

H-MBR: Hypervisor-Level Memory Bandwidth Reservation for Mixed Criticality Systems

Afonso Oliveira ✉ 

Centro ALGORITMI / LASI, Universidade do Minho, Portugal

Diogo Costa ✉ 

Centro ALGORITMI / LASI, Universidade do Minho, Portugal

Gonalo Moreira ✉ 

Centro ALGORITMI / LASI, Universidade do Minho, Portugal

Jos  Martins ✉ 

Centro ALGORITMI / LASI, Universidade do Minho, Portugal

Sandro Pinto ✉ 

Centro ALGORITMI / LASI, Universidade do Minho, Portugal

Abstract

Recent advancements in fields such as automotive and aerospace have driven a growing demand for robust computational resources. Applications that were once designed for basic Microcontroller Units (MCUs) are now deployed on highly heterogeneous System-on-Chip (SoC) platforms. While these platforms deliver the necessary computational performance, they also present challenges related to resource sharing and predictability. These challenges are particularly pronounced when consolidating safety-critical and non-safety-critical systems, the so-called Mixed-Criticality Systems (MCS) to adhere to strict Size, Weight, Power, and Cost (SWaP-C) requirements. MCS consolidation on shared platforms requires stringent spatial and temporal isolation to comply with functional safety standards (e.g., ISO 26262). Virtualization, mainly leveraged by hypervisors, is a key technology that ensures spatial isolation across multiple OSes and applications; however ensuring temporal isolation remains challenging due to contention on shared resources, such as main memory, caches, and system buses, which impacts real-time performance and predictability. To mitigate this problem, several strategies (e.g., cache coloring and memory bandwidth reservation) have been proposed. Although cache coloring is typically implemented on state-of-the-art hypervisors, memory bandwidth reservation approaches are commonly implemented at the Linux kernel level or rely on dedicated hardware and typically do not consider the concept of Virtual Machines that can run different OSes. To fill the gap between current memory bandwidth reservation solutions and the deployment of MCSs that operate on a hypervisor, this work introduces *H-MBR*, an open-source VM-centric memory bandwidth reservation mechanism. *H-MBR* features (i) VM-centric bandwidth reservation, (ii) OS and platform agnosticism, and (iii) reduced overhead. Empirical results evidenced no overhead on non-regulated workloads, and negligible overhead (<1%) for regulated workloads for regulation periods of 2 μ s or higher.

2012 ACM Subject Classification Computer systems organization → Real-time system specification; Computer systems organization → Embedded software

Keywords and phrases Virtualization, Multi-core Interference, Mixed-Criticality Systems, Arm, Memory Bandwidth Reservation

Digital Object Identifier 10.4230/OASICS.NG-RES.2025.4

Supplementary Material *Software (Source Code)*: <https://gitlab.com/ESRGv3/h-mbr>

Funding *Diogo Costa*: Supported by FCT grant 2022.13378.BD.

Jos  Martins: Supported by FCT grant SFRH/BD/138660/2018.



  Afonso Oliveira, Diogo Costa, Gonalo Moreira, Jos  Martins, and Sandro Pinto;
licensed under Creative Commons License CC-BY 4.0

Sixth Workshop on Next Generation Real-Time Embedded Systems (NG-RES 2025).

Editors: Patrick Meumeu Yomsi and Stefan Wildermann; Article No. 4; pp. 4:1–4:15

OpenAccess Series in Informatics



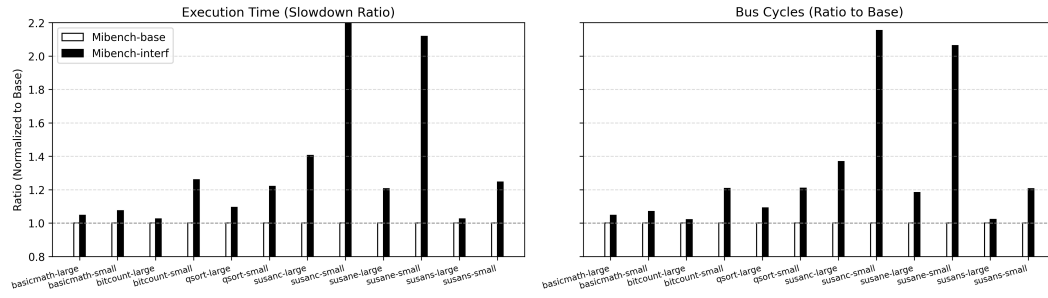
OASICS Schloss Dagstuhl – Leibniz-Zentrum f r Informatik, Dagstuhl Publishing, Germany

1 Introduction

Rapid advancements in industries such as automotive, aerospace, and industrial automation have led to increasingly demanding applications, driving a significant need for higher computational power [8, 5]. These applications now require systems capable of handling a diverse range of computationally intensive tasks, such as autonomous driving, flight control, and complex industrial automation workflows [16]. To address this computational needs, the once simple microcontroller units (MCUs) evolved to sophisticated heterogeneous architectures. Modern designs integrate multiple Central Processing Units (CPUs) alongside specialized accelerators, such as Graphics Processing Units (GPUs)[5], Field-Programmable Gate Arrays (FPGAs)[16], and novel AI accelerators like Tensor Processing Units (TPUs) and Neural Processing Units (NPU)s[21, 14]. This shift enables efficient processing of complex workloads but simultaneously introduces challenges related to resource sharing and system predictability, especially as systems grow in complexity and criticality[30]. In addition to computational demands, embedded systems increasingly face Size, Weight, Power, and Cost (SWAP-C) constraints, which are driving a trend towards system consolidation[8, 13], giving rise to the so called Mixed-Criticality Systems (MCS). However, integrating MCS on shared platforms faces a significant challenge: stringent requirements for both spatial and temporal isolation [40, 16, 1, 22, 12]. These systems must adhere to certification standards (e.g., ISO 26262 in automotive, CENELEC for railway systems, and ECSS for space applications) which place strict demands on system safety, reliability, and predictability [25, 6, 11, 34, 20].

Virtualization has become essential in the consolidation of MCS, with hypervisors playing a pivotal role. Hypervisors, particularly static partitioning hypervisors, must be minimal yet sufficiently robust to ensure reliable isolation while meeting the real-time (RT) performance requirements of modern applications [23, 11, 29, 26]. Despite advancements in spatial isolation through various partitioning techniques, temporal isolation remains an active area of research due to contention on shared resources [39, 40, 25]. These shared resources, including (i) main memory, (ii) caches, and (iii) the system bus, can introduce substantial unpredictability in response times if left unregulated [6, 30, 12]. Such variability in access latency poses significant risks to RT, safety-critical applications, where predictable performance is essential to meet strict timing constraints [17, 28, 20]. To address shared resource contention, considerable efforts have been made in both academia and industry. Techniques such as (i) cache coloring and (ii) MBR have emerged as widely recognized approaches to regulate and reduce contention induced by memory access [40, 39, 22]. However, unlike cache coloring, most MBR approaches are implemented at the OS-level, typically on the Linux kernel [12, 25], narrowing their flexibility and applicability, underscoring the need for a more versatile solution, one that can provide robust temporal isolation across diverse applications and platforms without sacrificing configurability or performance [7, 19, 6].

In this paper we present H-MBR, an open-source VM-centric MBR mechanism implemented in Bao hypervisor. Its primary contributions over traditional MBR methods stem from an innovative design: (i) VM-Centric bandwidth allocation, (ii) OS and platform agnosticism, and (iii) reduced overhead. Furthermore, H-MBR supports unbalanced distribution of bandwidth across a VM's vCPUs. As H-MBR's is implemented at the hypervisor level it enables the memory bandwidth regulation of VMs with different operating systems, (e.g., Linux, FreeRTOS and Zephyr). Additionally, H-MBR is designed to be platform-agnostic, allowing it to be ported across various architectures such as ARMv8-A, ARMv8-R, and RISC-V . Finally, as Bao hypervisor stands as the hypervisor with lowest interrupt latency among the state-of-the-art static partitioning hypervisors [24], this mechanism takes advantage of



■ **Figure 1** Mibench Execution Time and Bus Cycles (Ratio).

the fast interrupt handling, which minimizes the overhead of MBR. Empirical results show that *H-MBR* introduces no overhead for critical workloads, maintaining their performance and ensuring predictable behavior. For non-critical workloads, the mechanism introduces an overhead below 1% for regulation periods equal or higher than 2 μ s. It also fully eliminates critical workload performance degradation if strict budgets are enforced on non-critical workloads, demonstrating its reliability in MCSs.

2 Interference and Interference Mitigation

In multicore embedded systems, contention generated in shared resources is a significant challenge that impacts performance, predictability, and temporal isolation, particularly in RT and MCS [7, 40, 6]. Contention arises when multiple cores simultaneously access shared components [23, 34], (i) such as main memory [42, 33, 41, 20, 40], (ii) last-level caches [23, 22], (iii) system buses [33], and (iv) additional subsystems like the Generic Interrupt Controller [9, 10]. These shared resources can lead to delays as cores compete for access, resulting in increased response times and unpredictable behavior. This unpredictability is especially problematic in systems with strict timing requirements, where such delays affect the system’s temporal isolation, threatening reliability and safety [12, 22, 11]. As shown in Figure 1, this interference can arise up to 2.3x on memory-intensive benchmarks like in Mibench’s *susanc-small*. As system complexity grows, understanding and controlling memory utilization becomes essential to mitigate interference effects and ensure critical tasks’ predictability. To mitigate memory-induced contention and ensure temporal isolation of tasks, several techniques have been proposed.

Cache Partitioning. Cache partitioning divides LLC into distinct regions, or “partitions”, which are assigned to specific workloads to control cache access and reduce contention. Two primary approaches to cache partitioning are (i) cache locking and (ii) cache coloring. Cache locking relies on hardware assistance to restrict eviction from designated cache lines, securing specific portions of the cache for high-priority tasks. Cache coloring, on the other hand, uses the overlap between virtual page numbers and cache indices to partition cache sets without needing specialized hardware. Hybrid techniques, such as Colored Lockdown [22], combine coloring and locking, while other approaches propose dynamic re-coloring schemes to adapt to varying workloads [39, 38, 31]. Cache coloring has been successfully implemented in hypervisors like Bao [23], Jailhouse [20], XVisor [25], Xen [32] and KVM [35], effectively enhancing predictability in multicore systems with real-time requirements. However, the efficiency of cache coloring is workload-dependent: for high-demand cores, it significantly reduces cache interference and improves performance, but with lighter or variable loads, it

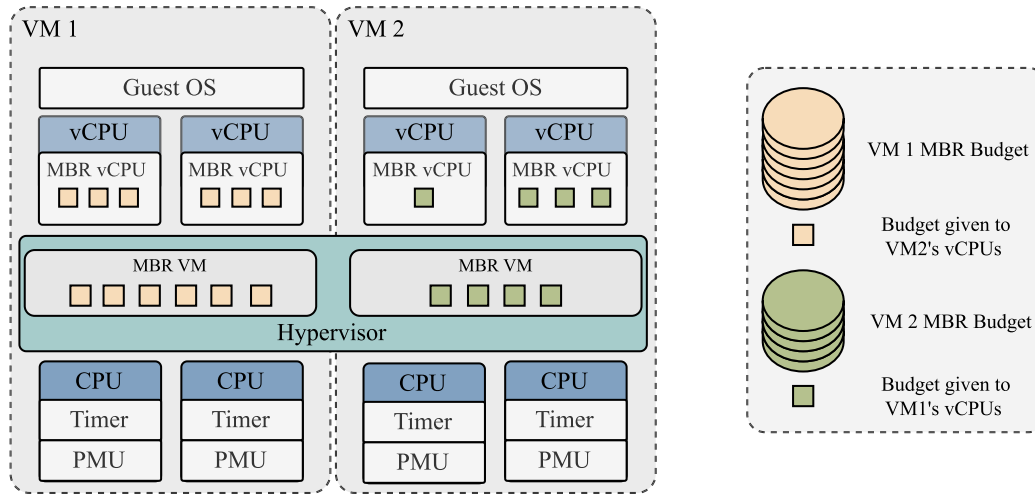
■ **Table 1** Gap analysis.

Paper	Dynamic	Static	Implementation Level	CPU-centric	VM-centric
A. Zuepke et al. [45]	●	○	Firmware	●	○
H.Yun et al. [41]	●	○	OS(Linux Kernel)	●	○
H.Yun et al. [43]	●	○	OS(Linux Kernel)	●	○
H.Yun et al. [40]	●	○	OS(Linux Kernel)	●	○
P. Sohal et al. [36]	●	○	Hardware	●	○
A. Agrawal et al. [2]	●	○	OS(Linux Kernel)	●	○
E. Seals [33]	●	○	OS(Linux Kernel)	●	○
D. Hoornaert et al. [18]	●	○	Hypervisor and Hardware	●	○
M. Pagani et al. [27]	●	○	Hardware	●	○
M. Xu et al. [37]	●	○	Hypervisor	●	○
P. Modica et al. [25]	○	●	Hypervisor	●	○
E. Gomes et al. [15]	●	○	Hypervisor	●	○
G. Brilli et al. [4]	●	○	Hypervisor	●	○
H-MBR	○	●	Hypervisor	○	●

may underutilize cache space, potentially impacting overall system efficiency. Additionally, cache coloring is hardware-specific, requiring an understanding of cache structure details such as size, associativity, and mapping policies, which can limit its portability across platforms.

Memory Bandwidth Reservation (MBR). MBR is a critical approach to reserving memory bandwidth, designed to prevent hardware contention on memory accesses, which is vital for achieving predictable performance in MCS [12, 40]. When memory is heavily utilized, increased memory contention leads to bus delays, subsequently extending execution times. MBR works by reserving memory bandwidth quotas per core or VM, either statically or dynamically, effectively reducing contention on shared resources and improving overall performance [7, 25]. Most MBR implementations operate at the OS level [41, 22, 40, 6, 34, 12, 44, 3] particularly within the Linux Kernel, where they manage bandwidth based on task requirements to ensure that critical tasks receive adequate access without overloading the memory controller [6, 34]. However, this OS-level integration limits MBR’s portability to other OSs, such as RTOSs like Zephyr and FreeRTOS, which are especially prevalent in real-time applications [17, 29]. On the other hand, some MBR implementations use FPGA-dedicated accelerators [44] to monitor and control memory bandwidth, which has the significant downside of being platform-specific. Additionally, some approaches implement MBR at the hypervisor level [25, 37, 15, 4]; however, existing methods are applied at the vCPU level and are integrated into hypervisors with slower interrupt handling compared to Bao [24]. This slower response introduces significant overhead, particularly for time-sensitive applications, where Bao’s efficient, low-latency interrupt handling offers a clear advantage. Additionally, some of these approaches are tied to specific architecture, reducing its portability. Table 1 compares various MBR techniques, highlighting that most Linux OS-based implementations are dynamically managed yet tied to the Linux kernel, restricting their flexibility and scalability across diverse embedded environments.

Gap Analysis. MemGuard [41] presents a dynamic allocation in Linux kernel while H.Yun et al. [43] presents an extension to previous implementation. Pallocc[40] implements a bank-aware memory allocation strategy. A. Zuepke et al. [45] adopt a distinct methodology



■ **Figure 2** System Overview.

by utilizing an MCU to control a larger APU. A. Agrawal et al. [2], and E. Seals [33] all implement dynamic allocation within the Linux kernel. Hardware-level implementations come from P. Sohal et al.[36] and M. Pagani et al.[27], both using dynamic allocation. At the hypervisor level, E. Gomes et al. [15] and G. Brilli et al. [4] leverage MemGuard on the hypervisor level, focusing on CPU-centric approaches. Moreover, M. Xu et al.[37] implement dynamic allocation, while P. Modica et al.[25] opt for static allocation. D. Hoornaert et al. [18] uniquely combine hypervisor and hardware approaches with dynamic allocation. While all these solutions focus on CPU-centric approaches, *H-MBR* introduces a VM-centric approach that aligns with MCSs' perspective by enabling bandwidth allocation based on VMs than individual cores, implementing static allocation at the hypervisor level.

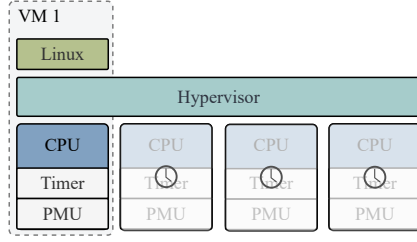
3 The mechanism

H-MBR Interface

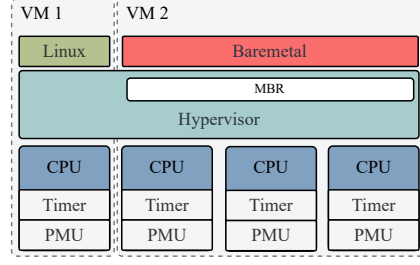
H-MBR focuses on a VM-centric perspective, where each VM requires two key MBR parameters: (i) *budget* and (ii) *period* configured at compile time. The MBR configuration operates on a per-VM basis, enabling customized budgets and regulation periods for individual VMs. Additionally, the VM budget distribution across CPUs can be handled in two distinct ways: through automatic balancing or via non-balanced distribution, where specific allocations can be manually set. This flexibility allows for uneven distribution of the VM budget across its vCPUs, as demonstrated by the green-colored VM in Figure 2.

H-MBR Run-Time

The MBR mechanism follows a Memguard [41] based approach on how to track this metrics in real-time, thereby the mechanism relies on two key peripherals: (i) PMU and (ii) generic timer, respectively. Both of this peripherals exist physically inside each CPU, as depicted in figures 3 and 4. The hypervisor-level approach leverages this physical peripherals common to most architectures and assigns them directly to each vCPU, which leads to CPU monitoring not impacting other CPUs performance and therefore ensuring spatial and temporal isolation.



■ **Figure 3** Standalone setup with single Linux VM.



■ **Figure 4** Co-existing setup with two VMs: (i) a Linux VM and (ii) a baremetal VM.

Timer. The timer is configured to overflow after a defined *Period*, setting the interval at which the memory bandwidth reserved for each core is reset. This periodic reset ensures that all cores start each *period* with a reset *budget*. Additionally, the timer re-activates any cores that were previously idled due to over-budget usage, allowing them to resume operation at the beginning of each new period. The overhead introduced by the timer interrupt can be quantified as follows: $\text{Overhead}_{\text{timer}} = \frac{D_{\text{timer}}}{\text{Period}_{\text{timer}}}$, where D_{timer} corresponds to the total execution time of the timer interrupt (i.e., the sum of interrupt injection latency with the interrupt callback execution time).

PMU. The PMU tracks memory access activities by counting *bus access* and triggers an overflow interrupt if a core exceeds its reserved *Budget* within the *Period*. When this overflow occurs, the PMU signals the MBR mechanism to temporarily idle the over-budget core. This action prevents excessive memory contention, helping to maintain consistent response times for critical tasks on other cores.

Furthermore, the MBR configuration which operates on a per-VM basis, enables customized budgets and regulation periods for individual VMs. The VM budget distribution across CPUs can be handled in two distinct ways: through automatic balancing or via non-balanced distribution, where specific allocations can be manually set.

4 Evaluation

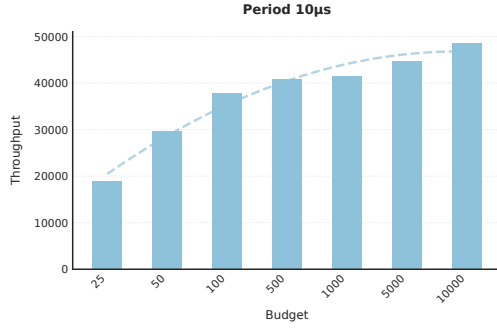
This section provides a detailed explanation of the methodology used to obtain the results, including descriptions of the guest OSs, configurations, and setups on the target hardware. Each component was carefully chosen to analyze memory interference under varying conditions.

■ Algorithm 1 Initial Assignment of MBR Parameters.

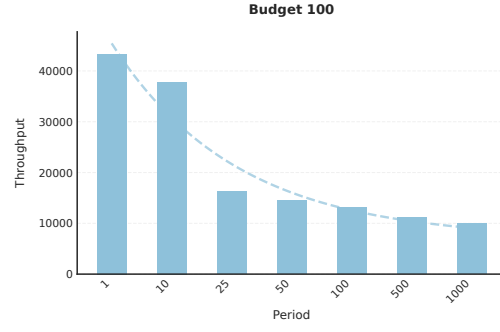
```

1: for each VM in VMs do
2:   VM.budget ← budget_vm
3:   VM.period ← period_vm
4:   for each vCPU in VM.vCPUs do
5:     vCPU.budget ← VM.has_custom_dist
6:       ? (VM.budget × vCPU_percentage[vCPU])
7:       : (VM.budget / VM.num_vCPUs)
8:   end for
9: end for

```



■ **Figure 5** *interf+mbr* NC throughput budget variation.



■ **Figure 6** *interf+mbr* NC throughput period variation.

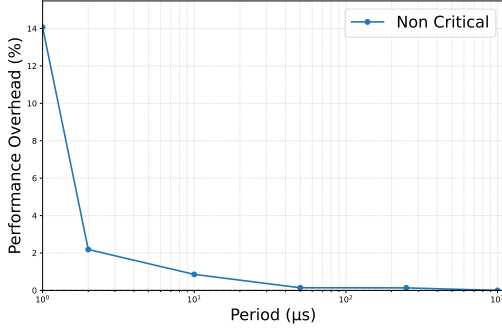
Target Platform & Measurement Tools. We conducted our experiments assessment on a Zynq UltraScale+ ZCU104 board (ARMv8-A-based architecture), which features a quad-core Arm Cortex-A53 processor. Each CPU has a private cache (data and instruction, 32KiB each), and the unified shared L2 cache (1MiB). The Cortex-A53 processor also features a ARM PMUv3, which has been leveraged to profile benchmark execution and gather key microarchitectural events. The selected events include bus cycles and execution cycles, providing insight into memory and system bus usage. Additionally, we used the `perf` tool on the Linux OS as an interface to the CPU’s PMU.

VM Workloads. We deployed two distinct guest environments for benchmarking: (i) a Linux-based VM, which will be denominated as Critical VM (C) and (ii) a baremetal VM, which will be denoted as Non-Critical VM (NC). For the critical VM, we deployed a Linux-based guest to enable the deployment of MiBench automotive suite [17], a widely-used benchmark for automotive embedded systems – to simulate real-world automotive workloads. For the NC VM, we deployed a baremetal guest that continuously writes to a buffer sized to match the LLC capacity, creating intentional memory contention to thoroughly assess interference effects on the memory hierarchy. We focused exclusively on write operations, rather than reads or a combination of both, since write operations stresses even more the memory shared hardware resources, when compared to read operations.

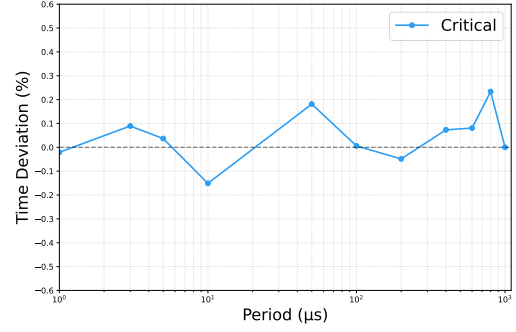
Setups and Configurations. We evaluated four distinct setups: (i) *solo*, (ii) *interf*, and (iii) *interf+mbr*. The *solo* configuration involves running Linux alone, whereas *interf* includes both Linux and Baremetal guests to introduce memory interference. The suffix *mbr* is used to identify setups with memory bandwidth reservation. As for hardware resources, we followed the CPU assignment identified in Figure 4: one CPU is allocated to Linux in all configurations, and three CPUs are assigned to the baremetal VM, when included.

4.1 Impact of MBR on Baremetal Throughput and Overhead

To assess the effect of MBR configurations on the non-critical baremetal VM, we measured the number of cache line writes achievable under different budget and period settings. This section will focus on the memory-intensive *susan-c* benchmark from the MiBench suite (the results of other benchmarks can be found in Appendix A) on the critical Linux-based VM. The *susan-c* benchmark was specifically chosen for its high memory contention characteristics, making it an ideal workload for evaluating the impact of MBR adjustments on memory access behavior in the non-critical VM.



■ **Figure 7** Performance overhead on NC.

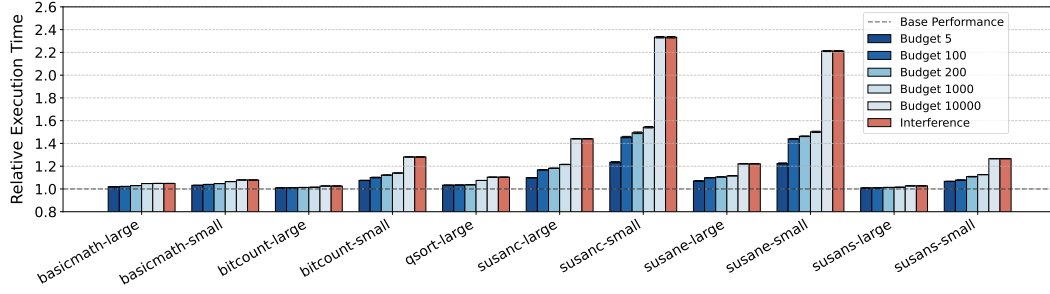


■ **Figure 8** Time deviation on C.

Budget Configuration. Increasing the MBR period while keeping the budget fixed at 100 reduces the cache line write capacity of the non-critical, as shown in the Figure 5. For instance, with a period of 25 μ s, cache line writes decrease by 2.6 times compared to a period of 1 μ s. This outcome reflects how longer periods effectively limit the NC's access to memory bandwidth, as the delay in budget reset constrains the number of operations the non-critical can perform within each interval. These findings align with our expectation that increasing the period value reduces memory bandwidth availability for the non-critical, thus lowering its potential to interfere with the critical VM.

Period Configuration. Figure 6 demonstrates that with a fixed period of 10 μ s, increasing the budget significantly increases the non-critical's cache line writes. For instance, a budget setting of 500 allows 2.2 times more cache line writes compared to a budget of 25. Higher budgets allocate more memory bandwidth to the non-critical VM, which in turn elevates its cumulative memory access capacity over time, maximizing the interference potential with the critical VM.

Overhead. Additionally, our empirical results show that *H-MBR* provides VM-level isolation with minimal overhead. To accurately assess the MBR mechanism's overhead, we deployed the baremetal VM in a solo setup, without the critical Linux VM. In this configuration, the PMU interrupt was disabled to prevent the baremetal from idling, as the PMU interrupt would otherwise trigger a negligible overhead due to the long idle wait times (which greatly exceed the combined interrupt latency and callback execution time). Consequently, the observed MBR overhead reflects a scenario where the budget never expires, and only the timer interrupt is active. As shown in Figure 7, the observed overhead for the non-critical VM remains minimal overall, with a noticeable spike only at a regulation period of 1 μ s, where it reaches 14.3%. This elevated overhead is anticipated due to the frequent timer interrupts within such a short interval. However, as the regulation period increases slightly, the overhead rapidly diminishes: at 2 μ s, it drops to 2%, and by 10 μ s, it is already below 1%. For periods beyond 10 μ s, the overhead becomes negligible, approaching zero. This trend demonstrates that the MBR mechanism maintains low interference impact on the non-critical VM, especially at moderate to higher regulation periods, reinforcing its effectiveness in providing controlled memory isolation.



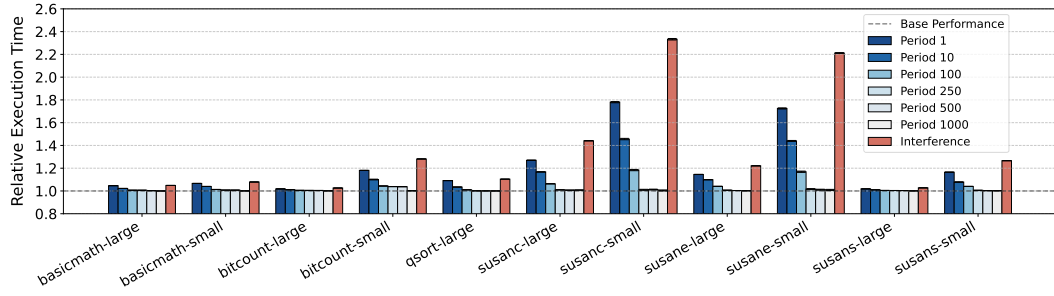
■ **Figure 9** Fixed 10µs period.

4.2 Effect of MBR on Linux Benchmark Performance

This section presents results for the *interf+mb* scenario, where MBR is applied to the non-critical baremetal VM, leaving the critical Linux VM unregulated. This approach enables a focused evaluation of MBR’s impact on the critical VM’s performance stability under varying MBR configurations. Two key aspects of MBR are analyzed: the influence of different budget values at a fixed period and the effect of varying periods at a fixed budget. These parameter variations directly affect memory bandwidth according to $\text{Bandwidth} = \frac{\text{Budget}}{\text{Period}}$.

Fixed Period, Variable Budgets. Figure 9 presents the effects of increasing the MBR budget for the non-critical VM with a fixed 10 µs period. Results show that as the MBR budget allocated to the non-critical VM increases, performance degradation in the critical VM also rises across all benchmarked workloads. This outcome aligns with expectations: a larger MBR budget allows the non-critical baremetal application increased memory bandwidth, thereby intensifying interference within shared memory resources. With budgets ranging from 50 to 10,000 memory accesses per period, lower budgets (e.g., a budget of 50) yield minimal interference, achieving performance levels close to the baseline for the critical VM. Higher budgets, such as 1,000 and 10,000, produce more substantial interference, notably in memory-intensive benchmarks like *susanc* and *susane*, where the relative execution time arises up to 2.3x and 2.2x, respectively. These benchmarks experience considerable performance degradation at higher budgets due to their sensitivity to memory contention. Interestingly, even with the short 10 µs period, small budgets (e.g., 100) still cause a measurable impact on the critical VM’s performance, increasing the execution time by 1.45x for the *susanc-small* benchmark.

Fixed Budget, Variable Periods. Figure 10 illustrates how different regulation periods, from 1 µs to 1000 µs, impact critical VM performance when the MBR budget for the non-critical VM is fixed at 100. Here, we observe that longer periods provide greater protection to the critical VM by reducing memory contention. This is especially evident in memory-sensitive applications, where the extended periods give the critical VM uninterrupted access to memory resources before the non-critical VM’s budget resets, thus improving performance. For instance, with a 1 µs period, critical benchmarks like *susanc* show a performance increase of 1.78x compared to the solo execution. However, as the period extends to 1000 µs, performance approaches solo levels, indicating a substantial reduction in interference. These results underscore that longer periods limit the non-critical VM’s memory bandwidth utilization, effectively mitigating its impact on the critical VM’s workloads.



■ **Figure 10** Fixed budget of 100.

Overhead. The results presented in Figure 8 confirm that applying MBR on the non-critical VM does not introduce any overhead on the critical VM, as each VM has dedicated, isolated CPUs. Additionally, the timer interrupts from MBR that regulate the non-critical VM’s memory access do not interfere with the critical VM, resulting in effectively zero overhead. This is evidenced by the consistency in the critical VM’s execution times across various MBR regulation periods, where variations remain minimal, within the range of 0.01% to 0.2%, equating to a minor fluctuation of around 20 μ s. These findings underscore MBR’s effectiveness for mixed-criticality systems, as it maintains reliable performance in the critical VM without impact from adjustments in the non-critical VM.

5 Discussion and Future Work

Currently, the MBR budget and period parameters are set statically at compile-time. Future work encompasses the introduction of dynamic configuration based on runtime monitoring of memory usage patterns could enable more adaptive and optimized bandwidth allocation. This could involve leveraging machine learning techniques to predict memory demands and proactively adjust MBR settings. Integration with additional isolation mechanisms, such as cache coloring, is another promising direction. By combining MBR with these techniques, contention could be mitigated at multiple levels of the memory hierarchy, providing even stronger temporal isolation guarantees. The interplay between these mechanisms and their cumulative impact on performance and predictability warrants further investigation. Benchmarking other guest OSs with MBR is also an important consideration. Since the mechanism targets real-time applications, and already supports other OSs, benchmarking a RTOS like Zephyr, would further enhance *H-MBR*’s potential. Demonstrating the portability and generalizability of *H-MBR* would broaden its applicability across a wider range of use cases. Finally, long-term efforts could explore extending *H-MBR* beyond CPUs to heterogeneous computing elements like GPUs, FPGAs, and AI accelerators. As embedded systems increasingly incorporate these specialized components, managing their shared memory resources becomes critical. Adapting MBR mechanisms to these contexts could unlock new possibilities for predictable acceleration in mixed-criticality environments.

6 Conclusion

This work introduced *H-MBR*, a VM-centric MBR mechanism for MCS on multicore platforms. Through extensive evaluations, the mechanism demonstrated its effectiveness in mitigating memory contention and enhancing isolation between VMs, while it presented a remarkably low overhead. The evaluation results showed that *H-MBR* significantly reduced the interference on critical tasks by carefully controlling memory bandwidth reservation. Under configurations

where MBR budget was kept low, critical workloads showed minimal to no performance degradation. Overhead analysis further confirmed that *H-MBR* imposes minimal impact on workloads, even when dealing with minor periods. In the worst-case scenario of a 1 μ s regulation period, the mechanism incurs an overhead of up to 14%; however, for regulation periods of 2 μ s or longer, this overhead drops to below 1%. Compared to existing solutions, this is the lowest interrupt overhead observed due to Bao's optimized interrupt handling and lightweight design. This capability makes *H-MBR* stand apart from other implementations, specially in embedded applications, where it maximizes available processing resources for critical tasks without compromising real-time performance.

References

- 1 Jaume Abella, Carles Hernández, Eduardo Quiñones, Francisco Cazorla, Philippa Ryan, Mikel Azkarate-askasua, Jon Perez-Cerrolaza, Enrico Mezzetti, and Tullio Vardanega. WCET analysis methods: Pitfalls and challenges on their trustworthiness. In *Proc. IEEE Int. Symp. on Industrial Embedded Systems*, 2015.
- 2 Ankit Agrawal, Renato Mancuso, Rodolfo Pellizzoni, and Gerhard Fohler. Analysis of Dynamic Memory Bandwidth Regulation in Multi-core Real-Time Systems. In *IEEE Real-Time Systems Symposium*, pages 230–241, 2018. doi:10.1109/RTSS.2018.00040.
- 3 Michael G Bechtel and Heechul Yun. Denial-of-Service Attacks on Shared Cache in Multicore: Analysis and Prevention. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2019.
- 4 Gianluca Brilli, Roberto Cavicchioli, Marco Solieri, Paolo Valente, and Andrea Marongiu. Evaluating Controlled Memory Request Injection for Efficient Bandwidth Utilization and Predictable Execution in Heterogeneous SoCs. *ACM Transactions on Embedded Computing Systems*, 22:1–25, 2022. doi:10.1145/3548773.
- 5 Paolo Burgio, Marko Bertogna, Nicola Capodieci, Roberto Cavicchioli, Michal Sojka, Přemysl Houdek, Andrea Marongiu, Paolo Gai, Claudio Scordino, and Bruno Morelli. A software stack for next-generation automotive systems on many-core heterogeneous platforms. In *Microprocess. Microsyst.*, volume 52, 2017. doi:10.1016/J.MICPRO.2017.06.016.
- 6 Daniel Casini, Alessandro Biondi, Geoffrey Nelissen, and Giorgio Buttazzo. A Holistic Memory Contention Analysis for Parallel Real-Time Tasks under Partitioned Scheduling. In *IEEE Real-Time and Embedded Technology and Applications Symposium*, pages 239–252, 2020. doi:10.1109/RTAS48715.2020.000–3.
- 7 Francisco J. Cazorla, Leonidas Kosmidis, Enrico Mezzetti, Carles Hernandez, Jaume Abella, and Tullio Vardanega. Probabilistic Worst-Case Timing Analysis: Taxonomy and Comprehensive Survey. In *ACM Comput. Surv.*, volume 52, 2019. doi:10.1145/3301283.
- 8 Jon Perez Cerrolaza, Roman Obermaisser, Jaume Abella, Francisco J. Cazorla, Kim Grüttner, Irune Agirre, Hamidreza Ahmadian, and Imanol Allende. Multi-Core Devices for Safety-Critical Systems: A Survey. In *ACM Comput. Surv.*, volume 53(4), 2020.
- 9 Diogo Costa, Luca Cuomo, Daniel Oliveira, Ida Savino, Bruno Morelli, José Martins, Alessandro Biasci, and Sandro Pinto. IRQ Coloring and the Subtle Art of Mitigating Interrupt-Generated Interference. In *IEEE 29th International Conference on Embedded and Real-Time Computing Systems and Applications*, 2023.
- 10 Diogo Costa, Luca Cuomo, Daniel Oliveira, Ida Maria Savino, Bruno Morelli, José Martins, Fabrizio Tronci, Alessandro Biasci, and Sandro Pinto. IRQ Coloring: Mitigating Interrupt-Generated Interference on ARM Multicore Platforms. In *Fourth Workshop on Next Generation Real-Time Embedded Systems*, pages 1–13, 2023. doi:10.4230/OASICS.NG-RES.2023.2.
- 11 Alfons Crespo, Patricia Balbastre, Jose Simo, Javier Coronel, Daniel Gracia Pérez, and Philippe Bonnot. Hypervisor-Based Multicore Feedback Control of Mixed-Criticality Systems. In *IEEE Access*, volume 6, 2018.

- 12 Farzad Farshchi, Qijing Huang, and Heechul Yun. BRU: Bandwidth Regulation Unit for Real-Time Multicore Processors. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2020.
- 13 Marisol García-Valls, Tommaso Cucinotta, and Chenyang Lu. Challenges in Real-Time Virtualization and Predictable Cloud Computing. In *booktitle of Systems Architecture*, volume 60, pages 726–740, 2014. doi:10.1016/J.SYSARC.2014.07.004.
- 14 Angelo Garofalo, Manuele Rusci, Francesco Conti, Davide Rossi, and Luca Benini. PULP-NN: Accelerating Quantized Neural Networks on Parallel Ultra-Low-Power RISC-V Processors. In *Philosophical Transactions of the Royal Society A*, volume 378, pages 01–55, 2020.
- 15 Everaldo P. Gomes, Gabriel De O. Aguiar, and Giovani Gracioli. Implementation and Evaluation of MemGuard in the Bao Hypervisor. In *XIV Brazilian Symposium on Computing Systems Engineering*, pages 1–6, 2024. doi:10.1109/SBESC65055.2024.10771909.
- 16 Giovani Gracioli, Rohan Tabish, Renato Mancuso, Reza Mirosanlou, Rodolfo Pellizzoni, and Marco Caccamo. Designing Mixed Criticality Applications on Modern Heterogeneous MPSoC Platforms. In *Proc. Euromicro Conf. Real-Time Syst.*, volume 133, 2019. doi:10.4230/LIPICS.ECRTS.2019.27.
- 17 M. R. Guthaus, J. S. Ringenberg, D. Ernst, T. M. Austin, T. Mudge, and R. B. Brown. Mibench: A free, commercially representative embedded benchmark suite. In *Proc. IEEE Int. Workshop on Workload Characterization*, pages 3–14, 2001.
- 18 Denis Hoornaert, Shahin Roozkhosh, and Renato Mancuso. A Memory Scheduling Infrastructure for Multi-Core Systems with Re-Programmable Logic. In *33rd Euromicro Conference on Real-Time Systems*, pages 1–22, 2021. doi:10.4230/LIPICS.ECRTS.2021.2.
- 19 Hyoseung Kim and Ragunathan (Raj) Rajkumar. Predictable Shared Cache Management for Multi-Core Real-Time Virtualization. In *ACM Trans. Embed. Comput. Syst.*, volume 17, 2017.
- 20 Tomasz Kloda, Marco Solieri, Renato Mancuso, Nicola Capodieci, Paolo Valente, and Marko Bertogna. Deterministic Memory Hierarchy and Virtualization for Modern Multi-Core Embedded Systems. In *IEEE Real-Time and Embedded Technology and Applications Symposium*, pages 1–14, 2019. doi:10.1109/RTAS.2019.00009.
- 21 Liangzhen Lai, Naveen Suda, and Vikas Chandra. CMSIS-NN: Efficient Neural Network Kernels for Arm Cortex-M CPUs. *arXiv preprint*, pages 1–10, 2018. arXiv:1801.06601.
- 22 Renato Mancuso, Roman Dudko, Emiliano Betti, Marco Cesati, Marco Caccamo, and Rodolfo Pellizzoni. Real-time cache management framework for multicore architectures. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2013.
- 23 José Martins, Adriano Tavares, Marco Solieri, Marko Bertogna, and Sandro Pinto. Bao: A Lightweight Static Partitioning Hypervisor for Modern Multi-Core Embedded Systems. In *Proc. Next Generation Real-Time Embedded Systems*, volume 77, 2020. doi:10.4230/OASICS.NG-RES.2020.3.
- 24 José Martins and Sandro Pinto. Shedding Light on Static Partitioning Hypervisors for Arm-based Mixed-Criticality Systems. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2023.
- 25 Paolo Modica, Alessandro Biondi, Giorgio Buttazzo, and Anup Patel. Supporting Temporal and Spatial Isolation in a Hypervisor for ARM Multicore Platforms. In *Proc. IEEE Int. Conf. on Industrial Technology*, 2018.
- 26 Afonso Oliveira, Gonçalo Moreira, Diogo Costa, Sandro Pinto, and Tiago Gomes. IA&AI: Interference Analysis in Multi-core Embedded AI Systems. In *International Conference on Data Science and Artificial Intelligence*, pages 181–193, 2024.
- 27 Marco Pagani, Enrico Rossi, Alessandro Biondi, Mauro Marinoni, Giuseppe Lipari, and Giorgio Buttazzo. A Bandwidth Reservation Mechanism for AXI-Based Hardware Accelerators on FPGAs. In *31st Euromicro Conference on Real-Time Systems*, pages 1–24, 2019. doi:10.4230/LIPICS.ECRTS.2019.24.
- 28 Rodolfo Pellizzoni, Emiliano Betti, Stanley Bak, Gang Yao, John Criswell, Marco Caccamo, and Russell Kegley. A Predictable Execution Model for COTS-Based Embedded Systems. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2013.

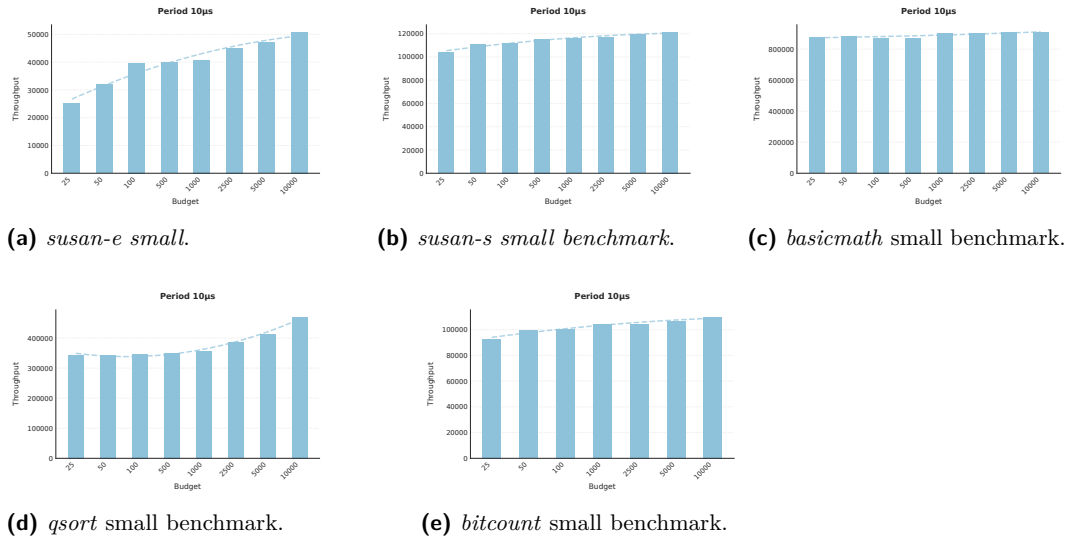
- 29 Sandro Pinto, Hugo Araújo, Daniel Oliveira, José Martins, and Adriano Tavares. Virtualization on TrustZone-enabled Microcontrollers? Voila! In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2019.
- 30 Simon Reder and Jürgen Becker. Interference-Aware Memory Allocation for Real-Time Multi-Core Systems. In *IEEE Real-Time and Embedded Technology and Applications Symposium*, 2020.
- 31 Shahin Roozkhosh and Renato Mancuso. The Potential of Programmable Logic in the Middle: Cache Bleaching. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2020.
- 32 Bernd Schulz and Björn Annighöfer. Evaluation of Adaptive Partitioning and Real-Time Capability for Virtualization With Xen Hypervisor. *IEEE Transactions on Aerospace and Electronic Systems*, 58(1):206–217, 2021. doi:10.1109/TAES.2021.3104941.
- 33 Eric Seals, Michael Bechtel, and Heechul Yun. BandWatch: A System-Wide Memory Bandwidth Regulation System for Heterogeneous Multicore. In *IEEE 29th International Conference on Embedded and Real-Time Computing Systems and Applications*, pages 38–46, 2023.
- 34 Alejandro Serrano-Cases, Juan M. Reina, Jaume Abella, Enrico Mezzetti, and Francisco J. Cazorla. Leveraging Hardware QoS to Control Contention in the Xilinx Zynq UltraScale+ MPSoC. In *Proc. Euromicro Conf. Real-Time Syst.*, 2021.
- 35 Prateek Sharma, Purushottam Kulkarni, and Prashant Shenoy. Per-vm page cache partitioning for cloud computing platforms. In *2016 8th International Conference on Communication Systems and Networks*, pages 1–8, 2016. doi:10.1109/COMSNETS.2016.7439971.
- 36 Parul Sohal, Rohan Tabish, Ulrich Drepper, and Renato Mancuso. Profile-driven memory bandwidth management for accelerators and CPUs in QoS-enabled platforms. In *Real-Time Systems*, pages 235–274, 2022. doi:10.1007/S11241-022-09382-X.
- 37 Meng Xu, Robert Gifford, and Linh Thi Xuan Phan. Holistic multi-resource allocation for multicore real-time virtualization. In *Proceedings of the 56th Annual Design Automation Conference*, 2019.
- 38 Meng Xu, Linh Phan, Hyon-Young Choi, and Insup Lee. vCAT: Dynamic Cache Management Using CAT Virtualization. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, pages 211–222, 2017.
- 39 Ying Ye, Richard West, Zhuoqun Cheng, and Ye Li. COLORIS: A Dynamic Cache Partitioning System Using Page Coloring. In *Proc. Int. Conf. on Parallel Architectures and Compilation Techniques*, pages 381–392, 2014. doi:10.1145/2628071.2628104.
- 40 Heechul Yun, Renato Mancuso, Zheng-Pei Wu, and Rodolfo Pellizzoni. PALLOC: DRAM Bank-Aware Memory Allocator for Performance Isolation on Multicore Platforms. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2014.
- 41 Heechul Yun, Gang Yao, R. Pellizzoni, Marco Caccamo, and Lui Sha. MemGuard: Memory Bandwidth Reservation System for Efficient Performance Isolation in Multi-core Platforms. In *Proc. IEEE Real-Time and Embedded Technology and Applications Symposium*, 2013.
- 42 Heechul Yun, Gang Yao, Rodolfo Pellizzoni, Marco Caccamo, and Lui Sha. Memory Access Control in Multiprocessor for Real-Time Systems with Mixed Criticality. In *Proc. Euromicro Conf. Real-Time Syst.*, 2012.
- 43 Heechul Yun, Gang Yao, Rodolfo Pellizzoni, Marco Caccamo, and Lui Sha. Memory Bandwidth Management for Efficient Performance Isolation in Multi-Core Platforms. In *IEEE Transactions on Computers*, volume 65(2), pages 562–576, 2016. doi:10.1109/TC.2015.2425889.
- 44 Matteo Zini, Giorgiomaria Cicero, Daniel Casini, and Alessandro Biondi. Profiling and controlling I/O-related memory contention in COTS heterogeneous platforms. In *Softw. Pract. Exper.*, volume 52(5), 2022. doi:10.1002/SPE.3053.
- 45 Alexander Zuepke, Andrea Bastoni, Weifan Chen, Marco Caccamo, and Renato Mancuso. MemPol: Policing Core Memory Bandwidth from Outside of the Cores. In *IEEE 29th Real-Time and Embedded Technology and Applications Symposium*, pages 235–248, 2023. doi:10.1109/RTAS58335.2023.00026.

A Appendix

Impact of MBR on Baremetal Throughput and Overhead

This appendix builds on the discussion in Section 4.1, which evaluated the memory-intensive *susan-c* benchmark running on the critical Linux-based VM (C) and focused on the throughput of the non-critical VM (NC). Here, we extend the analysis to additional benchmarks – *susan-e*, *susan-s*, *basicmath*, *qsort*, and *bitcount* – under the same conditions to provide a broader perspective on MBR’s impact on different workloads. By examining these diverse benchmarks, we gain deeper insights into the effectiveness of MBR in managing memory contention across different types of computational tasks.

- ***Susan-e* and *Susan-s* (Figures 11a and 11b):** The NC demonstrated significant improvements in throughput as the budget increased. These results underscore the high memory bandwidth usage of both benchmarks.
- ***Basicmath* (Figure 11c):** Unlike the other benchmarks, the NC running alongside *basicmath* displayed minimal sensitivity to budget changes, maintaining consistently high throughput across all configurations. This stability underscores the benchmark’s lower reliance on memory bandwidth for its computational tasks.
- ***Qsort* (Figure 11d):** The NC VM showed moderate throughput improvement as the budget increased, though not as dramatically as with the *susan* benchmarks, indicating moderate memory bandwidth dependency.
- ***Bitcount* (Figure 11e):** The NC VM’s throughput showed moderate improvement with increasing budgets, reflecting the benchmark’s medium memory bandwidth dependency, with a more gradual scaling compared to *susan-e* and *susan-s*.



■ **Figure 11** Throughput variation on NC VM with different budget configurations for various benchmarks. Benchmarks exhibit different levels of sensitivity to budget changes based on their memory intensity.

These findings underscore the adaptability and practical impact of MBR in managing memory bandwidth effectively. Memory-intensive benchmarks such as *susan-e*, *susan-s*, and *bitcount* display noticeable improvements in NC throughput when allocated higher budgets,

reflecting their heavy reliance on memory resources. Conversely, memory-lighter benchmarks like *basicmath* remain largely unaffected, proving that they are less constrained by memory bandwidth. This demonstrates the importance of workload-specific MBR tuning: adjusting budgets based on workload demands can maximize system performance while preserving isolation.