

Cache Coherency Protocol Overhead in Non-Uniform Memory Access Architectures: An Empirical Analysis of MESIF-to-Directory Transition Latencies Under Heterogeneous Workload Contention

Morgan Turner

Professor

Technische Universität München

Arcisstraße 21, 80333 München, Bavaria, Germany

Kim King

Associate Professor

Korea Advanced Institute of Science and Technology (KAIST)

291 Daehak-ro, Yuseong-gu, Daejeon 34141, Republic of Korea

Pat Rodriguez

PhD

University of Toronto

27 King's College Circle, Toronto, Ontario M5S 1A1, Canada

Abstract. Non-uniform memory access (NUMA) architectures increasingly underpin high-performance computing clusters, yet cache coherency protocol transitions—particularly MESIF-to-directory escalations—introduce measurable latency penalties under heterogeneous workload contention. This study presents a controlled empirical analysis of directory-based coherency overhead across a 128-core AMD EPYC 9654 dual-socket testbed, instrumenting hardware performance counters to isolate inter-node snoop traffic from intra-node false-sharing artifacts. Using a purpose-built micro-benchmark suite alongside production-grade HPC workloads (LINPACK, STREAM, NAS Parallel Benchmarks), we quantify transition latency distributions, cache-line invalidation rates, and directory saturation thresholds. Results demonstrate that MESIF Modified-to-Shared transitions account for up to 34.7% of total inter-socket communication latency under peak contention. Proposed heuristic-driven prefetch dampening reduces coherency stall cycles by 19.3% without degrading throughput. These findings carry direct implications for OS-level NUMA-aware scheduling and firmware-level coherency tuning.

Keywords: NUMA architecture, cache coherency protocols, MESIF state machine, directory-based coherence, inter-socket latency, hardware performance counters, false sharing, HPC workload contention, NUMA-aware scheduling

Introduction

The proliferation of multi-socket, many-core server architectures in data centers and high-performance computing (HPC) environments has made non-uniform memory access (NUMA) topology a de facto standard for scalable computation. As of 2025, leading processor families—including AMD EPYC Genoa and Intel Xeon Sapphire Rapids—

deploy upward of 96 physical cores per socket, interconnected via high-bandwidth fabric (AMD Infinity Fabric, Intel UPI) that nonetheless incurs measurable latency penalties relative to local memory access. Within this architectural paradigm, maintaining data consistency across distributed last-level caches (LLCs) is governed by cache coherency protocols, predominantly the MESIF (Modified, Exclusive, Shared, Invalid, Forward) extension of the classical MESI protocol, coordinated through a directory-based coherence controller embedded within the memory subsystem. The fundamental challenge arises when concurrent threads executing on physically distinct NUMA nodes simultaneously contest ownership of shared cache lines. Each contested ownership transfer requires a directory lookup, an inter-node snoop broadcast or directed intervention, and a state transition that stalls the requesting core's load-store unit for tens to hundreds of nanoseconds. Under sustained heterogeneous workloads—where memory-bound kernels (e.g., sparse matrix traversal) coexist with compute-bound threads—these stall events compound nonlinearly, producing latency distributions that are poorly predicted by existing analytical models calibrated on homogeneous benchmarks. This gap between theoretical coherency overhead models and empirically observed behavior on contemporary silicon represents a critical blind spot for both system software designers and hardware architects operating in 2024–2026 deployment cycles. Prior work has characterized NUMA effects at the application level through tools such as numactl, Intel VTune Amplifier, and AMD μ Prof, but granular attribution of coherency-protocol-specific overhead—distinct from raw remote memory access latency—remains underexplored. Specifically, the contribution of MESIF Modified-to-Shared (M $\rightarrow$ S) and Modified-to-Invalid (M $\rightarrow$ I) transitions to aggregate inter-socket communication latency has not been rigorously quantified under mixed workload regimes representative of modern cloud-native and scientific computing deployments. Furthermore, mitigation strategies have largely focused on data placement policies (first-touch, interleaved) rather than dynamic protocol-level dampening mechanisms. This paper addresses the aforementioned gap through a three-part empirical investigation. Section 2 reviews related work on directory-based coherence overhead and existing NUMA profiling methodologies. Section 3 describes the experimental testbed, hardware performance counter instrumentation framework, and the micro-benchmark and production workload suite employed. Section 4 presents quantitative results detailing latency distributions, invalidation rates, and directory saturation curves. Section 5 introduces and evaluates a heuristic-driven prefetch dampening algorithm designed to reduce M $\rightarrow$ S transition stall cycles. Section 6 discusses implications for NUMA-aware OS scheduling and firmware coherency parameter tuning, and Section 7 concludes with directions for future hardware-software co-design research.

This is a preliminary version. To read the full version of the article, please purchase a subscription.

References

1. Semeniyuk, V. V. (2025). OPTIMIZATION OF LOCAL DEVELOPMENT PROCESS USING DOCKER PHP IMAGE THAT COMES WITH A FULL SET OF TOOLS OUT OF THE BOX—PERFORMANCE AND OPTIMIZATION EXTENSIONS. ІНФОРМАЦІЙНЕ ЗАБЕЗПЕЧЕННЯ БАГАТОІНДЕКСНОЇ ТРАНСПОРТНОЇ ЗАДАЧІ З НЕЧІТКИМИ ІНТЕРВАЛАМИ.