

# Fat Address Translations

Akilan Selvacoumar\*

as251@hw.ac.uk

Heriot Watt

Edinburgh, Scotland, United Kingdom

## Abstract

The increasing disparity between application workloads and the capacity of translation lookaside buffers (TLBs) has prompted researchers to explore solutions to mitigate the extra clock cycles incurred during a TLB miss. One such approach involves leveraging physically contiguous memory and the use of huge pages. Concurrently, advancements in hardware-level system security—exemplified by the Capability Hardware Enhanced RISC Instructions (CHERI) architecture—offer additional opportunities for improving TLB performance. CHERI introduces capability-based addressing, a novel approach that enhances system security by associating capabilities with memory pointers. By leveraging capability-based addressing, we introduce a memory allocator that can integrate block-based allocations within huge pages. Through our evaluation using both micro and macro benchmarks, we show that our allocator can reduce TLB misses by up to 90%, leading to improvements in wall clock runtimes for memory-intensive applications.

## ACM Reference Format:

Akilan Selvacoumar. 2018. Fat Address Translations. In *Proceedings of Make sure to enter the correct conference title from your rights confirmation email (Conference acronym 'XX)*. ACM, New York, NY, USA, 9 pages. <https://doi.org/XXXXXXX.XXXXXXX>

## 1 Introduction

In computing, achieving high performance is an ongoing challenge, especially as applications handle increasingly complex workloads. Memory management is a key factor in performance, where efficient use of resources is essential. Translation Lookaside Buffers (TLBs) are crucial in this context, speeding up memory access by caching recent memory address translations. A TLB, a specialised cache in the memory management unit (MMU), reduces the time required to convert virtual addresses to physical ones. When a program accesses data in memory, the MMU first checks the TLB for a matching entry, avoiding the slower process of consulting page tables. However, as applications grow larger and more complex, the fixed size of TLBs often cannot keep up, leading to more TLB misses and performance slowdowns [1]. To tackle this issue, researchers have explored new solutions, including the use of huge pages [2].

Huge pages, also known as large pages, allow for the allocation of memory in significantly larger chunks compared to traditional small

pages. By reducing the number of TLB entries needed to access a given amount of memory, Huge pages offer a potential avenue for optimising TLB utilisation by reducing the number of entries needed to map large memory regions. This not only decreases the frequency of TLB misses but also lowers the overhead associated with address translation. By minimising these bottlenecks, huge pages can improve system performance in aspects such as speeding up memory-intensive applications, reducing latency in data access, and enhancing throughput for workloads that rely heavily on large datasets.

Simultaneously, advancements in hardware-level security, such as the Capability Hardware Enhanced RISC Instructions (CHERI) [3] architecture, present additional opportunities for performance enhancement. CHERI's capability-based addressing approach not only strengthens system security by tightly controlling memory access but also opens avenues for optimising memory management operations. By integrating CHERI's compressed encoded bounds [4] with the use of huge pages, We have shown it is possible to track and manage large, physically contiguous memory blocks without requiring numerous TLB entries. This combination reduces TLB pressure by minimising the number of entries required to map extensive memory regions, thereby decreasing TLB misses and improving address translation performance. Furthermore, it accelerates memory-intensive tasks by reducing the overhead associated with managing non-contiguous memory allocations. The contributions for the following paper are as follows:

- **Fat Addresses Translations:** Introduces fat-pointers that include memory bounds, allowing efficient tracking and management of physically contiguous memory regions (section 3).
- **CHERI's Capability-based Optimization:** Demonstrates how CHERI's architecture can be used to optimize memory allocation by encoding memory bounds directly within pointers, reducing TLB reliance (section 3.2).
- **Memory Allocation Algorithms:** Provides an algorithms for allocating, freeing physically contiguous memory and integrating huge pages with CHERI's capability-based bounds for enhanced memory management (section 4).

Through comprehensive evaluation, including micro and macro benchmarks, we demonstrate the allocator's ability to reduce TLB misses by up to 90%, yielding significant improvements in wall clock runtimes for memory-intensive applications. While its impact on larger, computation-heavy workloads is less pronounced, the proposed allocator shows strong potential for advancing memory management in scenarios requiring high memory throughput by reducing the address translation overhead.

Permission to make digital or hard copies of all or part of this work for personal or

Unpublished working draft. Not for distribution. This work is distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from [permissions@acm.org](mailto:permissions@acm.org).

Conference acronym 'XX, Woodstock, NY

© 2018 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-XXXX-X/18/06

<https://doi.org/XXXXXXX.XXXXXXX>

## 2 Related work

### 2.1 Huge Pages

Increasing TLB reach can be achieved by using larger page sizes, such as huge pages [2], which are common in modern computer systems. The x86-64 architecture supports huge pages of 2 MB and 1 GB, backed by OS mechanisms like Transparent Huge Pages (THP) [5] and HugeTLBFS in Linux. However, available page sizes in x86-64 are limited, leading to internal fragmentation issues.

For instance, allocating 1 MB with 4 KB base pages requires 256 PTEs, but using a 2 MB huge page would waste half of the memory space. Some architectures offer more page size choices, such as Intel Itanium, which allows different areas of the address space to have their own page sizes. Itanium uses a hash page table to organize huge pages, but without significant changes to the conventional page table, it only helps reduce page walk overheads. HP Tunable Base Page Size permits the OS to adjust the base page size, but still faces internal fragmentation problems, with HP recommending a base page size of no more than 16 KB. Shadow Superpage [6] introduces a new translation level in the memory controller to merge non-contiguous physical pages into a huge page in a shadow memory space, extending TLB coverage. However, this approach requires all memory traffic to be translated again in the memory controller, resulting in additional latency for memory accesses.

### 2.2 Direct Segment

Early processors often used segments to manage virtual memory, where a segment [7] essentially mapped contiguous virtual memory to contiguous physical memory. Unlike pages, which are relatively small, segments can be much larger, offering the potential for more efficient memory management in certain scenarios. This concept of segmentation has seen a resurgence in some modern approaches that aim to enhance translation coverage by designating specific areas in the virtual address space.

This method allows programmers to explicitly define a single segment for applications requiring significant memory. It introduces two new registers to the system, which indicate the start and end of this segment. Virtual addresses within this segment are translated by calculating the offset from the virtual start address and applying this offset to the physical start address. This straightforward method simplifies the translation process for large memory areas but requires significant modifications to the source code of applications.

### 2.3 Range Memory Mapping (RMM)

Redundant Memory Mappings (RMM) [8] enhance memory management by introducing an additional range table that pre-allocates contiguous physical pages for large memory allocations, creating ranges that are both virtually and physically contiguous. This approach simplifies address translation within these ranges by adding an offset, similar to Direct Segment, but RMM supports multiple ranges and operates transparently to programmers, requiring no source code modifications. The range table, separate from the conventional page table, holds the mappings for these large allocations. To determine which range an address belongs to, RMM compares

the address against all range boundaries, a process that is computationally expensive and therefore performed only after an L1 TLB miss. To optimize this, RMM uses a range TLB (RTLb) to quickly identify if an address falls within any pre-allocated range, facilitating efficient translation and reducing overhead. Range mapping works alongside the paging system by generating TLB entries on TLB misses and still performing TLB lookups for each virtual address translation. Unlike traditional segmentation mechanisms, range mapping activates a range lookaside buffer (RTLb) located with the last level TLB upon a miss. The hardware TLB miss handler then searches the RTLb for the miss address and, if found, generates a new TLB entry with the physical address derived from the base virtual address and range offset, along with permission bits. If the RTLb also misses, the system defaults to a standard page walk while a range table walker simultaneously loads the range into the RTLb in the background, avoiding delays in memory operations. The RTLb, functioning as a fully associative search structure, ensures that most last level TLB misses are handled efficiently by range mapping, reducing the need for costly page table walks.

### 2.4 CHERI

CHERI extends conventional processor Instruction-Set Architectures (ISAs) with architectural capabilities to enable fine-grained memory protection and highly scalable software compartmentalization. CHERI is a hybrid capability architecture that can combine capabilities with conventional MMU (Memory Management Unit) based systems. The contributions of CHERI include:

- ISA changes to introduce architectural capabilities.
- New microarchitecture proving that capabilities can be implemented efficiently in hardware, with support for efficient tagged memory to protect capabilities and compress capabilities to reduce memory overhead.
- A newly designed software construction model that uses capabilities to provide fine-grained memory protection and scalable software compartmentalization.
- Language and compiler extensions for using capabilities with C and C++.
- OS extensions to support fine-grained memory protection (spatial, referential, and (non-stack) temporal memory safety) and abstraction extensions for scalable software compartmentalization.

## 3 Fat Address Translations

This section talks about how Fat Address Translations (FAT) uses the CHERI architecture to bring about block based allocations in physically contiguous memory. FAT leverages techniques like Flex-Pointer [9] and Range Memory Mapping (RMM) [8] to achieve lesser pressure in the TLB. A key component in this implementation is the use of range addresses with CHERI CC [4].

Figure 1 illustrates a comparison between standard memory allocation (malloc()) and a proposed FAT method. The standard approach involves a C program interacting with a custom allocator, utilizing 48-bit free virtual addresses and a TLB walk (L1, L2 and L3 cache) to achieve non-contiguous allocation in physical memory. This typically results in more TLB entries and increased TLB misses increasing the reasoning to have more TLB walks. In contrast,

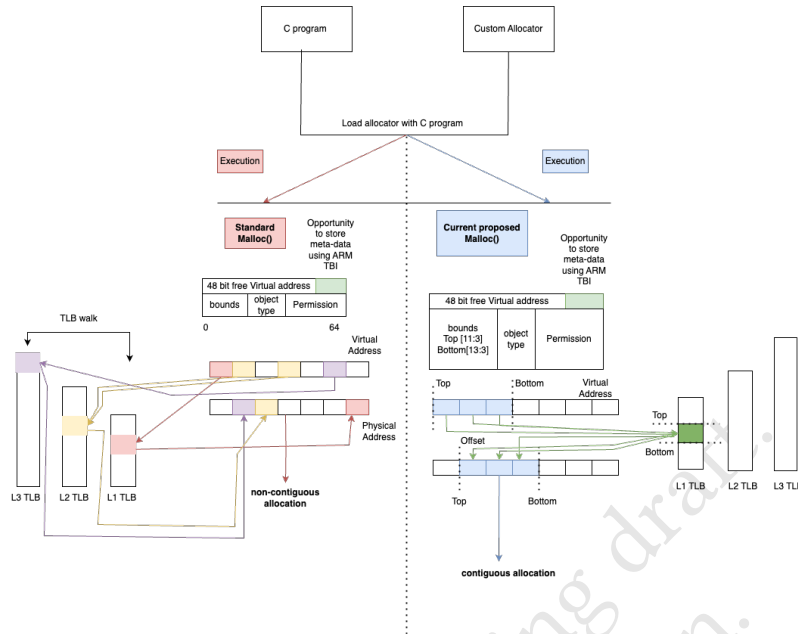


Figure 1: High overview architecture

the Fat-pointer Address Translations method employs a custom allocator leveraging physically contiguous memory by using CHERI to encode bounds within the pointers and as show in the figure 1 there is almost no reliance on walking the TLB hierarchy.

### 3.1 Encoding Ranges as Bounds to the Pointer

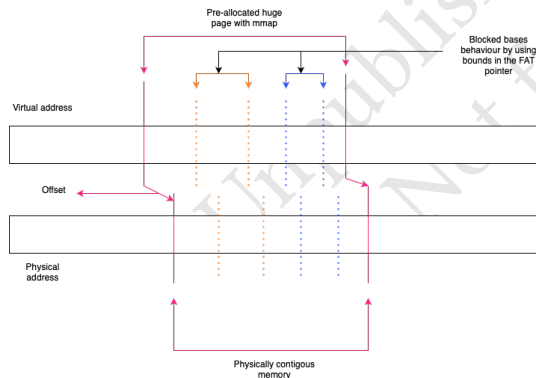


Figure 2: Range of memory

A memory range in FAT has 2 points to track memory in physical contiguous space which is the top and bottom. These 2 points are 2 virtual addresses and the range consists of addresses which lie within this and refers to addresses allocated by invoking malloc. In FAT memory ranges are established using bounds encoded within the pointer, adhering to CHERI CC [4] as referred in section 3.2.

Figure 2 illustrates a straightforward use-case in which the dark pink line represents a single, large contiguous memory area, or huge page. Within this huge page, the orange and blue lines indicate two

separate memory allocations equivalent to invoking malloc twice to allocate memory in distinct regions. This scenario simulates a block-based memory allocator operating within the confines of the huge page. The allocations use the bounds encoded in the FAT-pointer, ensuring tracking of the allocated memory regions. By using the CHERI bounds, this method maintains the contiguity of the allocated blocks within the huge page.

### 3.2 128 bit compressed bounds

We use CHERI CC [4] to track regions of memory in physically contiguous space. CHERI CC consists of compressed bounds that represent a 128-bit pointer within a 64-bit virtual address system. Our approach uses the work of CHERI CC by using a single cycle to decode bounds from the capability register and the bounds which are decoded are repurposed for tracking memory which is eagerly allocated.

Allocators like Jemalloc typically allocate objects under 512 bytes [4]. When an object's bounds cannot be precisely represented, padding is required to ensure memory safety. However, it has been observed that Jemalloc rarely needs more than 6 bits to store the exponent values within compressed bounds. This means that the default behavior of allocators such as Jemalloc, would allow precise representation of bounds within CHERI CC.

With FAT rather than using fixed size TLB entries page sizes (such as 4KB, 2MB and 1GB pages). We use CHERI CC bounds to be dynamic in size to be defined based on the size provided when calling malloc. Offering a more flexible alternative than fixed-size TLB entries.

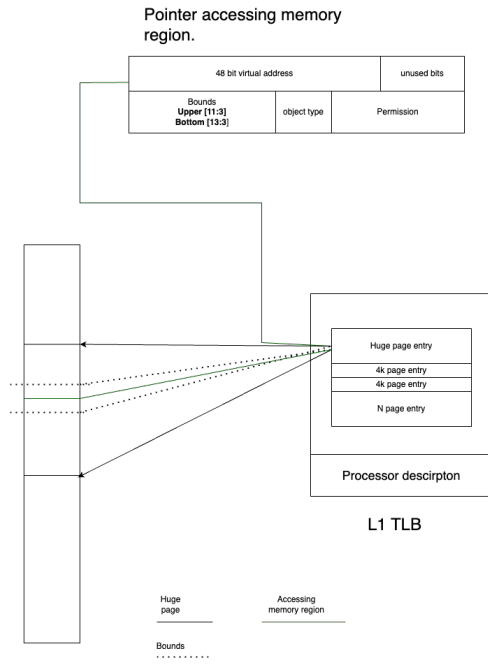


Figure 3: Fat-pointer Address Translations using huge pages

### 3.3 Instrumenting Block-Based Allocators with Physically Contiguous Memory

To build up based on Section 3.1 and 3.2. We are able to pre-allocate memory using huge pages and are able to mark smaller allocations using ranges by storing them as bounds within the pointer. Each of these memory ranges can be called as a block. Since we have numerous blocks inside a huge page we allow block based memory patterns within physically contiguous memory. As demonstrated with the allocator implementation in section 4.

Figure 3 illustrates a use-case of huge pages where the green line represents a sample access to read within a contiguous space of physical memory. The dotted lines represents the bounds for that particular pointer access. Using bounds stored on the pointer a block based pattern can be replicated on physically contiguous memory.

## 4 Memory allocator design

This section presents a straightforward memory allocator designed and implemented based on the principles outlined FAT 3. The allocator consists of three core functions: InitAlloc, malloc, and free. The InitAlloc function initializes the memory pool, setting up the necessary data structures and metadata required for efficient memory management. The malloc function is responsible for allocating a contiguous block of memory of a specified size, while the free function deallocates the memory, returning it to the pool for future use.

When the malloc function (Algorithm 1) is invoked, the algorithm employs an eager allocation strategy for physical memory. This is achieved through the use of the SetBounds mechanism, which constructs a FAT-pointer specialized pointer that encodes

### Algorithm 1 Malloc implementation

```

1: function MALLOC(sz)
2:   sz ← ALIGN_UP(sz, MAX_ALIGNMENT) ▷ Align size to
   max alignment
3:   MallocCounter ← MallocCounter – sz ▷ Update
   remaining memory
4:   ptrLink ← &ptr[MallocCounter] ▷ Calculate pointer
   address
5:   ptrLink ← SET_BOUNDS(ptrLink, sz) ▷ Set bounds for
   memory safety and to track the length of the pointer
6:   return ptrLink ▷ Return allocated memory pointer
7: end function

```

both the start and end addresses of the allocated memory region within the pointer itself. The start and end addresses correspond to the size of the memory block requested by malloc. This approach introduces a method of memory tracking, where the bounds of the allocated region are explicitly encoded in the address, enabling efficient monitoring and management of memory usage.

Furthermore, this design uses shared huge page TLB entries to map and track memory addresses. By encoding bounds directly into the address, the algorithm ensures that memory accesses remain within the allocated region, thereby reducing the risk of out-of-bounds errors. This use of FAT-pointers and shared TLB entries not only aligns with the principles of efficient memory management but also demonstrates a practical usecase of huge pages in CHERI.

### Algorithm 2 Free implementation

```

1: function FREE(ptr)
2:   len ← GET_LENGTH(ptr) ▷ Get length of memory block
   from the defined bounds
3:   UNMAP(ptr, len) ▷ Release memory block
4: end function

```

The memory deallocation (Algorithm 2) mechanism in the proposed allocator is facilitated by the FAT-pointer structure introduced in the malloc algorithm. When the free function is invoked, it uses the metadata embedded within the FAT-pointer to determine the range and size of the allocated memory region. Specifically, the start and end addresses encoded in FAT to provide the necessary information to identify the exact memory block to be deallocated. This allows the allocator to unmap the corresponding memory region from the address space.

By extracting the bounds and size directly from FAT, the free function eliminates the need for additional metadata lookups or complex data structures.

Algorithm 3 describes the initialization of physically contiguous memory through the use of huge pages, a mechanism supported by modern architectures to optimize memory management. The algorithm begins by allocating a fixed block of 1 GB of physically contiguous memory. This decision is driven by the architectural constraints of contemporary systems, particularly ARM-based CPUs, where 1 GB represents the largest supported page size. By leveraging huge pages, the algorithm reduces the overhead associated



**Algorithm 3** Init alloc function to create a initial 1 GB huge page

---

```

1: function INIT_ALLOC
2:   sz ← 1 GB           ▷ Define pre-allocated memory size
3:   fd ← CREATE_LARGE_PAGE_MEMORY(sz)  ▷ Create
   shared memory
4:   ptr ← MAP_MEMORY(sz)      ▷ Map memory region
5:   MallocCounter ← sz        ▷ Initialize memory counter
6: end function

```

---

with page table management and enhances memory access efficiency, which is critical for performance-sensitive applications and kernel-level operations.

## 5 Evaluation

We conducted tests of the FAT Pointer-based range addresses against Jemalloc [10], the default memory allocator for CHERI BSD [11], to assess the performance improvements enabled by a CHERI-based huge page-aware allocator. Specifically, we evaluated the reduction in TLB walks and misses and its impact on wall clock runtime.

To comprehensively analyze the proposed allocator, we categorized benchmarks into two classes which are micro and macro benchmarks. Micro benchmarks comprise smaller C programs designed to target specific allocator patterns, enabling us to evaluate detailed aspects of the allocator's behavior. Macro benchmarks, on the other hand, encompass larger, real-world C programs, allowing us to assess the allocator's performance in more practical, real-world scenarios.

The experiment setup (section 5.1) details the software stack used for evaluation. It includes the specific configurations, compiler options, and system environment tailored to benchmark the proposed allocator. This ensures consistency and repeatability in our results, providing a solid foundation for meaningful comparisons.

We further elaborated on the two classes of benchmarks executed. Micro benchmarks (section 5.2.1), focused on particular allocation and deallocation patterns, such as sequential and random memory accesses, to stress-test the allocator under controlled conditions. Macro benchmarks (section 5.2.2) involved real-world applications, offering insights into how the allocator performs with complex memory allocation demands, large datasets, and varying execution contexts.

The results (section 5.3) presents the outcomes of our benchmarks, highlighting key metrics such as TLB miss rates, memory usage, and runtime performance. We observed that the proposed allocator demonstrated significant improvements in reducing TLB misses, leading to noticeable enhancements in runtime efficiency for both micro and macro benchmarks. The behavior of specific allocation patterns and their impact on memory performance is detailed, providing a nuanced understanding of the allocator's effectiveness.

### 5.1 Experiment setup

The CHERI Morello [12] board was used to evaluate the proposed memory allocator. Morello implements the ARM A76 with enhanced server-class memory, featuring a quad-core ARM CPU with capability extensions. The L1 and L2 caches were modified to proliferate

the capability bit, ensuring compatibility with CHERI's capability-based memory model. When compiling the C programs for benchmarking, the Benchmark ABI was used as recommended by the CHERI community. This compilation mode was enabled using the Clang compiler.

The Benchmark ABI [13] was specifically designed because the Morello branch predictor was not expanded to predict bounds. Consequently, a capability-based jump introduces stalls in later PCC-dependent instructions until bounds are established. This issue is particularly significant during dynamically linked calls and returns between libraries, where bounds are changed to cover the called or returned-to library. Such stalls can negatively affect performance, making the Benchmark ABI an essential consideration for this evaluation.

Each C program was executed using two different memory allocators. The first was the modified C allocator, imported as a header file. This approach was necessary because the Benchmark ABI shared object file exhibited unexpected behavior, failing to overwrite the C program at runtime with the intended malloc functions. The second allocator was the standard OS memory allocator, which, in the case of CHERI BSD, is Jemalloc.

Performance measurements were carried out using ARM performance counters [14] to ensure accurate evaluation. These counters provided detailed metrics, allowing us to compare the performance of the two allocators and assess the impact of the proposed changes.

### 5.2 Benchmarks

The benchmarks [15] are classified into 2 classes:

#### 5.2.1 Micro benchmark.

- **GLIBC:** The Glibc benchmark evaluates the performance of malloc and free functions in single-threaded, multi-threaded, and emulated multi-threading scenarios using various block sizes and allocation patterns. It simulates real-world memory usage by partially deallocating blocks in FIFO order and fully deallocating them in LIFO order. Results are gathered across configurations to analyze performance variations.
- **MemAccess:** This benchmark by Alex Bordei evaluates the performance impact of memory access patterns by constructing and traversing a doubly linked list with varying working set sizes. It supports sequential or randomized structures, optional node operations, and multithreaded traversal using pthreads. The program dynamically allocates memory and systematically doubles the working set size to analyze memory hierarchy behavior.

#### 5.2.2 Macro benchmark.

- **Kmeans:** Kmeans implements a parallelized K-means clustering algorithm that assigns data points to clusters based on proximity to centroids, iteratively updating them until convergence. The computation is distributed across threads using the pthread library, dynamically assigning tasks to optimize performance. Parameters like data size and clusters are configurable, and the program ensures efficient memory management and synchronization.
- **Richards:** Richards is a task scheduling benchmark that simulates a multitasking environment with tasks of varying

types and priorities, communicating through queued packets. The schedule function manages task execution based on state and priority, tracking processed packets and held tasks for performance evaluation. Configurable iterations and timing help measure system performance and ensure correctness.

- BARNES: Implements the Barnes-Hut algorithm to efficiently simulate the interactions within an  $N$ -body system. A comprehensive overview of the Barnes-Hut method is provided by Singh in his doctoral dissertation [16]. The implementation we benchmark extends the original method by permitting multiple particles to be stored within each leaf cell of the spatial decomposition, enhancing performance and scalability. This extension is described by Holt and Singh [17].

### 5.3 Results

The graph (Figure 4) highlights the performance comparison between the modified memory allocator and Jemalloc, the default memory allocator. The FAT pointer memory allocator, specifically optimized for use with huge pages, demonstrates a clear advantage in scenarios where memory allocation patterns benefit from its design. The results align with expectations, showcasing the impact

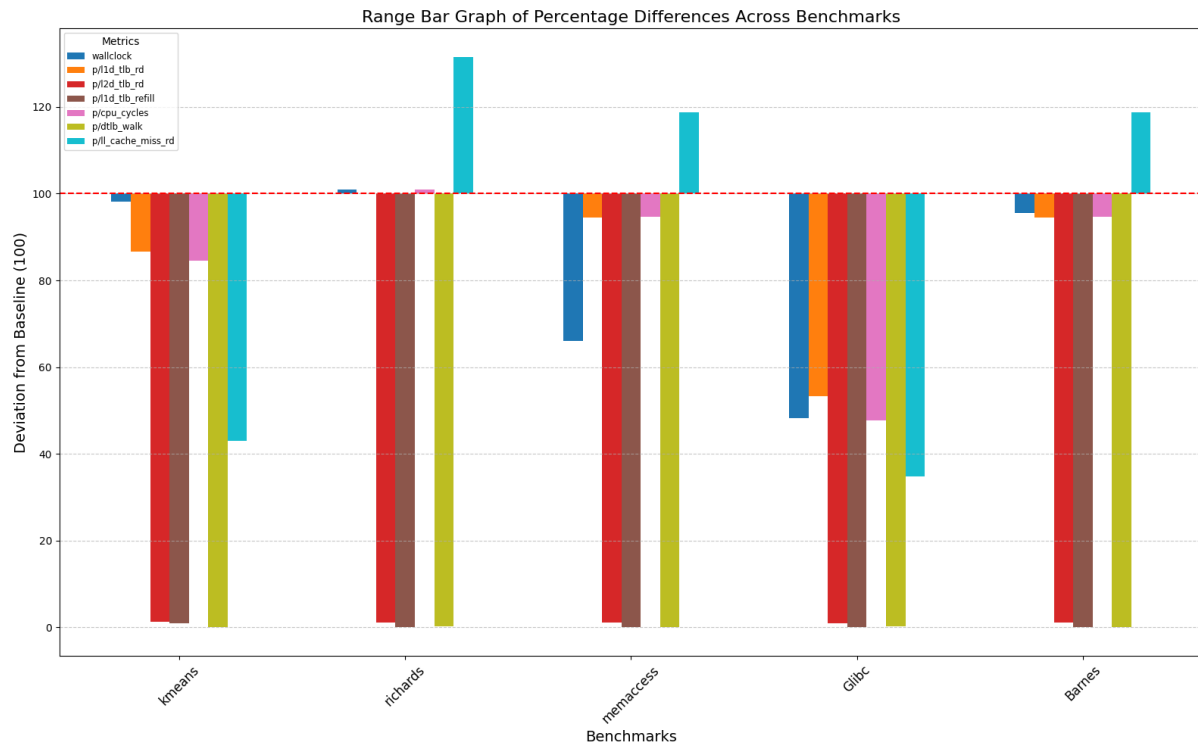
of its capability to handle memory more efficiently by leveraging huge pages.

- L1 DTLB reads: There was noticeable reduction of L1 DTLB reads for kmeans of about an average of 4% lesser and for Glibc there was significant reduction of 50% lesser than Jemalloc.
- L2 DTLB reads: For all the benchmarks on figure 4 there was on average 98% reduction on L2 DTLB reads. This demonstrates that all the TLB translations are read at the L1 TLB cache.
- DTLB walks: Due to most of the TLB entries getting hit at the L1 DTLB there is no need to walk the TLB cache hierarchy. This is shown by an average of 99% reduction in DTLB walks.
- L1 DTLB refills: Since there are lesser DTLB walks and most reads are done at the L1 DTLB layer there is no need for numerous TLB refills to take place. Our benchmarks show on average a 99% reduction on DTLB refills.

A particularly striking observation is the significant reduction in data TLB walks, L2 data TLB reads, and TLB refills—consistently showing a 90% decrease across all benchmarks compared to Jemalloc. This improvement is due to the modified allocator's use of a

Table 1: ARM performance counters

| Performance counter                              | Description                                                                                                                                                                                                                                                                                    |
|--------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Wall clock                                       | The actual time taken from the start of a computer program to the end.                                                                                                                                                                                                                         |
| (p/l1d_tlb_rd) L1 data TLB reads                 | Level 1 data TLB access, read                                                                                                                                                                                                                                                                  |
| (p/l2d_tlb_rd) L2 data TLB reads                 | Level 2 data TLB access, read                                                                                                                                                                                                                                                                  |
| (p/l1d_tlb_refill) L1 data TLB refills           | Level 1 data TLB refill.<br>The Level 1 data TLB refill counter tracks each access to the L1D_TLB that results in a refill of the Level 1 data or unified TLB. This includes any access that requires a memory lookup due to a translation table walk or accessing another level of TLB cache. |
| (p/cpu_cycles) CPU cycles                        | The CPU CYCLES counter increases with every clock cycle. However, it can be affected by changes in clock frequency, such as when WFI (Wait for Interrupt) or WFE (Wait for Event) instructions pause the clock.                                                                                |
| (p/dtlb_walk) Data TLB walks                     | Data TLB access with at least one translation table walk.                                                                                                                                                                                                                                      |
| (p/ll_cache_miss_rd) Last level cache miss reads | Last level cache miss, read<br>(This refers to every miss in the Last level cache that occurs during a memory read operation.)                                                                                                                                                                 |



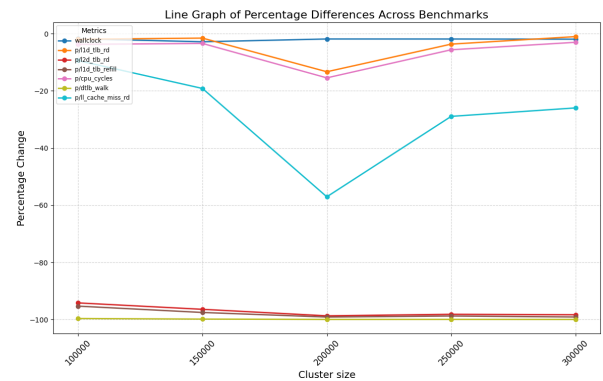
**Figure 4: Percentage difference between the modified memory allocator against the default system memory allocator**

single huge page entry at the L1 TLB layer. By enabling most address translations to be resolved directly at the L1 TLB, the need to walk through the deeper TLB hierarchy is largely eliminated. This reduction in translation overhead is a key factor in the allocator's performance for certain types of workloads.

The micro benchmarks, which are crafted to emphasize memory read operations, highlight the allocator's strengths. These tests simulate frequent and intensive memory access patterns, where the reduction in TLB misses directly translates into measurable performance gains. On average, the FAT allocator achieves a 50% reduction in wall clock runtimes for these workloads, underscoring its ability to optimize high-throughput memory operations.

On the other hand, macro benchmarks, which represent larger and more complex real-world applications, exhibit minimal differences in wall clock runtimes when using the FAT allocator. This outcome is expected, as macro benchmarks typically involve a broader range of operations beyond memory allocation, diluting the impact of the allocator's optimizations. Additionally, the benefits of huge pages may be less pronounced for these workloads, as they are often bottlenecked by factors such as computation or I/O rather than memory translation overhead.

The K-means algorithm was executed with varying cluster sizes to evaluate the performance difference between the FAT pointer allocator and the baseline allocator as the workload scales. This analysis aims to understand how the allocator's optimizations, particularly its ability to manage memory more efficiently with huge pages, impact performance under different workload conditions.



**Figure 5: Kmeans COZ benchmark executed against various cluster sizes**

For most cluster sizes tested, the percentage difference in performance remained relatively consistent. This indicates that the allocator's efficiency scales predictably with increasing workload sizes, suggesting a stable and uniform benefit across different configurations. The consistent performance gain is likely due to the allocator's ability to minimize TLB misses and efficiently manage memory allocations for the centroid and data point structures used in the K-means algorithm.

However, an anomaly was observed at a cluster size of 2000, where the percentage difference deviated significantly from the trend. At this cluster size, the memory access patterns and allocation behavior may align in a way that temporarily offsets the advantages of the FAT allocator. For example, the memory layout might interact with system-level caching mechanisms or TLB behavior differently, leading to an unexpected change in performance. Additionally, the increased complexity of managing a higher number of clusters might introduce computational overhead that overshadows the memory allocator's optimizations.

This observation highlights the importance of testing across a range of workload sizes and configurations to uncover edge cases or specific scenarios where performance deviates from the expected pattern. Understanding these anomalies can provide insights into the allocator's behavior and guide future improvements to address such outliers. Despite the deviation at a cluster size of 2000, the overall results reaffirm the allocator's capability to maintain consistent performance benefits across most scenarios.

## 5.4 Analysis

The FAT memory allocator demonstrates significant potential for enhancing memory management in systems that benefit from huge page optimizations. Its design effectively reduces TLB misses, achieving up to 90% fewer data TLB walks, L2 TLB reads, and TLB refills compared to Jemalloc. These improvements lead to noticeable performance gains, especially in micro benchmarks, where the allocator reduces wall clock runtimes by an average of 50%.

The allocator integrates seamlessly into memory-intensive workloads, as evidenced by its consistent performance across varying cluster sizes in the K-means benchmark, with only minor anomalies observed under specific conditions. These outliers provide valuable insights into the allocator's interaction with system-level caching and memory translation mechanisms.

While the allocator excels in scenarios emphasizing high memory throughput, its impact on macro benchmarks is less pronounced. This suggests that its benefits are most relevant for applications with frequent and intensive memory operations rather than those constrained by computation or I/O bottlenecks.

## 6 Future work

The current experimental setup on the ARM Morello board is constrained by the requirement that all memory reads must pass through the (TLB) for address translation. This necessitates frequent TLB lookups, potentially leading to performance bottlenecks. The planned future work aims to address this by leveraging CHERI (Capability Hardware Enhanced RISC Instructions) extensions on the RISC-V architecture, specifically using the Tooba implementation.

### 6.1 Storing Offsets Directly on Pointers

In the current ARM Morello setup, address translations rely on the TLB. The future approach on RISC-V Tooba involves storing the offset directly within the pointer. This is possible due to CHERI's capability model, which supports fine-grained memory protection and can encode bounds within pointers. Utilizing Bounds in CHERI for Block-Based Allocation:

CHERI capabilities allow pointers to carry metadata about memory bounds, providing hardware-enforced memory safety. By encoding the offset and bounds within the pointer, the system can directly access memory without needing intermediate translations via the TLB. This enables the implementation of a block-based allocator that can efficiently manage memory allocations and deallocations within defined bounds. Bypassing the TLB in RISC-V Tooba.

### 6.2 Hardware Modifications:

The Bluespec design of the RISC-V processor will be modified to allow certain memory operations to bypass the TLB. This means that when a pointer with encoded offset and bounds is used, the system can directly compute the physical address from the capability information. This modification reduces the dependency on the TLB, decreasing latency and improving performance, especially for frequent memory operations.

## 7 Conclusion

This paper addresses the growing disparity between application workloads and the capacity of TLBs. To mitigate this gap, we proposed leveraging physically contiguous memory with CHERI bounds to reduce TLB walks. We designed a memory allocator which uses huge pages with CHERI CC scheme to track allocations within the allocated huge page. This approach has helped reduce the number of TLB entries needed while using bounds to minimize fragmentation.

Comprehensive benchmarking demonstrates that the allocator reduces TLB misses by up to 90%, leading to substantial performance gains in memory-intensive workloads, though the improvements are less pronounced for larger, computation-heavy applications. These results highlight the allocator's potential to advance memory management by combining enhanced security and performance through CHERI's capability-based model with the use of huge pages.

## References

- [1] Sparsh Mittal. A survey of techniques for architecting TLBs. 29(10):e4061. [\\_eprint: https://onlinelibrary.wiley.com/doi/pdf/10.1002/cpe.4061](https://onlinelibrary.wiley.com/doi/pdf/10.1002/cpe.4061).
- [2] Ashish Panwar, Sorav Bansal, and K. Gopinath. HawkEye: Efficient fine-grained OS support for huge pages. In *Proceedings of the Twenty-Fourth International Conference on Architectural Support for Programming Languages and Operating Systems*, pages 347–360. ACM.
- [3] Jonathan Woodruff, Robert N.M. Watson, David Chisnall, Simon W. Moore, Jonathan Anderson, Brooks Davis, Ben Laurie, Peter G. Neumann, Robert Norton, and Michael Roe. The CHERI capability model: revisiting RISC in an age of risk. 42(3):457–468.
- [4] Jonathan Woodruff, Alexandre Joannou, Hongyan Xia, Anthony Fox, Robert M. Norton, David Chisnall, Brooks Davis, Khilan Gudka, Nathaniel W. Filardo, A. Theodore Markettos, Michael Roe, Peter G. Neumann, Robert N. M. Watson, and Simon W. Moore. CHERI concentrate: Practical compressed capabilities. 68(10):1455–1469.
- [5] Juan Navarro. Practical, transparent operating system support for superpages.
- [6] Cheol Ho Park and Daeyeon Park. Aggressive superpage support with the shadow memory and the partial-subblock tlb. *Microprocessors and Microsystems*, 25(7):329–342, 2001.
- [7] Arkaprava Basu, Jayneel Gandhi, Jichuan Chang, Mark D. Hill, and Michael M. Swift. Efficient virtual memory for big memory servers. *SIGARCH Comput. Archit. News*, 41(3):237–248, June 2013.
- [8] Vasileios Karakostas, Jayneel Gandhi, Furkan Ayar, Adrián Cristal, Mark D. Hill, Kathryn S. McKinley, Mario Nemirovsky, Michael M. Swift, and Osman Ünşal. Redundant memory mappings for fast access to large memories. In *Proceedings of the 42nd Annual International Symposium on Computer Architecture*, pages 66–78. ACM.



- [9] Dongwei Chen, Dong Tong, Chun Yang, Jiangfang Yi, and Xu Cheng. FlexPointer: Fast address translation based on range TLB and tagged pointers. 20(2):1–24.
- [10] JEMALLOC.
- [11] Benchmark ABI - CheriBSD 23.11 new features tutorial.
- [12] Department of computer science and technology – ChERI: The arm morello board.
- [13] Robert N. M. Watson, Jessica Clarke, Peter Sewell, Jonathan Woodruff, Simon W. Moore, Graeme Barnes, Richard Grisenthwaite, Kathryn Stacer, Silviu Baranga, and Alexander Richardson. Early performance results from the prototype Morello microarchitecture. Technical Report UCAM-CL-TR-986, University of Cambridge, Computer Laboratory, 15 JJ Thomson Avenue, Cambridge CB3 0FD, United Kingdom, phone +44 1223 763500, September 2023.
- [14] Arm architecture reference manual for a-profile architecture.
- [15] ChERI-allocator/benchmarks/benchmarks/StressTestMalloc/glibc-bench.c at main · akilan1999/ChERI-allocator.
- [16] Jaswinder Pal Singh. *Parallel Hierarchical N-body Methods and Their Implications for Multiprocessors*. PhD thesis, Stanford University, February 1993.
- [17] C. Holt and Jaswinder Pal Singh. Hierarchical n-body methods on shared address space multiprocessors. In *SIAM Conference on Parallel Processing for Scientific Computing*, February 1995. To appear.