

Though large-scale cluster systems remain the dominant solution for request and data-level parallelism [14], there have been a resurgence towards applying HPC techniques (e.g., DSM) for more efficient heterogeneous computation with more tightly-coupled heterogeneous nodes providing (hardware) acceleration for one another [7] [ADD MORE CITATIONS]. Orthogonally, within the scope of one motherboard, *heterogeneous memory management (HMM)* enables the use of OS-controlled, unified memory view across both main memory and device memory [13], all while using the same libc function calls as one would with SMP programming, the underlying complexities of memory ownership and data placement automatically managed by the OS kernel. On the other hand, while HMM promises a distributed shared memory approach towards exposing CPU and peripheral memory, applications (drivers and front-ends) that exploit HMM to provide ergonomic programming models remain fragmented and narrowly-focused. Existing efforts in exploiting HMM in Linux predominantly focus on exposing global address space abstraction to GPU memory – a largely non-coordinated effort surrounding both *in-tree* and proprietary code [10, 1]. Limited effort have been done on incorporating HMM into other variants of accelerators in various system topologies.

Orthogonally, allocation of hardware accelerator resources in a cluster computing environment becomes difficult when the required hardware accelerator resources of one workload cannot be easily determined and/or isolated as a “stage” of computation. Within a cluster system there may exist a large amount of general-purpose worker nodes and limited amount of hardware-accelerated nodes. Further, it is possible that every workload performed on this cluster asks for hardware acceleration from time to time, but never for a relatively long time. Many job scheduling mechanisms within a cluster *move data near computation* by migrating the entire job/container between general-purpose and accelerator nodes [27, 25]. This way of migration naturally incurs large overhead – accelerator nodes which strictly perform computation on data in memory without ever needing to touch the container’s filesystem should not have to install the entire filesystem locally, for starters. Moreover, must *all* computations be performed near data? [23], for example, shows that RDMA over fast network interfaces ($25\text{ Gbps} \times 8$), when compared to node-local setups, result in negligible impact on tail latencies but high impact on throughput when bandwidth is maximized.

This thesis paper builds upon an ongoing research effort in implementing a tightly coupled cluster where HMM abstractions allow for transparent RDMA access from accelerator nodes to local data and migration of data near computation, leveraging different consistency model and coherency protocols to amortize the communication cost for shared data. **Specifically, this thesis explores the effect of memory consistency model and coherency protocol on memory-sharing between cluster nodes**

The rest of the chapter is structured as follows...

1 Experiences from Software DSM

A majority of contributions to software DSM systems come from the 1990s [6, 9, 16, 15]. These developments follow from the success of the Stanford DASH project in the late 1980s – a hardware distributed shared memory (specifically NUMA) implementation of a multiprocessor that first proposed the *directory-based protocol* for cache coherence, which stores the ownership information of cache lines to reduce unnecessary communication that prevented previous multiprocessors from scaling out [20].

While developments in hardware DSM materialized into a universal approach to cache-coherence in contemporary many-core processors (e.g., *Ampere Altra*[2]), software DSMs in clustered computing languished in favor of loosely-coupled nodes performing data-parallel computation, communicating via message-passing. Bandwidth limitations with the network interfaces of the late 1990s was insufficient to support the high traffic incurred by DSM and its programming model [31, 21].

New developments in hardware network interfaces provides much improved bandwidth and latency compared to ethernet in the 1990s. RDMA-capable NICs have been shown to improve the training efficiency sixfold compared to distributed *TensorFlow* via RPC, scaling positively over non-distributed training [18]. Similar results have been observed for *APACHE Spark*[22] and *what?*. Consequently, there have been a resurgence of interest in software DSM systems and programming models [24, 8].

1.1 Munin: Multi-Consistency Protocol

Munin[9] is one of the older developments in software DSM systems. The authors of Munin identify that *false-sharing*, occurring due to multiple processors writing to different offsets of the same page triggering invalidations, is strongly detrimental to the performance of shared-memory systems. To combat this, Munin exposes annotations as part of its programming model to facilitate multiple consistency protocols on top of release consistency. An immutable shared memory object across readers, for example, can be safely copied without concern for coherence between processors. On the other hand, the *write-shared* annotation explicates that a memory object is written by multiple processors without synchronization – i.e., the programmer guarantees that only false-sharing occurs within this granularity. Annotations such as these explicitly disables subsets of consistency procedures to reduce communication in the network fabric, thereby improving the performance of the DSM system.

Perhaps most importantly, experiences from Munin show that *restricting the flexibility of programming model can lead to more performant coherence models*, as *corroborated* by the now-foundational *Resilient Distributed Database* paper [33] – which powered many now-popular scalable data processing frameworks such as *Hadoop MapReduce*[3] and *APACHE Spark*[4]. “To achieve fault tolerance efficiently, RDDs provide a restricted form of shared memory [based on]...transformations rather than...updates to shared state” [33]. This al-

lows for the use of transformation logs to cheaply synchronize states between unshared address spaces – a much desired property for highly scalable, loosely-coupled clustered systems.

1.2 Treadmarks: Multi-Writer Protocol

Treadmarks[6] is a software DSM system developed in 1996, which featured an intricate *interval*-based multi-writer protocol that allows multiple nodes to write to the same page without false-sharing. The system follows a release-consistent memory model, which requires the use of either locks (via **acquire**, **release**) or barriers (via **barrier**) to synchronize. Each *interval* represents a time period in-between page creation, **release** to another processor, or a **barrier**; they also each correspond to a *write notice*, which are used for page invalidation. Each **acquire** message is sent to the statically-assigned lock-manager node, which forwards the message to the last releaser. The last releaser computes the outstanding write notices and piggy-backs them back for the acquirer to invalidate its own cached page entry, thus signifying entry into the critical section. Consistency information, including write notices, intervals, and page diffs, are routinely garbage-collected which forces cached pages in each node to become validated.

Compared to *Treadmarks*, the system described in this paper uses a single-writer protocol, thus eliminating the concept of “intervals” – with regards to synchronization, each page can be either in-sync (in which case they can be safely shared) or out-of-sync (in which case they must be invalidated/updated). This comes with the following advantage:

- Less metadata for consistency-keeping.
- More adherent to the CPU-accelerator dichotomy model.
- Much simpler coherence protocol, which reduces communication cost.

In view of the (still) disparate throughput and latency differences between local and remote memory access [8], the simpler coherence protocol of single-writer protocol should provide better performance on the critical paths of remote memory access.

1.3 Hotpot: Single-Writer & Data Replication

Newer works such as *Hotpot*[28] apply distributed shared memory techniques on persistent memory to provide “transparent memory accesses, data persistence, data reliability, and high availability”. Leveraging on persistent memory devices allow DSM applications to bypass checkpoints to block device storage [28], ensuring both distributed cache coherence and data reliability at the same time [28].

We specifically discuss the single-writer portion of its coherence protocol. The data reliability guarantees proposed by the *Hotpot* system requires each

shared page to be replicated to some *degree of replication*. Nodes who always store latest replication of shared pages are referred to as “owner nodes”, which arbitrate other nodes to store more replications in order to reach the degree of replication quota. At acquisition time, the acquiring node asks the access-management node for single-writer access to shared page, who grants it if no other critical section exists, alongside list of current owner nodes. At release time, the releaser first commits its changes to all owner nodes which, in turn, commits its received changes across lesser sharers to achieve the required degree of replication. These two operations are all acknowledged back in reverse order. Once all acknowledgements are received from owner nodes by commit node, the releaser tells them to delete their commit logs and, finally, tells the manager node to exit critical section.

The required degree of replication and logged commit transaction until explicit deletion facilitate crash recovery at the expense of worse performance over release-time I/O. While the study of crash recovery with respect to shared memory systems is out of the scope of this thesis, this paper provides a good framework for a **correct** coherence protocol for a single-writer, multiple-reader shared memory system, particularly when the protocol needs to cater for a great variety of nodes each with their own memory preferences (e.g., write-update vs. write-invalidate, prefetching, etc.).

1.4 MENPS: A Return to DSM

MENPS[12] leverages new RDMA-capable interconnects as a proof-of-concept that DSM systems and programming models can be as efficient as *partitioned global address space* (PGAS) using today’s network interfaces. It builds upon *TreadMark*’s [6] coherence protocol and crucially alters it to a *floating home-based* protocol, based on the insight that diff-transfers across the network is comparatively costly compared to RDMA intrinsics – which implies preference towards local diff-merging. The home node then acts as the data supplier for every shared page within the system.

Compared to PGAS frameworks (e.g., MPI), experimentation over a subset of *NAS Parallel Benchmarks* shows that MENPS can obtain comparable speedup in some of the computation tasks, while achieving much better productivity due to DSM’s support for transparent caching, etc. [12]. These results back up their claim that DSM systems are at least as viable as traditional PGAS/message-passing frameworks for scientific computing, also corroborated by the resurgence of DSM studies later on[23].

2 PGAS and Message Passing

While the feasibility of transparent DSM systems over multiple machines on the network has been made apparent since the 1980s, predominant approaches to “scaling-out” programs over the network relies on the message-passing approach [30]. The reasons are twofold:

1. Programmers would rather resort to more intricate, more predictable approaches to scaling-out programs over the network [30]. This implies manual/controlled data sharding over nodes, separation of compute and communication “stages” of computation, etc., which benefit performance analysis.
2. Enterprise applications value throughput and uptime of relatively computationally inexpensive tasks/resources [14], which requires easy scalability of tried-and-true, latency-inexpensive applications. Studies in transparent DSM systems mostly require exotic, specifically-written programs to exploit global address space, which is fundamentally at odds in terms of reusability and flexibility required.

2.1 PGAS

Partitioned Global Address Space (PGAS) is a parallel programming model that (1) exposes a global address space to all machines within a network and (2) explicates distinction between local and remote memory [11]. Oftentimes, message-passing frameworks, for example *OpenMPI*, *OpenFabrics*, and *UCX*, are used as backends to provide the PGAS model over various network interfaces/platforms (e.g., Ethernet and Infiniband)[29, 26].

Examples of PGAS programming languages and models include Notably, implementation of a *global* address space across machines on top of machines already equipped with their own *local* address space (e.g., cluster nodes running commercial Linux) necessitates a global addressing mechanism for shared/shared data objects. DART[34], for example, utilizes a 128-bit “global pointer” to encode global memory object/segment ID and access flags in the upper 64 bits and virtual addresses in the lower 64 bits for each (slice of) memory object allocated within the PGAS model. A *non-collective* PGAS object is allocated entirely local to the allocating node’s memory, but registered globally. Consequently, a single global pointer is recorded in the runtime with corresponding permission flags for the context of some user-defined group of associated nodes. Comparatively, a *collective* PGAS object is allocated such that a partition of the object (i.e., a subarray of the repr) is stored in each of the associated node – for a k -partitioned object, k global pointers are recorded in the runtime each pointing to the same object, with different offsets and (naturally) independently-chosen virtual addresses. Note that this design naturally requires virtual addresses within each node to be *pinned* – the allocated object cannot be re-addressed to a different virtual address i.e., the global pointer that records the local virtual address cannot be auto-invalidated.

Similar schemes can be observed in other PGAS backends/runtimes, albeit they may opt to use a map-like data structure for addressing instead. In general, PGAS backends differ from DSM systems in that, despite providing memory management over remote nodes, they provide no transparent caching and transfer of remote memory objects accessed by local nodes. The programmer is still expected to handle data/thread movement manually when working with

shared memory over network to maximize performance metrics of interest.

...
Improvement in NIC bandwidth and transfer rate benefits DSM applications that expose global address space, and those that leverage single-writer capabilities over hierarchical memory nodes. **[GAS and PGAS (Partitioned GAS) technologies for example Openshmem, OpenMPI, Cray Chapel, etc. that leverage specially-linked memory sections and /dev/shm to abstract away RDMA access].**

Contemporary works on DSM systems focus more on leveraging hardware advancements to provide fast and/or seamless software support. Adrias [23], for example, implements a complex system for memory disaggregation over multiple compute nodes connected via the *ThymesisFlow*-based RDMA fabric, where they observed significant performance improvements over existing data-intensive processing frameworks, for example APACHE Spark, Memcached, and Redis, over no-disaggregation (i.e., using node-local memory only, similar to cluster computing) systems.

2.2 Programming Model

2.3 Move Data to Process, or Move Process to Data?

(TBD – The former is costly for data-intensive computation, but the latter may be impossible for certain tasks, and greatly hardens the replacement problem.)

3 Replacement Policy

In general, three variants of replacement strategies have been proposed for either generic cache block replacement problems, or specific use-cases where contextual factors can facilitate more efficient cache resource allocation:

- General-Purpose Replacement Algorithms, for example LRU.
- Cost-Model Analysis
- Probabilistic and Learned Algorithms

3.1 General-Purpose Replacement Algorithms

Practically speaking, in the general case of the cache replacement problem, we desire to predict the re-reference interval of a cache block [17]. This follows from the Belady’s algorithm – the optimal case for the *ideal* replacement problem occurs when, at eviction time, the entry with the highest re-reference interval is replaced. Under this framework, therefore, the commonly-used LRU algorithm could be seen as a heuristic where the re-reference interval for each entry is predicted to be immediate. Fortunately, memory access traces of real computer systems agree with this tendency due to spatial locality [source]. (Real systems are complex, however, and there are other behaviors...) On the other hand, the

hypothetical LFU algorithm is a heuristic that captures frequency. [...] While the textbook LFU algorithm suffers from needing to maintain a priority-queue for frequency analysis, it was nevertheless useful for keeping recurrent (though non-recent) blocks from being evicted from the cache [\[source\]](#).

Derivatives from the LRU algorithm attempts to balance between frequency and recency. [\[Talk about LRU-K, LRU-2Q, LRU-MQ, LIRS, ARC here ...\]](#)

Advancements in parallel/concurrent systems had led to a rediscovery of the benefits of using FIFO-derived replacement policies over their LRU/LFU counterparts, as book-keeping operations on the uniform LRU/LFU state proves to be (1) difficult for synchronization and, relatedly, (2) cache-unfriendly [32]. [\[Talk about FIFO, FIFO-CLOCK, FIFO-CAR, FIFO-QuickDemotion, and Dueling CLOCK here ...\]](#)

Finally, real-life experiences have shown the need to reduce CPU time in practical applications, owing from one simple observation – during the fetch-execution cycle, all processors perform blocking I/O on the memory. A cache-unfriendly design, despite its hypothetical optimality, could nevertheless degrade the performance of a system during low-memory situations. In fact, this proves to be the driving motivation behind Linux’s transition away from the old LRU-2Q page replacement algorithm into the more coarse-grained Multi-generation LRU algorithm, which has been mainlined since v6.1.

3.2 Cost-Model Analysis

The ideal case for the replacement problem fails to account for invalidation of cache entries. It also assumes for a uniform, dual-hierarchical cache-store model that is insufficient to capture the heterogeneity of today’s massively-parallel, distributed systems. High-speed network interfaces are capable of exposing RDMA interfaces between computer nodes, which amount to almost twice as fast RDMA transfer when compared to swapping over the kernel I/O stack, while software that bypass the kernel I/O stack is capable of stretching the bandwidth advantage even more (source). This creates an interesting network topology between RDMA-enabled nodes, where, in addition to swapping at low-memory situations, the node may opt to “swap” or simply drop the physical page in order to lessen the cost of page misses.

[\[Talk about GreedyDual, GDSF, BCL, Amortization\]](#)

Traditionally, replacement policies based on cost-model analysis were utilized in content-delivery networks, which had different consistency models compared to finer-grained systems. HTTP servers need not pertain to strong consistency models, as out-of-date information is considered permissible, and single-writer scenarios are common. Consequently, most replacement policies for static content servers, while making strong distinction towards network topology, fails to concern for the cases where an entry might become invalidated, let alone multi-writer protocols. One early paper [19] examines the efficacy of using page fault frequency as an indicator of preference towards working set inclusion (which I personally think is highly flawed – to be explained). Another paper [5] explores

the possibility of taking page fault into consideration for eviction, but fails to go beyond the obvious implication that pages that have been faulted *must* be evicted.

The concept of cost models for RDMA and NUMA systems are relatively underdeveloped, too. (Expand)

3.3 Probabilistic and Learned Algorithms for Cache Replacement

Finally, machine learning techniques and low-cost probabilistic approaches have been applied on the ideal cache replacement problem with some level of success. [Talk about LeCaR, CACHEUS here].

4 Cache Coherence and Consistency in DSM Systems

(I need to read more into this. Most of the contribution comes from CPU caches, less so for DSM systems.) [Talk about JIAJIA and Treadmark's coherence protocol.]

Consistency and communication protocols naturally affect the cost for each faulted memory access ...

[Talk about directory, transactional, scope, and library cache coherence, which allow for multi-casted communications at page fault but all with different levels of book-keeping.]

References

- [1] URL: <https://www.phoronix.com/search/Heterogeneous%20Memory%20Management>.
- [2] URL: https://uawartifacts.blob.core.windows.net/upload-files/Altra_Max_Rev_A1_DS_v1_15_20230809_b7cdce449e_424d129849.pdf.
- [3] URL: <https://hadoop.apache.org/>.
- [4] URL: <https://spark.apache.org/>.
- [5] J. Aguilar and E.L. Leiss. "A Coherence-Replacement Protocol For Web Proxy Cache Systems". In: *International Journal of Computers and Applications* 28.1 (2006), pp. 12–18. DOI: 10.1080/1206212X.2006.11441783. eprint: <https://doi.org/10.1080/1206212X.2006.11441783>. URL: <https://doi.org/10.1080/1206212X.2006.11441783>.
- [6] Cristiana Amza et al. "Treadmarks: Shared memory computing on networks of workstations". In: *Computer* 29.2 (1996), pp. 18–28.

- [7] Javier Cabezas et al. “GPU-SM: shared memory multi-GPU programming”. In: *Proceedings of the 8th Workshop on General Purpose Processing using GPUs*. 2015, pp. 13–24.
- [8] Qingchao Cai et al. “Efficient distributed memory management with RDMA and caching”. In: *Proceedings of the VLDB Endowment* 11.11 (2018), pp. 1604–1617.
- [9] John B Carter, John K Bennett, and Willy Zwaenepoel. “Implementation and performance of Munin”. In: *ACM SIGOPS Operating Systems Review* 25.5 (1991), pp. 152–164.
- [10] Jonathan Corbet. *Heterogeneous memory management meets EXPORT_SYMBOL_GPL()*. 2018. URL: <https://lwn.net/Articles/757124/>.
- [11] Mattias De Wael et al. “Partitioned global address space languages”. In: *ACM Computing Surveys (CSUR)* 47.4 (2015), pp. 1–27.
- [12] Wataru Endo, Shigeyuki Sato, and Kenjiro Taura. “MENPS: a decentralized distributed shared memory exploiting RDMA”. In: *2020 IEEE/ACM Fourth Annual Workshop on Emerging Parallel and Distributed Runtime Systems and Middleware (IPDRM)*. IEEE. 2020, pp. 9–16.
- [13] Mark Harris. *Unified memory for cuda beginners*. 2017. URL: <https://developer.nvidia.com/blog/unified-memory-cuda-beginners/>.
- [14] John L Hennessy and David A Patterson. *Computer architecture: a quantitative approach*. Elsevier, 2011.
- [15] Weiwu Hu, Weisong Shi, and Zhimin Tang. “JIAJIA: A software DSM system based on a new cache coherence protocol”. In: *High-Performance Computing and Networking: 7th International Conference, HPCN Europe 1999 Amsterdam, The Netherlands, April 12–14, 1999 Proceedings* 7. Springer. 1999, pp. 461–472.
- [16] Ayal Itzkovitz, Assaf Schuster, and Lea Shalev. “Thread migration and its applications in distributed shared memory systems”. In: *Journal of Systems and Software* 42.1 (1998), pp. 71–87.
- [17] Aamer Jaleel et al. “High performance cache replacement using re-reference interval prediction (RRIP)”. In: *ACM SIGARCH computer architecture news* 38.3 (2010), pp. 60–71.
- [18] Chengfan Jia et al. “Improving the performance of distributed tensor-flow with RDMA”. In: *International Journal of Parallel Programming* 46 (2018), pp. 674–685.
- [19] Richard P. LaRowe and Carla Schlatter Ellis. “Page placement policies for NUMA multiprocessors”. In: *Journal of Parallel and Distributed Computing* 11.2 (1991), pp. 112–129. ISSN: 0743-7315. DOI: [https://doi.org/10.1016/0743-7315\(91\)90117-R](https://doi.org/10.1016/0743-7315(91)90117-R). URL: <https://www.sciencedirect.com/science/article/pii/074373159190117R>.
- [20] Daniel Lenoski et al. “The stanford dash multiprocessor”. In: *Computer* 25.3 (1992), pp. 63–79.

- [21] Honghui Lu et al. “Message passing versus distributed shared memory on networks of workstations”. In: *Supercomputing’95: Proceedings of the 1995 ACM/IEEE Conference on Supercomputing*. IEEE. 1995, pp. 37–37.
- [22] Xiaoyi Lu et al. “Accelerating spark with RDMA for big data processing: Early experiences”. In: *2014 IEEE 22nd Annual Symposium on High-Performance Interconnects*. IEEE. 2014, pp. 9–16.
- [23] Dimosthenis Masouros et al. “Adrias: Interference-Aware Memory Orchestration for Disaggregated Cloud Infrastructures”. In: *2023 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*. IEEE. 2023, pp. 855–869.
- [24] Jacob Nelson et al. “{Latency-Tolerant} software distributed shared memory”. In: *2015 USENIX Annual Technical Conference (USENIX ATC 15)*. 2015, pp. 291–305.
- [25] SeungYong Oh and JongWon Kim. “Stateful Container Migration employing Checkpoint-based Restoration for Orchestrated Container Clusters”. In: *2018 International Conference on Information and Communication Technology Convergence (ICTC)*. 2018, pp. 25–30. DOI: 10.1109/ICTC.2018.8539562.
- [26] *Platform-Specific Notes*. 2023. URL: <https://chapel-lang.org/docs/platforms/index.html#>.
- [27] Manuel Rodríguez-Pascual et al. “Job migration in hpc clusters by means of checkpoint/restart”. In: *The Journal of Supercomputing* 75 (2019), pp. 6517–6541.
- [28] Yizhou Shan, Shin-Yeh Tsai, and Yiyang Zhang. “Distributed Shared Persistent Memory”. In: *Proceedings of the 2017 Symposium on Cloud Computing*. SoCC ’17. Santa Clara, California: Association for Computing Machinery, 2017, pp. 323–337. ISBN: 9781450350280. DOI: 10.1145/3127479.3128610. URL: <https://doi.org/10.1145/3127479.3128610>.
- [29] *upcc.1*. 2022. URL: <https://upc.lbl.gov/docs/user/upcc.html>.
- [30] Maarten Van Steen and Andrew S Tanenbaum. *Distributed systems*. Maarten van Steen Leiden, The Netherlands, 2017.
- [31] Paul Werstein, Mark Pethick, and Zhiyi Huang. “A performance comparison of dsm, pvm, and mpi”. In: *Proceedings of the Fourth International Conference on Parallel and Distributed Computing, Applications and Technologies*. IEEE. 2003, pp. 476–482.
- [32] Juncheng Yang et al. “FIFO can be Better than LRU: the Power of Lazy Promotion and Quick Demotion”. In: *Proceedings of the 19th Workshop on Hot Topics in Operating Systems*. 2023, pp. 70–79.

- [33] Matei Zaharia et al. “Resilient Distributed Datasets: A Fault-Tolerant Abstraction for In-Memory Cluster Computing”. In: *9th USENIX Symposium on Networked Systems Design and Implementation (NSDI 12)*. San Jose, CA: USENIX Association, Apr. 2012, pp. 15–28. ISBN: 978-931971-92-8. URL: <https://www.usenix.org/conference/nsdi12/technical-sessions/presentation/zaharia>.
- [34] Huan Zhou et al. “DART-MPI: An MPI-based implementation of a PGAS runtime system”. In: *Proceedings of the 8th International Conference on Partitioned Global Address Space Programming Models*. 2014, pp. 1–11.