



Profile-driven memory bandwidth management for accelerators and CPUs in QoS-enabled platforms

Parul Sohal¹ · Rohan Tabish² · Ulrich Drepper³ · Renato Mancuso¹

Accepted: 31 March 2022 / Published online: 26 April 2022

© The Author(s), under exclusive licence to Springer Science+Business Media, LLC, part of Springer Nature 2022

Abstract

The proliferation of multi-core, accelerator-enabled embedded systems has introduced new opportunities to consolidate real-time systems of increasing complexity. But the road to build confidence on the temporal behavior of co-running applications has presented formidable challenges. Most prominently, the main memory subsystem represents a performance bottleneck for both CPUs and accelerators. And industry-viable frameworks for full-system main memory management and performance analysis are past due. In this paper, we propose our *Envelope-aWare Predictive model*, or E-WarP for short. E-WarP is a methodology and technological framework to: (1) analyze the memory demand of applications following a profile-driven approach; (2) make realistic predictions on the temporal behavior of workload deployed on CPUs and accelerators; and (3) perform saturation-aware system consolidation. This work aims at providing the technological foundations as well as the theoretical grassroots for truly workload-aware analysis of real-time systems. This work combines traditional CPU-centric bandwidth regulation techniques with state-of-the-art hardware support for memory traffic shaping via the ARM QoS extensions. We make three key observations. First, our profile-driven methodology achieves, on average, 6% over-prediction on the runtime of bandwidth-regulated applications. Second, we experimentally validate that the calculated bounds hold system-wide if the main memory subsystem operates below saturation. Third, we show that the E-WarP methodology is practical even when applications exhibit input-dependent memory access patterns. We provide a full implementation of our techniques on a commercial platform (NXP S32V234).

Keywords Profiling · Accelerators and CPU benchmarks · Memory profiles · Cumulative memory transaction · Bandwidth regulation · WCET predictions · Traffic-shaping · ARM QoS regulator · Temporal isolation · DDR saturation and utilization · Input-dependent memory pattern · Multi-core real-time systems

✉ Parul Sohal
psohal@bu.edu

Extended author information available on the last page of the article

1 Introduction

The proliferation of inexpensive and high-performance multi-core embedded platforms has been enthusiastically embraced by the industry. These are seen as an opportunity to migrate away from system designs with many interconnected single-core chips; to consolidate all the application workload onto a few systems-on-chip (SoC) with multiple CPUs and accelerators. And while the transition has been smooth for general-purpose workloads, the same cannot be stated for safety-critical systems.

It is well known that contention over shared hardware resources leads to substantial violation of temporal properties when workload developed and tested in isolation is consolidated on the same multi-core platform. Effects like shared cache contention (Ward et al. 2013; Ye et al. 2014), DRAM bank conflicts (Yun et al. 2014; Kim et al. 2014), and contention at the DDR controller (Pellizzoni and Yun 2016) have significantly slowed down the adoption of multi-core solutions in the safety-critical domain. The presence of performance *interference channels* has been acknowledged by certification authorities (C. A. S. Team 2016), which have mandated methodologies to “account and bound” the temporal effect of interference channels for the certification of avionic systems. The last decade has produced seminal results (Maiza et al. 2019) on techniques to manage contention at the different levels of the memory hierarchy. But unfortunately, there is a substantial lack of frameworks and methodologies that can be applied system-wide to: (1) take into account realistic applications, (2) consider that processing workload does not occur only on CPUs; accelerators (e.g. DMAs, video-encoders, GPUs) are also fundamental components in real systems, and (3) that can be deployed on existing platforms while ensuring that the models assumed to derive analytical guarantees are in match with the true behavior of the hardware.

In realistic systems that harness the power of multiple CPUs and accelerators, the main memory subsystem represents the performance bottleneck (Pellizzoni and Yun 2016; Yun et al. 2013). Thus, we focused on the problem of contention in the DRAM subsystem. DRAM bandwidth management (Yun et al. 2013, 2017) is a promising grassroots technique to control main memory contention. Many works have studied the behavior of applications in multi-core systems under main memory bandwidth regulation (Agrawal et al. 2018; Hassan and Pellizzoni 2020; Kim et al. 2014). But the overwhelming focus of these works has been put on formulating an increasingly more refined model of the DRAM subsystem (Hassan and Pellizzoni 2020; Kim et al. 2014) to reduce the pessimism in the timing analysis. On the other hand, the behavior of applications is abstracted away with only a few parameters, for instance, to summarize the worst-case end-to-end number of cache misses (Mancuso et al. 2015; Kim et al. 2014; Pellizzoni and Yun 2016).

In this paper, we propose a focus shift. We introduce a comprehensive framework of techniques called *Envelope-aWare Predictive model*, or E-WarP for short. In E-WarP, accurate predictions on the worst-case execution time (WCET) of co-running applications are made following a profile-driven approach. Profiling represents a substantial refinement of measurement-driven approaches, where

fine-grained knowledge of the interaction between applications and the platform is collected and leveraged. Conversely, as much as possible, we treat the DRAM subsystem as a black-box. By shifting our emphasis on a more precise representation of memory bandwidth requirements of applications and by ensuring that the DRAM subsystem operates below its saturation threshold, we demonstrate that highly accurate predictions on the behavior of tasks operating on CPUs and accelerators can be made.

We stress upfront that we do not construct a formal model of the DRAM subsystem nor formulate provable guarantees. The correctness of our approach is corroborated by a full-system evaluation, which provides evidence that the work presented is practical for industrial applications. Furthermore, our profile-driven approach enables a better understanding of the important aspects that have traditionally received little attention. Precise regulation overheads, the impact of burst size on DRAM utilization, and the unexpected presence of memory instructions that bypass regulation are some examples. The proposed E-WarP framework can be used to integrate multi-core, accelerator-enabled real-time systems in all those domains where a measurement-based approach was deemed acceptable for single-core systems.

This journal paper, which is an extension of “E-WarP: A System-wide Framework for Memory Bandwidth Profiling and Management (Sohal et al. 2020)” makes the following contributions. (1) It introduces the E-WarP model where the time-varying demand for main memory resources is characterized via envelopes. (2) It introduces key requirements and design principles for profile-driven approaches. (3*) In depth discussion of how the ARM QoS infrastructure operates along with the traffic controls that it enables. (4) It considers the integration of broadly implementable techniques for DRAM bandwidth regulation of CPUs and accelerators. (5) It describes how to leverage memory enveloping to perform accurate WCET predictions under regulation for both CPU and accelerator workloads. (6*) It further expands on how to apply the proposed methodology to applications that exhibit input-dependent memory access patterns. (7) It provides a technique to reason on the saturation level of the DRAM subsystem. (8) Lastly, it proposes a full-system implementation and evaluation that includes a low-overhead profiler and an augmented partitioning hypervisor.¹

2 Related work

There have been many research works (Agrawal et al. 2018; Kim et al. 2014; Pellizzoni and Yun 2016; Yao et al. 2016) that aimed to provide hard real-time guarantees for tasks running on multi-core systems. A common denominator in these works is that they consider the worst-case number of main memory transactions (LLC misses) for tasks in isolation (Agrawal et al. 2018; Pellizzoni and Yun 2016; Mancuso et al. 2015; Hassan and Pellizzoni 2020); then compute an upper-bound on memory interference when multiple applications run in parallel. This type of

¹ Contributions indicated with a * are new additions in the journal extension.

analysis has been proposed with various degrees of refinement on different DRAM/CPU models. For instance, in (Kim et al. 2014) the authors assume that there is only one outstanding request per CPU; while (Pellizzoni and Yun 2016) focus on the First-Ready First-Come First-Served (FR-FCFS) DRAM scheduling policy. Compared to this line of work, E-WarP is substantially different because its premise is to rely on high-accuracy observations of the memory demands of applications while treating the DRAM subsystem mostly as a black-box.

Other works, such as (Yun et al. 2013, 2017; Modica et al. 2018) focus on implementable mechanisms to regulate/throttle the bandwidth of other low criticality tasks with the goal of reducing contention and improving performance isolation. The first work in this sense was (Yun et al. 2013), where budget-based bandwidth enforcement is proposed. The work in (Yun et al. 2017) builds on this technique by allowing high-priority tasks to acquire a “bandwidth lock” on the memory controller. These techniques have also been shown to be implementable at the hypervisor level (Modica et al. 2018; Kim and Rajkumar 2016). Recently, there have been important efforts to control, account, and ultimately integrate the behavior of accelerators into real-time systems. The work in (Pagani et al. 2017) lays the groundwork for managing hardware accelerators defined in FPGA, while (Houdek et al. 2017) touches on the topic of non-CPU components regulated via platform-specific throttling mechanisms. In many ways, E-WarP builds on top of the seminal results achieved in this context and complements the CPU-centric management by integrating traditional accelerators (e.g., DMAs, GPUs) in the picture.

Finally, the need for a DRAM controller capable of enforcing bandwidth partitioning and traffic prioritization has been expressed in multiple papers (Li et al. 2016; Bui et al. 2011; Valsan and Yun 2015; Akesson et al. 2007). We acknowledge the important design principles proposed in said works. However, as we strive for immediate industrial applicability, we restrict ourselves to commercial-off-the-shelf platforms.

In summary, our work sets itself apart because it proposes a novel profile-driven methodology to characterize the behavior of applications that execute on CPUs and accelerators. It then combines (1) CPU-centric bandwidth regulation techniques with (2) broadly available hardware support for the regulation of non-CPU masters. In doing so, key relationships between extracted bandwidth and saturation of the DRAM subsystem are derived. Finally, a full-system integration is proposed where we demonstrate that E-WarP is practical in real systems.

3 System model and assumptions

We consider a heterogeneous multi-core system with accelerators and traditional CPUs. A hardware accelerator can be any module capable of initiating transactions to the main memory. DMA engines, GPUs, video encoders/decoders, audio sequencers, network interfaces, are some notable examples. We use m to indicate the number of CPUs present in the system and the index $k \in \{1, \dots, m\}$ to refer to a given CPU_k . The system also features a accelerators indexed using l , with $l \in \{1, \dots, a\}$. The l -th

accelerator is indicated with ACC_l . We use *processing element* (PE) whenever what stated applies to both CPU and accelerator.

We make a restriction, namely the *single driver assumption*, on how accelerators are used in our system. We assume that there exists a single CPU task that acts as the driver for a given accelerator. I.e., it must hold that for ACC_l there exists at most one CPU task acting as the driver. The assumption allows us to abstract away the differences in the preemption model of accelerators. The single driver assumption is accurate only in a subset of possible system designs, but it allows us to keep our focus on how accelerators interact with the main memory. For the same reason, we make the assumption that shared caches (Mancuso et al. 2013; Kim and Rajkumar 2016) and DRAM banks (Yun et al. 2014; Kim et al. 2014) are statically partitioned on a per-core basis to ensure that the load generated by each application toward main memory does not change when multiple applications execute in parallel.

We assume that only one main memory controller is used by all the tasks under analysis. This is referred to as the “DDR controller” or the “DRAM controller”. If more than one controller exists, the techniques presented in this work can be extended by partitioning tasks to memory controllers, and then considering each sub-system independently. We assume that the traffic originated by CPUs and accelerators towards main memory can be regulated. We use budget-based periodic regulation (MemGuard (Yun et al. 2013)) to manage traffic from the CPU; we leverage standard ARM QoS support that is broadly available in modern ARM-based SoCs to regulate traffic from accelerators (see Sect. 6.2). Lastly, the bandwidth at the interconnect should be greater than the bandwidth of both memory controllers.

4 E-WarP Task Model

The E-WarP task model incorporates the relationship between a task’s progress and its demand for main memory. This relationship, expressed via *cumulative memory envelopes*, is captured for each task in isolation. It is leveraged to derive precise predictions on the behavior of the task under regulation. Section 8 is dedicated to constructing memory envelopes following a profile-driven approach.

We consider a set of *n* sporadic, deadline-constrained real-time tasks scheduled according to fixed-priority. The generic task τ_i is statically assigned to execute on a given CPU_k —partitioned fixed-priority scheduling. A task τ_i is a tuple of the following form: $\tau_i = \{T_i, D_i, C_i, \mathcal{M}_i\}$. T_i represents the minimum inter-arrival time between two jobs of the same task, D_i is the relative deadline of task τ_i , and C_i captures the worst-case execution time (WCET) of τ_i in isolation and without memory bandwidth regulation. The $\mathcal{M}_i$ parameter is a super-set of memory envelopes, one per each PE that the task uses. Each memory envelope $M_j \in \mathcal{M}_i$ is of the form $\{R_j, \sigma_j(1), \dots, \sigma_j(L_i)\}$. Here, L_i is simply the number of $\sigma_j(h)$ elements that compose the envelope, while each $\sigma_j(h)$ captures the activity of the task over a fixed small time interval δ . It follows that $L_i = \lceil C_i / \delta \rceil$. The generic $\sigma_j(h)$ has the structure $\{x_j^+(h), x_j^-(h)\}$, where $x_j^+(h)$ (resp., $x_j^-(h)$) captures the upper-bound (resp., lower-bound) on the cumulative number of memory transactions at time $h \cdot \delta$ for the task

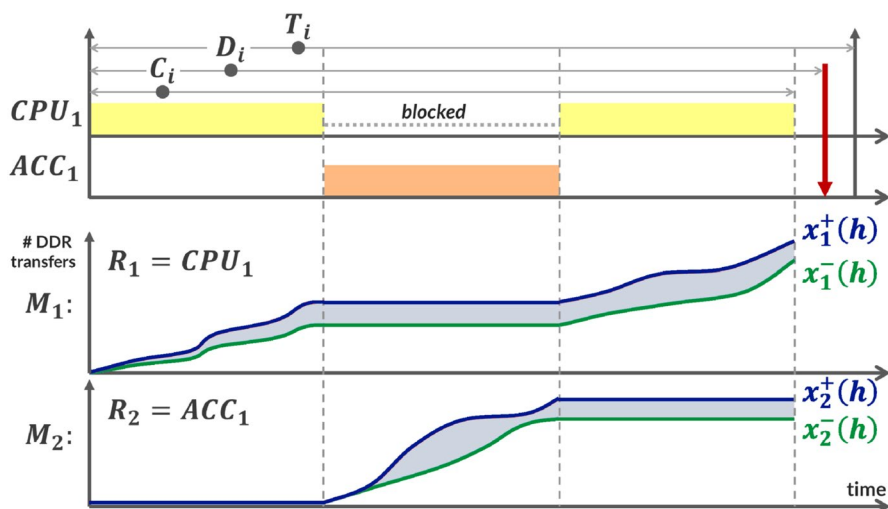


Fig. 1 Overview of main parameters in the E-WarP model for a generic task τ_i executing on CPU_1 and Acc_1

in isolation and without regulation. Lastly, R_j keeps track of the PE (CPU_k or ACC_l) on which the transactions are executed. Figure 1 provides a visualization of the main parameters in the E-WarP model.

Note that, in principle, the E-WarP model is capable of encapsulating the behavior of input-dependent applications. For simplicity, we initially assume that the variations in terms of memory access patterns exhibited by an application under analysis are negligible when a different input vector is provided. We provide a discussion of how such an assumption can be lifted in Sect. 8.3.

5 Transparent profiling

In the E-WarP methodology, the profiler plays a key role in (1) defining the memory envelopes of applications, (2) studying the saturation point of the DDR subsystem, and (3) differentiating between the behavior in main memory of different PEs. To precisely measure these quantities, the profiler must be designed to satisfy the *transparency requirement*. Under this requirement, the task under observation is not (or minimally) impacted by the activity of the profiler. On the other hand, the higher the granularity of the profiler, the smaller will be the pessimism on WCET estimations. This corresponds to the *fine-granularity requirement*. Indeed, an extremely coarse profiler degenerates into a model where the entire activity of the task is summarized by a single value for the worst-case number of transactions.

The definition of a good profiler is challenging because the two requirements are opposing objectives. Indeed, to achieve fine granularity, the profiler needs to frequently sample DDR performance and to keep a complete history of the acquired samples in memory. Thus, a fine-grained profiler acts as a *memory bomb*. We briefly

outline the principles according to which a profiler that satisfies both requirements can be designed and implemented. We describe our software-only implementation in Sect. 9.

To satisfy the fine-granularity requirement, the target platform must provide a performance monitoring interface to sample key metrics of the DDR activity. The closer the interface is to where the transactions are served, the higher the accuracy of the resulting profile. Modern embedded platforms include extensive facilities for performance monitoring (Neill et al. 2017). Some of these interfaces, such as the ARM Performance Monitoring Unit (PMU), are broadly known and supported in software. Often, however, there exist better interfaces that operate much closer to main memory. A few notable examples are discussed in the following paragraph.

The P- and T-series of NXP embedded platforms (NXP 2020a, b) have been extensively studied in the literature (Mancuso et al. 2015; Freitag et al. 2018; Yun et al. 2014; Agrawal et al. 2017, 2018). These platforms include an Event Processing Unit (EPU) and a DDR debug subsystem. The DDR debug subsystem can be configured to generate a trace of events at the DDR controller(s) that includes performed reads, writes, DRAM refreshes, DRAM row hit/miss events, etc. The trace can be processed on chip to create custom event counters. The trace can be also exported to be stored in any block of addressable memory (NXP 2015, 2016).

The Xilinx Zynq UltraScale+ family platforms (Xilinx 2016) that are surging in popularity in the recent years (Roozkhosh et al. 2020; Gracioli et al. 2019; Serrano et al. 2021) include an AXI Performance Monitor (APM) that is interposed between the interconnect and the DDR and that well fits our requirement. The APM can measure the exact number of bytes read/written, as well as their max/min latency over a user-specified sampling interval (Xilinx 2016). Unsurprisingly, support for fine-grained monitoring close to memory resources is not limited to embedded platforms. Intel has recently introduced its family of memory monitoring and management techniques under the name of Resource Director Technology (RDT) (Intel 2019). RDT includes support to monitor the memory bandwidth extracted by CPUs via the Memory Bandwidth Monitoring (MBM) interface (Nguyen 2016). On top of the families of platforms mentioned above, yet another example is the NXP S32V234 (NXP S32V family) platform (NXP 2020c) targeted in our implementation.

To ensure transparency, the platform must allow storing the profiled samples without introducing spurious DRAM traffic. Fortunately, modern embedded platforms feature heterogeneous memory subsystems, with two common features that can be leveraged. (1) The presence of fast on-chip scratchpad memories (SPM); and (2) the existence of multiple DDR controllers. Both are valid alternatives. But the limited size of SPMs restricts the length of profiled application, and/or the granularity of the profile. The NXP P- and T-series family of platforms, the Xilinx Zynq UltraScale+ SoCs, the NXP S32V and S32G family of platforms all define both multiple DDR controllers and on-chip SPMs. A key takeaway is that fine-grained transparent profiling is possible today on a range of modern platforms. A sound implementation requires careful consideration of platforms-specific features and the flow of data within the memory hierarchy. Nonetheless, this level of knowledge of the underlying hardware is not uncommon in the development cycle of safety-critical systems.

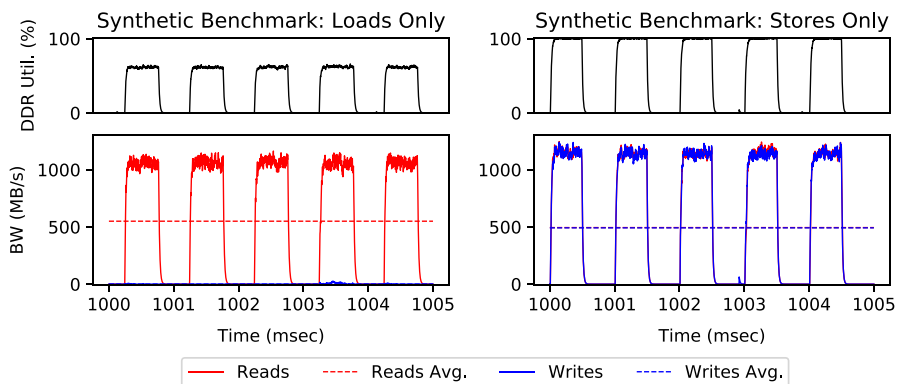


Fig. 2 Profile of a MemGuard-regulated synthetic task performing only loads (left) or only stores (right) under the same regulation budget

6 System-wide bandwidth regulation

To achieve system-wide control over memory bandwidth allocation, with the goal of keeping the DRAM subsystem below its saturation threshold, we combine two different mechanisms for CPUs and accelerators, respectively.

6.1 Regulating CPU memory traffic

The first mechanism is budget-based periodic regulation following the MemGuard (Yun et al. 2013) approach to regulate CPU memory traffic. MemGuard defines a global *regulation period* P and a per-core budget of last-level cache line refills Q_k . The performance measurement unit (PMU) is used to keep track of per-core cache line refills since the beginning of the current P . A local interrupt is delivered by the PMU to stop the core until the next P interval if Q_k refills have occurred. All the cores have their budget synchronously replenished at the beginning of the next regulation interval.

Two important considerations need to be made. Consider a single CPU_k active in the system. First, until CPU_k hits Q_k , it alone can potentially drive the DDR controller to 100% utilization. In other words, MemGuard guarantees no regulation at a time scale that is smaller than P . However, when Q_k is selected appropriately if CPU_k performs back-to-back transactions, it will eventually be regulated/stopped until the next regulation period. Thus, the DDR utilization observed over a time period, not shorter than P can be kept below 100%. We statically set $P = 1$ ms.

Second, the same number of cache refills can result in very different data exchanges with main memory because of write-back CPU caches. When a load or store instruction causes a cache miss, a read transaction is initiated to load the cache block from main memory. If the line being evicted from cache was dirty, then it is written back to memory with a write transaction. Thus, MemGuard only directly regulates read transactions. This phenomenon is shown in Fig. 2 which was produced by our profiler. The figure depicts the behavior in terms of read/write

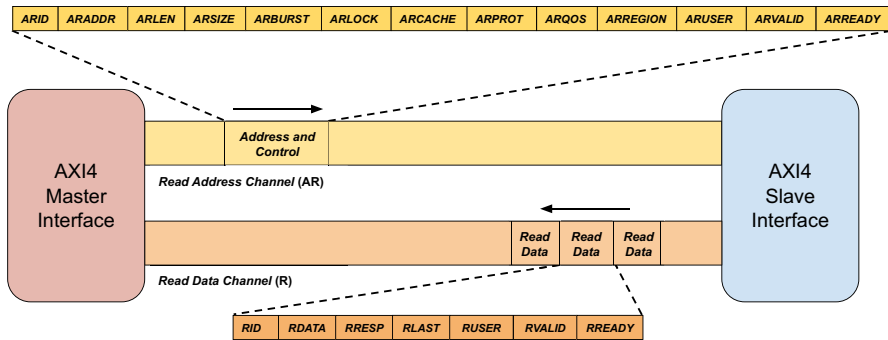


Fig. 3 AXI4 Read address and data channel

bandwidth extracted by a synthetic load- (left) and store-intensive (right) benchmark regulated with the same budget. Note how the same budget can impact the DDR utilization differently, depending on the exact state of the cache. In this particular instance, the same MemGuard budget leads to double the DDR utilization when the benchmark performs writes instead of only reads. Once again, this is due to MemGuard only regulating cache refills. It follows that a safe approach is to construct memory envelopes of read transactions (parameters $x_j^+(h)$ and $x_j^-(h)$), and always assume that the impact on DDR utilization is that of read+write transactions—counted through the profiler.

6.2 Regulating the memory traffic of accelerators

To enforce regulation on accelerators, we leverage traffic shaping via ARM QoS support. The ARM QoS support encompasses a rich set of functionalities implemented in many popular high-performance embedded platforms to manage memory traffic at the level of bus masters. We first provide some necessary background required to understand how the ARM QoS infrastructure works in combination with the Advanced eXtensible Interface 4.0 (AXI4) (Xilinx 2017).

6.2.1 QoS signals in AXI4

Modern ARM-based platforms rely on the AXI4 protocol to establish on-chip data communication channels between processing elements, I/O devices, and memory modules. In AXI4, component-to-component interactions occur through one-to-one full-duplex AXI segments. Each AXI segment defines five channels and the interfaces at the two ends of a segment are referred to as *master* and *slave* interfaces. The master interface is responsible for initiating any read/write transaction. The slave interface produces responses to master-initiated requests.

Figure 3 depicts the signals used between the master and slave interfaces when handling read transactions. A similar view is provided in Fig. 4 for write transactions. AXI4 reads are carried out through two channels: (1) the read address channel (**AR**), and (2) the read data channel (**R**). The **AR** channel carries, among

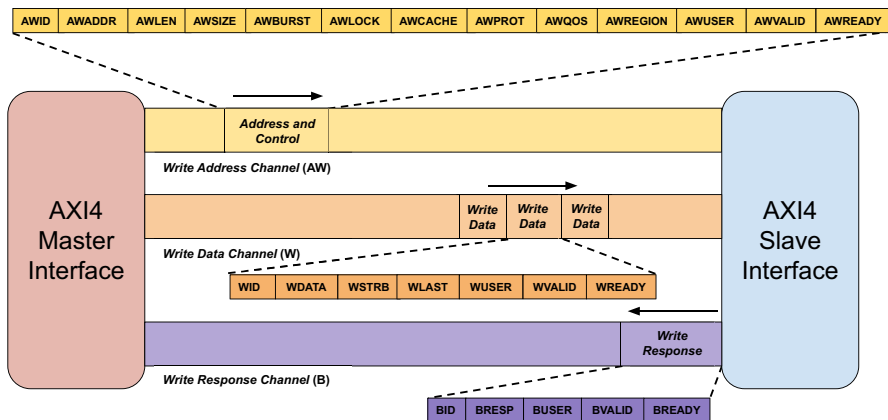


Fig. 4 AXI4 write address (**AW**), data (**W**) and write response (**B**) channels

other signals, the address of the data to be read (**ARADDR**). The amount of data to retrieve and the number of subsequent data beats that should comprise the response are also carried on the **AR** channel by the signals **ARSIZE**—data to be transmitted in each data beat—and **ARLEN**—number of data beats to transfer, i.e. burst length. The resulting data width of each request can be computed as $w = \text{ARLEN} \times \text{ARSIZE}$. The value of w is an important parameter that can change when analyzing the memory behavior of a CPU as opposed to that of an accelerator. Indeed, a single request might transfer more or less data (and thus extract higher bandwidth) depending on the value of w used by the master interface. CPUs addressing cacheable memory generally set w equal to the cache line size when performing cache refills and write-backs. In DMA engines, the burst length, and thus the w parameter, can be configured. The slave responds with a sequence of **ARLEN** data beats where the requested data is carried by the **RDATA** signals through the **R** channel.

To carry out write transactions, a similar interaction occurs between master and slave interfaces. In this case, as depicted in Fig. 4, the width w of the data to be written is determined by the **AWLEN** and **AWSIZE** signals on the **AW** channel. The master then produces **AWLEN** data beats on the **W** channel with the memory content to be written. The slave uses the **B** channel to transmit acknowledgments and for error signaling.

Both read and write interfaces supply additional signals (**ARCACHE**/**AWCACHE**) to propagate cacheability attributes and permission attributes (**ARPROT**/**AWPROT**) through the memory hierarchy. Furthermore, the AXI4 standard includes a set of signals, namely **ARQOS** and **AWQOS**, to relay traffic prioritization information for the purpose of on-chip memory QoS enforcement. Unfortunately, these signals are not meaningful on their own, but require (1) upstream PEs and interconnects to appropriately set these values; and (2) downstream memory components to appropriately perform QoS-aware request handling.

6.2.2 Transaction prioritization at the memory controller

QoS signaling through the AXI4 interfaces described above is only useful if slave components that serve transactions are QoS-aware. An excellent example of QoS-awareness is provided by the DRAM controller implemented in the NXP S32V234 platform that we consider for our evaluation. The controller defines separate read and write queues. Requests arriving at the two queues are admitted to a single *reordering queue* (NXP 2020c). From the reordering queue, the access to be dispatched to the DDR depends on a dynamic priority calculated for any pending transaction in the reordering queue.

The dynamic score for individual transactions is computed as the sum of four values. (1) A 3-bit *static page hit score* that is applied if the pending transaction will result in a DRAM row hit; (2) a 3-bit static access hit score that is applied if the transaction is a read (resp. a write) and the previously forwarded transaction was also a read (resp., write); (3) a 4-bit dynamic jump score which tracks the number of times the access was not chosen during arbitration, and (4) the 4-bit QoS value reported in the ARQOS/AWQOS signals associated with the transaction. The sum of contributions (1)–(4) is then used to compute a 4-bit score with the priority assigned to each transaction in the reordering queue. If an overflow occurs, then the transaction is treated at the highest priority. It follows that setting the value encoded through the ARQOS/AWQOS signals allows direct control over the prioritization of transactions at the DRAM controller.

6.2.3 Traffic control at QoS-enabled interconnects

If the platform includes an interconnect with QoS extensions, such as the ARM QoS-301 (ARM 2011) or ARM QoS-400 (ARM 2013), then three primitives for the control of memory traffic that traverse the interconnect are available. Traffic control is supported for individual master ports attached to the main interconnect. In all the instances of QoS-enabled platforms we have studied, all the CPUs are treated as a single master. Hence, QoS support cannot be used to regulate the traffic of individual CPUs but only to control the aggregated traffic of all the CPUs. Conversely, QoS-based regulation is well suited to regulate traffic from accelerators.

The three primitives for traffic control are (1) *transaction rate regulation*, (2) *outstanding transactions regulation*, and (3) *latency regulation*. Transaction rate regulation enables traffic shaping at the granularity of individual memory requests. In this mode, the network interconnect (NIC) enforces a minimum gap between any two requests originated by a given master as they are forwarded by the interconnect. Read and write traffic can be treated separately or jointly, depending on a configuration switch. The calculation of the inter-transaction gap depends on three configurable parameters, namely the peak rate (ar_p/aw_p), burstiness allowance (ar_b/aw_b), and average rate (ar_r/aw_r). Transactions are forwarded at the peak rate until the number of outstanding transactions reaches the burstiness allowance, after which they are forwarded at the average rate.

Next, in the outstanding transactions regulation mode, the NIC keeps track of the number of transactions that have been forwarded by the interconnect and for which

a response has not been received yet—i.e., outstanding. The NIC allows setting a maximum value of outstanding transactions that can be issued by a master. It stops forwarding additional transactions from said master when the current limit has been reached. The rationale of this approach is that the memory component downstream is a lossless queuing system. Therefore, at steady-state, one can describe the relationship between the average memory latency T_q , arrival rate λ , and number of outstanding transactions q as $q = \lambda T_q$ —i.e. using Little’s Law (Gustafson 2011). Therefore, by controlling q , one can assert implicit control over the average issuance rate λ and average response latency T_q .

Finally, in the latency regulation mode, the NIC tracks the latency of the transactions forwarded by a given master. With that, it allows setting three main parameters. (1) A target latency expressed in clock cycles; (2) the minimum and (3) the maximum value of QoS to be used when forwarding traffic from the considered master. Next, the NIC manipulates the value of the ARQOS/AWQOS signals of forwarded transactions to try and meet the configured target latency. This regulation mode is effective on the considered platform because, as described in the previous section, the downstream memory controller is capable of appropriately prioritizing traffic based on the emitted ARQOS/AWQOS signals.

6.2.4 Our implementation

As mentioned earlier, in transaction rate regulation mode, one can specify a set of parameters to enforce traffic shaping separately for traffic on the **AR** (read requests) and **AW** (write requests) channels. Focusing on the **AR**, the `ar_r` parameter controls the rate of read requests; the `ar_b` the accepted burst size, and `ar_p` the peak rate of read transactions within the allowed burst size. As MemGuard works on a similar regulation technique of controlling the number of transactions within a period, using this QoS mechanism ensures the calculation for DDR utilization is compliant for both accelerators and CPUs. Also, the pipeline for outstanding transaction request regulation did not provide the finer granularity control we needed to achieve regulation for accelerator traffic.

Setting `ar_b` = 1 and `ar_p` = 0 enforces strict regulation at the rate selected by `ar_r`, which is the way QoS is used in this work. The value of `ar_r` can be any 12-bit value greater than 0 (ARM 2011). The resulting inter-transaction gap can be computed as: $t_{ar} = (2^{12}/ar_r)/f_{clk}$, where f_{clk} is the reference clock. In our platform, f_{clk} corresponds to the DDR clock (0.5 GHz). The same applies to write requests. In this work, we always set `ar_r` = `aw_r`, and refer to this value as “QoS level” with the notation Q_1 for each regulated ACC_i . Equation 1 can be used to compute the bandwidth in MB/s given Q_1 , and the size w in bytes of each transaction initiated by ACC_i .

$$b_{qos,w} = \frac{w \cdot Q_1 \cdot f_{clk}}{2^{32}} \quad (1)$$

Figure 5 was obtained by programming the Enhanced DMA (eDMA) on our platform at different QoS levels, and by performing only reads (left), only writes

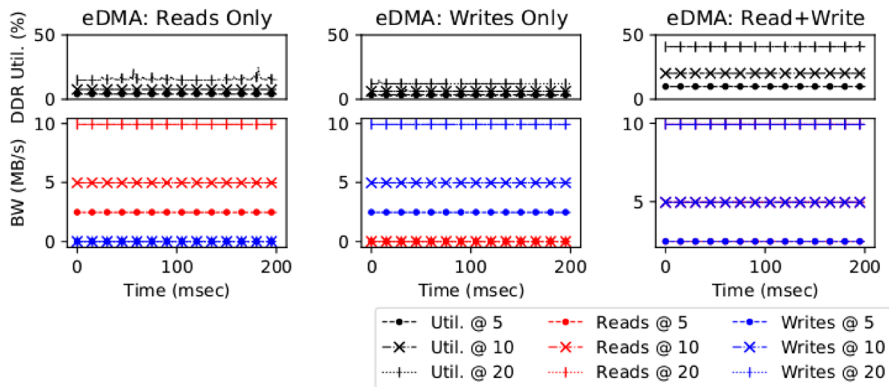


Fig. 5 Profile of QoS-regulated DMA traffic performing only reads (left), only writes (center), or reads+writes (right). Each plot depicts regulation at three enforced rates: 5, 10, and 20

(center), or reads+writes (right). For the eDMA, $w = 4$ bytes. Firstly, unlike MemGuard, QoS regulation operates with a transaction-level granularity (flat lines). Secondly, the read+write traffic once again achieves higher DDR utilization at the same QoS level.

6.3 DDR saturation model under regulation

We propose a linear model that describes how the different active PEs subject to (MemGuard or QoS) regulation impact the utilization U_{tot} of the DDR subsystem. The proposed model is simple and the actual values of the parameters are described in Sect. 10.3. The model is given in Eq. 2:

$$U_{tot} = \sum_{CPU_k} \left(U_{mg}^{\alpha} \cdot \frac{Q_k \cdot L_s}{2^{20} \cdot P} + U_{mg}^{\beta} \right) + \sum_{ACC_l} \left(U_{qos,w}^{\alpha} \cdot Q_1 + U_{qos,w}^{\beta} \right) \quad (2)$$

In the equation, CPUs and accelerators are treated differently because they are differently regulated. For CPUs, we first convert the MemGuard budget to the corresponding bandwidth in MB/s—where P is expressed in seconds and L_s represents the size of a cache line in bytes. Then a linear slope U_{mg}^{α} is applied and the initial offset U_{mg}^{β} is added to find the contribution of each CPU_k to the total utilization. For accelerators, instead, we convert directly from QoS level to contribution in utilization with similar parameters $U_{qos,w}^{\alpha}$ and $U_{qos,w}^{\beta}$. These parameters depend on the transfer size in bytes, w , that masters are capable of transferring with each read/write request—recall that ARM QoS only enforces a minimum inter-arrival time on memory requests, regardless of their size.

7 From profiles to E-WarP tasks

In order to instantiate the E-WarP model, the starting point is the profiles acquired on the task under analysis in isolation. Indeed, a large number of runs and corresponding profiles are required to build confidence on the worst-case behavior, like in traditional single-core measurement-based WCET estimation. The profiles are then integrated to build the task envelopes M_i for the task under analysis. If a task executes on multiple processing elements, then multiple sets of profiles need to be acquired, one per each processing element R_j used by the task. We hereby focus on the definition of the generic M_j for processing element R_j .

Let us first consider a single run and resulting acquired raw profile. A profile is an ordered collection of samples $\{s_r(1), s_r(2), \dots\}$. Each sample collected by the profiler captures the activity of the task under analysis during an interval of length δ . The smaller the parameter, the more accurate the E-WarP model will be. Moreover, for the model to produce valid predictions on the task's WCET under regulation, it must hold that $\delta < P$. We hereby consider that $\delta \ll P$ and evaluate how to find a suitable lower limit for δ in Sect. 10.2.

We use the notation $s_r(h)$ to refer to the h -th sample in the r -th run. Each sample collected by the profiler carries the following information. (1) s_r^r number of bytes read during δ ; (2) the s_r^w number of bytes written during δ . The profile also contains (3) the $s_r^u \in [0, 1]$ utilization of the DDR controller during the δ time window. The latter information is not stored in the task envelopes, but it is useful to study the saturation point of the DDR controller, as studied in Sect. 10.3.

Algorithm 1 Envelope M_j from profiler runs

```

1: function GETENVELOPE( $\tau_i, R_j, \mathcal{R}$ )
2:    $L_i \leftarrow 0$ 
3:    $M_j \leftarrow \{R_j\}$  ▷ The first element is the proc. element
4:   for  $r \leftarrow 1, |\mathcal{R}|$  do ▷ Consider each run
5:      $x_r \leftarrow 0$  ▷ Cumulative num. of transfers in run  $r$ 
6:      $h \leftarrow 1$  ▷ Current sample index
7:      $L_r \leftarrow 0$ 
8:     for  $\exists s_r(h), h \leftarrow h + 1$  do
9:        $L_r \leftarrow L_r + 1$  ▷ Track length of the run
10:       $x_r \leftarrow x_r + s_r^r(h)$ 
11:      if  $L_r > L_i$  then
12:         $L_i \leftarrow L_r$  ▷ Remember longest run
13:         $x_j^+(h) \leftarrow \max(x_j^+(h-1), x_r)$ 
14:         $x_j^-(h) \leftarrow x_r$ 
15:         $\sigma_j(h) \leftarrow \{x_j^+(h), x_j^-(h)\}$ 
16:         $M_j \leftarrow M_j + \{\sigma_j(h)\}$ 
17:      else
18:        if  $x_r > x_j^+(h)$  then  $x_j^+(h) \leftarrow x_r$ 
19:        if  $x_r < x_j^-(h)$  then  $x_j^-(h) \leftarrow x_r$ 
20:   return  $M_j, L_i \cdot \delta$  ▷ Return envelope and WCET

```

Algorithm 1 constructs the envelope M_j and also returns the observed task's WCET in isolation from an arbitrary set of runs $\mathcal{R}$ sorted by shortest-to-longest. The algorithm's logic is simple: (1) we expand the length of the envelope M_j if longer

runs are observed (Lines 11–17); and (2) we keep track of the highest and lowest cumulative number of transactions in each run (Lines 18–19). Note that Algorithm 1 only considers read traffic in the profiles, which is the correct way of deriving envelopes when R_j is a CPU. To apply the algorithm to accelerator tasks, it is enough to replace Line 10 with: $x_r \leftarrow \max(s_r^r(h), s_r^w(h))$, to only keep track of the type of traffic that constitutes the bottleneck. As shown in Fig. 11, the gap between the upper and the lower memory curve for an accelerator task is less due to the presence of finer controls at the level of the bus master.

8 Predicting WCETs from regulation levels

In this section, we describe how to predict the WCET of tasks for which a memory envelope has been constructed according to Sect. 7. The key idea is to mimic the behavior of budget-based regulation (for CPU envelopes) or QoS-based regulation (for accelerator envelopes) as we move through the envelope.

Let us first consider CPU envelopes. Given a generic envelope M_j where $R_j = CPU_k$, we use Algorithm 2 to predict the WCET of the task when CPU_k is assigned MemGuard budget Q_k . To be correct in practice, an extra overhead introduced by MemGuard needs to be taken into account. There are two types of overhead involved. The first, namely t_{ovh} is the upper-bound on the extra time overhead introduced by each periodic budget replenishment. Each activation of MemGuard might also pollute some of the cache partition of the application under analysis, leading to extra memory transactions x_{ovh} being budgeted to the task, compared to when it operates without regulation. We incorporate this overhead as a restriction on the budget given to the core under analysis. Hence, Algorithm 2 considers $Q'_k = Q_k - x_{ovh}$.

8.1 Intuition

Algorithm 2 returns the predicted WCET by keeping track of the additional time t_{add} due to regulation at quota Q'_k . During every regulation period of length P , the algorithm performs multiple steps through the profile samples. At each step, from a memory bandwidth perspective, the worst-case is when (1) the behavior of the application has followed the lower envelope, i.e. when at the generic sample $h - 1$ its cumulative number of memory transactions is exactly $x_j^-(h - 1)$ (Line 16); and (2) at sample h the cumulative number of memory transactions jumps to $x_j^+(h)$. If this difference is greater than Q'_k , (Line 12) then we increase the overall regulation stall. But in doing that, we remember that at least Q'_k transactions were performed by increasing the value of x_{off} which is always considered instead of $x_j^-(\cdot)$ ($\cdot$ refers to an arbitrary sample) when $x_{off} > x_j^-(\cdot)$. This prevents the algorithm from being overly pessimistic. Indeed, by tracking x_{off} , the algorithm captures the worst-case progress of the application as a trajectory somewhere between $x_j^+(h)$ and $x_j^-(h)$.

8.2 Correctness

To understand why the algorithm is safe, let's take a closer look. Consider the easy case where the upper-envelope is equal to the lower envelope, i.e. $\forall h, x_j^+(h) = x_j^-(h)$. In this case, it is enough to keep tracking the progression of transactions. If within a regulation period P we observe more transactions than Q'_k , then the extra regulation time is added to the WCET (Lines 12–15). Conversely, if the budget is not exceeded, it is replenished and counting transactions restarts (Lines 9–11). In this case, transactions might suffer a t_{stall} time due to contention, which is accounted (Line 10). This parameter makes the calculation generic and applicable to in-order micro-architectures. In our observations, no visible stall was measured when the saturation of the DDR is kept below 100%; hence we considered $t_{stall} = 0$. Moreover, any carry-in due to misalignments between δ and P needs to be taken into account—see Line 11.

In the more general case, i.e. when $x_j^+(h) \neq x_j^-(h)$, one must consider the case where the task might have been idle (in terms of DDR activity) and then suddenly performs $x_j^+(h) - x_j^-(h - 1)$ transactions. If the jump incurs regulation, we add the regulation time but also shift up the lower envelope by Q'_k , always preventing it from exceeding $x_j^+(h)$ —see Lines 15–16.

Algorithm 2 Predict WCET for CPU Envelope

```

1: function GETWCET_CPU( $\tau_i, M_j, CPU_k$ )
2:    $t_{add} \leftarrow P$                                 ▷ Track time added by regulation, add tail
3:    $x_{off} \leftarrow 0$                                 ▷ Tracks offset of lower envelope
4:    $t_s \leftarrow 0$                                   ▷ Start time of regulation period
5:    $x_s \leftarrow 0$                                   ▷ Transactions at beginning of regul. period
6:    $h \leftarrow 1$ 
7:   for  $\exists \sigma_j(h), h \leftarrow h + 1$  do
8:      $t \leftarrow \delta \cdot h$                                 ▷ Advance time
9:     if  $t - t_s \geq P$  then                                ▷ No regulation
10:       $t_{add} \leftarrow t_{add} + t_{stall} \cdot x_s + t_{ovh}$     ▷ Add stall due to contention
11:       $t_s \leftarrow t - ((t - t_s) - P)$                 ▷ New beginning of regulation period.
12:      if  $x_j^+(h) - x_s \geq Q'_k$  then                                ▷ Budget exceeded
13:         $t_{add} \leftarrow t_{add} + P - (t - t_s) + t_{ovh}$     ▷ Add regulation stall
14:         $t_s \leftarrow t$ 
15:         $x_{off} \leftarrow \max(x_{off}, x_j^-(h)) + Q'_k$         ▷ Track offset on lower env.
16:         $x_s \leftarrow \min(x_j^+(h), \max(x_j^-(h), x_{off}))$     ▷ New initial number of trans.
17:   return  $t_{wcet} = t + t_{add}$                                 ▷ Predicted WCET

```

To compute WCET predictions on envelopes defined on accelerators, i.e. when $R_j = ACC_l$, we follow a similar yet different strategy because regulation performed by MemGuard differs from QoS traffic shaping. For this purpose, we formulate Algorithm 3. In this case, instead of tracking if Q_k has been surpassed, we track the number of transactions executed in sample period δ . The number of transactions (N) allowed per δ is calculated based on the inter-transaction gap (explained in Sect. 6.2) i.e. N should be less than δ/t_{trans} for regulation to not take place—see Line 8.

If enough transactions to induce regulation fall in the current interval, the resulting regulation-induced inflation to the total runtime is computed at Line 9.

Since we do not know in principle how close will be the first transaction in the subsequent sampling interval, we conservatively assume that an extra regulation gap will be inserted. Next, the algorithm tracks the transactions that have already been considered for regulation (Line 10).

Conversely, when regulation does not take place, we increase the time by δ only. Lastly, for the next iteration of the algorithm, we update the new initial number of transactions in x_s that will be used in iteration $h + 1$, while ensuring we do not exceed the upper envelope. At the same time, the total number of transactions that must have occurred at interval index $h + 1$ must remain above (or in match with) the lower envelope $x_j^-(h)$.

Algorithm 3 Predict WCET for Accelerator Envelope

```

1: function GETWCET_ACC( $\tau_i, M_j, ACC_l$ )
2:    $t_{trans} \leftarrow (2^{12}/Q_l)/f_{clk}$  ▷ Compute inter-transactions spacing at  $Q_l$ 
3:    $t_{wcet} \leftarrow 0$ 
4:    $x_{off} \leftarrow 0$  ▷ Offset for lower envelope
5:    $x_s \leftarrow 0$  ▷ Transaction at the beginning of the considered interval
6:   for  $\exists \sigma_j(h), h \leftarrow h + 1$  do
7:      $N \leftarrow x_j^+(h) - x_s$  ▷ Update number of transactions in this sample period
8:     if  $(N + 1) * t_{trans} \geq \delta$  then ▷ Transactions can cause regulation
9:        $t_{wcet} \leftarrow t_{wcet} + ((N + 1) * t_{trans})$  ▷ Add regulation gaps
10:       $x_{off} \leftarrow \max(x_{off}, x_j^-(h)) + N$  ▷ Track regulated transactions
11:     else ▷ No Regulation
12:        $t_{wcet} \leftarrow t_{wcet} + \delta$  ▷ Add time spent without regulation
13:       $x_s \leftarrow \min(x_j^+(h), \max(x_j^-(h), x_{off}))$  ▷ New initial number of transactions
14:   return  $t_{wcet}$  ▷ Predicted WCET

```

If a task τ_i runs only on a CPU_k , then the new WCET $C_i(Q_k)$ under regulation with budget Q_k can be computed by invoking Algorithm 2. In this case, schedulability can be checked using the traditional partitioned fixed-priority scheduling with preemptions, as long as preemptions are restricted to occur only at the boundaries of regulation periods. However, if preemptions can occur, care must be taken in adding the additional overhead in terms of extra memory transactions performed by τ_i due to cache-related preemption delay (CRPD) (Altmeyer and Burguière 2011; Altmeyer et al. 2010).

However, if a task assigned to CPU_k also uses an accelerator ACC_l , then we assume it will be blocked on CPU_k while it executes on ACC_l . From a CPU scheduling point of view, the time it takes for ACC_l to return control to CPU_k is a self-suspension interval. τ_i 's worst-case response time can be computed leveraging the results in (Nelissen et al. 2015). To compute the overall WCET of τ_i $C_i(Q_k, Q_1)$ subject to regulation on CPU_k with budget Q_k and with ACC_l subject to QoS-based regulation at level Q_1 , the following needs to be computed. First, we compute the stall due to regulation on CPU_k as $t_k^{stall} = C_i(Q_k) - C_i$ computed using Algorithm 2; next, we compute $t_l^{stall} = C_i(Q_1) - C_i$ using the equivalent of Algorithm 2 for QoS regulation. Finally, $C_i(Q_k, Q_1) = C_i + t_k^{stall} + t_l^{stall}$.

8.3 Multiple input vectors for E-WarP

We further discuss how the proposed E-WarP approach can be employed to handle applications whose memory access pattern is impacted by changes in the supplied inputs. The key idea is to consider the behavior produced by multiple input vectors to create a *global* upper and lower envelope that encapsulates the individual $x_j^\pm(h)$ of all the different input vectors. For the remainder of this paper, we will refer to the *global upper envelope* and to the *global lower envelope* with the notation $X_j^+(h)$ and $X_j^-(h)$, respectively.

Creating the global envelope can be done in two steps. First, finding a set of inputs $V = \{v_1, v_2, \dots, v_j\}$ that exhaustively exercise multiple execution paths. Input generation for definitive code coverage is an open challenge for complex applications. Important seminal results have been achieved in this area with the combination of symbolic and concrete execution (Dinges and Agha 2014). The problem of meaningful input generation is beyond the scope of this paper. For all purposes we assume that the set of inputs V for a given application under analysis can be constructed. Second, the global envelope created via the set of input vectors V , should illustrate the worst-case memory execution pattern for the benchmark. For each of the input vectors v_j we have a memory envelope M_j which captures the $x_j^-(h)$ and $x_j^+(h)$ over the runtime of the application with a fixed profiling time interval δ . The global upper and lower envelope is computed by invoking Algorithm 1; the algorithm accepts a set of upper and lower cumulative curves M_j instead of individual single-input profiles. Note that the $X_j^\pm(h)$ is a strictly non-decreasing function as it is a cumulative curve.

Once the $X_j^\pm(h)$ is created over the input set V , we need a model to compute predictions under different CPU and accelerator regulations. There is a probability of each input vector v_j has a drastically different memory envelope making the global memory curve explode. Hence, there is a need to quantify how close the global envelope is to the original memory envelope M_j for input v_j . In our work, we evaluate this phenomenon by understanding the proportion of overlapping area of $x_j^\pm(h)$ to $X_j^\pm(h)$ i.e. comparing the area of M_j to the global envelope created over the input vector set $V = \{v_1, v_2, \dots, v_j\}$ for each j in V . The global envelope will be calculated including all the input vectors and the closeness will be measured by recording the overlapping areas of each input in V with the global envelope. The smaller the overlapping between individual and global envelopes, the more the application's memory behavior can be deemed to be input-dependent. For global envelopes that result much wider/longer than individual input envelopes, non-negligible over-prediction in the resulting WCET estimates should be expected. In this case, an online mechanism to detect the class of inputs supplied to each job and then use an envelope constructed only on the restricted input class can be devised. The latter approach is currently out of the scope of this work, but it lays in our current future work roadmap. Finally, to compute WCET predictions that can apply to an input-dependent application, Algorithm 2 and Algorithm 3 can be reused as is, with the difference that instead of individual envelopes the algorithms are provided in input global

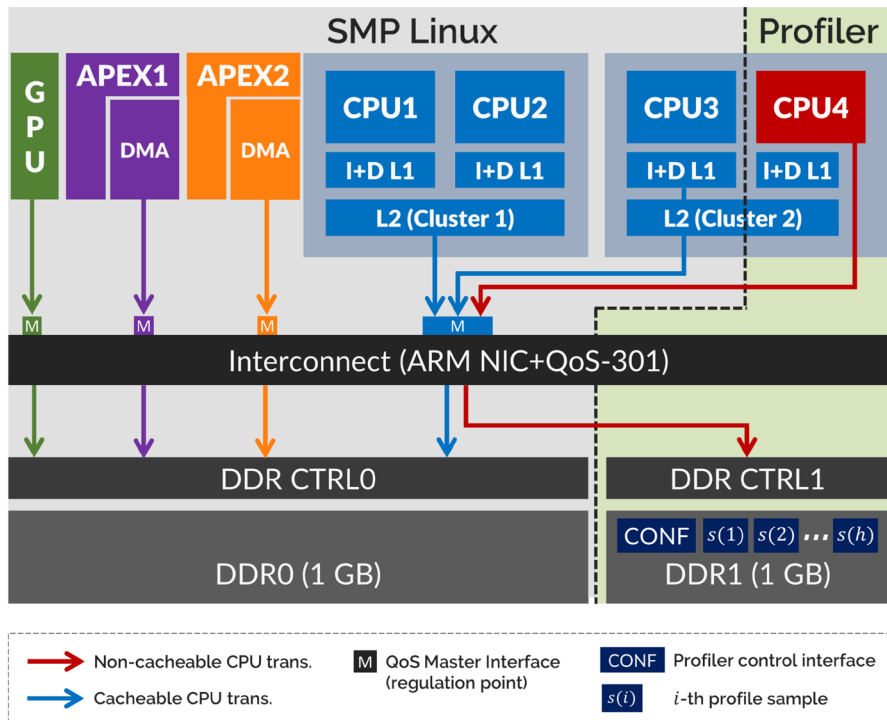


Fig. 6 Block diagram of the main PEs and memory modules in the NXP S32V234 platform. The division between computation and profiling sub-shell is highlighted

envelopes. We refer the reader to Sect. 8 for the construction of global memory envelopes in the case of input-dependent applications.

9 Implementation

We have performed a full-system implementation that includes a low-overhead, high-accuracy profiler, and a partitioning hypervisor augmented to support ARM QoS features. Our implementation was carried out on the NXP S32V234 (NXP 2020c) embedded platform. The main hardware blocks are presented in Fig. 6. The SoC features 4 ARM Cortex A53 CPUs operating at a clock frequency of 1 GHz ($f_{cpu} = 1 \times 10^9$ Hz) divided into two clusters. Each core has a private 32 KB+32 KB I+D L1 cache, and a 256 KB L2 cache is present in each cluster. Because this platform is designed for vision applications, it also integrates two accelerators. The first is a programmable GC3000 GPU (Vivante 2020) and the second is the APEX-CL Image Cognition Processor, or APEX for short. The device contains two identical instances of the APEX engine, namely APEX0 and APEX1. This accelerator promises to deliver “high-performance parallel processing capability” (NXP 2020c).

The APEX are highly complex processing subsystems that include scalar and vector processing units, local scratchpad memories, and DMA engines. We focus on the APEX in our evaluation as a realistic instance of a high-performance accelerator.

The platform features two DDR controllers, namely DDR0 and DDR1, that operate independently on two separate portions of DRAM memory of 1 GB each. The controllers operate at $f_{clk} = 0.5$ GHz and have a bus width of 32 bits. Importantly, each controller exposes a set of memory-mapped performance counters that report: (1) the number of DDR cycles elapsed `tot_ddr_cyc`; (2) the number of busy DDR cycles `busy_ddr_cyc`; (3) the total number of bytes transferred in read (`rd_bytes`) and (4) in write (`wr_bytes`) transactions. The DDR profiling interface also allows defining a filter on the source of traffic (e.g. CPU cluster 1, APEX1, etc.) that is applied when counting read/write bytes. To differentiate between the traffic coming from different masters, counters (3) and (4) can be programmed to only filter the traffic coming from a specific master(s) based on their AXI-ID.

The last component that requires some introduction is the interconnect. The S32V234 system uses a standard ARM NIC-301 (ARM 2010) with ARM QoS-301 (ARM 2011) extensions. The QoS extension of the NIC is where traffic regulation is performed on traffic that traverses the interconnect towards DDR. ARM QoS extensions are surprisingly, broadly available in many current-generation ARM-based platforms. When we started this work, we were surprised to discover that little-to-no software support or research literature was available on these modules. So we had to implement our own to carry out this research. The NIC+QoS-301 provides a memory-mapped interface to control the regulation parameters on a per-master basis. Regulation interfaces are depicted as colored squares on top of the NIC in Fig. 6. Because the traffic from all the CPUs arrives through the same master interface, QoS regulation cannot be used to regulate individual CPUs, but only the total traffic from all the CPUs. Conversely, it allows one to set individual regulation regimes for each of the APEX, for the GPU (see Fig. 6), for the DMAs, for the network interface and the I/O sub-shell (not shown).

We use the Jailhouse partitioning hypervisor (Kiszka et al. 2020) to partition resources in our system. Jailhouse is the ideal choice for this type of implementations because it does not perform scheduling of virtual CPUs (VCPUs), it is lightweight and easy to port/modify, includes support for cache coloring and DRAM bank partitioning (Kloda et al. 2019), and is open-source. It also includes libraries to define bare-metal guest-OS that can be launched directly on a subset of the CPUs. Unfortunately, Jailhouse was not ported to the NXP S32V234 platform at the time we started this work. Our first implementation tasks concerned writing a layer of SoC-dependent code to port Jailhouse onto the target platform. Doing so required a few modifications to the stock boot-loader(u-boot), and to the CPU hotplug support in the Linux kernel.² It also involved writing a driver for the LINFlexD device in the S32 that controls the console outputs.

² This was required to overcome the lack of a PSCI firmware provided by the vendor to control CPU shutdown.

Next, we integrated into our porting an implementation of MemGuard originally proposed in the context of the HERCULES project (Scirdino et al. 2018). We also implemented from scratch a platform-independent support for ARM QoS features, along with the platform-specific code to setup QoS regulation in the S32V234 system. With the implemented support, system designers can set multiple QoS parameters for multiple masters in a single hypercall, making the interface suitable for efficient online dynamic QoS management.

Finally, we implemented a profiler comprised of two parts: a low-level profiler, `profvm`, and a user-space control toolkit, `profctl`. First, `profvm` is a small-footprint bare-metal guest-OS that can be loaded by Jailhouse. To meet the stringent accuracy and transparency requirements of our profiler, we proceeded as follows. When loaded, `profvm` takes exclusive ownership of a single CPU (CPU4), and of an entire DDR controller (DDR1). Our `profvm` uses the dedicated 1 GB of DRAM memory for two purposes. (1) It exposes a shared command&control interface; and (2) when active, stores a sequence of samples of DDR0 activity. The other three CPUs are assigned to Linux in SMP mode and are used to run the user-space applications that need to be profiled. When active, `profvm` performs periodic sampling of DDR0 at a configurable sampling rate expressed in CPU clock cycles. Each sample collected in DDR1 contains the values, and the time of sampling, of: (1) CPU cycles counter, (2) value of `tot_ddr_cyc`, (3) value of `busy_ddr_cyc`, (4) value of `rd_bytes` and (5) `wr_bytes`. The ratio between (3) and (2) provides the instantaneous utilization of the DDR subsystems. Some porting was also required to ensure that the APEX driver does not attempt to use any memory space in DDR1. This is because the out-of-the-box drivers execute APEX code from the memory space of DDR1 controller.

Second, to facilitate profile acquisition, the `profctl` toolkit is provided. It takes care of all the low-level coordination with the `profvm` module; launches the benchmark(s) to be profiled; and at the end of the experiment gathers samples from DDR1 to save them to disk for later analysis.³ Multiple parameters can be configured directly from `profctl`, most prominently sampling period, and filter on individual masters.

Despite all the changes mentioned above, two important features are needed to port E-WarP to another hardware platform. (1) Profiling: The requirements for such a profiling tool are discussed in detail in Sect. 5. (2) Bandwidth Control: MemGuard is a widely-implementable technique and ARM QoS extensions are drop-in modules (ARM QoS-310/QoS-400) bound to increase in popularity. Another tool, ARM Memory System Resource Partitioning and Monitoring (MPAM) (Arm 2018–2020) combines shared cache, memory, and interconnect bandwidth management.

³ <https://github.com/rntmancuso/jailhouse-rt>

10 Validation and evaluation

In this section, we first build a set of experiments to identify key parameters in our system. Next, we discuss how we instantiated the E-WarP model on real-world applications and evaluate the WCET predictions under regulation. Then an in depth analysis of QoS-based controls for accelerators is provided. Finally, we present a full-system integration where all the applications are analyzed in isolation on the CPU and the accelerators are deployed to run in parallel.

10.1 Experimental setup

We used the NXP S32V234 (NXP 2020c) platform introduced in Sect. 9. A combination of synthetic and real benchmarks are used to gain insight into the platform. The synthetic benchmarks used to stress/evaluate specific parameters of our platform are described in the corresponding subsections. For our real benchmarks, we use a subset of the benchmarks in the San-Diego Vision Benchmarks (SD-VBS) suite (Venkata et al. 2009). Because we are interested in applications that are DRAM-bound, the selection was performed by taking all the benchmarks that operate on images. These come with different input sizes, but we have excluded the FULLHD inputs, which lead to impractically long runtimes. We instead focus on the next two largest sizes, i.e. VGA and CIF. The complete list of selected benchmarks is reported in Table 2.

In terms of accelerators, we focus on the APEX engine included in the S32 platform. The S32 features two independent APEX accelerators. Each accelerator is fully programmable and includes a high-performance parallel processing unit (APU) for vector and scalar operations, a DMA, and internal scratchpad memories to operate on data/image tiles. The ARM QoS control interface instantiated on this platform allows setting regulation parameters on the main bus independently for the two APEX engines. The selection of benchmarks available for this unit is limited to the examples released by the manufacturer. We were able to fully integrate the APEX within our profiling infrastructure. But the benchmarks we observed insisted on the processing capabilities of the engine as opposed to generating a lot of DDR traffic. We focus our evaluation on the most DRAM-intensive one we found, i.e., the “Region of Interest” (RoI) benchmark. The RoI benchmark processes different parts of the image on APEX.

For consistency, we always activate the Jailhouse hypervisor. As most of our experiments involve the use of the presented profiler, the `profvm` bare-metal VM is generally loaded (unless specified otherwise) and pinned to core 4. Linux v4.19 is deployed on the other 3 CPUs. Some minor modifications to the kernel were performed to port Jailhouse and to enable support for the APEX. The kernel is compiled in full-tickless (`NO_HZ_FULL`) mode. All the benchmarks are always deployed using the `SCHED_FIFO` scheduler and with explicit pinning to CPUs. We use the `proftcl` to synchronously launch multiple benchmarks in parallel and to coordinate profiling and collection of execution times. All the min/max/

avg statistics were calculated on 30 runs for each configuration to remain statistically significant.

10.2 Profiler transparency and accuracy

As a first experiment, we evaluate how well the proposed profiler satisfies the transparency and accuracy requirements.

The accuracy was evaluated along two sub-dimensions. First, we evaluated how closely the obtained profile matches the expected number of read/write bytes in a synthetic benchmark of known characteristics. To limit the number of spurious DDR transactions in the experiment, we (1) program the platform DMA (eDMA) engine to transfer a known number of bytes; (2) leverage the filtering capabilities of our profiler to only capture eDMA transactions. The resulting profiles cumulative number of read/write bytes were in perfect match with the synthetic benchmark.

Next, we want to find a suitable value for δ that directly relates to the profiler's accuracy. To do so, we varied the configuration of `profvm`'s sampling period and selected the smallest number of CPU clock cycles that leads to a measurement error no larger than ± 2 clock cycles⁴ with 99.99% confidence over 100,000 consecutive measurements. Setting 1,500 clock cycles as the sampling period of `profvm` satisfies this specification. This value was used in all the experiments. With this setting, each acquired sample captures the behavior of the DDR subsystem within a $1.5\mu s$ window. The profiler operates $1500\times$ faster than MemGuard, so it holds that $\delta \ll P$.

Lastly, we evaluated the impact of the profiler on all the selected SD-VBS applications. We first capture the runtime of a benchmark executing without the `profvm` loaded in the system. The runtime is then compared to the case where `profvm` is loaded and configured to collect the profile of the application under analysis. On average across all cases, we observed a runtime increase of 0.33%, with a maximum of 1.67%. Since the profiler is designed to bypass the shared cache and only interact with a private DDR controller, the overhead necessarily arises at the shared interconnect. Because the profiler is not required in production, this overhead will not affect the final applications and all the WCET predictions will still be safe.

10.3 DRAM controller saturation

In this section, we study the saturation of the DDR controller under MemGuard and QoS regulation with the goal of establishing appropriate values for the U_{mg}^a , U_{mg}^b , $U_{qos,w}^a$, $U_{qos,w}^b$ parameters discussed in the previous sections.

10.3.1 MemGuard regulation

We first establish a relationship between MemGuard budget assigned to a CPU, the resulting bandwidth extracted from the DRAM, and the measured DRAM

⁴ The DRAM operates at half the frequency of the CPUs.

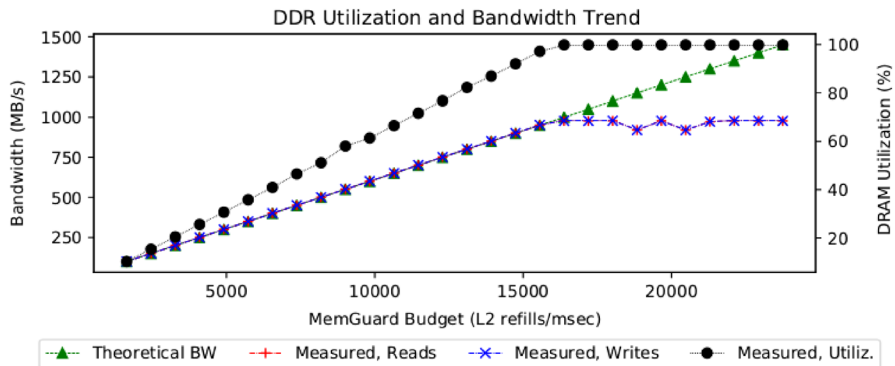


Fig. 7 DDR controller saturation analysis. Bandwidth grows linearly for reads and writes and close to theoretical value as long as the utilization remains below 100%

utilization. Because we are interested in an upper-bound on the utilization, it is important to design an experiment where the DDR utilization is maximized at the selected budget. It is already clear from Fig. 2 that performing stores achieves higher utilization at the same budget level. Furthermore, following the analysis in Hassan (2019) we want to make sure that each DRAM transaction performed by our benchmark results in a DRAM row miss.

With this in mind, we consider the mapping of physical addresses to DRAM coordinates (banks/rows/columns), and design the *USTRESS* synthetic benchmark. *USTRESS* allocates in user-space a 2 MB buffer that is contiguous in physical memory leveraging standard support for huge-pages (`MAP_HUGETLB`). It then performs the first store on column 0 and row 0. The next store is performed 2^{15} bytes away—because the first row bit is bit 15. This pattern keeps all the accesses on bank 0. Once we reach the last accessible row, we set the column offset to 64 bytes and restart from row 0 to fetch the second cache line in the first row. We proceed by scanning all the rows (inner loop) and then increasing the column offset (outer loop) until reaching the last accessible column of the last row. This pattern not only always accesses a closed row in the same bank, but it also bypasses the cache and ensures that no prefetching is performed because subsequent accesses cross the 4 KB page boundary.

We then profile *USTRESS* subject to variable regulation enforced with MemGuard. We compare the theoretical bandwidth that should be extracted with what is observed in the profiles. Simultaneously studying the trend of DDR utilization as returned by the profiles. The results are shown in Fig. 7. As predicted by our model in Eq. 2 for cache line size $L_s = 64$ bytes, the utilization grows linearly as the extracted bandwidth increases. At bandwidth 950 MB/s (budget = 15,565) the controller is running at 97% utilization. At the next budget value we considered (budget = 16,384), 100% utilization is reached, and the observed bandwidth starts to level-off and deviate from the linear trend. Hence we consider 950 MB/s to be a safe bound on the cumulative budget that can be extracted by the CPUs without saturating the DDR controller. By finding the angular coefficient and y-intercept

Table 1 MemGuard budgets (Q_k) and QoS levels (Q_i) with the corresponding bandwidth and utilization

Setting		Bandwidth (MB/s)		DDR Utilization (%)	
MemGuard	QoS	MemGuard	QoS	MemGuard	QoS
492	5	30.03	74.51	3.14	15.68
819	10	49.99	149.01	5.18	30.73
1475	20	90.03	298.02	9.27	60.83
2130	40	130.00	596.05	13.36	121.02
4096	80	250.00	1192.09	25.62	241.41
5734	100	349.98	1490.12	35.84	301.61
7373	160	450.01	2384.19	46.06	482.2
9830	320	599.98	4768.37	61.39	963.76

of the utilization trend before saturation, we can set $U_{mg}^{\alpha} = 6.23856 \times 10^{-3}$ and $U_{mg}^{\beta} = 6.68742 \times 10^{-2}$.

10.3.2 QoS-based regulation

We conducted a similar analysis to the previous case, but this time we use two accelerators to study the relationship between extracted bandwidth and DDR utilization with 4-byte and 128-byte transfers, respectively.

First, we use the eDMA module to generate read+write access patterns at various levels of QoS regulation. The eDMA is configured to generate 32-bit-wide transfers. Unfortunately, the eDMA is not very efficient and hence it cannot bring the DDR to 100% utilization. Indeed, the eDMA does not incur any slowdown when regulated at QoS level 40 and above. Hence, we use the regulation levels between 5 and 20 to establish the linear relationship between QoS levels and utilization. The behavior of the eDMA at these regulation levels was already shown in Fig. 5—see the rightmost subplot for the read + write case. Once again, we observe a linear trend between QoS levels and DDR utilization, and we identify the following parameters: $U_{qos,4}^{\alpha} = 2.05867$ and $U_{qos,4}^{\beta} = -0.383333$. We repeat the same type of experiment using the APEX engine, which transfers 128-bytes with each transaction. This allows us to set the parameters $U_{qos,128}^{\alpha} = 3.00978$ and $U_{qos,128}^{\beta} = 0.632288$.

We provide in Table 1 a recap on the relationship between MemGuard budgets, QoS values, and upper-bounds on extracted bandwidth and the impact on DDR utilization.

10.4 MemGuard regulation: practical quirks

Before moving on to instantiating our E-WarP on real applications, a couple of aspects need to be clarified. These are CPU regulation overheads and limitations to what can be regulated. Both these aspects have been overlooked in previous works because they were hard to evaluate. We were able to use our profiler to evaluate these.

In terms of overhead, we mentioned that MemGuard introduces two types of overheads, i.e. t_{ovh} and x_{ovh} . We designed two synthetic tasks to evaluate these quantities. To evaluate t_{ovh} , we implemented a task that defines a buffer smaller than the L1 cache size, and that continuously samples the CPU cycle counter, storing the difference between two successive samples in the buffer. Because the benchmark does not generate DDR traffic, it will not be regulated by MemGuard. It will, however be interrupted by the periodic interrupt used for budget replenishment. To discover the end-to-end overhead, we then look for discontinuities in the sampled time deltas. With this, we measured the overhead of our Jailhouse implementation of MemGuard to be up to 450 cycles. This is also in line with Dall et al. (2016) and we set $t_{ovh} = 450/1.0 \text{ GHz} = 4.5 \times 10^{-4} \text{ ms}$.

To compute x_{ovh} , we rely on the profiler. We created a benchmark that allocates a buffer of the same size as the last-level cache—256 KB. Like in USTRESS, the buffer is placed contiguously in physical memory to control cache-set conflicts. When this benchmark is profiled, we observe small spikes of memory transactions at the periodicity of MemGuard activations. By counting these transactions on a per-period basis, we computed $x_{ovh} = 35$ transactions.

Another phenomenon we observed by analyzing the profiles of some of our benchmarks is unregulated CPU activity. MemGuard, as well as later implementations like the one in Yun et al. (2017), rely on the `L2_CACHE_REFILL` event to count transactions. Clearly, a CPU can perform DRAM transactions that are not counted by this event by accessing non-cacheable memory, or when performing cache maintenance operations—e.g. a cache flush. Fortunately, these operations are not common in user-space applications. But there exists a class of instruction routinely used in user-space applications that behave in a similar way. Instructions like `STM` (in ARM aarch32) and `STP` (in ARM aarch64) that might be treated as write-no-allocate, which bypass the cache and generate DRAM write transactions. Common operations such as `memset` are implemented using these instructions. We have modified our benchmarks to avoid the use of the problematic instructions.

10.5 QoS regulation: adherence to single-bottleneck model

As we mentioned in Sect. 6.2, we assume that at any point in time, the workload executing on an accelerator is bottle necked by either read or write operations. This is trivially true if the accelerator first copies a block of inputs, computes the result locally, and then writes back the result to memory.

But some accelerators might perform stream-processing, with interleaved reads and writes, where this assumption is less obvious. If the assumption holds, when the traffic that represents the bottleneck is regulated, the other type of traffic will also slow down. This assumption allows us to reason on a single upper-envelope curve and to predict the effect of QoS regulation on the combined read/write traffic. We hereby validate this assumption.

We study the profile of a simple benchmark, namely `ADD`, deployed on the APEX accelerator, and performing streaming vector additions. Because the benchmark reads in input, two operands for each unit of output, we expect the read bandwidth

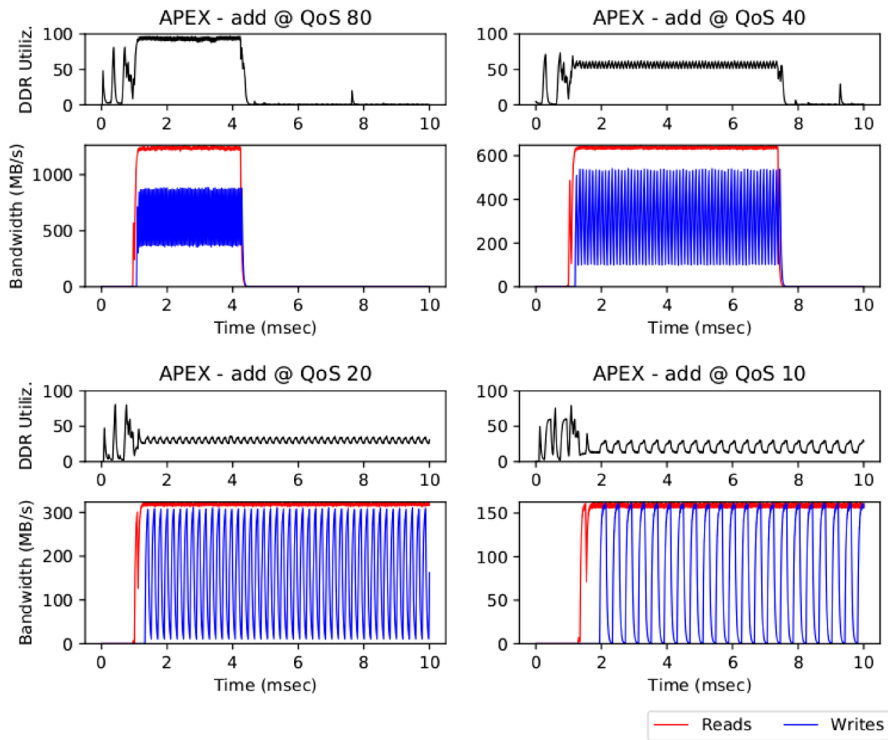


Fig. 8 Profile of QoS-regulated APEX computation over the Add benchmark. From left-to-right, top-to-bottom, the regulation levels are 80, 40, 20, and 10

to be the bottleneck. Figure 8 displays the activity in DDR of this benchmark at different levels of QoS regulation. First, we note that at QoS 80 (top-left) the APEX is able to saturate the DDR and that, as predicted, it is bottlenecked by read operations. As we lower the QoS level to 40, 20, and 10 it can be noted that (1) the extracted write bandwidth also drops; and (2) that the overall length of the operation is dictated by the bottleneck traffic.

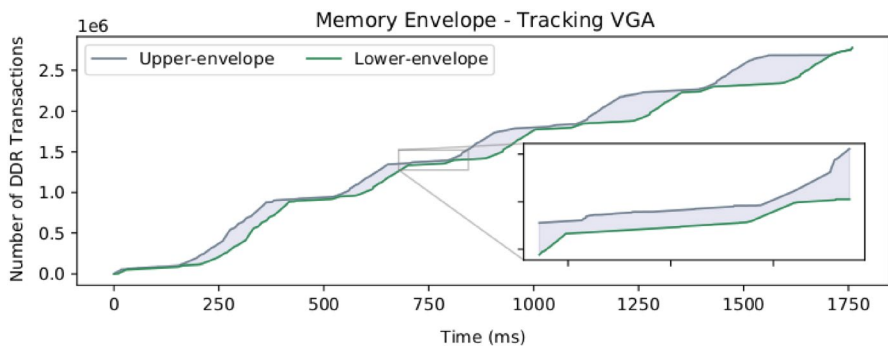
While all the benchmarks we observed on the target platform behaved in a similar way, we do not exclude that workload violating this assumption could be defined. The model presented in this paper does not directly apply to such cases.

10.6 E-WarP instantiation and prediction: CPU tasks

Having identified key system parameters and having assessed the validity of fundamental assumptions, we present the results obtained by instantiating the E-WarP model on the CPU-only SD-VBS applications in this section. In all the predictions presented in this section, 30 runs/profiles of execution were obtained for each benchmark in isolation and without regulation. Then, we produce a prediction for each budget in Table 1 using Algorithm 2. We then run the benchmark under regulation

Table 2 CPU benchmarks—prediction and overestimation

B.mark	In	Budgets									
		492		983		1475		1966		2458	
		Pr. (s)	+ %	Pr. (s)	+ %	Pr. (s)	+ %	Pr. (s)	+ %	Pr. (s)	+ %
sift	cif	5.28	22.51	3.19	13.32	2.6	8.53	2.31	7.44	2.15	6.46
	vga	13.91	9.45	8.74	4.55	7.23	3.08	6.5	2.16	6.08	2.9
disparity	cif	10.68	5.53	5.3	2.57	3.57	1.74	2.72	2.07	2.21	2.69
	vga	29.11	4.49	14.39	1.65	9.65	1.07	7.32	0.74	5.93	1.13
mser	cif	1.79	1.72	0.91	0.98	0.63	0.94	0.49	0.38	0.42	0.5
	vga	8.23	18.82	4.14	15.41	2.83	13.49	2.19	12.76	1.82	13.24
tracking	cif	2.13	8.9	1.16	4.81	0.85	3.46	0.71	2.47	0.63	2.01
	vga	7.2	23.69	3.92	18.9	2.89	16.51	2.39	14.92	2.13	14.34
localiz.	cif	1.16	10.9	1.02	3.82	0.99	2.12	0.97	1.5	0.97	1.73
	vga	4.48	5.28	3.7	2.02	3.5	1.8	3.41	1.1	3.38	0.83
tex_synth	cif	0.43	14.16	0.3	7.27	0.26	5.13	0.24	2.92	0.24	2.62
	vga	1.42	10.35	1.1	3.04	1.01	1.51	0.98	0.62	0.97	0.55
stitch	cif	1.42	9.69	1.09	3.38	1	2.11	0.96	1.27	0.93	0.74
svm	cif	1.73	6.19	1.58	1.65	1.55	1.07	1.53	0.97	1.53	0.87
						MAX	23.69%	MIN	0.38%	AVG	5.71%

**Fig. 9** Upper- and lower-envelope computed over 30 runs of the tracking benchmark

at each of the selected values, and compare our predictions against the maximum runtime observed under regulation.

The full list of results is summarized in Table 2. For each benchmark/input size, we report the predicted time in seconds (“Pred. (s)”) column and the overestimation percentage (“Incr. (%)”) compared to the longest run observed under the considered budget. We only report numbers for the lower values of budgets because they are where predictions become worse. However, we have carried out our predictions on the full range of considered budgets and confirmed that the obtained WCET always upper-bounds the experimental observations.

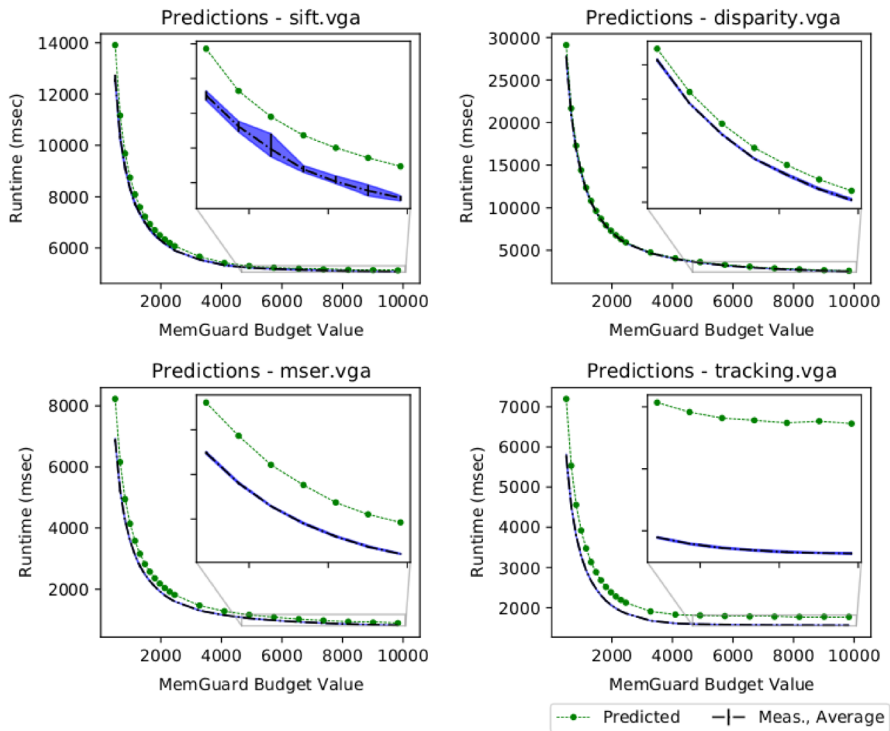


Fig. 10 Prediction v/s measured runtime for SIFT, DISPARITY, MSR, and TRACKING benchmarks with VGA input size. Zoomed in where curve flattens

We visualize the memory envelope obtained on the application that led to higher overestimation, namely **TRACKING** with input **VGA** in Fig. 9. Unlike most of the other applications we studied, the upper and lower envelopes create a visible gap, which forces the prediction to be more pessimistic. As part of future work, we want to explore how much these curves diverge, and how that affects predictions, when different inputs are provided in different runs.

In order to understand how precisely the runtime of CPU tasks can be predicted following the E-WarP methodology, we refer to Fig. 10. Here, we plot the trend of our predictions against the timing observed in the actual runs. The blue area represents the min-max error range around the average. We provide insets in each sub-plot to zoom in on the portions that are otherwise harder to appreciate.

All in all, what stands out is that our prediction remains extremely close to the observed runtimes, with an average over-prediction below 6% for budgets between 492 and 2415 (Table 2); and with max an average over-prediction of 13.6% and 3%, respectively, for budgets in the mid-range 4915–9830 (not shown in Table 2).

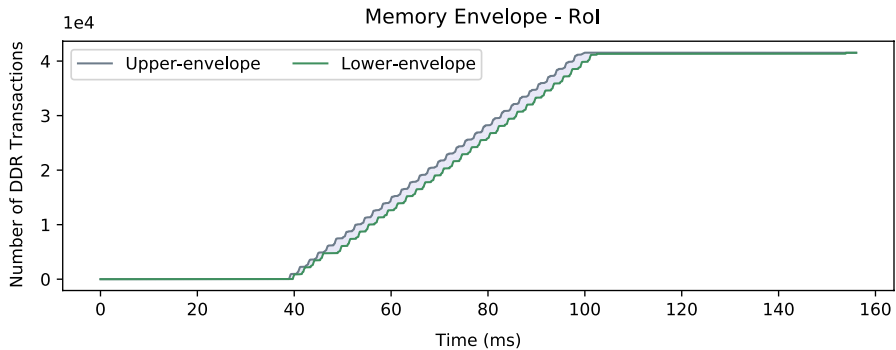


Fig. 11 Upper- and lower-envelope computed over multiple runs of RoI benchmark

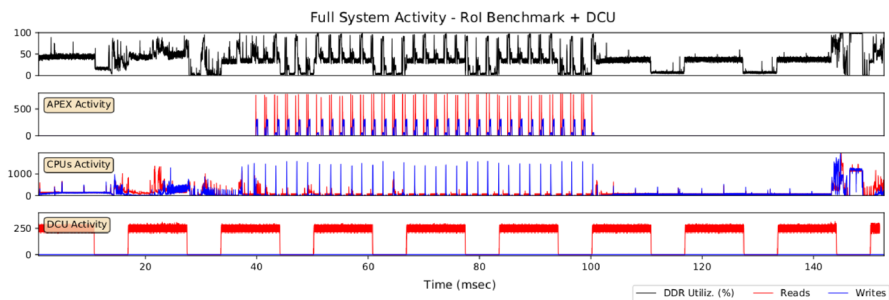


Fig. 12 Full system DDR activity. RoI benchmark in execution. Utilization on top; APEX, CPUs and DCU activity in the other sub-plots

10.7 E-WarP instantiation and prediction: accelerator tasks

We now consider the most memory-intensive benchmark, i.e. RoI, available for the APEX accelerator. With RoI, we perform the combined analysis for a task that runs in intermittent phases on the CPU and the APEX. The upper- and lower-envelopes for the activity of the RoI benchmark on the APEX accelerator is provided in Fig. 11. Compared to the CPU tasks studied in the previous section, the APEX exhibits a much more deterministic behavior that translates into envelopes with very similar shape and small separation.

To better contextualize the behavior of the RoI benchmark, the full profile of CPU and APEX activity in DRAM when RoI is executed is depicted in Fig. 12. The figure was obtained by acquiring three profiles of RoI, in each profile, we filter traffic by the AXI-ID either APEX, the CPU, or the Display Control Unit (DCU). More details about the DCU are provided in the next section.

We first perform prediction of the runtime of the benchmark when no QoS-based regulation is enforced on transactions from the APEX block, while the CPU activity of the benchmark is regulated at different levels considered thus far. The results of this experiment are reported in the left plot of Fig. 13. Once again,

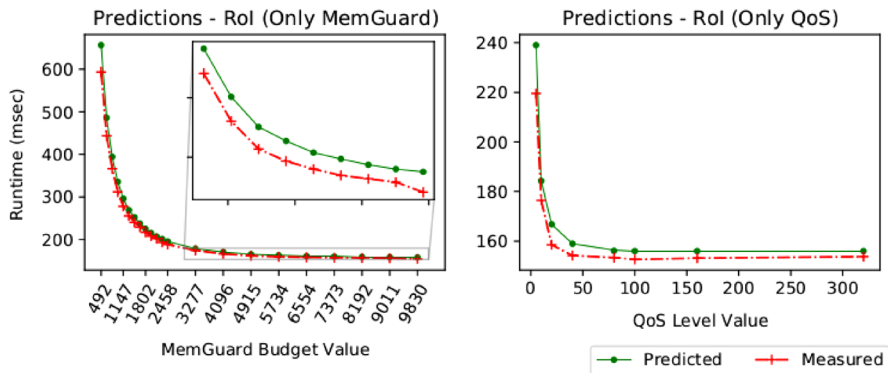


Fig. 13 Comparison of measured runtime of RoI benchmark under MemGuard-only regulation for CPUs (left) and QoS-only regulation for APEX (right)

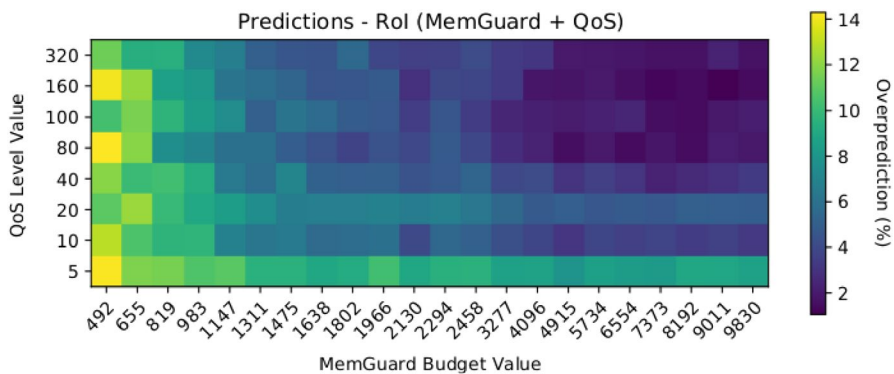


Fig. 14 Comparison of measured runtime of RoI benchmark under both MemGuard- (CPUs) and QoS-based regulation (APEX)

we observed that the predictions remain very close to the maximum measured runtimes, with 10.58% max over-prediction at the lowest budget and 4.15% on average across all the considered budget values.

We then studied our prediction on the same benchmark but when only the APEX is regulated using QoS, while the CPU is not throttled. The results for the QoS values considered in Table 1 are reported in the right sub-plot in Fig. 13. Once again, with E-WarP we are able to predict with good accuracy the total execution time of the benchmark. The maximum over-prediction was 8.82%, while the average sat at 3.62%.

Next, we evaluate the accuracy of E-WarP predictions over all the possible pairs of QoS values and MemGuard budgets in Table 2. We visualize the results in terms of over-prediction as a heat-map in Fig. 14, where CPU regulation levels and APEX regulation levels are varied on the x- and y-axis, respectively. As expected, the highest levels of throttling for both CPU and APEX lead to the largest

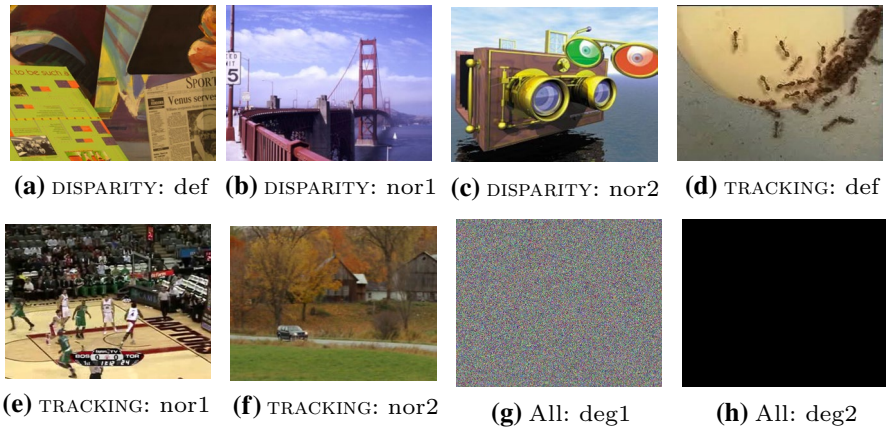


Fig. 15 Default and additional input images to understand global envelopes with changing input vectors

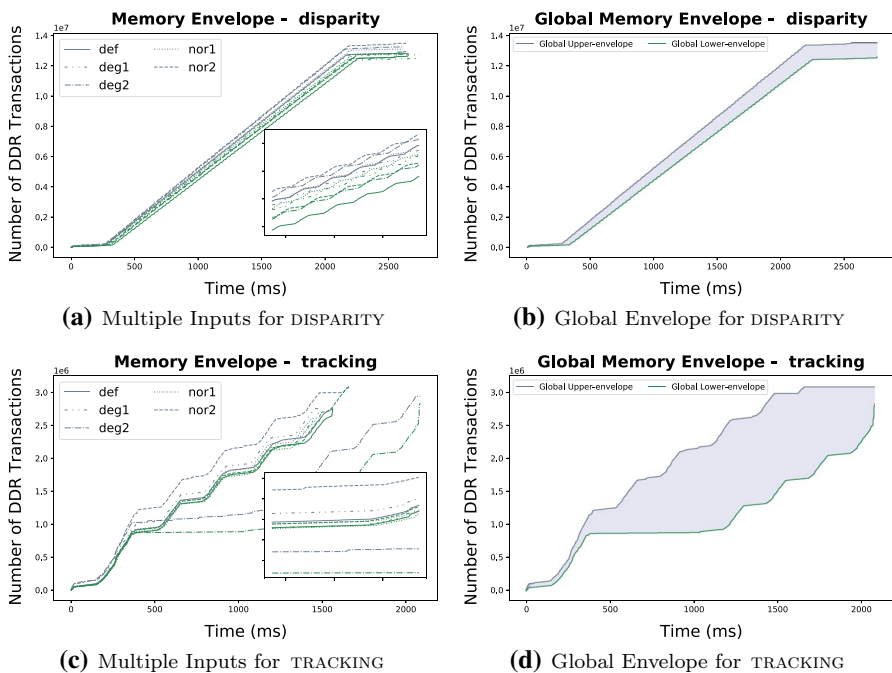


Fig. 16 Representation of multiple inputs of same size for DISPARITY and TRACKING

over-estimations—around 14%, bottom-left corner. Conversely, over-estimation is below 2% for high QoS and budget values (top-right corner).

10.8 E-WarP instantiation and prediction: multiple input vectors

For the input-dependent global envelopes, we focused on two different scenarios: (1) where the various inputs deviate slightly from each other, and (2) where certain inputs increase the size of the global envelope tremendously. Hence, we focused on these two benchmarks, as visually illustrated in Fig. 16. DISPARITY (top) with the five input vectors (explained in further detail next) creates a compact global memory envelope. TRACKING (Fig. 16) on the other hand has an inflated global envelope. In particular, note that one input vector (*deg2*) causes the application to run longer with a final total number of memory transactions in line with what is observed with the other inputs. This causes the computed envelope to stretch over a longer time interval. When the *deg2* envelope is used to create the global envelope, a wide gap is introduced between the $X_j^+(h)$ and $X_j^-(h)$ curves, as depicted in Fig. 16d. In light of the differences that exist between DISPARITY and TRACKING when considering variable input vectors, we take a closer look at the WCET (over-)prediction that can be achieved when reasoning on global envelopes.

SD-VBS benchmark suite has the default (“def”) input images it comes with and we have described previously. The selected input size is VGA for DISPARITY and TRACKING. For each benchmark, we have produced four additional input images. The first two called “nor1” and “nor2” are meaningful (*normal*) scenes, while the last two, namely “deg1” and “deg2” are scenes that correspond to corner (*degenerative*) cases. Specifically, “deg1” corresponds to random noise while “deg2” to a solid-color frame.

Figure 15⁵ provides a visualization of the considered input vectors. In the case of DISPARITY, the input is a stereoscopic scene. We depict one of the two images in Fig. 15; finally, TRACKING takes as input four subsequent video frames. We provide the first frame of the sequence.

We divided the experiments for each benchmark into three separate cases to generate the global memory envelopes: $V_1 = \{deg1, deg2, def\}$, $V_2 = \{nor1, nor2, def\}$, and $V_3 = \{deg1, deg2, nor1, nor2, def\}$. The results from this experiment are summarized in Table 3. For example, let's take $V_1 = \{deg1, deg2, def\}$ i.e. Row 1 and Row 4 in Table 3. These rows show, for each columns, the overlapping area between the global memory envelope obtained with V_1 and the envelope obtained with only one specific input image (i.e., *deg1*, *deg2*, or *def*). We also show, on the right-hand side of the table, the WCET over-prediction percentage for three MemGuard budgets when an envelope constructed using V_1 is considered. We calculate

⁵ Figure 15b, c: original photos by Alexander Klein and Stefan Wernthaler, respectively, from <https://www.stereoscopy.com/>; Figure 15e, f: original video frames from the Visual Tracker Benchmark, respectively Basketball and CarScale data sets available at http://cvlab.hanyang.ac.kr/tracker_benchmark/datasets.html. The original photos have been scaled and/or cropped to match the same resolution and aspect ratios as the default SD-VBS image files.

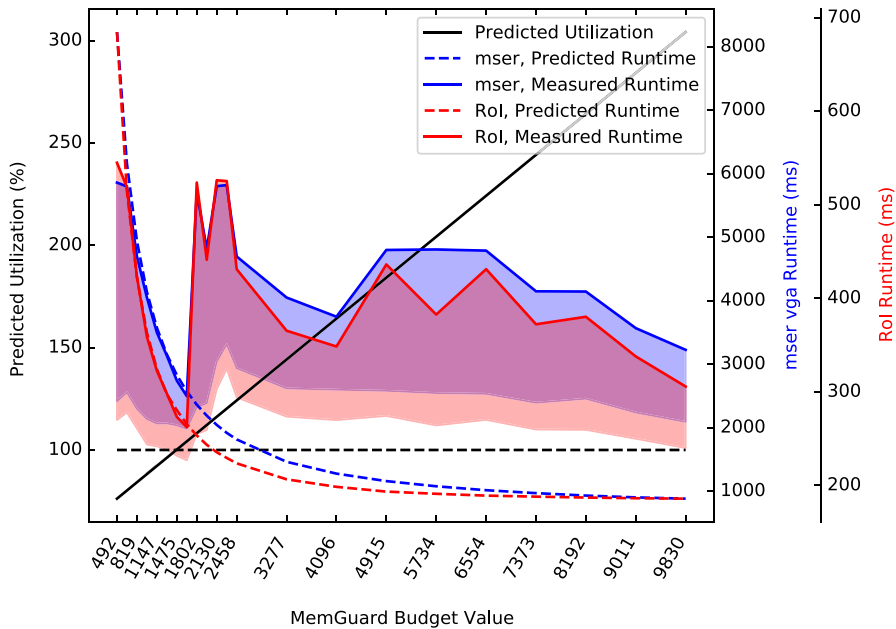


Fig. 18 Full system setup with RoI, MSER and two bandwidth intensive synthetic applications applications on CPU 1, 2, 3, and 4, respectively

the MSER (or TRACKING) benchmark on (CPU 3) with VGA input; and instantiate two memory bombs continuously performing DDR transactions on CPUs 2 and 4.

First, we need to determine suitable budget and QoS levels, since the unconstrained system easily drives the DDR to 100% utilization. In this system, when all the drivers from the manufacturer have been loaded, the DCU becomes active. The role of the DCU is to transfer display frames from the frame-buffer in DDR, to the display port. The DCU is active even without a display connected to the I/O port. Our profiler was able to reveal the presence of this spurious activity, which is visible in the bottom plot of Fig. 12, which underlines the importance of having a tool like the one proposed in this work when working on modern complex embedded platforms. While it is possible to disable the DCU, we believe an active DCU makes for a more realistic setup. Hence, we conduct our experiments by simply accounting for its impact on the utilization of the DDR subsystem.

We measure the upper-bound on DDR utilization caused by the DCU at 36%. Unfortunately, the DCU cannot be regulated using QoS. To reduce the number of parameters, we also set the QoS level for the APEX at 10, so that the APEX can increase the DDR utilization by at most 30.73%. From Sect. 6.3, we know that a safe utilization is 97%. Hence the cores need to be assigned MemGuard budgets so that they increase the DDR utilization by no more than 30.27%, which corresponds to a total budget of 4915 (about 300 MB/s). With 4 active CPUs, and by performing an even division of this quota, we expect that the DDR remains below the saturation threshold as long as the individual CPU budgets remain below 1228.

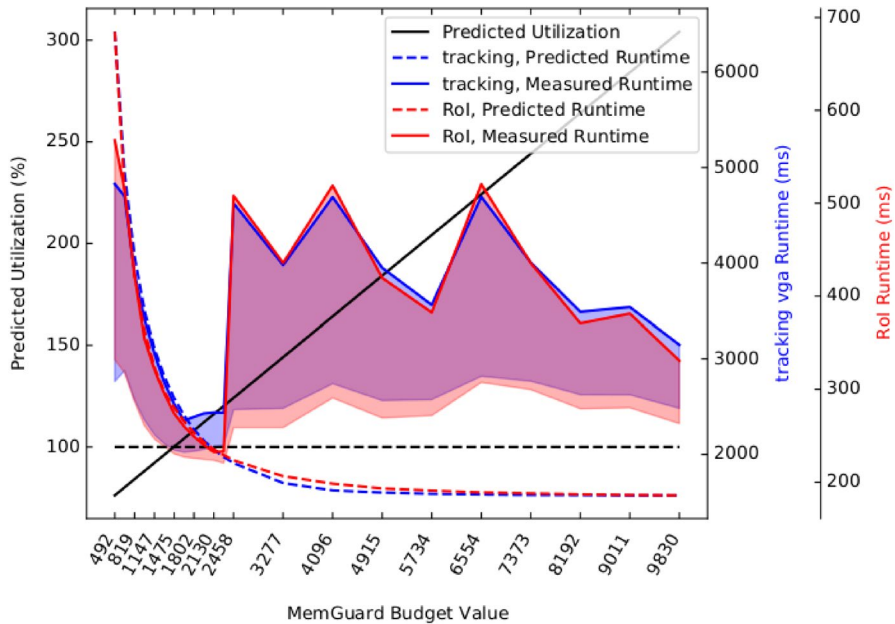


Fig. 19 Full system setup with RoI, TRACKING and two bandwidth intensive synthetic applications applications on CPU 1, 2, 3, and 4, respectively

In Fig. 18 (resp., Fig. 19), we plot what happens to the runtime of the CPU tasks, i.e. RoI and MSER (resp., TRACKING) with VGA input as we increase the budgets on the CPUs. The black solid line tracks the predicted DDR utilization, with the 100% threshold marked with the dashed line. Solid blue lines are used to plot the maximum observed runtime of the MSER (resp., TRACKING), with our predictions depicted in the same color and dashed lines. The same convention using red lines is used to plot the runtime of the RoI benchmark. The areas under the blue/red curves captures the difference between observed maximum and average runtimes. Three main characteristics stand out in the figures. (1) In both, the maximum runtimes correctly remain below the predicted WCETs until 100% DDR utilization is reached, which confirms the validity of the E-WarP approach. (2) Once the saturation point is exceeded, the behavior of the system is highly unpredictable, with our benchmarks experiencing large swings in execution times that are not mitigated by increasing the CPU budgets. (3) In the system with TRACKING, the benchmarks behave erratically slightly later than the predicted saturation point. This is possible because the proposed utilization model has to be conservative to be safe.

11 Conclusion and future work

This work presented E-WarP, a framework of technologies to profile and bound the temporal behavior of workload on CPUs and accelerators. E-WarP achieves full-system memory bandwidth management by integrating two broadly available regulation mechanisms. We design and implement a fine-granularity, transparent profiler. We show how to build relationships between regulation levels and DDR saturation. Finally, we experimentally demonstrate that the formulated WCET predictions hold as long as the main memory subsystem remains below its saturation threshold.

E-WarP is meant to be a stepping stone for profile-driven real-time application analysis with realistic upper-bounds on application runtimes. It enables important future research avenues in directions that include: (1) optimally setting regulation parameters leveraging the convexity of the E-WarP's predictions; (2) performing WCET impact-aware dynamic regulation control in the OS; and (3) integrating our profile-driven approach with formal DRAM models for provable performance guarantees.

Acknowledgements The material presented in this paper is based upon work supported by the National Science Foundation (NSF) under Grant Number CCF-2008799. The work was also supported through the Red Hat Research program. Any opinions, findings, and conclusions or recommendations expressed in this publication are those of the authors and do not necessarily reflect the views of the NSF.

References

- Agrawal A, Fohler G, Freitag J, Nowotsch J, Uhrig S, Paulitsch M (2017) Contention-aware dynamic memory bandwidth isolation with predictability in COTS multicores: an avionics case study. In: 29th Euromicro conference on real-time systems (ECRTS 2017). Schloss Dagstuhl-Leibniz-Zentrum fuer Informatik
- Agrawal A, Mancuso R, Pellizzoni R, Fohler G (2018) Analysis of dynamic memory bandwidth regulation in multi-core real-time systems. *IEEE Real-Time Syst Symp (RTSS)* 2018:230–241
- Akesson B, Goossens K, Ringhofer M (2007) Predator: a predictable SDRAM memory controller. In: 2007 5th IEEE/ACM/IFIP international conference on hardware/software codesign and system synthesis (CODES + ISSS). pp 251–256
- Altmeyer S, Burguière CM (2011) Cache-related preemption delay via useful cache blocks: survey and redefinition. *J Syst Architect* 57(7):707–719
- Altmeyer S, Maiza C, Reineke J (2010) Resilience analysis: tightening the CRPD bound for set-associative caches. *ACM Sigplan Notices* 45(4):153–162
- ARM (2010) AMBA network interconnect (NIC-301) technical reference manual. accessed 07 Jan 2020
- ARM (2011) *ARM[®] CoreLink[™] QoS-301 network interconnect advanced quality of service*. accessed 07 Jan 2020
- ARM (2013) *ARM[®] CoreLink[™] QoS-400 network interconnect advanced quality of service*. accessed 07 Jan 2020
- Arm (2018–2020) arm architecture reference manual supplement memory system resource partitioning and monitoring (MPAM), for Armv8-A. accessed 16 Oct 2020
- Bui D, Lee E, Liu I, Patel H, Reineke J (2011) Temporal isolation on multiprocessing architectures. In: 2011 48th ACM/EDAC/IEEE Design Automation Conference (DAC). pp 274–279
- C. A. S. Team (2016) Multi-core processors position paper. accessed 07 Jan 2020
- Dall C, Li S-W, Lim JT, Nieh J, Kolovrentzos G (2016) ARM virtualization: performance and architectural implications. In: (2016) ACM/IEEE 43rd annual international symposium on computer architecture (ISCA). IEEE, pp 304–316

- Dinges P, Agha G (2014) Targeted test input generation using symbolic-concrete backward execution. In: Proceedings of the 29th ACM/IEEE international conference on automated software engineering, ser. ASE '14. New York, NY, USA: Association for Computing Machinery, pp 31–36. <https://doi.org/10.1145/2642937.2642951>
- Freitag J, Uhrig S, Ungerer T (2018) Virtual timing isolation for mixed-criticality systems. In: 30th Euro-micro conference on real-time systems (ECRTS 2018) ser. Leibniz. In: Altmeyer S (ed) International proceedings in informatics (LIPIcs), vol. 106. Dagstuhl, Germany: Schloss Dagstuhl–Leibniz-Zentrum fuer Informatik, pp 13:1–13:23. <http://drops.dagstuhl.de/opus/volltexte/2018/8990>
- Gracioli G, Tabish R, Mancuso R, Miroslanlou R, Pellizzoni R, Caccamo M (2019) Designing mixed criticality applications on modern heterogeneous MPSoC platforms. In: 31st Euromicro conference on real-time systems (ECRTS 2019). Schloss Dagstuhl–Leibniz-Zentrum fuer Informatik
- Gustafson JL (2011) Little's law. Springer, Boston, pp 1038–1041
- Hassan M (2019) Reduced latency DRAM for multi-core safety-critical real-time systems. *Real-Time Syst* 56:1–36
- Hassan M, Pellizzoni R (2020) Analysis of memory-contention in heterogeneous COTS MPSoCs (ECRTS2020)
- Houdek P, Sojka M, Hanzálek Z (2017) Towards predictable execution model on ARM-based heterogeneous platforms. In: 2017 IEEE 26th international symposium on industrial electronics (ISIE). IEEE pp. 1297–1302
- Intel (2019) Resource director technology reference manual. accessed 07 Jan 2020
- Kim H, Rajkumar R (2016) Real-time cache management for multi-core virtualization. *Int Conf Embedded Softw (EMSOFT)* 2016:1–10
- Kim H, De Niz D, Andersson B, Klein M, Mutlu O, Rajkumar R (2014) Bounding memory interference delay in COTS-based multi-core systems. In: 2014 IEEE 19th real-time and embedded technology and applications symposium (RTAS). pp 145–154
- Kiszka J, Sinitsin V, Schild H, contributors, Jailhouse Hypervisor. accessed 07 Jan 2020 <https://github.com/siemens/jailhouse>
- Kloda MST, Mancuso R, Capodiceci N, Valente P, Bertogna M (2019) Deterministic memory hierarchy and virtualization for modern multi-core embedded systems. In: 25th IEEE real-time and embedded technology and applications symposium (RTAS 2019), Montreal, Canada, conference, pp 1–14
- Li Y, Akesson K, Goossens K (2016) Architecture and analysis of a dynamically-scheduled real-time memory controller. *Real-Time Syst* 52(5):675–729
- Maiza C, Rihani H, Rivas JM, Goossens J, Altmeyer S, Davis RI (2019) A survey of timing verification techniques for multi-core real-time systems. *ACM Comput. Surv.* 52(3):1–38. <https://doi.org/10.1145/3323212>
- Mancuso R, Dudko R, Betti E, Cesati M, Caccamo M, Pellizzoni R (2013) Real-time cache management framework for multi-core architectures. In: 19th IEEE real-time and embedded technology and applications symposium (RTAS 2013), Philadelphia, PA, USA. pp 45–54
- Mancuso R, Pellizzoni R, Caccamo M, Sha L, Yun H (2015) WCET(m) estimation in multi-core systems using single core equivalence. In: 2015 27th Euromicro conference on real-time systems, pp 174–183
- Modica P, Biondi A, Buttazzo G, Patel A (2018) Supporting temporal and spatial isolation in a hypervisor for ARM multicore platforms. *IEEE Int Conf Ind Technol (ICIT)* 2018:1651–1657
- Neill R, Drebes A, Pop A (2017) Fuse: accurate multiplexing of hardware performance counters across executions. *ACM Trans Archit Code Optim (TACO)* 14(4):1–26
- Nelissen G, Fonseca J, Ravari G, Nélis V (2015) Timing analysis of fixed priority self-suspending sporadic tasks. In: 2015 27th Euromicro conference on real-time systems. pp 80–89
- Nguyen KT (2016) Introduction to memory bandwidth monitoring in the Intel® Xeon® processor. accessed 07 Jan 2020
- NXP (2015) P4080 multicore communication processor reference manual. accessed 07 Jan 2020
- NXP (2016) QorIQ T2080 reference manual. accessed 07 Jan 2020
- NXP (2020a) P-series in QorIQ processing platforms
- NXP (2020b) T-series in QorIQ processing platforms
- NXP (2020) S32V234 reference manual. accessed 07 Jan 2020
- Pagani M, Balsini A, Biondi A, Marinoni M, Buttazzo G (2017) A Linux-based support for developing real-time applications on heterogeneous platforms with dynamic FPGA reconfiguration. In: 2017 30th IEEE international system-on-chip conference (SOCC). pp 96–101

- Pellizzoni R, Yun H (2016) Memory servers for multicore systems. In: IEEE Real-time and embedded technology and applications symposium (RTAS). pp. 1–12
- Roozkhosh S, Mancuso R (2020) The potential of programmable logic in the middle: cache bleaching. In: 2020 IEEE real-time and embedded technology and applications symposium (RTAS). IEEE pp 296–309
- Scirdino C, Cuomoand L, Solieri M, Sojka M (2018) HERCULES: high-performance real-time architectures for low-power embedded systems. accessed 07 Jan 2020
- Serrano-Cases A, Reina JM, Abella J, Mezzetti E, Cazorla FJ (2021) Leveraging hardware QoS to control contention in the Xilinx Zynq UltraScale+ MPSoC
- Sohal P, Tabish R, Drepper U, Mancuso R (2020) E-WarP: a system-wide framework for memory bandwidth profiling and management. In: 2020 IEEE real-time systems symposium (RTSS), pp 345–357
- Valsan PK, Yun H (2015) MEDUSA: a predictable and high-performance DRAM controller for multicore based embedded systems. In: 2015 IEEE 3rd international conference on cyber-physical systems, networks, and applications. pp 86–93
- Venkata SK, Ahn I, Jeon D, Gupta A, Louie C, Garcia S, Belongie S, Taylor MB (2009) SD-VBS: the san diego vision benchmark suite. In: 2009 IEEE international symposium on workload characterization (IISWC). IEEE, pp 55–64
- Vivante, Vega Cores for 3D. accessed 07 Jan 2020. <http://www.vivantecorp.com/index.php/en/technology/3d.html>
- Ward BC, Herman JL, Kenna CJ, Anderson JH (2013) Outstanding paper award: making shared caches more predictable on multicore platforms. In: 2013 25th Euromicro conference on real-time systems. IEEE, pp 157–167
- Xilinx (2016) ZCU102 user guide. accessed 07 Jan 2020
- Xilinx (2017) AXI4 reference guide. accessed 07 Jan 2020
- Yao G, Yun H, Wu ZP, Pellizzoni R, Caccamo M, Sha L (2016) Schedulability analysis for memory bandwidth regulated multicore real-time systems. IEEE Trans Comput 65(2):601–614
- Ye Y, West R, Cheng Z, Li Y (2014) Coloris: a dynamic cache partitioning system using page coloring. In: 2014 23rd international conference on parallel architecture and compilation techniques (PACT). IEEE. pp 381–392
- Yun H, Yao G, Pellizzoni R, Caccamo M, Sha L (2013) MemGuard: memory bandwidth reservation system for efficient performance isolation in multi-core platforms. In: 2013 IEEE 19th real-time and embedded technology and applications symposium (RTAS). pp 55–64
- Yun H, Mancuso R, Wu Z-P, Pellizzoni R (2014) PALLOC: DRAM bank-aware memory allocator for performance isolation on multicore platforms. In: IEEE 19th real-time and embedded technology and applications symposium (RTAS). IEEE pp 155–166
- Yun H, Ali W, Gondi S, Biswas S (2017) BWLOCK: a dynamic memory access control framework for soft real-time applications on multicore platforms. IEEE Trans Comput 66(7):1247–1252

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Parul Sohal is a Ph.D. student at Boston University, advised by Professors Orran Krieger and Renato Mancuso. Her interests lie in providing temporal isolation for applications sharing resources at the different levels of the memory hierarchy. Her approach to co-locating applications on multi-core platforms involves in-depth profiling of applications online and offline to ensure soft performance guarantees. Her work is in collaboration with Red Hat. She is also a student member of the ACM and the IEEE.

Rohan Tabish received his Ph.D. from the Department of Computer Science at the University of Illinois at Urbana-Champaign, IL, USA. His research interests are in real-time and embedded systems with a focus on the use of OS-level techniques in multi-core processors to achieve predictability and strong timing guarantees such that they can be deployed in safety-critical automotive and avionics applications. Recently, he has also been exploring neural-networks based software stacks to build real-time and safe AI. He is also interested in wireless communications, HPC, and sensor networks. He is a student member of the IEEE.

Ulrich Drepper is a distinguished engineer at Red Hat, working in the research group. His interests extend from hardware design, reprogrammable hardware, OS and runtime design, to the tools to make all this possible and better.

Renato Mancuso is an assistant professor in the Computer Science Department at Boston University (BU). Before joining BU in 2017, Renato received his Ph.D. from the Department of Computer Science at the University of Illinois at Urbana-Champaign. He is interested in high-performance cyber-physical systems, with a specific focus on techniques to enforce strong performance isolation and temporal predictability in multi-core, accelerator-enabled systems. He has published more than 50 papers in major conferences and journals. His papers received multiple awards in top conferences in the field of real-time and embedded systems. His research is supported by the NSF, Bosch, Cisco Systems, and Red Hat. He is a member of the IEEE.

Authors and Affiliations

Parul Sohal¹  · Rohan Tabish² · Ulrich Drepper³ · Renato Mancuso¹

Rohan Tabish
rtabish@illinois.edu

Ulrich Drepper
drepper@redhat.com

Renato Mancuso
rmancuso@bu.edu

¹ Boston University, Boston, USA

² University of Illinois at Urbana-Champaign, Champaign, USA

³ Red Hat, Raleigh, USA