effectiveness of reorder*: We use *search* operations in this experiment. When the reorder algorithm is disabled, the keys in the input query list remain unchanged. As shown in Fig. 7, given different expected lengths, enabling the reorder algorithm always gives better performance. The average speed up is around $4.33\times$. It is noteworthy that our reorder algorithm achieves small speed ups only when the expected length is short. This is because when it is short (e.g., less than 0.2), the number of buckets is too many, the total number of data that each row can hold is small (i.e., $M = \lceil N/B \rceil$ in Algorithm 1). In this case, since the key distribution across buckets is not perfectly uniform, the reorder algorithm has to find available locations for some keys in other rows, which decreases the performance of the reorder algorithm.

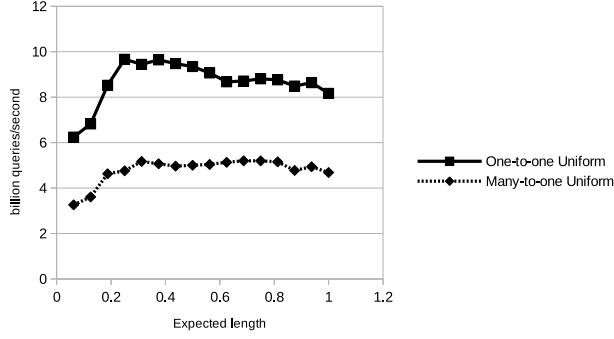
2) *Cache performance*: We also conducted a search experiment to measure the effectiveness of our reorder algorithm on L2 cache performance (data provided by the NVIDIA Visual Profiler) under different situations. From Table I, one can see with reordering enabled, our proposed DACHash boosts L2 cache bandwidth by $9.18\times$, and L2 cache hit rate by $2.68\times$. In short, as mentioned early in Section IV, the effect of memory operations on GPU performance cannot be ignored. By applying our proposed reorder algorithm, we can improve data locality because adjacent threads will possibly map to the same or nearby buckets. Therefore, the overall performance is improved by increasing L2 cache bandwidth and hit rate.

	Reorder enabled	Reorder disabled
L2 Cache Bandwidth	901.25 GB/s	98.196 GB/s
L2 Cache Hit Rate	91.00 %	34.00 %

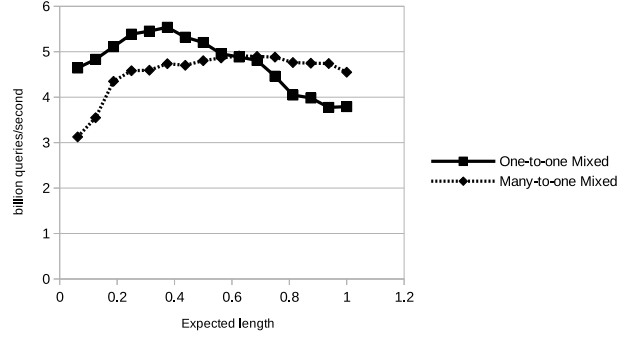
Table I: L2 cache performance comparison.

C. Build Rate Comparison

We compared the build rate of DACHash to those of SlabHash. As shown in Fig. 8, our proposed DACHash performs better than the state-of-the-art SlabHash when the expected length roughly ranges from 0.1 to 0.6. The peak building rate of DACHash is *3.42 billion elements/second*, compared to *2.9 billion elements/second* in SlabHash with the total number of buckets of 2^{22} . However, we noticed that our hash table suffers when the expected length is very low (e.g., 0.0625). We think this is due to the limitation of the reorder algorithm as we mentioned in Section VI-B1. The reorder algorithm has to update indexes for some data in other rows. Also, our build rates are lower when expected length is high (e.g., 0.6). We believe there are two reasons for it. The first reason is that when we have larger expected length, we also expect the time for traversing the super nodes to find empty spots longer. Increased traversal times decrease



(a) Uniform operations



(b) Mixed operations

Figure 6: The performance impact of threshold across different settings. The total number of elements is fixed at 2^{22} and the expected length varies.

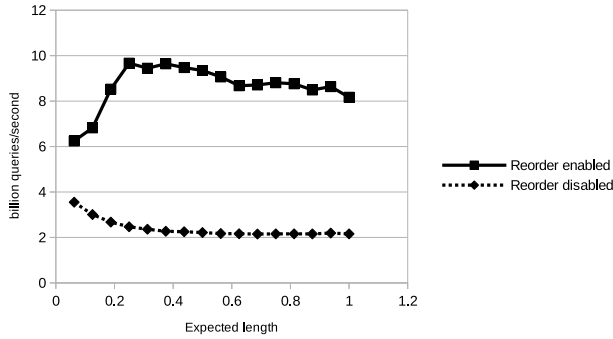


Figure 7: The performance impact of the proposed reorder algorithm. The total number of elements stored in the table is 2^{22} .

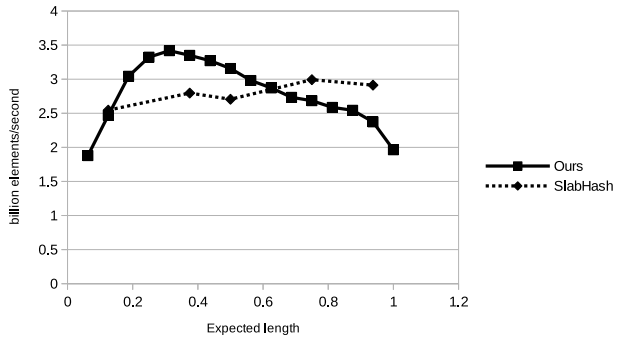


Figure 8: Build rate comparison. The total number of elements inserted in the table is 2^{22} .

the overall build rate. The second reason is that when we have fewer buckets, the contention on the same memory unit (line 8 in Algorithm 1) intensifies so that the overall throughput decreases.

D. Static Search Comparison

For static search comparison, we have two different setups as shown in Fig. 9. For both setups, we first build a hash table with randomly generated key-value pairs, then create a query list with the same keys but re-arranged to perform *search* operations. The total number of elements varies in Fig. 9a, while the expected length varies in Fig. 9b.

Fig. 9a shows the result of DACHash compared to SlabHash when the total number of queries varies. The peak performance of DACHash is *8.66 billion elements/second*, while the peak performance of SlabHash is only *7.55 billion elements/second*. On average, DACHash is improved by 7.1% compared to SlabHash. It is noteworthy that DACHash under-performs when the total number of elements is small or large while the expected length is fixed (e.g., 0.5). Recall that we synchronize the reorder algorithm by adopting a host-side CUDA API. So when the total number of elements is small (e.g., 2^{19}), our reorder algorithm suffers from the extra kernel launch and its extra synchronization time, which increases the overall finishing time. When the total elements ranges from 2^{25} to 2^{28} , our reorder algorithm suffers from the situation where writing to different buckets in the reorder algorithm results in an inefficient memory access pattern. Also, one can see that DACHash outperforms SlabHash when the total elements ranges from 2^{20} to 2^{24} . The reason is that our reorder algorithm indeed finds a balance point at which the benefits of reordering outweigh its drawbacks. In addition, one may notice that in Fig. 9a, we have an extra base line. This is a baseline measuring the performance of purely *one-to-one mapping* scheme without any other techniques involved. Note that when the total number is small (e.g., 2^{19}), either DACHash or SlabHash under-performs the baseline. We think this is due to the GPU characteristics of latency hiding. When the total number of elements is small, GPUs actually hide memory latency fairly well. This also implies why we prefer to design a reorder algorithm instead

of using a sorting algorithm directly because sorting may perform well on a small data scale but will suffer on a large data scale.

Fig. 9b also demonstrates the searching performance where the expected length varies. The experiment shows that DACHash outperforms SlabHash and achieves a searching throughput above *8.65 billion elements/second* and improves SlabHash by *41.53%* on average. However, there are several observations: First, the performance of our hash table decreases when the expected length increases. This is because, in this experiment, we rely on the *one-to-one mapping* scheme so that when the expected length increases, the bucket traversal time to find the matching key also increases, and the contention on the same memory unit will increase as well along with the increased expected length. Second, when the expected length is small (e.g., <0.1), SlabHash could perform better than DACHash. We believe the reason is that our reorder algorithm is less efficient when mapping to other rows.

E. Concurrent Operations Comparison

DACHash also supports concurrent execution of the *search*, *update*, and *delete* operations. In the experiment, we tested the performance in two groups. One is a mix of *80% search* operations and *20% updates* operations (*10% update* and *10% delete* operations respectively). The other is a mix of *60% search* operations and *40% updates* operations (*20% update* and *20% delete* operations respectively). Fig. 10 shows the DACHash and SlabHash results. It is clear that DACHash outperforms Slabhash by *19.92%* on average, and the peak performance of DACHash with the dynamic mapping schemes is *5.54 billion operations/second*, while SlabHash has the peak performance of *4.41 billion operations/second*. We believe there are mainly three reasons that account for the difference. First, the optimized structure of DACHash enables each thread to traverse its target bucket no matter what operations it has. This helps reduce thread divergence and improve warp execution efficiency. Second, our proposed reorder algorithm is more cache-friendly. Even though GPUs hide memory latency by scheduling available thread blocks, frequent memory operations still deteriorates the efficacy of latency hiding. So in our reordered input list, keys with the same or nearby hash values are grouped together. This increases the likelihood of threads in a warp being mapped to the same or nearby super nodes, so that the cache hit rate improves. Third, our dynamic mapping schemes, especially the *many-to-one mapping* scheme, minimizes the effect of thread divergence by making threads in a group process operations cooperatively with the help of CUDA primitives such as *shfl()* and *ballot()*, thus improves the overall performance. Also, note that when there are more update operations such as *20% update* and *20% delete*, the throughput for them is lower than that of *10% update* and *10% delete*. We believe that updates operations (e.g., *update*

and *delete*) may trigger more memory operations so that they are more expensive than *search*.

VII. CONCLUSIONS

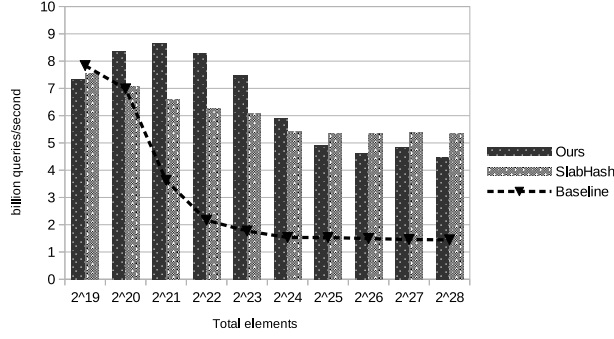
In the past decade, rapidly growing data in numerous fields such as computational geometry and bio-informatics have given rise to research in high throughput hash tables. Hash table suffers from poor memory performance and thread divergence on GPUs. We present a dynamic, high throughput, GPU architecture-aware hash table in this paper. Our proposed reorder algorithm and dynamic mapping schemes help improve the performance of hash table significantly, and beats the state-of-the-art implementation reported in the literature. We conducted experiments in three different categories to compare against the state-of-the-art solution SlabHash to verify our proposed DACHash. We also demonstrated the effectiveness of our proposed reorder algorithm by presenting the performance counters for L2 cache hit rate and bandwidth from the NVIDIA Visual Profiler.

ACKNOWLEDGMENT

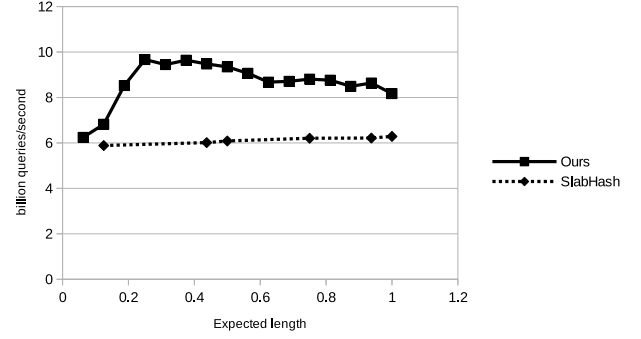
This research is supported by the National Science Foundation under Grant 1907838.

REFERENCES

- [1] Y. Xu, L. Gao, R. Wang, Z. Luan, W. Wu, and D. Qian, "Lock-based synchronization for gpu architectures," in *Proceedings of the ACM International Conference on Computing Frontiers*, 2016, pp. 205–213.
- [2] M. M. Michael, "High performance dynamic lock-free hash tables and list-based sets," in *Proceedings of the fourteenth annual ACM symposium on Parallel algorithms and architectures*, 2002, pp. 73–82.
- [3] P. Misra and M. Chaudhuri, "Performance evaluation of concurrent lock-free data structures on gpus," in *2012 IEEE 18th International Conference on Parallel and Distributed Systems*. IEEE, 2012, pp. 53–60.
- [4] L. Verkleij, "Boosting shared hash tables performance on gpu," Ph.D. dissertation, University of Twente, Enschede, The Netherlands, 2016.
- [5] B. Lessley and H. Childs, "Data-parallel hashing techniques for gpu architectures," *IEEE Transactions on Parallel and Distributed Systems*, vol. 31, no. 1, pp. 237–250, 2019.
- [6] S. Ashkiani, M. Farach-Colton, and J. D. Owens, "A dynamic hash table for the gpu," in *2018 IEEE International Parallel and Distributed Processing Symposium (IPDPS)*. IEEE, 2018, pp. 419–429.
- [7] D. A. Alcantara, V. Volkov, S. Sengupta, M. Mitzenmacher, J. D. Owens, and N. Amenta, "Building an efficient hash table on the gpu," in *GPU Computing Gems Jade Edition*. Elsevier, 2012, pp. 39–53.



(a) The expected length is 0.5.



(b) The total number of elements is 2²².

Figure 9: Static search comparison.

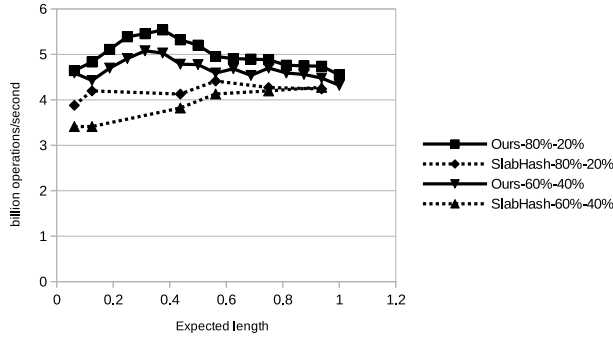


Figure 10: Concurrent operation comparison. The total number of operations is 2²² and node size is 16*4 bytes. Note that we use the dynamic mapping schemes in this experiment.

Available: <https://www.nvidia.com/en-us/data-center/ampere-architecture/>

[14] D. Guide, “Cuda c programming guide,” *NVIDIA*, July, 2013.

[8] R. Pagh and F. F. Rodler, “Cuckoo hashing,” *Journal of Algorithms*, vol. 51, no. 2, pp. 122–144, 2004.

[9] I. García, S. Lefebvre, S. Hornus, and A. Lasram, “Coherent parallel hashing,” *ACM Transactions on Graphics (TOG)*, vol. 30, no. 6, pp. 1–8, 2011.

[10] F. Khorasani, M. E. Belviranli, R. Gupta, and L. N. Bhuyan, “Stadium hashing: Scalable and flexible hashing on gpus,” in *2015 International Conference on Parallel Architecture and Compilation (PACT)*. IEEE, 2015, pp. 63–74.

[11] L. Gao, Y. Xu, C. Xu, R. Wang, H. Yang, Z. Luan, and D. Qian, “Towards a general and efficient linked-list hash table on gpus,” in *2019 IEEE 21st International Conference on High Performance Computing and Communications; IEEE 17th International Conference on Smart City; IEEE 5th International Conference on Data Science and Systems (HPCC/SmartCity/DSS)*. IEEE, 2019, pp. 1452–1460.

[12] D. Jünger, R. Kobus, A. Müller, C. Hundt, K. Xu, W. Liu, and B. Schmidt, “Warpcore: A library for fast hash tables on gpus,” *arXiv preprint arXiv:2009.07914*, 2020.

[13] NVIDIA. (2021) Nvidia ampere architecture. [Online].