

FAT allocator

Akilan Selvacoumar*

as251@hw.ac.uk

Heriot Watt

Edinburgh, Scotland, United Kingdom

Abstract

Modern processors increasingly suffer from Translation Lookaside Buffer (TLB) pressure, where frequent misses lead to costly page table walks and degraded performance. While huge pages can mitigate some of this overhead, allocator-level inefficiencies persist.

We present Fat Address Translations (FAT), a memory allocator for capability-aware architectures such as CHERI. FAT embeds allocation metadata directly within pointer capabilities and manages memory in block-based allocations inside huge pages, reducing reliance on page table lookups and enabling more efficient pointer dereferencing. Evaluation shows that FAT reduces TLB walks by up to 99% and L2 TLB reads by 98%. Runtime performance improves by up to 11% in macro-benchmarks (e.g. XSBench), 1.8–4% in Kmeans and Barnes, and 4% in micro-benchmarks such as Memaccess, while Richards remains near baseline. These results demonstrate that capability-based addressing can be repurposed to mitigate TLB pressure, offering a scalable and efficient strategy for memory allocation in future capability-aware systems.

ACM Reference Format:

Akilan Selvacoumar. 2018. FAT allocator. In *Proceedings of Make sure to enter the correct conference title from your rights confirmation email (Conference acronym 'XX)*. ACM, New York, NY, USA, 10 pages. <https://doi.org/XXXXXX.XXXXXXX>

1 Introduction

In computing, achieving high performance is an ongoing challenge, especially as modern applications handle increasingly memory intensive workloads. Memory management is a key factor in reducing the time it takes to access a memory region. A Translation Lookaside Buffer (TLB) is a specialised cache in the memory management unit (MMU). It reduces the time required to convert virtual addresses to physical addresses. When a program accesses data in memory, the MMU first checks the TLB for a matching entry and avoids the slow process of accessing the page tables situated in memory. TLBs are crucial in speeding up memory access by caching recent memory address translations. The TLB hierarchy [1] comprises of a smaller, faster L1 TLB situated close to the processor core and a larger slightly slower L2 TLB, both of which are usually private to each core thereby reducing the latency associated with page table lookups. However, as applications grow more complex the TLBs

often cannot keep up with the number of translation entries leading to more TLB misses which requires walking the TLB hierarchy and this leads to performance slowdowns [2]. To tackle this issue, researchers have explored new solutions, including the use of huge pages [3].

Huge pages also known as large pages allow for the allocation of memory in significantly larger chunks compared to traditional small pages which leads to fewer number of TLB entries needed to access a given amount of memory. Huge pages offer a potential avenue for optimising TLBs which are used 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 in regards to traversing the TLB cache hierarchy and therefore minimising these bottlenecks.

Simultaneously, advancements in hardware-level security, such as the Capability Hardware Enhanced RISC Instructions (CHERI) [4] which are standard 64-bit pointers that are replaced with capabilities with 128-bit or 256-bit encoding scheme which consists of both the address and metadata. The metadata includes bounds, permissions and validity of the pointer which forms a capability. This transformation of larger pointers enforces strict memory safety as capabilities prevent arbitrary pointer arithmetic and unauthorised memory access. 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 [5] with the use of huge pages, one of the contribution is demonstrating that it is possible to track, manage large and physically contiguous memory blocks without requiring numerous TLB entries. This combination reduces the TLB pressure by minimising the number of entries required to map extensive memory regions thereby decreasing TLB misses and improving address translation performance. Furthermore, this paper introduces the Fat Address Translations (FAT) allocator which accelerates memory-intensive tasks by reducing the overhead associated with managing non-contiguous memory allocations by emulating block allocations on physically contiguous memory to reduce the TLB pressure. The contributions for the following paper are as follows:

- **FAT Addresses Translations:** Introduces FAT that include memory bounds, allowing efficient tracking and management of physically contiguous memory regions (Section 3).
- **CHERI's Capability-based Optimisation:** Demonstrates how CHERI's architecture can be used to optimise memory allocation by encoding memory bounds directly within pointers, reducing TLB reliance (Section 4).
- **Memory Allocation Algorithms (FAT allocator):** Presents an algorithm for allocating, freeing physically contiguous memory, and integrating huge pages with CHERI's capability-based bounds for enhanced memory management (Section 5).

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.

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/XXXXXX.XXXXXXX>

2025-08-30 13:00. Page 1 of 1–10.

- **Modified Jemalloc Memory Mapping with FAT allocator:** Modification to Jemalloc with the FAT allocator to support physically contiguous memory allocation using a block-based strategy with capability-based addressing (Section 6).

Through evaluating micro and macro benchmarks the FAT allocator through the use of CHERI's capabilities and huge pages demonstrates the allocator's ability to reduce TLB walks by up to 99% which yields to improvements of wall clock runtimes by upto 11% for memory-intensive applications. While its impact on larger and 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.

2 Related work

2.1 Huge Pages

Increasing TLB reach [6] can be achieved by using larger page sizes such as huge pages [3] which are common in modern computer systems. The x86-64 architecture supports huge pages of 2 MB and 1 GB which are backed by OS mechanisms like Transparent Huge Pages (THP) [7]. 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 (Page Table Entries) and in contrast using a 2 MB huge page would waste half of the memory space. Some architectures offer more page size choices, such as Intel Itanium [8] which allows different areas of the address space to have their own page sizes. Itanium uses a hash page table to organise huge pages and without significant changes to the conventional page table. It only helps reduce page walk overheads. Huge page tunable base page size permits the OS to adjust the base page, but still faces internal fragmentation problems with huge page recommending a base page size of no more than 16 KB. Shadow Superpage [9] introduces a new translation level in the memory controller to merge non-contiguous physical pages into a huge page in a shadow memory space by 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 [10] 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. 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 Redundant Memory Mapping (RMM)

Redundant Memory Mappings (RMM) [11] enhance memory management by introducing an additional range table that pre-allocates contiguous physical pages for large memory allocations thus creating ranges that are both virtually and physically contiguous. This approach simplifies address translation by adding ranges. RMM is similar to Direct Segment with the support of multiple ranges and operates transparently to programmers. No source code modifications are required on the programmers side. The range table which is 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. Since this process is computationally expensive hence it is therefore performed only after an L1 TLB miss. To optimise this, RMM uses a Range TLB (RTLb) to quickly identify if an address falls within any pre-allocated range. Range mapping works alongside the paging system by generating TLB entries on TLB misses and still performs TLB lookups for each virtual address translation. Unlike traditional segmentation mechanisms, range mapping activates the RTLb located with the last level TLB upon a miss. The hardware TLB miss handler then searches the RTLb for the missed address. If found, generates a new TLB entry with the physical address derived from the base virtual address and range offset along with the 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 on the background. This avoids delays for in-memory operations. The RTLb functions as a fully associative search structure ensuring that most last-level TLB misses are handled efficiently by range mapping which reduces the need for costly page table walks.

2.4 FlexPointer

The key insight behind FlexPointer [12] is that large memory objects are relatively uncommon. This allows memory ranges to be constructed around them and assigned unique identifiers. These range IDs are embedded within the unused bits of pointers, enabling direct indexing of the range TLB and simplifying its design. Since the range ID is unaffected by address generation, range lookups can occur earlier in the pipeline, concurrently with address computation. Simulation results show that FlexPointer significantly reduces L1 TLB misses and the need for page walks in a range of memory-intensive workloads. When compared with a traditional 4KB-page system, FlexPointer delivers an average performance improvement of 14%, with peak gains of up to 2.8 $\times$, and introduces no performance regressions in less demanding scenarios.

2.5 CHERI

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

to introduce architectural capabilities and is a new microarchitecture which shows that capabilities can be implemented efficiently in hardware. CHERI provides support for efficient tagged memory to protect capabilities and compresses them to reduce memory overhead. The CHERI ecosystem provides language and compiler extensions for using capabilities with C and C++. An OS extensions is also provided to support fine-grained memory protection (including spatial, referential, and non-stack temporal memory safety) and abstraction extensions for scalable software compartmentalisation.

2.6 CHERI CC

CHERI CC(Concentrate Compressed Capabilities) [5] introduces a compression scheme for CHERI that aims to address the performance and compatibility challenges associated with capability pointers. Capability pointers enhance memory safety by embedding bounds and permissions directly within pointers. Traditional implementations of CHERI double their size leading in bounds to increased memory usage. CHERI CC proposes a compression strategy that preserves security while reducing size and inefficiencies. Key contributions include a floating-point bound encoding technique with an internal exponent mechanism that offers greater precision for smaller objects and optimised space usage for larger ones.

3 Fat Address Translations

Fat Address Translations (FAT) is our new memory allocator for CHERI. It uses the CHERI architecture to bring about block-based allocations in physically contiguous memory. FAT leverages techniques like FlexPointer [12], RMM [11] in context of introducing a range table to store custom address translation memory ranges to reduce pressure on the TLB. A key component in this implementation is the use of range addresses with CHERI CC [5].

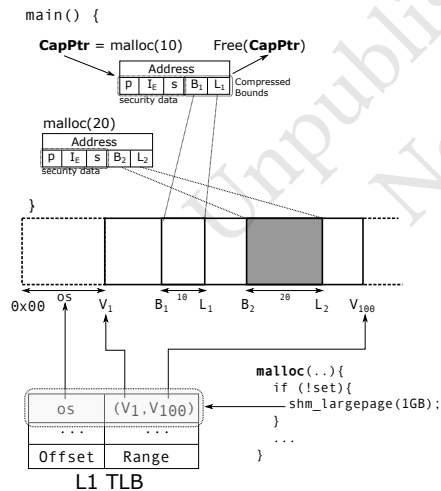


Figure 1: High overview architecture

Figure 1 illustrates a high-level overview of the FAT allocator. The dotted rectangular box represents addresses in memory, and the box from v1 to v100 represents a physically contiguous region of memory. This physically contiguous region is typically allocated using huge pages. When malloc is called, a capability pointer is

returned with metadata related to the bounds information encoded within the pointer. The bounds encoded in the pointer are reused for tracking memory blocks within a huge page. As shown, this in turn reduces TLB pressure.

3.1 Encoding Ranges as Bounds to the Pointer

Range is defined from work based on RMM [11] and FlexPointer [12]. In RMM a range refers to a contiguous region of memory defined by a base address and a bound. This information is stored in a hardware managed table called Range Table which is called when there is a L1 TLB miss. FlexPointer builds up on the work of RMM and stores the ID value between 48th bit and 64th bit which is the value to check the range table on parallel to the L1 TLB lookup.

The FAT allocator builds up on the concept of range from RMM and FlexPointer. Instead of using a hardware range table using CHERI range information can be encoded within a capability pointer. A memory range in FAT has two points to track memory in physically contiguous space which is the top and bottom. These two points are two virtual addresses and the range consists of addresses that lie within this and refers to addresses allocated by invoking *malloc*. FAT memory ranges are established using bounds encoded within the pointer, adhering to CHERI CC [5] compression scheme.

4 128 bit compressed bounds

FAT uses CHERI CC [5] 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 that are decoded are repurposed for tracking memory which is eagerly allocated.

Allocators like Jemalloc typically allocate objects under 512 bytes. 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 (as shown in [5]). This means that the default behavior of allocators such as Jemalloc, would allow precise representation of bounds within CHERI CC.

Instead of relying on fixed-size TLB entries with set page sizes (such as 4KB, 2MB, or 1GB), FAT uses CHERI CC to define dynamic bounds based on the size requested at allocation time (e.g., during a *malloc* call). This approach offers flexible translation entry sizes compared to the traditional fixed-size TLB model.

This means that the default behavior of most allocators, such as Jemalloc, would allow precise representation of bounds within a FAT. These pointers can then be repurposed as memory ranges in custom memory allocators, offering a more flexible alternative to fixed-size TLB entries.

4.1 Instrumenting Block-Based Allocators with Physically Contiguous Memory

FAT is able to pre-allocate memory using huge pages and is able to mark smaller allocations using ranges by storing them as bounds within the pointer. Each of these memory ranges can be called a block. Since there are numerous blocks inside a huge page, This

allows for abbreviated block-based memory patterns within physically contiguous memory.

By consolidating address translations into a single TLB entry, this method cuts down on the overhead of managing many entries. It also takes advantage of the bounds encoded within FAT to track and access memory within physically contiguous memory.

5 Memory allocator design

This section presents a memory allocator design which is implemented based on the principles outlined in FAT (Section 3). The allocator consists of three core functions: *InitAlloc*, *malloc*, and *free*. The *InitAlloc* function initialises 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. When the *free* function deallocates memory it is returned to the pool for future use.

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

```

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. This constructs a FAT-specialised pointer that encodes 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 is explicitly encoded in the address which enables efficient monitoring and management of memory usage.

Furthermore, this design uses 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 and shared TLB entries not only align with the principles of efficient memory management but also demonstrate a practical use case 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 structure introduced in the *malloc* algorithm. When the *free* function is invoked, it uses the metadata embedded within FAT to determine the range and size of the allocated memory region. Specifically, FAT encodes the start and end addresses of each allocation and provides the information needed to identify the memory block to be deallocated. This enables the allocator to accurately 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 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

```

Algorithm 3 describes the initialisation of physically contiguous memory through the use of huge pages which is a mechanism supported by modern architectures to optimise memory management. The algorithm begins by allocating a fixed block of 1 GB 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 with page table management and enhances memory access which is critical for performance-sensitive applications and kernel-level operations.

6 Embedding FAT allocator inside Jemalloc

This section describes the FAT allocator implementation (Section 5) embedded inside Jemalloc. The objective here is to describe the changes needed for a block based allocator to use physically contiguous memory with a block based strategy with the help of capability based addresses. In the case of Jemalloc the only changes required was to replace the *mmap* with the *malloc* function (Algorithm 1) and for *munmap* with the *free* function (Algorithm 2).

6.1 Mmap replaced with MALLOC

The *os_pages_map* function (Algorithm 4) simulates Jemalloc's low-level page mapping routine. It first checks if a specific address is requested in case relevant to CheriABI [13] where such behavior is disallowed and returns NULL. It then allocates memory using the custom MALLOC(size) (Algorithm 1) function and validates whether the returned pointer matches the requested address if one was provided. If there's a mismatch, it unmmaps calling FREE() (Algorithm 2) the memory and returns NULL; otherwise, it returns the allocated pointer.

This approach mentioned above is embedded inside Jemalloc's strategy of managing memory through arenas [14] and size [14] classes. In Jemalloc, memory is divided into chunks, which are

Algorithm 4 Modified Jemalloc Memory Mapping Routines

Function: OS_PAGES_MAP(ADDR, SIZE, ALIGNMENT, COMMIT)
Require: addr aligned to os_page, size aligned to os_page, size $\neq 0$
Ensure: If addr is non-NULL (CheriABI), return NULL

```

1: if addr  $\neq$  NULL then
2:   return NULL
3: end if
4: if os_overcommits then
5:   commit  $\leftarrow$  true
6: end if
7: ret  $\leftarrow$  MALLOC(size)
8: assert(ret  $\neq$  NULL)
9: if addr  $\neq$  NULL and ret  $\neq$  addr then
10:  os_pages_unmap(ret, size)
11:  ret  $\leftarrow$  NULL
12: end if
13: return ret

```

Function: PAGES_MAP(ADDR, SIZE, ALIGNMENT, COMMIT)
Require: alignment $\geq$ PAGE, addr aligned to alignment

```

14: ret  $\leftarrow$  MALLOC(size)
15: return ret

```

Function: PAGES_COMMIT_IMPL(ADDR, SIZE, COMMIT)
Require: addr aligned to PAGE, size aligned to PAGE

```

16: result  $\leftarrow$  MALLOC(size)
17: if result  $\neq$  addr then
18:  os_pages_unmap(result, size)
19:  return true
20: end if
21: return false

```

further subdivided into runs and regions to handle allocations of various sizes efficiently. By aligning sizes and managing allocations within predefined structures, Jemalloc minimizes fragmentation [15].

Algorithm 5 os_pages_unmap

Require: addr aligned to os_page, size aligned to os_page
Ensure: Memory region at addr is unmapped

```

1: assert(ALIGNMENT_ADDR2BASE(addr, os_page) = addr)
2: assert(ALIGNMENT_CEILING(size, os_page) = size)
3: FREE(addr)

```

6.2 Munmap replaced with FREE

The os_pages_unmap (Algorithm 5) represents a customized abstraction of Jemalloc's memory unmapping routine, designed to integrate with the previously defined simplified free(ptr) implementation. In conventional Jemalloc configurations, os_pages_unmap would invoke low-level system calls such as munmap to release virtual memory pages back to the operating system. However, in this adapted version, the function instead delegates the deallocation to a higher-level FREE(addr) (Algorithm 2).

The function begins by enforcing two invariants through assertions: first, that the input address addr is aligned to the operating system's page size and second that the size of the memory region is also a multiple of os_page. These alignment checks are critical for maintaining consistency with Jemalloc's internal page-based memory management semantics and ensuring compatibility with the allocator's expectations. Following these checks the memory at the specified address is deallocated via the FREE(addr) (Algorithm 2) operation.

The following changes done to free is embedded inside Jemalloc's deallocation mechanism, where metadata associated with each allocation (such as size and location) is used to efficiently return memory to the appropriate arena or pool. Jemalloc maintains separate metadata structures to track allocations, allowing for quick deallocation and reuse of memory blocks without significant overhead [15].

7 Evaluation

Benchmarks of the FAT memory allocator and the FAT allocator embedded within Jemalloc against the standard Jemalloc [14] allocator was conducted. Jemalloc is the default memory allocator for CHERI BSD [16]. The objective was to evaluate the reduction of TLB walks and misses and its impact on the wall clock runtime.

To comprehensively analyse the implemented allocator, the benchmarks [17] were categorised into two classes which are micro and macro benchmarks. Micro benchmarks comprise smaller C programs designed to target specific allocator patterns such as the memory read operations which enables us to evaluate detailed aspects of the allocators behavior. Macro benchmarks, on the other hand, encompass larger real-world C programs allowing us to assess the allocators performance in a more practical and real-world scenarios.

7.1 Experiment setup

The CHERI Morello [18] board was used to evaluate the proposed memory allocator. Morello implements the ARMv8 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 which ensures compatibility with CHERI's capability-based memory model. When compiling the C programs for benchmarking, the Benchmark ABI [19] was used as recommended by the CHERI community. This compilation mode was enabled using the Clang compiler.

The Benchmark ABI 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 (Program Counter Capability) 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 FAT allocator (Section 5) which is imported as a header file. This approach was necessary because the Benchmark ABI shared object file exhibited unexpected behavior by 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 CHERIBSD is Jemalloc.

Performance measurements (table 1) were carried out using ARM hardware performance counters [20] 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.

7.2 Benchmarks

Elaborated here are two classes of benchmarks. Micro benchmarks focus on particular allocation, deallocation patterns such as sequential and random memory accesses. This is to stress-test the allocator under controlled conditions. Macro benchmarks involves real-world applications offering insights into how the allocator performs with complex memory allocation demands such as large datasets with varying execution contexts.

7.2.1 Micro benchmark.

- **MemAccess:** This benchmark 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 randomised structures with optional node operations and multithreaded traversal using pthreads. The program dynamically allocates memory and systematically doubles the working set size to analyse memory hierarchy behavior.

- **Richards:** Richards is a task scheduling benchmark that simulates a multitasking environment with tasks of varying types and priorities which is communicated through queued packets. The schedule function manages task execution based on the state, priority and tracks processed packets which are held tasks for performance evaluation. Configurable iterations and timing help measure system performance to ensure correctness.

7.2.2 Macro benchmark.

- **Kmeans:** Kmeans implements a parallelised K-means clustering algorithm that assigns data points to clusters based on proximity to the centroids. This iteratively updates them until convergence. The computation is distributed across threads using the pthread library, dynamically assigning tasks to optimise performance. Parameters like data size and clusters are configurable and the program ensures efficient memory management and synchronisation.
- **XSbench:** XSbench [21] is a simplified version of the OpenMC Monte Carlo code that focuses only on the most time-consuming part: calculating macroscopic neutron cross sections. Instead of running a full particle transport simulation it uses a benchmark reactor model, a large unionised energy grid to reproduce the same types of calculations and memory use patterns but with far less complexity. The key point is that cross-section lookups involve many random accesses to large data tables in memory, rather than a small amount of repeated arithmetic. This makes performance heavily dependent on

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. 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/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 (This refers to every miss in the Last level cache that occurs during a memory read operation.)

how quickly data can be read from and written to different levels of the memory hierarchy (caches, main memory). XSBench is therefore closely tied to studying memory reads and writes, since it allows researchers to see how modern processors handle the irregular, data-intensive behaviour that dominates in real Monte Carlo simulations.

- 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 [22]. This implementation 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 [23].

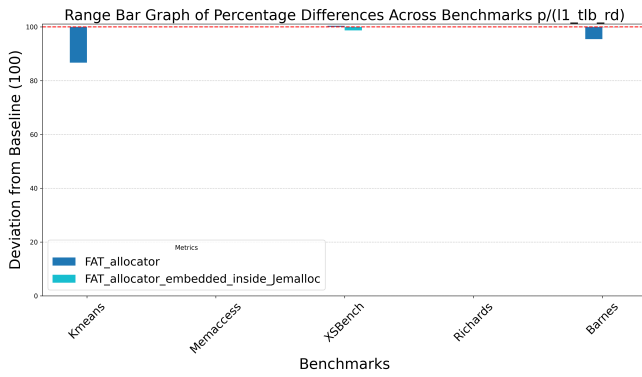
7.3 Results and discussion

Figure 2 highlights the performance comparison between the modified memory allocator and Jemalloc, the default memory allocator. The FAT memory allocator is specifically optimised for use with huge pages and 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.

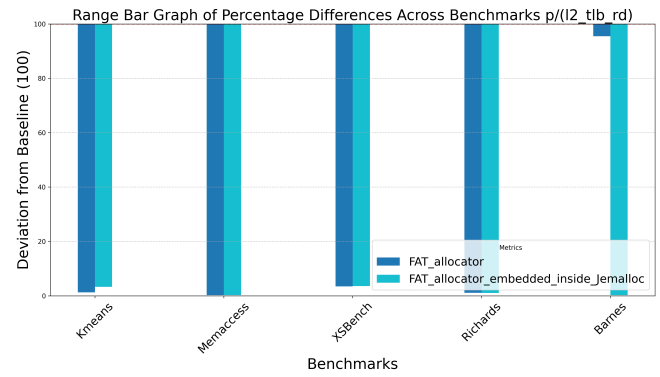
- L1 DTLB reads (Figure 2a): L1 TLB reads are critical for achieving fast memory access; therefore, a reduction in events that could signify misses or lead to further lookups is generally beneficial. In the Kmeans benchmark, the FAT allocator demonstrated 13% fewer L1 TLB reads than the baseline allocator. However, when the FAT allocator was embedded within Jemalloc the L1 TLB reads were the same as the baseline. For memaccess, XSBench and Richards benchmark the L1 TLB were the same as the baseline for both allocators. For the Barnes benchmark, the FAT allocator exhibited 5% fewer L1 TLB reads, and the result was the same for the FAT allocator embedded within Jemalloc.
- L2 DTLB reads (Figure 2b): L2 Data TLB reads act as a secondary cache for address translations. The FAT allocator consistently reduced L2 TLB reads by around 98% across all benchmarks, with the exception of Barnes, where the reduction was 4%. When embedded within Jemalloc, the FAT allocator also achieved substantial improvements, lowering L2 TLB reads by more than 86% for Kmeans, XSBench, Memaccess, and Richards. These results indicate that the majority of address translation requests are successfully resolved at the L1 TLB level, thereby avoiding the need to access deeper levels of the memory hierarchy.
- DTLB walks (Figure 2c): which occur when a virtual-to-physical address translation is not found in the DTLB and a page table traversal is necessary, represent a performance cost. thus, fewer walks are preferable. In the observed tests neither FAT allocator nor FAT allocator embedded inside Jemalloc demonstrated significant deviation from the baseline performance of 99% lesser walks. This consistent behavior was noted across all benchmarks evaluated: Kmeans, Memaccess, XSBench, Richards, and Barnes. This indicates that almost all address translations were resolved directly at

the L1 DTLB level without triggering expensive traversals through the page table. More generally, this finding demonstrates that both allocator designs make efficient use of the hardware translation system and maintain consistently low page walk overheads across a range of workloads.

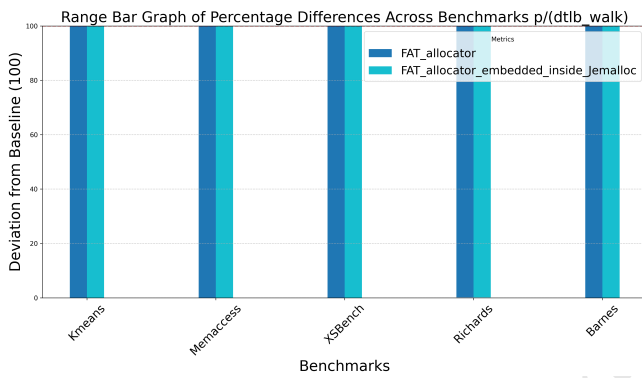
- L1 DTLB refills (Figure 2d): L1 Data DTLB refills are a direct consequence of L1 DTLB misses; therefore, fewer refills indicate better performance. Interestingly, despite the variations observed in L1 DTLB reads for the FAT allocator embedded within Jemalloc, the L1 DTLB refills metric showed a significant reduction of 99% for both the standalone FAT allocator and the FAT allocator embedded within Jemalloc. This consistent and substantial improvement over the baseline was observed across all tested benchmarks: Kmeans, Memaccess, XSBench, Richards, and Barnes.
- Last-level cache (Figure 2e): Cache read misses in the Last-Level Cache (LLC) are important performance indicators, as each miss typically results in slower data retrieval from main memory. Thus, fewer misses are generally preferable. The results for this metric show notable variation across both allocators and benchmarks. With the FAT allocator, kmeans achieved a 57% reduction in LLC misses, while memaccess showed no meaningful change. In contrast, XSBench experienced 6% more misses, Richards 32% more, and Barnes 19% more compared to the baseline. When the FAT allocator was embedded within Jemalloc, the behaviour shifted: kmeans recorded 19% more misses, memaccess again showed negligible change, XSBench improved slightly with 3% fewer misses, Richards had a marginal 1% increase, and Barnes exhibited a substantial 69% increase.
- Wall clock (Figure 2f): Wallclock time serves as the definitive metric for evaluating overall execution performance. When using the FAT allocator, XSBench demonstrated the most significant improvement, with a 11% reduction in wallclock runtime compared to the baseline allocator. This was followed by memaccess, barnes with a 4% decrease and kmeans with a modest 1.8% improvement. Richards exhibited a slight increase of 1% in runtime. When the FAT allocator was integrated into Jemalloc, the performance impact varied. Memaccess and XSBench demonstrated similar improvements compared to the FAT allocator. in wallclock time, while barnes showed a 7% rise. Both Richards and kmeans maintained the same runtime as the baseline allocator. In regards to memaccess (Figure 3) the comparative performance of two allocators FAT allocator embedded inside jemalloc (depicted by the solid blue line with circular markers) and FAT allocator (represented by the dashed orange line with square markers) across progressively increasing memory access iterations ranging from 500 to 64,000. The results indicate that the standalone FAT allocator exhibits superior performance at smaller memory sizes, maintaining an initial lead of approximately 0.78% at size 500. However, its trajectory is characterised by minor fluctuations as memory size increases. Conversely, the jemalloc embedded allocator demonstrates a more consistent and steadily rising performance pattern. The annotated red values denote the performance differential between the two allocators. Although the FAT allocator retains



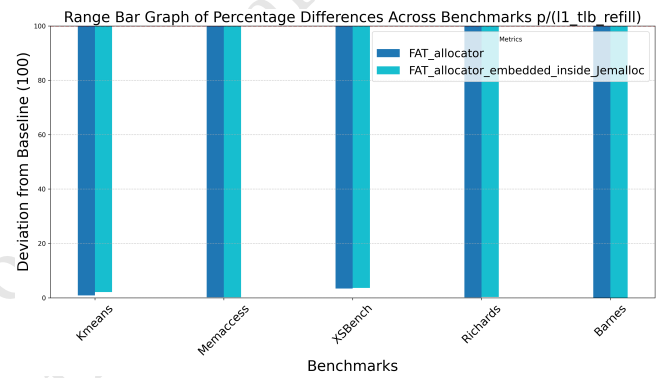
(a) L1 TLB Reads



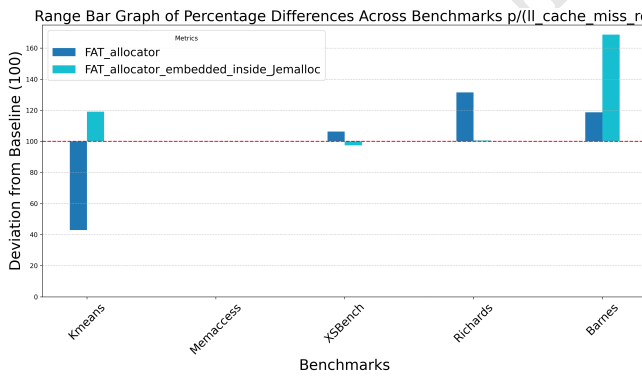
(b) L2 TLB Reads



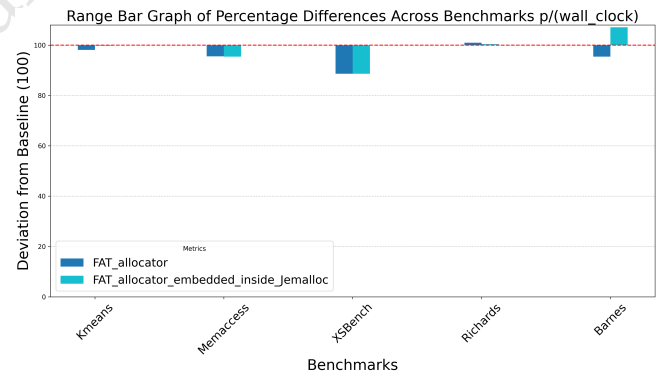
(c) DTLB Walks



(d) L1 DTLB Refill



(e) LL Cache Reads



(f) Wall Clock Time

Figure 2: Benchmarks comparing the percentage difference between FAT.

an advantage in the lower ranges, this advantage progressively diminishes with increasing memory reads. Notably, at the largest tested size of 64,000, the difference marginally reverses (-0.07), with the jemalloc-embedded allocator slightly outperforming. Overall, the findings suggest that while the FAT allocator is more effective at smaller scales, the jemalloc

embedded FAT allocator exhibits superior scalability and reliability across larger memory reads.

Overall, the evaluation shows that the FAT allocator provides consistent improvements in TLB efficiency, with near elimination of L1 refills and L2 reads across all benchmarks, thereby reducing pressure on the memory translation subsystem. Its effect on last-level

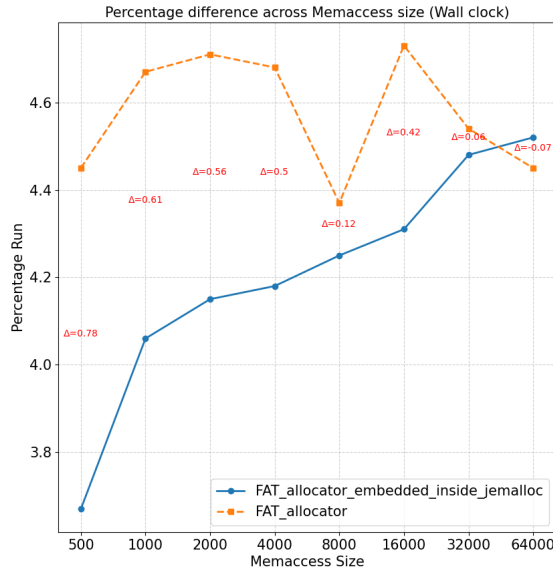


Figure 3: Memaccess percentage increase in wallclock run times

cache behaviour is more variable, with notable gains in Kmeans but regressions in other workloads such as Richards and Barnes. Wall-clock performance reflects these mixed results: while XSBench and Memaccess benefit from measurable reductions in execution time, benchmarks with irregular access patterns (e.g., Richards) show little to no improvement. A closer comparison of Memaccess further reveals that the standalone FAT allocator performs best at smaller memory scales, whereas the Jemalloc-embedded FAT allocator demonstrates superior scalability and stability as memory sizes increase. These findings suggest that FAT’s huge-page-oriented design is particularly effective for workloads with predictable or structured access patterns, whilst integration with Jemalloc enhances adaptability for larger or more irregular memory usage.

8 Conclusion

This paper has presented FAT, a memory allocator designed to address the growing mismatch between application memory demands and the limited reach of TLBs in modern processors. By leveraging physically contiguous memory through huge pages and embedding allocation metadata within CHERI’s compressed capability-based pointers, FAT significantly reduces the overhead associated with virtual-to-physical address translation.

FAT achieves block-based allocation within huge pages, enabling memory tracking without relying heavily on traditional page table mechanisms. Benchmark evaluations demonstrate that FAT can reduce TLB walks by up to 99%, resulting in substantial improvements in memory-intensive workloads. When applied to micro-benchmarks, FAT reduced wall-clock runtime by up to 4% for

Memaccess and maintained near-baseline performance for Richards, while also achieving up to 98% fewer L2 TLB reads. In contrast, macro-benchmarks such as XSBench, Kmeans, and Barnes showed runtime improvements of 11%, 1.8%, and 4%, respectively, reflecting the allocator’s stronger benefits in memory-bound scenarios compared to compute-heavy ones.

Overall, the evaluation shows that FAT provides consistent improvements in TLB efficiency, with near elimination of L1 refills and L2 reads across all benchmarks, thereby reducing pressure on the memory translation subsystem. Its effect on last-level cache behaviour is more variable, with notable gains in Kmeans but regressions in workloads such as Richards and Barnes. Wall-clock performance reflects these mixed outcomes: while XSBench and Memaccess benefit from measurable reductions in execution time, workloads with irregular access patterns (e.g., Richards) show little to no improvement. A closer comparison of Memaccess further reveals that the standalone FAT allocator performs best at smaller memory scales, whereas the Jemalloc-embedded FAT allocator demonstrates superior scalability and stability as memory sizes increase.

Although performance gains are less significant for larger or computation-heavy applications, the results underscore the allocator’s potential to enhance memory management in high-throughput environments. More broadly, this work demonstrates how capability-based architectures, originally designed to ensure memory safety, can be effectively repurposed to optimise address translation. FAT thus represents a promising direction for developing more scalable and efficient memory allocation strategies in systems adopting capability-aware architectures.

References

- [1] Daniel Lustig, Abhishek Bhattacharjee, and Margaret Martonosi. Tlb improvements for chip multiprocessors: Inter-core cooperative prefetchers and shared last-level tlbs. *ACM Trans. Archit. Code Optim.*, 10(1), April 2013. ISSN 1544-3566. doi: 10.1145/2445572.2445574. URL <https://doi.org/10.1145/2445572.2445574>.
- [2] Sparsh Mittal. A survey of techniques for architecting TLBs. 29(10):e4061. ISSN 1532-0634. doi: 10.1002/cpe.4061. URL <https://onlinelibrary.wiley.com/doi/abs/10.1002/cpe.4061>. _eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1002/cpe.4061>.
- [3] 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, ISBN 978-1-4503-6240-5. doi: 10.1145/3297858.3304064. URL <https://dl.acm.org/doi/10.1145/3297858.3304064>.
- [4] 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. ISSN 0163-5964. doi: 10.1145/2678373.2665740. URL <https://doi.org/10.1145/2678373.2665740>.
- [5] 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. ISSN 0018-9340, 1557-9956, 2326-3814. doi: 10.1109/TC.2019.2914037. URL <https://ieeexplore.ieee.org/document/8703061/>.
- [6] Binh Pham, Abhishek Bhattacharjee, Yasuko Eckert, and Gabriel H. Loh. Increasing tlb reach by exploiting clustering in page translations. In *2014 IEEE 20th International Symposium on High Performance Computer Architecture (HPCA)*, pages 558–567, 2014. doi: 10.1109/HPCA.2014.6835964.
- [7] Juan Navarro, Sitararn Iyer, Peter Druschel, and Alan Cox. Practical, transparent operating system support for superpages. *SIGOPS Oper. Syst. Rev.*, 36(SI):89–104, December 2003. ISSN 0163-5980. doi: 10.1145/844128.844138. URL <https://doi.org/10.1145/844128.844138>.
- [8] Marius Cornea, John Harrison, and Ping Tak Peter Tang. Intel® itanium® floating-point architecture. In *Proceedings of the 2003 Workshop on Computer Architecture Education: Held in Conjunction with the 30th International Symposium on Computer Architecture*, WCAE '03, page 3–es, New York, NY, USA, 2003. Association for

- Computing Machinery. ISBN 9781450347327. doi: 10.1145/1275521.1275526. URL <https://doi.org/10.1145/1275521.1275526>.
- [9] 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. ISSN 0141-9331. doi: [https://doi.org/10.1016/S0141-9331\(01\)00125-9](https://doi.org/10.1016/S0141-9331(01)00125-9). URL <https://www.sciencedirect.com/science/article/pii/S0141933101001259>.
- [10] 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. ISSN 0163-5964. doi: 10.1145/2508148.2485943. URL <https://doi.org/10.1145/2508148.2485943>.
- [11] Vasileios Karakostas, Jayneel Gandhi, Furkan Ayar, Adrián Cristal, Mark D. Hill, Kathryn S. McKinley, Mario Nemirovsky, Michael M. Swift, and Osman Ünsal. Redundant memory mappings for fast access to large memories. In *Proceedings of the 42nd Annual International Symposium on Computer Architecture*, pages 66–78. ACM, ISBN 978-1-4503-3402-0. doi: 10.1145/2749469.2749471. URL <https://dl.acm.org/doi/10.1145/2749469.2749471>.
- [12] Dongwei Chen, Dong Tong, Chun Yang, Jiangfang Yi, and Xu Cheng. Flexpointer: Fast address translation based on range tlb and tagged pointers. *ACM Trans. Archit. Code Optim.*, 20(2), March 2023. ISSN 1544-3566. doi: 10.1145/3579854. URL <https://doi.org/10.1145/3579854>.
- [13] Brooks Davis, Robert N. M. Watson, Alexander Richardson, Peter G. Neumann, Simon W. Moore, John Baldwin, David Chisnall, Jessica Clarke, Nathaniel Wesley Filardo, Khilan Gudka, Alexandre Joannou, Ben Laurie, A. Theodore Markettos, J. Edward Maste, Alfredo Mazzinghi, Edward Tomasz Napierala, Robert M. Norton, Michael Roe, Peter Sewell, Stacey Son, and Jonathan Woodruff. Cheriabi: Enforcing valid pointer provenance and minimizing pointer privilege in the posix c run-time environment. In *Proceedings of the Twenty-Fourth International Conference on Architectural Support for Programming Languages and Operating Systems*, ASPLOS '19, page 379–393, New York, NY, USA, 2019. Association for Computing Machinery. ISBN 9781450362405. doi: 10.1145/3297858.3304042. URL <https://doi.org/10.1145/3297858.3304042>.
- [14] Jason Evans. A scalable concurrent malloc (3) implementation for freebsd. In *Proc. of the BSDcan conference, ottawa, canada*, 2006.
- [15] Jason Evans. A Scalable Concurrent malloc(3) Implementation for FreeBSD.
- [16] Benchmark ABI - CheriBSD 23.11 new features tutorial. URL <https://www.cheribsd.org/tutorial/23.11/benchmark/index.html>.
- [17] ChERI-allocator/benchmarks/benchmarks/StressTestMalloc/glibc-bench.c at main · akilan1999/ChERI-allocator. URL <https://github.com/Akila1999/ChERI-Allocator/blob/main/benchmarks/benchmarks/StressTestMalloc/glibc-bench.c>.
- [18] Department of computer science and technology – ChERI: The arm morello board. URL <https://www.cl.cam.ac.uk/research/security/ctsr/cheri/cheri-morello.html>.
- [19] 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.
- [20] Arm architecture reference manual for a-profile architecture. URL <https://developer.arm.com/documentation/ddi0487/latest>.
- [21] John R Tramm, Andrew R Siegel, Tanzima Islam, and Martin Schulz. XSBench - the development and verification of a performance abstraction for Monte Carlo reactor analysis. In *PHYSOR 2014 - The Role of Reactor Physics toward a Sustainable Future*, Kyoto, 2014. URL <https://www.mcs.anl.gov/papers/P5064-0114.pdf>.
- [22] Jaswinder Pal Singh. *Parallel Hierarchical N-body Methods and Their Implications for Multiprocessors*. PhD thesis, Stanford University, February 1993.
- [23] Jaswinder Pal Singh, John L. Hennessy, and Anoop Gupta. Implications of hierarchical n-body methods for multiprocessor architectures. *ACM Trans. Comput. Syst.*, 13(2):141–202, May 1995. ISSN 0734-2071. doi: 10.1145/201045.201050. URL <https://doi.org/10.1145/201045.201050>.