

Allocation Policies Matter for Hybrid Memory Systems

Adnan Maruf
amaru009@fiu.edu
Florida International University
Miami, Florida, USA

Daniel Carlson
dcarl026@fiu.edu
Florida International University
Miami, Florida, USA

Ashikee Ghosh
ghashike@amazon.com
Amazon Web Services
New York City, New York, USA

Manoj Saha
msaha002@fiu.edu
Florida International University
Miami, Florida, USA

Janki Bhimani
jbhimani@fiu.edu
Florida International University
Miami, Florida, USA

Raju Rangaswami
raju@cs.fiu.edu
Florida International University
Miami, Florida, USA

ABSTRACT

Existing tiered memory systems all use DRAM-Preferred as their allocation policy whereby pages get allocated from higher-performing DRAM until it is filled after which all future allocations are made from lower-performing persistent memory (PM). The novel insight of this work is that the right page allocation policy for a workload can help to lower the access latencies for the newly allocated pages. We design, implement, and evaluate three page allocation policies within the real system deployment of the state-of-the-art dynamic tiering system. We observe that the right page allocation policy can improve the performance of a tiered memory system by as much as 17x for certain workloads.

KEYWORDS

hybrid memory systems, memory allocation, memory tiering

ACM Reference Format:

Adnan Maruf, Daniel Carlson, Ashikee Ghosh, Manoj Saha, Janki Bhimani, and Raju Rangaswami. 2018. Allocation Policies Matter for Hybrid Memory Systems. In *Proceedings of The 32nd International Symposium on High-Performance Parallel and Distributed Computing (HPDC'23)*. ACM, New York, NY, USA, 2 pages. <https://doi.org/XXXXXXX.XXXXXXX>

1 INTRODUCTION

In tiered memory systems, optimizing dynamic page migration has received substantial attention. The state-of-the-art tiered memory systems such as MULTI-CLOCK, AutoTiering, AMP, Nimble, and others [?] dynamically reorganize the pages across memory tiers based upon the accesses to the pages. These systems improve the overall performance of the dynamic workloads by periodically scanning/sampling the page accesses to determine page importance and perform *page selection* followed by *page migration* to move the page(s) to an appropriate tier. However, to our surprise, we noticed that all the existing tiered memory systems allocate the new pages in DRAM until full, and then allocate the remaining new pages throughout the workload from the lower tier (i.e., the tier with

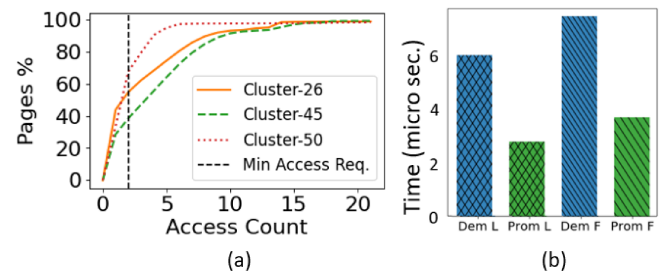


Figure 1: Page Selection (a) and Migration (b) Overhead.

higher access latency compared to DRAM). In the rest of this paper, we refer to this as the DRAM-Preferred allocation policy. Hence, in the steady state of any workload when DRAM is fully utilized, the important pages get accessed from the slower tier until the page is selected and promoted to the faster tier. *The novel insight of this research is that with the right page allocation policy in addition to tiering algorithm, the workload performance can be further improved by the lower access latencies as well as reduced page selection and migration overheads for the newly allocated pages.*

The *page allocation policy* of a tiered memory system determines which tier new page allocations are made from prior to engaging page migration mechanisms.

Figure ??(a) shows the page selection overhead of the dynamic tiered systems, with the number of page access on the X-axis and the percentages of total pages on the Y-axis for Twitter cluster trace [?]. The dashed line presents the required number of page accesses (i.e., 2) for Multi-Clock to select the page to migrate. *From Figure ??(a), we can see that the accesses to the 67% of the total pages for cluster 50, 57% for cluster 26, and 38% for cluster 45 depend on the initial placement of the pages due to the page allocation policy used.* These pages would not even be considered for migrations as these pages have a total access count less than the required threshold. Figure ??(b) shows the time taken to demote and promote a page in most of the dynamic tiering systems [?]. *With the default DRAM-Preferred allocation policy, hot pages that are allocated after the DRAM is filled will trigger migration. Thus, the number of initial demotions and promotions can be hundreds of millions for workloads like the Twitter clusters where the working set size of a cluster can reach terabytes [?]. These demotions and promotions would require a significant amount of time just for the migrations. The above overheads can be significantly reduced by choosing the right tier to allocate new pages.*

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

HPDC'23, June 20–23, 2023, Orlando, FL

© 2018 Association for Computing Machinery.

ACM ISBN 978-1-4503-XXXX-X/18/06...\$15.00

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

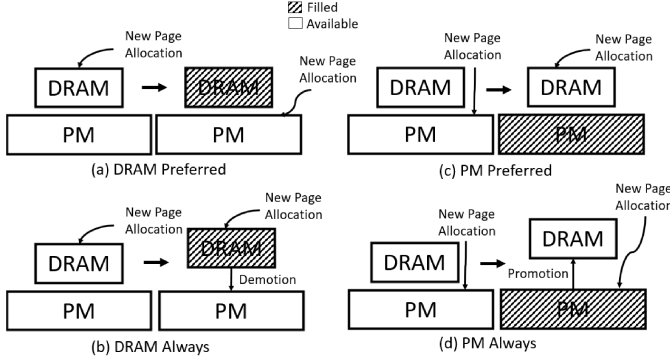


Figure 2: Page allocation policies.

In this work, first, we design and implement three different allocation policies. Second, we integrate the above allocation policies with the state-of-the-art dynamic tiering systems deployed using Linux kernel to study the impact of allocation policies within tiering systems using various real workload traces and several popular benchmarks.

We find that the page allocation policy is as important as the dynamic tiering policy. Allocation policies can impact the tiering system performance by up to 17x. While all of the state-of-the-art tiering mechanisms use the *DRAM-Preferred* allocation policy, we find that no single allocation policy performs the best for all workloads.

2 ALLOCATION POLICIES

DRAM-Preferred: In the *DRAM-Preferred* allocation policy, all new pages are allocated from the DRAM tier as long as it has free space. Once the DRAM is filled, new pages are allocated from the PM tier. Figure 2(a) illustrates this allocation policy.

DRAM-Always: In *DRAM-Always* allocation, all new pages are allocated from the DRAM, even if the DRAM is already filled. As Figure 2(b) shows, when DRAM is filled, new page allocations force the demotion of cold pages to the PM tier.

PM-Preferred: In the *PM-Preferred* allocation policy, new pages are allocated from the PM tier as long as the PM has free space. Once it is filled, new pages are then allocated from the DRAM tier. Figure 2(c) shows how page allocations get made with the *PM-Preferred* policy.

PM-Always: The final allocation policy, *PM-Always*, always allocates new pages from the PM tier as shown in Figure 2(d). With *PM-Always*, in the Linux kernel, we observe that allocating new pages from the swap space while DRAM space is free often causes an Out Of Memory (OOM) error which kills the running application. Hence, in the rest of the paper, we focus our experiments on the remaining three allocation policies.

3 EXPERIMENT SETUP

All experiments are performed using an Intel Xeon Gold 5218 dual-socket processor with 16 cores per socket, i.e., 32 cores in total. The system is configured with twelve DDR4 DIMMs, totaling 192GB of DRAM, and 4 Intel Optane DC Persistent Memory (DCPM) DIMMs totaling 512GB of PM. We implemented the allocation policies in Linux kernel version 5.3.1. We used SPEC [?], NAS [?], YCSB [?

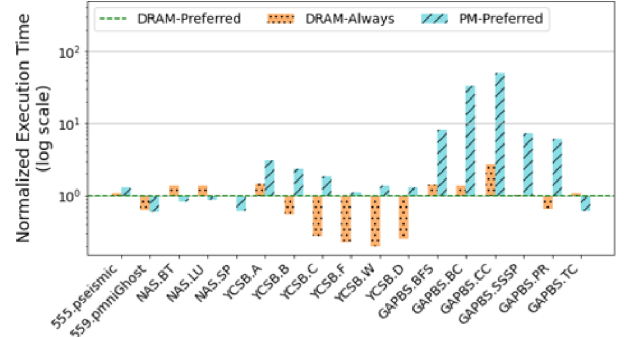


Figure 3: Impact of allocation policy on the dynamic tiered memory system.

], and GAPBS [?] benchmarks to evaluate the performance of the allocation policies.

4 IMPACT OF ALLOCATION POLICIES

Figure 3 shows the large variation in performance that we can obtain by changing the allocation policy for various workloads. For example, in Fig 3, the performance of the dynamic tiering can significantly improve if the default allocation policy of *DRAM-Preferred* is replaced by the *DRAM-Always* allocation policy for the YCSB.D workload. In YCSB workload D, new items are added, and the most recent items are the most popular items [?]. With *DRAM-Preferred*, since new pages are allocated from the PM once the DRAM is filled, the newer popular pages would get accessed from the PM. On the other hand, with *DRAM-Always* allocation policy, newer pages containing popular items are allocated and accessed from the DRAM. Thus, *the allocation policies can significantly impact the performance of the tiered memory systems*. Furthermore, we observed that *no single allocation policy always performs best for different types of workloads*.

5 CONCLUSION

In this work, we investigated the problem of page allocation in tiered memory systems which was previously unexplored. We introduced several allocation policies and evaluated the performance of these allocation policies in dynamic tiered memory systems by using a variety of workloads. We observed that allocation policies can significantly impact the performance of tiered memory systems.