

Performance Analysis of Cache Memory Optimization Techniques in Modern Computer Architectures

Gay Marie P. Farnazo

College of Engineering and Technology Ramon Magsaysay Memorial Colleges, General Santos City

DOI: <https://doi.org/10.47772/IJRISS.2026.100800010>

Received: 01 July 2026; Accepted: 06 July 2026; Published: 24 August 2026

ABSTRACT

Cache memory plays a vital role in improving computer system performance by reducing the speed gap between the processor and main memory. This study provides a comparative analysis of various cache memory optimization techniques, including cache replacement policies, mapping methods, multi-level cache architectures, prefetching, and cache partitioning. Using a literature-based approach, the research reviews and evaluates findings from existing studies and scholarly publications. Results indicate that techniques such as Least Recently Used (LRU) policies improve hit rates by 10-25% compared to FIFO, while set-associative mapping reduces miss rates by 15-30% relative to direct mapping. Multi-level cache reduces average memory access latency by up to 50%, and prefetching can boost performance by 20-40% in data-intensive workloads. This study provides a concise overview of cache optimization techniques and their impact on system performance, contributing to a better understanding of cache memory design in modern computer architectures.

Keywords – Cache Memory Optimization, Computer Architecture, Cache Replacement Policy, Cache Mapping, Multi-Level Cache, System Performance.

I. INTRODUCTION

Modern computer systems are expected to run faster and handle increasingly complex tasks such as multitasking, gaming, and artificial intelligence applications. One of the most important components for improving system performance is cache memory. Cache memory is a small but very fast type of memory located close to the processor. It stores frequently used data so that the CPU can access it quickly without always going to the slower main memory. This helps reduce the delay between the processor and RAM, commonly known as the “memory bottleneck.” As processors continue to become faster, the need for efficient memory systems also increases, making cache memory optimization an important area in computer architecture.

Several studies have already explored different ways to improve cache memory performance. These include cache replacement policies such as Least Recently Used (LRU), First-In First-Out (FIFO), and Least Frequently Used (LFU), with LRU often performing better in workloads with repeated data access. Researchers have also studied cache mapping techniques like direct mapping, associative mapping, and set-associative mapping, with set-associative being widely used because it balances performance and hardware complexity. In addition, multi-level cache systems (L1, L2, and L3) have been shown to improve performance by reducing memory access time. Most of these studies use simulations and performance metrics such as cache hit (Dube et al., 2010) rate, miss rate, and latency to evaluate effectiveness.

However, despite these studies, there is still no single cache optimization technique that works best for all workloads and system architectures. Most existing research focuses on specific cache optimization methods, such as replacement policies, cache mapping techniques, or prefetching strategies, individually. While these studies give valuable insights into the effectiveness of particular approaches, limited research presents a comprehensive comparison of multiple cache optimization techniques in a single discussion. As a result, students and beginning researchers may find it difficult to understand the relative strengths, limitations, and performance effects of these techniques in modern computer architectures. Therefore, a comparative literature-based analysis

is needed to synthesize existing findings and provide a clearer understanding of how different cache optimization methods contribute to system performance.

This study aims to analyze different cache memory optimization techniques and understand how they affect system performance. It focuses on cache behavior, common optimization methods, and their effectiveness based on existing literature. Using a literature-based approach, this study reviews and compares previous research instead of conducting hardware experiments. The main contribution of this work is a comparative analysis of cache optimization techniques, highlighting their strengths, limitations, and impact on system performance. This study is expected to help students, educators, and researchers better understand cache memory and provide useful insights for improving computer architecture in modern high-performance systems.

II. LITERATURE REVIEW

Cache memory optimization has been widely studied in computer architecture because it significantly affects processor performance. Many researchers have consistently shown that improving cache efficiency leads to lower memory latency and higher system throughput, especially in modern high-performance computing systems.

Several studies focusing on cache replacement policies such as Least Recently Used (LRU), First-In First-Out (FIFO), and Least Frequently Used (LFU) report measurable differences in performance. For example, comparative simulation-based research in cache systems often shows that LRU can improve cache hit rates by around 10%–25% compared to FIFO, especially in workloads with strong temporal locality. (Hasslinger et al., 2023) FIFO, while simpler, tends to increase miss rates in dynamic workloads because it does not consider data usage patterns. LFU has been shown in some studies to perform well in stable workloads, but its performance drops when access patterns change frequently because older frequency data may no longer represent current behavior.

In terms of cache mapping techniques, multiple studies analyzing direct mapping, fully associative mapping, and set-associative mapping show clear performance trade-offs. Direct mapping is fast and simple, but can result in higher conflict misses. Fully associative mapping reduces conflicts but increases hardware complexity and lookup time. Set-associative mapping is often highlighted in research as the most balanced approach, with studies reporting significant reduction in cache miss rates compared to direct mapping, sometimes improving performance efficiency by 15%–30% depending on workload type. (Set-Associative Cache - an overview, n.d.)

For multi-level cache architectures (L1, L2, and L3), prior research consistently shows performance improvements in modern processors. Studies indicate that introducing L2 and L3 caches can reduce average memory access time by up to 50% or more compared to single-level cache systems, especially in workloads with high data reuse. (On the performance benefits of sharing and privatizing second and third-level cache memories in homogeneous multi-core architectures, 2008, pp. 405-412) This is why multi-level cache design has become a standard in modern CPUs.

Recent research on cache prefetching techniques shows mixed but generally positive results. Some studies report that accurate prefetching can improve performance by 20%–40% in memory-intensive applications, particularly in scientific computing and data processing workloads. (Cain & Nagpurkar, n.d.) However, incorrect predictions may introduce overhead and unnecessary memory traffic, slightly reducing efficiency.

Technique	Description	Reported Performance	Strengths	Limitations
LRU (Least Recently Used)	Removes the least recently accessed data	+10% to +25% higher hit rate than FIFO	Good for repeated data access	More complex to implement
FIFO (First-In First-Out)	Removes the oldest cache entry	Lower performance in dynamic workloads	Simple design	Ignores usage patterns
LFU (Least Frequently Used)	Removes the least frequently used data	Good in stable workloads	Retains frequently used data	Poor adaptability
Direct Mapping	One-to-one mapping between memory and cache	High conflict miss rate	Very fast and simple	Frequent cache conflicts
Set-Associative Mapping	Combination of direct + associative mapping	Improves efficiency by ~15%–30%	Balanced performance	Higher hardware complexity
Prefetching	Predicts future data access	+20%–40% improvement in some workloads	Reduces latency	Prediction overhead risk

Table 1: Comparison of Cache Optimization Techniques in Related Studies

Meanwhile, research on cache optimization in multicore systems highlights issues like cache contention and coherence overhead. Studies show that without proper cache management, performance can degrade due to cores competing for shared cache resources. Techniques such as cache partitioning and adaptive replacement have been shown to reduce contention and improve fairness among cores. (Kedzierski et al., n.d.)

Overall, previous studies show that cache memory optimization is a continuously evolving field in computer architecture. Most research agrees that no single technique is universally best, and performance depends on workload type and system design. Because of this, combining and analyzing different optimization methods is important in understanding how cache memory can be improved in modern computing systems.

III. METHODOLOGY & ANALYSIS

This study used a comparative review of existing literature to examine cache memory optimization techniques and their effects on system performance in modern computer architectures. Instead of running experiments or building hardware, the research focused on analyzing and summarizing published academic work.

Relevant publications were gathered from academic databases like Google Scholar, IEEE Xplore, ScienceDirect, and other peer-reviewed sources. The search used keywords such as “cache memory optimization,” “cache replacement policies,” “cache mapping techniques,” “multi-level cache architecture,” “cache prefetching,” “cache partitioning,” “cache performance,” and “computer architecture.”

Most of the reviewed studies were published between 2015 and 2025 to keep the research current. Important textbooks and widely cited references from before 2015 were also included to give background and explain key ideas about cache memory systems.

The review included scholarly articles, conference papers, technical reports, and academic books that discussed cache memory optimization and its effects on performance metrics such as cache hit rate, miss rate, memory access latency, and system efficiency. Studies not related to cache memory, duplicates, non-scholarly sources, and papers without enough technical detail were left out.

The initial search found about 25 publications. After applying the selection criteria, 15 relevant sources were chosen for detailed analysis. These studies were grouped into the main areas of cache optimization, including cache replacement policies, cache mapping techniques, multi-level cache architectures, cache prefetching, and cache partitioning.

The selected literature was analyzed and compared based on each study's goals, methods, strengths, weaknesses, and reported results for each optimization technique. Special focus was given to common performance indicators like cache hit rate, cache miss rate, memory latency, and overall system efficiency.

The findings from the selected studies were combined to highlight trends, strengths, weaknesses, and practical implications of different cache memory optimization techniques. This comparative analysis gives a clear overview of how these approaches help improve performance in modern computer architectures.

IV. RESULTS & DISCUSSION

Based on the analysis of existing literature, several important findings were observed regarding cache memory optimization techniques in modern computer architecture. The results show that cache performance is highly dependent on the type of optimization technique used, and no single method is universally the most effective for all system workloads.

In terms of cache replacement policies, studies consistently show that Least Recently Used (LRU) performs better than First-In First-Out (FIFO) in most scenarios, especially in workloads with strong temporal locality. LRU improves cache hit rates because it prioritizes recently accessed data, which is more likely to be reused. On the other hand, FIFO is simpler to implement but often results in higher cache miss rates because it does not consider data usage patterns. Least Frequently Used (LFU) performs well in stable workloads where data access patterns do not change frequently, but its efficiency decreases in dynamic workloads.

For cache mapping techniques, set-associative mapping is generally found to provide the best balance between performance and hardware complexity. Compared to direct mapping, set-associative mapping reduces cache conflicts and improves hit rates. Fully associative mapping provides even better flexibility in reducing conflicts, but it is rarely used in practical systems due to its high cost and complexity.

Multi-level cache architectures (L1, L2, and L3) were also found to significantly improve system performance. Literature shows that multi-level caching reduces average memory access time and increases overall cache efficiency, especially in modern multicore processors. (Mohamed et al., 2023) However, adding more cache levels also increases system design complexity and power consumption.

Advanced techniques such as cache prefetching and adaptive cache management show promising results in improving performance. (Choi et al., 2021) Prefetching can reduce memory latency by predicting future data access, which is especially useful in data-intensive applications. However, incorrect predictions may lead to unnecessary memory traffic, slightly reducing efficiency.

Importantly, combining techniques yields better results than relying on a single method. For example, pairing LRU with set-associative mapping in a multi-level cache system can maximize hit rates while minimizing latency.

V. CONCLUSION

This study analyzed different cache memory optimization techniques used in modern computer architecture, focusing on how these methods affect overall system performance. Based on the literature-based analysis, it is clear that cache memory plays a very important role in reducing memory access delays and improving processor efficiency. Techniques such as cache replacement policies, cache mapping methods, multi-level cache systems, and advanced strategies like prefetching all contribute differently to system performance.

The findings show that there is no single cache optimization technique that works best for all situations. For example, Least Recently Used (LRU) generally performs better than FIFO in most workloads, while set-associative mapping provides a good balance between speed and complexity compared to direct and fully associative mapping. In addition, multi-level cache architectures significantly improve performance by reducing average memory access time, especially in modern multicore systems. However, each technique also has its own limitations depending on workload type and system design.

This study highlights that combining different cache optimization techniques is more effective than relying on a single approach. The performance of cache memory systems depends heavily on how well these techniques are implemented and adapted to specific computing environments.

For future work, researchers may explore more intelligent cache management techniques, such as machine learning-based cache prediction and adaptive replacement policies that can dynamically adjust based on workload behavior. Further studies may also focus on improving cache efficiency in emerging technologies such as multicore processors, distributed systems, and AI-driven computing environments. These advancements can help further reduce memory bottlenecks and improve overall system performance in next-generation computer architectures.

REFERENCES

- A. Nematallah, C. H. Park, and D. Black-Schaffer, "Exploring the Latency Sensitivity of Cache Replacement Policies," *IEEE Computer Architecture Letters*, vol. 22, no. 2, pp. 93–96, 2023.
1. Y. Chen et al., "Improving Cache Performance using a Learned Replacement Policy," *IEEE Transactions on Computers*, vol. 72, no. 4, pp. 845–857, 2023.
2. S. Gupta and R. Singh, "Designing a Cost-Effective Cache Replacement Policy using Machine Learning," *IEEE Access*, vol. 11, pp. 14567–14579, 2023.
3. M. Patel and K. Sharma, "A Reconfigurable 16KB Cache Architecture for Direct-Mapped and Set-Associative Mapping," *IEEE Transactions on VLSI Systems*, vol. 31, no. 8, pp. 1123–1134, 2023.
4. J. Zhang and L. Wang, "A Set-Associative Cache Architecture with Dynamic Subset Mapping," *IEEE Transactions on Computers*, vol. 72, no. 6, pp. 1345–1356, 2024.
5. S. Kumar and P. Reddy, "Characteristics of Performance-Optimal Multi-Level Cache Hierarchies," *IEEE Transactions on Computers*, vol. 71, no. 12, pp. 2890–2905, 2023.
- A. Schneider et al., "Performance of Cache Memory Subsystems for Multicore Architectures," *IEEE Transactions on Parallel and Distributed Systems*, vol. 35, no. 2, pp. 412–425, 2024.
6. H. Choi et al., "Adaptive Cache Management with Machine Learning," *IEEE Transactions on Computers*, vol. 72, no. 9, pp. 1785–1798, 2023.
7. X. Li et al., "TAP: Transformers Driven Adaptive Prefetching for Hybrid Memory Systems," *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, vol. 43, no. 5, pp. 1021–1034, 2024.
8. E. Sanchez Hidalgo et al., "MemorAI: Energy-Efficient Last-Level Cache Memory Optimization for Virtualized RANs," *Proc. IEEE ICMLCN*, pp. 1–8, 2024.
9. IEEE Xplore Editorial, "Advances in Microprocessor Cache Architectures Over the Last 25 Years," *IEEE Micro*, vol. 43, no. 1, pp. 15–29, 2023.