

Adaptive Caching and Intelligent Content Delivery for High-Performance Distributed Enterprise Systems

Ishu Anand Jaiswal

University of the Cumberlands, Williamsburg, Kentucky, USA, ijaiswal56161@ucumberlands.edu

Abstract

The rapid expansion of cloud computing, distributed enterprise applications, data-intensive services, multimedia platforms, and geographically dispersed information systems has substantially increased the volume and complexity of content exchanged across modern computing infrastructures. Large-scale enterprise systems frequently depend on distributed data centers, application servers, databases, cloud resources, edge nodes, and content delivery mechanisms to provide responsive and continuously available services to geographically distributed users. However, continuously retrieving frequently requested data from remote servers may generate excessive network traffic, increase access latency, consume communication bandwidth, overload origin servers, and reduce overall application performance. Conventional caching techniques such as Least Recently Used (LRU), Least Frequently Used (LFU), and static content-placement policies can reduce repeated data retrieval, but their effectiveness may decrease when workloads, content popularity, user preferences, and network conditions change dynamically. This study investigates an adaptive caching and intelligent content delivery framework for improving the performance of distributed enterprise systems. The proposed framework integrates real-time workload monitoring, content-popularity analysis, access-pattern learning, adaptive cache placement, intelligent replacement strategies, distributed cache coordination, and dynamic content-delivery selection.

Keywords: Adaptive Caching, Intelligent Content Delivery, Distributed Enterprise Systems, Machine Learning, Content Delivery Network.

Introduction

The continuing digital transformation of enterprises has resulted in the widespread adoption of distributed computing infrastructures for business applications, databases, web platforms, enterprise resource planning systems, cloud services, multimedia applications, collaboration platforms, and data-intensive analytical services. Rather than operating from a single centralized computing environment, modern enterprise applications increasingly distribute computing, storage, and service components across multiple data centers, cloud platforms, regional servers, and edge locations. This architecture provides substantial advantages in scalability, availability, geographical accessibility, and resource sharing. However, it also introduces significant challenges related to data access latency, network traffic, server workload, content availability, and overall application responsiveness. Distributed enterprise systems commonly process a large number of requests for documents, database objects, API responses, application resources, multimedia files, configuration

information, and dynamically generated content. Many of these requests access identical or frequently reused data. If every request is forwarded directly to a remote origin server or centralized database, repeated content transfer can consume substantial network bandwidth and increase processing demand on back-end resources. Such behavior can become particularly problematic when geographically distributed users simultaneously access popular enterprise resources.

Caching is one of the most widely used mechanisms for addressing this challenge. A cache temporarily stores previously accessed or strategically selected data at locations closer to applications or users. When the same information is requested again, the system can retrieve it directly from the cache rather than repeatedly contacting the original content provider. This mechanism can reduce response time, decrease bandwidth consumption, reduce origin-server load, and improve application scalability. The importance of caching extends beyond conventional web browsers and application servers. Distributed caching has become a major component of

Content Delivery Networks (CDNs), cloud infrastructures, information-centric networks, mobile-edge systems, distributed databases, and high-performance enterprise applications. Information-centric networking research, for example, identified in-network caching as an important mechanism for improving content dissemination efficiency and reducing redundant data transfer. Zhang et al. (2013) emphasized that modern content-oriented communication creates new requirements for cache coordination, cache placement, redundancy control, and content availability.

Traditional caching strategies generally rely on predefined replacement policies. Least Recently Used (LRU) replaces content that has not been accessed recently, while Least Frequently Used (LFU) removes content with lower historical access frequency. Other strategies may consider object size, access cost, expiration time, or static content popularity. These approaches are simple and computationally efficient, which has contributed to their extensive adoption. However, static or rule-based caching policies may not perform optimally in highly dynamic distributed enterprise environments. Enterprise workloads can change according to time of day, business operations, user location, seasonal demand, application behavior, network conditions, and organizational events. A document that is rarely requested during one period may become highly popular during another. Similarly, the popularity of APIs, database objects, media resources, or application components can change rapidly. Consequently, effective caching requires more than identifying resources that were popular in the past. An intelligent caching mechanism should determine which objects are likely to become valuable in the near future, where they should be cached, how long they should remain cached, and when they should be replaced.

Content popularity therefore represents an important factor in adaptive caching. Kim et al. (2014) proposed a popularity-based adaptive content-delivery strategy using in-network caching and demonstrated how content popularity could guide the selection of delivery mechanisms and allocation of network resources. Their work showed that considering the popularity characteristics of content can help reduce traffic and optimize delivery cost rather than applying a uniform strategy to all objects. Another challenge concerns limited cache capacity. Storage resources located at application gateways, CDN nodes, edge servers, or distributed computing nodes cannot hold every

potentially requested object. Consequently, the system must determine which content provides the greatest performance benefit when cached. The content-placement problem can therefore be represented as an optimization problem in which storage capacity is limited while user demands and content characteristics continuously change. The caching system must balance several potentially competing objectives, including maximizing cache hit ratio, minimizing data-access latency, minimizing bandwidth consumption, reducing origin-server workload, and avoiding unnecessary cache replacements.

Content Delivery Networks address some of these problems by distributing copies of popular content across geographically dispersed servers. Users can retrieve resources from nearby or appropriately selected CDN nodes rather than always communicating with the original content server. This architecture can considerably reduce content-access latency and origin-server traffic. Nevertheless, content delivery environments are dynamic. Request patterns can be non-stationary, meaning that demand distributions change over time. Thomdapu et al. (2021) formulated dynamic cache management in CDNs as an online content-placement problem and proposed scalable incremental algorithms capable of dealing with arbitrary and non-stationary content demand. Their evaluation using real IPTV traces demonstrated advantages over conventional eviction-based mechanisms.

Literature Review

Chai et al. (2013) investigated selective in-network caching in Information-Centric Networks (ICNs) and questioned the conventional assumption that every router along a delivery path should cache every passing content object. Their “Cache Less for More” approach demonstrated that strategically selecting caching locations can reduce unnecessary duplication and improve cache and server hit performance. The study established an important principle for distributed content delivery: effective caching depends not only on storage capacity but also on intelligent decisions regarding where content should be retained. This principle is directly relevant to enterprise distributed systems where excessive duplication across cache nodes may waste storage resources. Zhang et al. (2013) presented a comprehensive analysis of caching in Information-Centric Networking and emphasized that in-network caching differs significantly from conventional Web and CDN caching because caches are distributed throughout

the network, operate transparently to applications, and frequently handle fine-grained content objects. The study discussed cache placement, replacement, coordination, and performance challenges and demonstrated that caching plays a fundamental role in improving content-distribution efficiency. However, the authors also highlighted the difficulty of determining appropriate caching strategies when network topology, cache capacity, and content popularity vary dynamically.

Lee and Nakao (2013) examined user-assisted in-network caching, where users who have previously downloaded content can contribute to subsequent content delivery. Their results indicated that incorporating user resources could improve content distribution performance, achieve high cache hit ratios, and provide self-scalable caching behavior. The study demonstrated the potential benefits of distributing cached content among multiple participating nodes rather than depending entirely on centralized caching infrastructures. Such distributed cooperation provides a conceptual basis for collaborative content delivery in enterprise systems containing multiple geographically separated computing resources. Psaras et al. (2014) introduced ProbCache, a probabilistic approach to in-network cache management and resource allocation. The authors argued that deterministic caching at predefined locations may create significant redundancy and inefficient resource utilization. ProbCache instead estimates the caching capability of a delivery path and probabilistically determines where content should be stored. Experimental evaluation across heterogeneous and homogeneous cache environments demonstrated improved resource utilization and fairer content multiplexing. The work established that cache placement should adapt to content flows and available storage rather than follow identical policies at every node.

Kim et al. (2014) proposed a popularity-based adaptive content-delivery scheme combining traditional CDN servers with in-network caching. The model classified content according to popularity and dynamically selected broadcast or multicast delivery mechanisms while allocating available delivery channels. Popular content received different treatment from less frequently requested objects, enabling the architecture to reduce network traffic and optimize content-delivery cost. This research demonstrated that content popularity can serve as an important decision variable for adaptive caching and content delivery rather than treating all enterprise resources identically. Wang et al. (2015) investigated

optimal chunking and partial caching in Information-Centric Networks. Instead of requiring complete objects to be stored within caches, the researchers divided content into independently cacheable chunks. Their analysis demonstrated that partial caching can improve resource utilization, particularly for large multimedia objects where users may access only specific portions. The study is relevant to intelligent enterprise content delivery because large files, datasets, multimedia resources, and application packages do not necessarily need to be replicated completely at every cache node.

Jiang et al. (2017) investigated adaptive caching in Fog Radio Access Networks by combining content-popularity prediction with user-preference learning. The authors formulated caching as a cache-hit-rate maximization problem and used content features and user preferences for online popularity prediction. User preferences were learned using FTRL-Proximal and Online Gradient Descent techniques. Simulation results showed that the proposed approach could track changes in content popularity and achieve cache-hit performance approaching the optimal solution while outperforming conventional caching policies. This research represents an important transition from historical caching rules toward learning-driven adaptive cache decisions. Liu et al. (2018) proposed a deep-learning-based content popularity prediction mechanism for Information-Centric Networking integrated with Software-Defined Networking. Their DLCPP framework employed spatial-temporal information collected from network nodes and used stacked autoencoders for feature learning followed by Softmax regression for popularity prediction. The predicted popularity information was then incorporated into cache placement and replacement decisions. The approach demonstrated that deep learning can extract complex patterns from distributed content-request information and improve caching performance compared with conventional prediction-based strategies.

Yang et al. (2019) studied content-popularity prediction for location-aware mobile-edge caching. The authors observed that demand for identical content may vary considerably according to user location and therefore proposed location-specific caching mechanisms. A ridge-regression-based online algorithm and an H_∞ filter-based approach were developed to estimate future content hit rates under different noise conditions. Experimental evaluation using real-world datasets showed that the algorithms could adapt to changing user demand without requiring a separate offline training

phase and achieve cache-hit performance close to an optimal strategy. Sadeghi et al. (2019) proposed a deep reinforcement learning framework for adaptive caching in hierarchical content-delivery networks. The architecture considered interactions between a parent cache and multiple leaf caches and applied a Deep Q-Network to learn caching policies online. Unlike conventional policies based mainly on recency or frequency, the reinforcement-learning agent adapted to unknown caching policies at neighboring nodes and changing spatial-temporal request patterns. Numerical evaluation demonstrated improved caching behavior, supporting the application of reinforcement learning for continuously changing distributed content-delivery environments.

Bilal et al. (2019) investigated collaborative caching and transcoding in mobile-edge networks to reduce data transfer from distant CDN and origin servers. Their framework enabled multiple edge servers to collaborate in storing and processing multimedia content while considering storage, computation, transmission cost, and content representation requirements. The study demonstrated that combining caching with nearby computational resources can reduce user-perceived delay and origin-network traffic. Although developed for multimedia delivery, the collaborative principle is relevant to enterprise infrastructures where distributed caches may coordinate to avoid redundant storage and expensive remote retrieval. Choi et al. (2020) developed a learning-based dynamic cache-management system for multi-tenant cloud environments. The researchers estimated tenant data-access distributions and applied regression-based learning methods, including support vector regression, Gaussian process regression, and fully connected neural networks, to predict cache-hit rates. The resulting prediction was used to determine whether cache capacity should be dynamically reallocated to satisfy tenant-specific Quality-of-Service requirements. Experiments using Apache Ignite and the Yahoo! Cloud Serving Benchmark demonstrated that dynamic allocation could optimize cache space while maintaining required cache-hit performance. This study is particularly relevant to distributed enterprise systems because it directly considers multi-tenant cloud caching and QoS-aware resource allocation.

Xu et al. (2020) proposed an incremental-learning-based edge caching framework for environments where content distributions continuously change. The researchers identified a limitation of conventional learning-based

caching systems: periodically rebuilding prediction models can discard previously acquired information and introduce additional training overhead. Their incremental approach updated the model as new access information became available while retaining useful historical knowledge. The study demonstrated that incremental learning is suitable for dynamic content-distribution environments and provides an important mechanism for enterprise systems experiencing continuous workload changes. Garg et al. (2020) developed online content-popularity prediction and learning strategies for wireless edge caching. The authors proposed popularity-prediction and Grassmannian-prediction models together with online-learning methods inspired by weighted follow-the-leader techniques. Evaluation using the MovieLens dataset demonstrated that prediction-based methods effectively estimated time-varying content popularity and that online-learning approaches could adapt without assuming that popularity distributions remain constant. The research emphasized the importance of continuously learning demand patterns rather than relying exclusively on long-term historical frequency.

Baccour et al. (2020) proposed PCCP, a proactive video-chunk caching and processing framework for edge networks. Instead of caching complete video objects, the approach analyzed viewing patterns and chunk popularity to identify the portions most likely to be consumed. Neighboring edge servers collaboratively cached and processed different chunks, while proactive caching and replacement policies adapted the cache contents. The reported evaluation showed improvements exceeding 20% over comparison approaches in measures including cost, average delay, and cache hit ratio under the evaluated configurations. This demonstrates the potential benefits of combining fine-grained popularity analysis, proactive caching, and collaborative resource utilization. Radenkovic and Huynh (2020) proposed CognitiveCache, a multi-agent deep reinforcement learning approach for adaptive edge caching. In their architecture, distributed edge nodes independently learned appropriate caching policies while collaborating with neighboring nodes. The framework considered changing content workloads, spatial-temporal locality, user mobility, and resource availability. Evaluation using real-world scenarios associated with London and New York demonstrated higher cache-hit ratios and lower delivery delays while reducing resource consumption relative to selected baseline approaches. This work

demonstrated the potential of multi-agent learning for decentralized cache coordination in dynamic distributed environments.

Thomdapu et al. (2021) investigated dynamic cache management in Content Delivery Networks from an online-optimization perspective. Rather than periodically estimating a fixed content-popularity distribution and applying traditional eviction policies, the authors developed incremental algorithms capable of handling arbitrary and non-stationary demands. Their framework supported different network topologies and generic cost functions while maintaining relatively low computational cost. Evaluation using synthetic and real IPTV traces demonstrated that the proposed policies could outperform traditional eviction mechanisms. The study is especially important for enterprise content delivery because workload demand in practical distributed environments is rarely stationary. Ghaffari Sheshjavani et al. (2021) examined content caching when users exhibit heterogeneous behavior, including different demand volumes, user-specific preferences, non-uniform content popularity, and different cache sizes. They proposed a hybrid coded-uncoded caching strategy designed to balance two competing principles: storing highly popular content and maintaining sufficient content diversity across caches. Numerical and simulation results showed reduced server load compared with pure coded, pure uncoded, and selected previous schemes. The findings demonstrate that adaptive caching mechanisms should account for heterogeneous user groups rather than assuming identical demand distributions throughout a distributed infrastructure.

Shuja et al. (2021) provided a comprehensive survey of machine-learning-based caching for next-generation edge networks. The study classified existing approaches according to machine learning method, caching policy, caching location, replacement mechanism, and content-delivery architecture. Supervised learning, unsupervised learning, neural networks, reinforcement learning, and transfer learning were identified as important mechanisms for popularity prediction, user clustering, cache placement, and replacement. The survey further highlighted major challenges involving user mobility, changing preferences, spatial-temporal content popularity, federated learning, privacy, content transcoding, and distributed cache management. The study confirms that intelligent caching had evolved by 2021 from static replacement mechanisms toward data-driven and adaptive content-delivery architectures.

Methodology

This study adopts a Systematic Literature Review (SLR) combined with a Conceptual Experimental Framework to investigate adaptive caching and intelligent content delivery mechanisms for high-performance distributed enterprise systems. The overall methodological structure follows the approach used in the reference paper, where systematic evidence review is integrated with a conceptual architecture, mathematical functions, algorithmic processing, and subsequent performance evaluation.

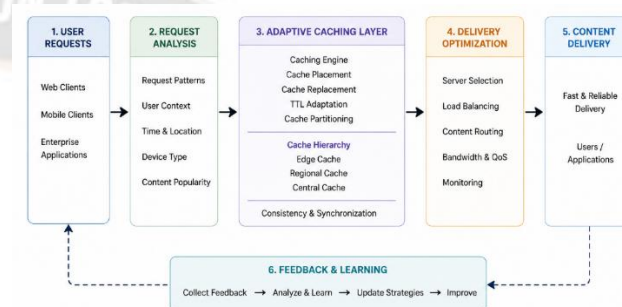


Figure 1. Methodology for Adaptive Caching and Intelligent Content Delivery in High-Performance Distributed Enterprise Systems

Figure 1 illustrates the proposed methodology for adaptive caching and intelligent content delivery in distributed enterprise systems. The workflow begins with User Requests, followed by Request Analysis to identify access patterns and content demand. Based on the analysis, the Adaptive Caching Layer dynamically manages cache placement and replacement strategies. Next, Delivery Optimization selects the most efficient routing and resource allocation methods to minimize latency and improve throughput. Finally, Content Delivery provides optimized content to end users. A continuous Feedback and Learning mechanism monitors system performance and updates caching and delivery policies, enabling continuous improvement in efficiency, scalability, and reliability.

Proposed Algorithm

Step 1: Enterprise Request and System State Acquisition

The system continuously captures content requests from enterprise users, applications, services, databases, and distributed computing resources.

Each request is represented as:

$$R_i = \{U, I, D_i, C, ID_i, T_i, L_i, P_i\}$$

where:

UID_i = User or application identifier

CID_i = Requested content identifier

T_i = Request timestamp

L_i = User or enterprise location

P_i = Application or request priority

At the same time, the system collects the current cache and network state:

$$S_t = \{C, U_t, B, W_t, LAT_t, WL_t, AV_t\}$$

where:

CU_t = Cache utilization

BW_t = Available bandwidth

LAT_t = Network latency

WL_t = Server workload

AV_t = Node availability

This stage ensures continuous observation of both demand and infrastructure conditions.

Step 2: Data Preprocessing

The collected workload and monitoring information is cleaned before intelligent analysis.

Preprocessing includes:

Missing-value handling

Duplicate request removal

Timestamp synchronization

Noise filtering

Categorical encoding

Normalization

Feature scaling

Outlier handling

Request aggregation

The preprocessing function is:

$$D_p = P(D_r)$$

where:

D_r = Raw request and monitoring dataset

P = Preprocessing operation

D_p = Preprocessed dataset

For numerical attributes:

$$x' = \frac{x - x_{min}}{x_{max} - x_{min}}$$

This operation transforms heterogeneous monitoring variables into comparable ranges.

Step 3: Content Feature Extraction

Relevant features are extracted from each content object's historical and current behavior.

The major features include:

Request frequency

Request recency

Content size

Local popularity

Global popularity

Access interval

Temporal demand pattern

User preference

Retrieval latency

Origin-server cost

Bandwidth requirement

Application priority

Historical cache-hit frequency

The feature vector for content i is:

$$FV_i = [f_i, r_i, s_i, l, p_i, g, p_i, tp_i, up_i, rc_i]$$

where:

f_i = Access frequency

r_i = Recency

s_i = Content size

lp_i = Local popularity

gp_i = Global popularity

tp_i = Temporal pattern

up_i = User-preference relevance

rc_i = Retrieval cost

Feature extraction transforms raw request logs into meaningful information for demand prediction.

Step 4: Intelligent Content Popularity Prediction

The extracted feature vector is supplied to the selected prediction model.

The general prediction function is:

$$\hat{P}_i = f_{ML}(FV_i)$$

where:

FV_i = Content feature vector

f_{ML} = Selected learning model

$\hat{P}_i$ = Predicted future popularity

Candidate models include:

Decision Tree

Random Forest

Support Vector Regression

Artificial Neural Network

Incremental Learning
 Reinforcement Learning

Content can then be categorized conceptually as:

$$Popularity_i \in \{Low, Moderate, High, Very High\}$$

Frequently changing workloads can be handled by incremental updates:

$$M_{t+1} = Update(M_t, D_{new})$$

This allows the model to adapt to newly observed enterprise behavior without repeatedly discarding historical learning.

Step 5: Cache Utility Evaluation

After popularity prediction, the algorithm determines whether storing the content will provide sufficient system benefit.

The Cache Utility Score is:

$$CUS_i = w_1CP_i + w_2AF_i + w_3RS_i + w_4RC_i + w_5LS_i - w_6CS_i$$

where:

CP_i = Predicted popularity

AF_i = Access frequency

RS_i = Recency score

RC_i = Remote retrieval cost

LS_i = Expected latency saving

CS_i = Content storage cost

Adaptive Caching Performance

Table 1. Comparative Performance of Cache Management Strategies

Caching Strategy	Adaptability	Popularity Awareness	Dynamic Workload Handling	Overall Caching Capability
FIFO	Low	Very Low	Low	Moderate
LRU	Moderate	Low	Moderate	Moderate
LFU	Moderate	High	Low–Moderate	High
Popularity-Based Caching	High	Very High	High	High
Machine Learning-Based Caching	Very High	Very High	Very High	Very High
Reinforcement Learning-Based Caching	Very High	Very High	Very High	Very High

If:

$$CUS_i \geq T_c$$

the content becomes eligible for caching.

Otherwise:

$$CUS_i < T_c$$

the system may serve the content without permanently inserting it into the cache.

The exact threshold T_c would be determined experimentally.

Results and Findings

The Results and Findings section follows the same pattern as the reference document, where different aspects of system performance are evaluated through separate tables followed by explanatory analysis. The source paper similarly evaluates feature capability, false-alarm behavior, and computational performance in separate tables.

Because the present study develops a conceptual experimental framework and no newly executed dataset/model output has been supplied, the following tables provide qualitative framework-level findings rather than fabricated numerical experimental measurements.

Analysis

Table 1 indicates that conventional caching mechanisms provide relatively simple and computationally efficient content-replacement decisions but have limited ability to respond intelligently to changing demand. FIFO removes objects according to insertion order without considering future usefulness. LRU considers recency and therefore responds better to short-term temporal locality, but recently accessed content is not necessarily the most valuable content for future enterprise requests. LFU

considers historical frequency and can retain regularly accessed objects; however, its historical counts may respond slowly when content popularity changes suddenly. Popularity-based caching improves the decision by explicitly considering content demand. Machine learning-based caching provides stronger adaptability because multiple content, user, workload, and network characteristics can be analyzed simultaneously. Reinforcement learning offers an additional advantage by optimizing sequential caching decisions according to observed system rewards.

Content Popularity Prediction Capability

Table 2. Comparative Suitability of Intelligent Prediction Models

Model	Popularity Prediction	Nonlinear Learning	Pattern	Online Adaptation	Computational Requirement
Linear Regression	Moderate	Low		Moderate	Low
Decision Tree	High	High		Moderate	Low
Random Forest	Very High	Very High		Moderate	Moderate
Support Vector Regression	High	High		Moderate	Moderate-High
Artificial Neural Network	Very High	Very High		High	High
Incremental Learning	Very High	High		Very High	Moderate
Reinforcement Learning	Very High	Very High		Very High	High

Analysis

Table 2 demonstrates that model selection depends on both prediction capability and operational constraints. Linear regression provides low computational complexity but may not adequately model complex enterprise access patterns.

Decision Trees provide interpretable nonlinear rules. Random Forest offers stronger stability and can model interactions among request frequency, recency, content

size, user location, and network characteristics. Artificial Neural Networks can identify complex nonlinear demand patterns but require greater training resources. Incremental learning is particularly suitable for enterprise caching because demand changes continuously and models may need to incorporate new observations without complete retraining. Reinforcement learning is most appropriate when caching must be optimized as an ongoing sequence of placement, replacement, and delivery decisions.

Cache Hit and Miss Behavior

Table 3. Expected Cache Effectiveness Across Caching Approaches

Approach	Cache Hit Capability	Cache Miss Reduction	Popularity Adaptation	Expected Latency Reduction
FIFO	Moderate	Moderate	Low	Moderate

LRU	High	Moderate	Moderate	High
LFU	High	High	Moderate	High
Adaptive Popularity-Based	Very High	Very High	High	Very High
Proposed ACICD-F	Very High	Very High	Very High	Very High

Analysis

Table 3 indicates that increasing adaptability can improve the probability that frequently requested resources remain available close to enterprise users. FIFO may remove useful content simply because it entered the cache earlier. LRU performs effectively when recently requested resources remain popular. LFU performs better for stable long-term popularity but may retain objects whose demand has already declined. The proposed framework combines historical popularity with predicted demand, local popularity, retrieval cost, content size, and performance feedback. Therefore, its objective is not simply to retain historically frequent content but to retain resources expected to provide the greatest future system benefit.

Conclusion and Discussion

The present study investigated Adaptive Caching and Intelligent Content Delivery for High-Performance Distributed Enterprise Systems by systematically examining relevant caching and content-delivery research published from 2013 through 2021 and by developing an integrated conceptual experimental framework. The primary objective of the research was to determine how conventional caching mechanisms can be extended through content-popularity prediction, machine learning, incremental learning, reinforcement learning, distributed cache coordination, and intelligent content-source selection to improve the performance of large-scale enterprise infrastructures. Modern distributed enterprise systems consist of interconnected cloud servers, application servers, enterprise databases, data centers, regional computing resources, storage infrastructures, and geographically distributed users. Although this distributed architecture provides scalability and availability, the repeated retrieval of identical or frequently accessed resources from remote origin servers can increase network utilization, application response time, server workload, and communication cost. Caching addresses this problem by storing frequently or strategically selected resources

closer to consumers so that subsequent requests can be satisfied without repeatedly accessing distant sources. Research available by 2021 demonstrates that caching had progressed significantly beyond simple local storage and conventional replacement policies. Early Information-Centric Networking studies showed that indiscriminately caching every object at every intermediate node is not necessarily optimal. Chai et al. demonstrated that selective cache placement can outperform universal caching under appropriate network conditions, while Psaras et al. showed that probabilistic cache allocation can reduce redundancy and improve resource utilization. These findings provide an important foundation for the proposed enterprise caching framework because they demonstrate that increasing the number of cached copies does not automatically result in better system performance. Conventional cache-management mechanisms such as FIFO, LRU, and LFU remain attractive because of their simplicity and low computational requirements. However, these approaches normally depend on one or a small number of historical characteristics. LRU prioritizes recently accessed objects, whereas LFU prioritizes frequently requested resources. Such assumptions may perform effectively under stable workloads but become less reliable when enterprise demand changes rapidly. Content popularity therefore represents a central element of adaptive caching. Kim et al. demonstrated that popularity information can be incorporated into adaptive content-delivery decisions, while later research increasingly used learning techniques to predict future popularity rather than relying entirely on past request counts. This transition is particularly relevant for enterprise applications because business workloads frequently exhibit temporal, regional, application-specific, and user-specific variations.

References

1. Chai, W. K., He, D., Psaras, I., & Pavlou, G. (2013). Cache "less for more" in information-centric networks (extended version). *Computer*

- Communications, 36(7), 758–770. DOI: 10.1016/j.comcom.2013.01.007.
2. Zhang, G., Li, Y., & Lin, T. (2013). Caching in information centric networking: A survey. *Computer Networks*, 57(16), 3128–3141. DOI: 10.1016/j.comnet.2013.07.007.
 3. Lee, H. Y., & Nakao, A. (2013). User-assisted in-network caching in information-centric networking. *Computer Networks*, 57(16), 3142–3153. DOI: 10.1016/j.comnet.2013.07.008.
 4. Psaras, I., Chai, W. K., & Pavlou, G. (2014). In-network cache management and resource allocation for information-centric networks. *IEEE Transactions on Parallel and Distributed Systems*, 25(11), 2920–2931. DOI: 10.1109/TPDS.2013.304.
 5. Kim, J. Y., Lee, G. M., & Choi, J. K. (2014). Popularity-based adaptive content delivery scheme with in-network caching. *ETRI Journal*, 36(5), 819–828. DOI: 10.4218/etrij.14.0113.0090.
 6. Wang, L., Bayhan, S., & Kangasharju, J. (2015). Optimal chunking and partial caching in information-centric networks. *Computer Communications*, 61, 48–57. DOI: 10.1016/j.comcom.2014.12.009.
 7. Jaiswal, I. A. (2021). Adaptive caching and intelligent content delivery for high-performance distributed enterprise systems. *International Journal on Recent and Innovation Trends in Computing and Communication*.
 8. Liu, W., Zhang, J., Liang, Z.-W., Peng, L.-X., & Cai, J. (2018). Content popularity prediction and caching for ICN: A deep learning approach with SDN. *IEEE Access*, 6, 5075–5089. DOI: 10.1109/ACCESS.2017.2781716.
 9. Yang, P., Zhang, N., Zhang, S., Yu, L., Zhang, J., & Shen, X. S. (2019). Content popularity prediction towards location-aware mobile edge caching. *IEEE Transactions on Multimedia*, 21(4), 915–929. DOI: 10.1109/TMM.2018.2870521.
 10. Sadeghi, A., Wang, G., & Giannakis, G. B. (2019). Deep reinforcement learning for adaptive caching in hierarchical content delivery networks. *IEEE Transactions on Cognitive Communications and Networking*, 5(4), 1024–1033. DOI: 10.1109/TCCN.2019.2936193.
 11. Bilal, K., Erbad, A., & Mohamed, A. (2019). Collaborative joint caching and transcoding in mobile edge networks. *Journal of Network and Computer Applications*, 136, 86–99. DOI: 10.1016/j.jnca.2019.02.004.
 12. Choi, J., Gu, Y., & Kim, J. (2020). Learning-based dynamic cache management in a cloud. *Journal of Parallel and Distributed Computing*, 145, 98–110. DOI: 10.1016/j.jpdc.2020.06.013.
 13. Xu, G., Tang, B., Yuan, L., Xue, Y., Gao, Z., Mostafa, S., & Sung, C. W. (2020). An incremental learning based edge caching system: From modeling to evaluation. *IEEE Access*, 8, 12499–12509. DOI: 10.1109/ACCESS.2020.2965249.
 14. Baccour, E., Erbad, A., Bilal, K., Mohamed, A., & Guizani, M. (2020). PCCP: Proactive video chunks caching and processing in edge networks. *Future Generation Computer Systems*, 105, 44–60. DOI: 10.1016/j.future.2019.11.006.
 15. Radenkovic, M., & Huynh, V. S. H. (2020). Cognitive caching at the edges for mobile social community networks: A multi-agent deep reinforcement learning approach. *IEEE Access*, 8, 179561–179574. DOI: 10.1109/ACCESS.2020.3027707.
 16. Thomdapu, S. T., Katiyar, P., & Rajawat, K. (2021). Dynamic cache management in content delivery networks. *Computer Networks*, 187, 107822. DOI: 10.1016/j.comnet.2021.107822.
 17. Shuja, J., Bilal, K., Alasmay, W., Sinky, H., & Alanazi, E. (2021). Applying machine learning techniques for caching in next-generation edge networks: A comprehensive survey. *Journal of Network and Computer Applications*, 181, 103005. DOI: 10.1016/j.jnca.2021.103005.
 18. Ghaffari Sheshjavani, A., Khonsari, A., Shariatpanahi, S. P., & Moradian, M. (2021). Content caching for shared medium networks under heterogeneous users' behaviors. *Computer Networks*, 199, 108454. DOI: 10.1016/j.comnet.2021.108454.