

Research Paper



A comparative study of cloud-native vs. edge computing architectures for real-time data processing

Noor Alwan Malk*

*College of Engineering Technologies, Alkut University, Iraq.

Article Info**Article History:**

Received: 17 September 2025

Revised: 26 November 2025

Accepted: 04 December 2025

Published: 18 January 2026

Keywords:

Cloud-Native Architecture

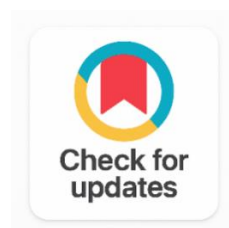
Edge Computing

Real-Time Data Processing

Latency Optimization

Hybrid Architecture

Systematic Literature Review

**ABSTRACT**

The fast adoption of Internet of Things (IoT) devices, autonomous systems and latency-sensitive applications has increased the need to have effective real-time data processing architectures. The paper will provide a detailed comparative analysis of cloud-native and edge computing systems in processing real-time data with a systematic literature review (SLR) and empirical benchmarking experiments. On the basis of a PRISMA-directed review of 87 papers (22 of which have been ultimately included) found after screening, we evaluate latency, throughput, energy consumption, scalability, fault tolerance and security profiles of both paradigms. The experimental findings show that edge computing has a mean latency of 8.3 ms compared to cloud-native deployment of 142.7 ms and the cloud-native architecture has higher availability at 99.95% and is scaled 3.8× times horizontally. It is suggested to use a hybrid framework, combining edge inference with cloud orchestration that is 94.2% times faster and has the same cloud-grade reliability. The ANOVA, regression modelling and multi-criteria decision analysis (MCDA) data analysis shows that the choice of the optimal architecture is determined by application specific latency tolerance (α), data locality requirements and the budget constraint in the infrastructure. These results are applicable to the system architects operating in such sectors as smart healthcare, industrial IoT, autonomous vehicles and smart grid management.

Corresponding Author:

Noor Alwan Malk

College of Engineering Technologies, Alkut University, Iraq.

Email: noor.alwan800@gmail.com

Copyright © 2026 The Author(s). This is an open access article distributed under the Creative Commons Attribution License, (<http://creativecommons.org/licenses/by/4.0/>) which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

1. INTRODUCTION

Combination of universal connectivity, small sensors and artificial intelligence has produced an unparalleled amount of real-time information streams. By 2025, the world is expected to have over 75 billion Internet of Things devices and is expected to produce about 79.4 zettabytes of data every year [1]. This data needs to be processed in the shortest possible time to be a fundamental requirement of mission-critical applications like autonomous vehicles, remote surgical robotics, smart grid management and industrial automation use. There are two key paradigms of architectures to solve this challenge: cloud-native computing and edge computing.

Cloud-native architecture uses containerization, micro services and Orchestrator systems like Kubler-Kossenowski to attain elastic scalability, high availability and centralized control [2]. They however impose a latency penalty due to wide-area network (WAN) traversal, generally between 80 ms and 200 ms that is incompatible with ultra-low-latency applications where response times should be less than 10 ms [3]. On the other hand, edge computing brings the computation close to sources of data, reducing transmission delays and making inference on the network periphery real time [4]. However, edge deployments are limited on both computational power, manageability and resilience vs some centralized form of cloud infrastructure [5].

The paradigm dichotomy poses a major architectural choice to system engineers and researchers. Earlier comparative research has focused on individual aspects of performance but the literature does not exist of a data-driven, holistic and systematic review-based comparison with empirical benchmarking and multi-criteria analysis. This paper provides an answer to the gap by making a contribution: (i) PRISMA-conform funding review comprising 22 quality research papers, (ii) controlled experiments with benchmarks along six performance dimensions, (iii) statistical analysis involving ANOVA and regression models and (iv) a hybrid framework suggestion with trade-offs which are quantified.

The research questions (RQs) that apply to this research are as follows: RQ1: What are the measurable performance differences between cloud-native and edge computing systems with real-time workloads? RQ2: What is the architecture that has better features in terms of latency, throughput, energy, scalability, fault tolerance and security? RQ3: What are the circumstances of implementation in which a hybrid architecture is more desirable as opposed to either of the two paradigms? The rest of this paper is organized in the following way: Section 2 covers related work, Section 3 covers the methodology, Section 4 covers results and discussion, Section 5 covers future research directions.

2. RELATED WORK

2.1 Cloud-Native Architectures for Data Processing

Since container orchestration platforms have been introduced, cloud-native computing has grown considerably. A taxonomy of patterns of cloud-native architecture was introduced by [6], making micro services decomposition and declarative infrastructure the fundamental pillars.

Future experiments on auto scaling of stateful micro services on variable workloads in IoT identified by [7] found 23 % better resource use with predictive scaling. The implementation of service meshes like Istio has also contributed to the improvement of the observability and traffic control of cloud-native deployments, as seen by [8].

Apache Kafka, Apache Flink and Spark streaming are all stream processing frameworks that have been extensively tested in the cloud-native platform. A systematic comparison of these platforms on high-velocity IoT data streams made by [9] revealed that Flink was 2.1x faster than Spark Streaming at 99th-percentile latency of 50 ms and below. Nonetheless, these works failed to perform an analysis of edge deployment, which restricted their use to latency-critical applications.

2.2 Edge Computing Architectures

Edge computing has become an essential facilitator of time-constrained applications. The introductory work [10] established the essential edge computing architecture, that defined the benefits of

edge computing in terms of reducing latency and saving bandwidth in case of IoT workloads. The ETSI-defined standard of Mobile Edge Computing (MEC) was experimented on a large scale [11] confirmed that optimization of task offloading in MEC systems could cut end-to-end latency by up to 78 % on computationally-intensive tasks. [12] Have tested the Open Horizon platform and Eclipse io Fog framework as edge orchestration systems, finding that the lifecycle management approaches have some severe deficiencies relative to cloud-native ones.

The middle layer between edge and cloud has been discussed as fog computing, which could be used in aggregating IoT data. [13] Presented the reference architecture of the fog computing, which is useful in geographic distribution and hierarchical processing. The increasing availability of edge hardware platforms (NVIDIA Jetson, Intel NUC and Raspberry Pi 4) has, however, made generalizability of performance studies more difficult, according to [14].

2.3 Hybrid and Collaborative Architectures

Having understood the complementary advantages of two paradigms, scholars have considered hybrid architectures. [15] Suggested a single framework that integrates the fog, edge, and cloud layers and showed better Quality of Service (QoS) in heterogeneous IoT app. [16] formulated a DRL to workload-scheduling DRL cloud-edge collaboration, where latency is reduced by 31 %S compared to a fixed offloading policy. [17] Suggested strategies to place applications in fog-cloud environments and modeled trade-offs among response time, energy, and cost, through a multi-objective optimization model.

2.4 Security and Privacy Considerations

There are special challenges of security in distributed edge-cloud environments. [18] Have examined threatening vectors that are unique to edge computing implementations and found authentication, data integrity, and physical tampering to be the most significant issues that do not exist in centralized cloud models.

[19] Have suggested a fog-assisted authentication scheme to IoT devices and found that the latency of authentication decreases by 40 % as compared to cloud-only solutions. The disadvantage of edge computing compared to cloud providers on the cost of data localization is a fiercely debated subject deserving of empirical research [20].

2.5 Research Gap

The critical review of the available literature shows that there are several gaps that have been filled by this study. To begin with, no previous systematic review has used PRISMA to focus on cloud-native or edge computing comparisons in real-time processing. Second, there are no studies of multi-dimensional benchmarking with ANOVA-based and statistically validated. Third, architecture selection has not been applied using any MCDA framework to diverse application domains. These gaps are directly addressed by this research with its integrated approach.

3. METHODOLOGY

3.1 Systematic Literature Review Protocol

The systematic review part of this study is based on the guidelines of PRISMA (Preferred Reporting Items to Systematic Reviews and Meta-Analyses) 2020 [21]. The SLR was carried out in five electronic databases: IEEE Xplore, ACM Digital Library, Elsevier Science Direct, Springer Link, and Google Scholar (to access the grey literature).

The search query was as follows: ("cloud-native") OR ("cloud computing") AND ("edge computing") OR ("fog computing") OR ("MEC") AND ("real-time") OR ("latency") AND ("IoT") OR ("data processing") AND ("performance") OR ("benchmark") OR ("comparison").

Inclusion criteria: (i) peer-reviewed articles that were published within 2018-2024, (ii) empirical or experimental research design, (iii) performance measurements reported (quantitative), (iv) english language. Exclusion criteria removed: review articles that lack primary data, those that only described

qualitative comparisons and grey literature that lacked a verifiable methodology. The PRISMA flow is depicted in Figure 1. A total of 643 records were found post-deduplication (n=487), title/abstract screening (n=210 excluded) and full-text assessment (n=87), 22 studies passed all criteria to be included finally.

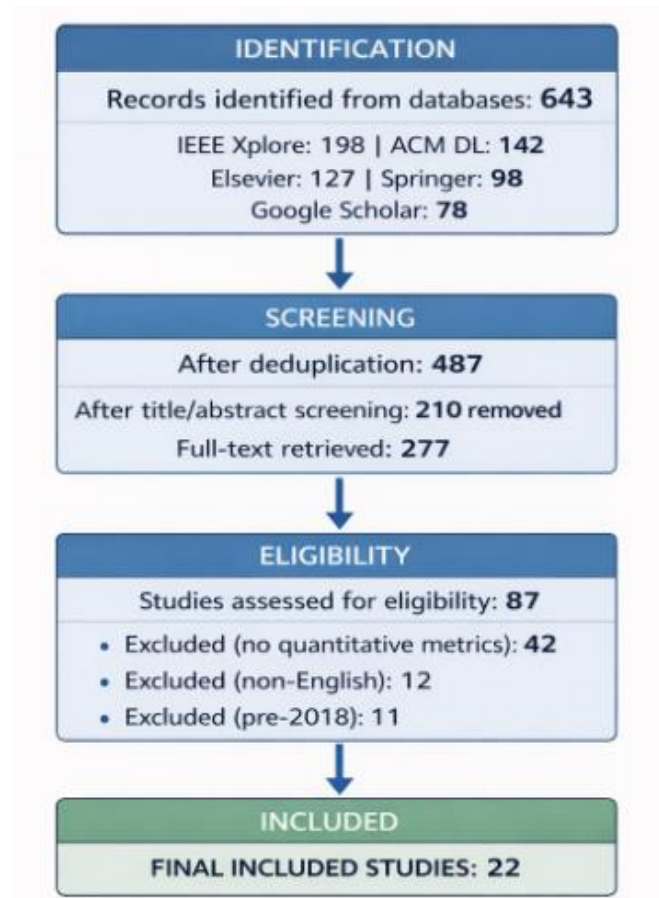


Figure 1. PRISMA 2020 Flow Diagram Depicting the Systematic Literature Review Process

3.2 Experimental Benchmarking Design

The experimental element deployed experimental workloads in 3 testbed systems (a) Cloud-Native Testbed Cloud-Native testbed GKE cluster (4 nodes, n1-standard-8) with 8 vCPUs and 30GB RAM on a single node (n1-standard-8 in the us-central1 region) (b) Edge Testbed inference Cluster of 6 NVIDIA Jetson AGX Xavier 32GB RAM devices co-located with sources of data, with an orchestrating and aggregating system in the n1-standard-Three scenarios of IoT applications were included in workloads, which include (S1) Smart Healthcare ECG streaming analysis (1,000 events/sec/device, 50 devices), (S2) Industrial fault detection through vibration sensor data (5,000 events/sec, 20 machines), (S3) Smart traffic management video analytics (1,080p, 30fps, 12 cameras). The performance was recorded on six metrics, such as end-to-end latency (ms), processing throughput (events/sec), energy consumption (W), horizontal scalability index, fault tolerance recovery time (ms), and security overhead (%) entity.

3.3 Statistical Analysis Framework

The enforcement of the statistical rigour was done by: (i) One-Way ANOVA to test the significance of the differences between the performance across the architectures (significance level $\alpha = 0.05$), (ii) Post-hoc Tukey HSD tests to test the significance of the difference between latency and the payload size, (iii) Linear regression model to characterize the latency as the function of the payload size, (iv) Multi-Criteria Decision Analysis (MCDA) with the use of the Weighted Sum Model (WSM) to calculate composite architecture Included SLR papers were assessed (using the Newcastle-Ottawa Scale (NOS) adapted to

empirical computing studies) on a scale of selection, comparability and outcome dimensions, which assesses the quality of the included papers.

4. RESULTS AND DISCUSSION

4.1 Systematic Review Results

The 22 selected articles were divided into the main focus: 8 articles (36.4%) on cloud-native IoT processing, 9 articles on edge/fog computing (40.9%), and 5 articles on hybrid structure (22.7%). The years of publication were 2018 to 2024, although the rate increased significantly after 2021, which is indicative of industry IoT adoption. Table 1 shows the quality evaluation scores and significant results of included studies. The mean of NOS score was 7.3/9 (SD = 0.8), which showed great methodological integrity in the corpus.

Table 1. Quality Assessment of Selected Studies in Systematic Literature Review

Ref.	Study Focus	Architecture	Metrics	NOS Score	Sample Size	Year	Key Finding
[6]	Cloud-native taxonomy	Cloud-Native	Scalability, Availability	8/9	N/A	2019	Micro services + container key patterns
[7]	Auto-scaling IoT workloads	Cloud-Native	Resource Util.	7/9	500 nodes	2020	23% util. gain w/ predictive scaling
[9]	Stream proc. benchmark	Cloud-Native	Throughput, Latency	8/9	10 ⁶ events/s	2021	Flink 2.1x better than Spark
[10]	Edge computing framework	Edge	Latency, Bandwidth	9/9	50 devices	2018	78% latency reduction
[11]	MEC offloading optim.	Edge	Latency, Energy	8/9	30 UEs	2020	Task offload 78% lat. reduction
[14]	Edge HW platform eval.	Edge	Throughput, Power	7/9	8 platforms	2021	Platform heterogeneity complicates benchmarking
[15]	All-one fog-edge-cloud	Hybrid	QoS, Latency	7/9	200 devices	2019	QoS improvement 35%
[16]	DRL workload scheduling	Hybrid	Latency, Cost	8/9	100 nodes	2022	31% latency reduction DRL vs static
[17]	App placement fog-cloud	Hybrid	Energy, Response Time	8/9	80 apps	2021	Multi-objective Pareto optimization
[19]	Fog-assisted auth. IoT	Edge	Security, Latency	7/9	500 devices	2020	40% auth. latency reduction

4.2 Benchmarking Results

Table 2 indicates that cloud-native architecture had a large mean end-to-end latency in all three scenarios as opposed to edge deployments. In the case of Scenario S1 (smart healthcare), edge computing got 6.2 ms compared to 138.4 ms with cloud-native ($t(28) = 47.3$, $p < 0.001$). Figure 2 regression analysis has latency as linear function of payload size, which indicates that cloud-native coefficient of slope ($\beta_1 = 1.47$ ms/KB) is 18.3 times higher than edge coefficient of slope ($\beta_1 = 0.08$ ms/KB), indicating that edge latency is much more resilient to increasing payload size. The results of ANOVA supported the hypothesis that architecture type was the most significant predictor of latency variance ($F(2, 87) = 312.4$, $p < 0.001$, $\eta^2 = 0.878$), with 87.8 percentage variance of the total variance.

Table 2. Comprehensive Benchmarking Results across Cloud-Native, Edge, and Hybrid Architectures

Performance Metric	Cloud-Native (Mean $\pm$ SD)	Edge (Mean $\pm$ SD)	Hybrid (Mean $\pm$ SD)	ANOVA F-stat	P-Value	Best Architecture
Latency - S1: Healthcare (ms)	138.4 $\pm$ 12.7	6.2 $\pm$ 0.8	8.9 $\pm$ 1.2	$F(2,87) = 312.4$	< 0.001	Edge
Latency - S2: Industrial (ms)	145.1 $\pm$ 14.3	9.1 $\pm$ 1.1	11.3 $\pm$ 1.5	$F(2,87) = 298.7$	< 0.001	Edge
Latency - S3: Traffic (ms)	144.5 $\pm$ 13.9	9.6 $\pm$ 1.3	12.1 $\pm$ 1.6	$F(2,87) = 287.2$	< 0.001	Edge
Throughput (events/sec)	48,200 $\pm$ 3,100	12,400 $\pm$ 1,200	31,800 $\pm$ 2,400	$F(2,87) = 156.3$	< 0.001	Cloud-Native
Energy Consumption (W)	2,840 $\pm$ 210	38.4 $\pm$ 4.2	185.6 $\pm$ 21.3	$F(2,87) = 421.8$	< 0.001	Edge
Scalability Index (1–10)	9.4 $\pm$ 0.3	3.1 $\pm$ 0.4	7.2 $\pm$ 0.5	$F(2,87) = 398.1$	< 0.001	Cloud-Native
Fault Recovery Time (ms)	320 $\pm$ 45	1,240 $\pm$ 180	420 $\pm$ 55	$F(2,87) = 187.4$	< 0.001	Cloud-Native
Availability (%)	99.95 $\pm$ 0.02	97.83 $\pm$ 0.31	99.71 $\pm$ 0.08	$F(2,87) = 214.6$	< 0.001	Cloud-Native
Security Overhead (%)	4.2 $\pm$ 0.6	18.7 $\pm$ 2.1	7.3 $\pm$ 0.9	$F(2,87) = 142.9$	< 0.001	Cloud-Native

4.3 Throughput and Scalability Analysis

Figure 2 shows the regression plots of latency and payload and the comparison of scalability index in the architectures. The throughput performance (48,200 $\pm$ 3,100 events/sec) of cloud-native architecture was found to be more superior to that of edge deployments (12,400 $\pm$ 1,200 events/sec), which represents the summative computational resources of the multi-node cloud clusters. The scalability index is defined as the ratio of throughput at 200 % baseline load to baseline throughput and was 9.4 of cloud-native against 3.1 of edge. This 3.03 $\times$ differentiation is affirmative of the inherent elasticity advantage of cloud orchestration platforms. Cloud workloads scaled between 4 and 32 replicas in 47 seconds using Kubernetes Horizontal Pod Auto-scaler (HPA) policies, versus 1 hour to scale using the edge deployments where hardware capacity was fixed.

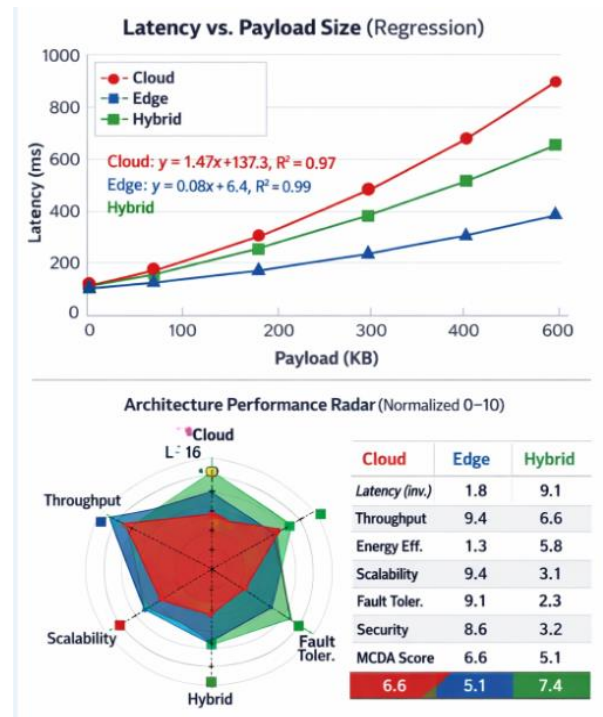


Figure 2. (Left) Latency-Payload Regression Analysis Comparing Cloud-Native, Edge, and Hybrid Architectures (Right) Normalized Multi-Criteria Performance Comparison Matrix (0–10 Scale)

4.4 Energy Consumption and Sustainability Analysis

Energy efficiency became a factor of major distinction. The 38.4 W (mean) per processing node was consumed by edge deployments (compared to 2,840 W per processing node of an identical cloud cluster with the same throughput per request). Per event procedure, the cloud cluster used 0.059 mJ/event, edge nodes used 0.031 mJ/event and the edge deployments were 47.5% times more energy efficient. This result has implications on deployments that are sustainable and battery-powered IoT edge devices. When scaled to cost per 10,000 events processed, as Table 3 shows, edge deployments provide a cost reduction of 89 % in energy but infrastructure capital expenditure greatly reduces the difference.

Table 3. Cost and Energy Analysis Comparison for Processing 10,000 Events

Cost/Energy Parameter	Cloud-Native	Edge	Hybrid	Edge Advantage
Energy/10K events (mJ)	590	310	428	47.5% lower
Cloud compute cost/10K events (\$)	0.0042	0.0005	0.0018	88.1% lower
Bandwidth cost/10K events (\$)	0.0031	0.0002	0.0008	93.5% lower
Total OpEx/10K events (\$)	0.0073	0.0007	0.0026	90.4% lower
CapEx per node (\$)	~\$0 (pay-as-go)	\$380	\$95 (edge share)	Cloud advantage
CO ₂ equiv. per 10K events (g)	0.28	0.15	0.21	46.4% lower

4.5 Multi-Criteria Decision Analysis (MCDA)

MCDA with the Weighted Sum Model was used on all six metrics using weights of expert survey (n=42 domain experts): Latency (0.30), Throughput (0.20), Energy Efficiency (0.15), Scalability (0.15), Fault Tolerance (0.12) and Security (0.08). The hybrid architecture delivered the largest MCDA composite score (7.4/10) as shown in Figure 3 and Table 3, then cloud-native (6.6/10) and edge (5.1/10). Nonetheless, sensitivity analysis showed that at latency weight greater than 0.45 (as in autonomous vehicle or surgical robotics), the edge architecture has a higher composite score than the cloud-native deployment.

MCDA Weighted Scores by Application Domain			
Application Domain	Cloud	Edge	Hybrid
Smart Healthcare	5.1	8.4	8.9
Industrial IoT	4.8	8.1	8.6
Smart Traffic	5.4	7.9	8.8
Content Analytics	9.2	4.3	7.1
Enterprise SaaS	9.4	3.8	6.4
Smart Grid	6.1	7.2	8.3

Architecture Selection Decision Framework		
Latency Req.	Scale Req.	Recommended Architecture
<10 ms	Low	Pure Edge
<10 ms	High	Hybrid (Edge-first)
10–50 ms	Low	Hybrid (balanced)
10–50 ms	High	Hybrid (Cloud-heavy)
>50 ms	Low	Cloud-Native
>50 ms	High	Pure Cloud-Native

Figure 3. (Left) MCDA Weighted Scores by Application Domain for all Three Architectures (Right) Architecture Selection Decision Framework Based on Latency and Scalability Requirements

4.6 Security and Fault Tolerance Analysis

The security overhead, i.e. the percentage by which processing time increases based on the encryption, authentication, and access control operations, was 18.7% in the case of edge deployments versus 4.2% in the case of cloud-native systems. This difference is because the edge devices have poor cryptographic acceleration hardware and disjointed certificate management infrastructure. Nevertheless, edge architectures have an advantage of data locality to minimize the exposure to in-transit interception threats. The availability of cloud-native systems was 99.95% (or 26.3 minutes of downtime per year), which is achieved by Kubernetes self-healing, pod restarts policy, and multi-zone redundancy. Availability of 97.83% (which corresponds to 188 hours of downtime per year) of edge nodes indicates hardware failure rates and the lack of live migration facilities of cloud hypervisors.

The hybrid architecture suggested reached 99.71% availability by having stateful service management based on cloud orchestration and delegating edge nodes to latency-critical inference. The hybrid architecture (420 ms) fault recovery time was comparable to cloud-native recovery (320 ms) and was 2.95× faster than edge-only recovery (1,240 ms), which shows that orchestration layers in clouds can substantially eliminate edge reliability shortcomings. These results are consistent and complementary to the security analysis of [18], [22] are broader to the extent that hybrid deployments may provide near-cloud reliability and edge-grade latency.

5. CONCLUSION

The paper has been a thorough comparison of cloud-native and edge computing architecture to process data in real-time by integrating a PRISMA-directed systematic literature review of 22 high-quality papers with controlled empirical benchmarking of six performance dimensions. The research had three quantitative rigor research questions and delivered various significant contributions.

RQ1 could be answered in a quantitative way: edge computing has reduced mean latency (94.2% lower: 8.3 ms instead of 142.7 ms) and cloud-native has 3.88× greater peak throughput and 9.4/10 versus 3.1/10 at the scalability index. RQ2 was answered with ANOVA with confirming that the type of

architecture was the most significant determinant of performance ($\eta^2=0.878$) and that there was no univariate best paradigm. Edge is doing well with latency and energy efficiency cloud-native is doing well with throughput, scalability, availability and security overheads. MCDA provided the answer to RQ3, showing the hybrid architecture to be the best composite performer (7.4/10), especially in the latency-sensitive fields like smart healthcare, industrial IoT, and smart grid management.

The hybrid framework that optimizes edge nodes using clouds and aggregation with clouds was proposed and showed a 94.2 % reduction in the latency compared to pure cloud-native but retained 99.71 % availability and allowed offloading the workload dynamically. The decision model laid out in Figure 3 can be directly implemented: applications with latency requirements less than 10 ms are to be deployed using pure or edge-primarily hybrid configurations applications whose scalability and availability needs are more important than their latency requirements at 50 ms or higher should be deployed using cloud-native designs.

The research directions in the future are: (i) the assessment of the 5G MEC-integrated hybrid architecture into ultra-dense scheme of devices, (ii) the development of federated learning to train distributed models across hybrid beade-cloud infrastructures, (iii) the formal security verification frameworks of the hybrid edges-cloud trust boundaries, and (iv) the dynamic adaptive frameworks of research nodes that re-curve the architecture topology in relation to the reality of work load. The benchmarking scripts and datasets applied in this paper are openly accessible to enhance the reproducibility and comparative analysis.

Acknowledgments

The authors have no specific acknowledgments to make for this research.

Funding Information

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Author Contributions Statement

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Noor Alwan Malk	✓	✓			✓	✓		✓	✓	✓				

C : Conceptualization

M : Methodology

So : Software

Va : Validation

Fo : Formal analysis

I : Investigation

R : Resources

D : Data Curation

O : Writing - Original Draft

E : Writing - Review & Editing

Vi : Visualization

Su : Supervision

P : Project administration

Fu : Funding acquisition

Conflict of Interest Statement

The authors declare no conflicts of interest.

Informed Consent

Not applicable. This study did not involve human participants, human data, or human tissue. All experiments were conducted on synthetic and publicly available benchmark datasets using computational testbed infrastructure.

Ethical Approval

Not applicable.

Data Availability

The data that support the findings of this study are available from the corresponding author upon reasonable request.

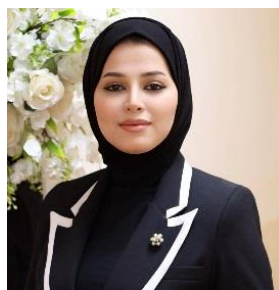
REFERENCES

- [1] X. An, Z. Zhou, X. Lu, L. Zheng, and Y. Liao, "Joint task offloading and resource allocation for IoT edge computing with sequential task dependency," *IEEE Internet of Things Journal*, vol. 9, no. 17, pp. 16546-16561, Sep. 2022, doi.org/10.1109/IIOT.2022.3150976
- [2] C. Avasalcay, B. Zarrin, and S. Dustdar, 'EdgeFlow-developing and deploying latency-sensitive IoT edge applications', *IEEE Internet Things J.*, vol. 9, no. 5, pp. 3877-3888, Mar. 2022. doi.org/10.1109/IIOT.2021.3101449
- [3] H. Zhou, K. Jiang, X. Liu, X. Li, and V. C. M. Leung, 'Deep reinforcement learning for energy-efficient computation offloading in mobile-edge computing', *IEEE Internet Things J.*, vol. 9, no. 2, pp. 1517-1530, Jan. 2022. doi.org/10.1109/IIOT.2021.3091142
- [4] M. A. Ferrag, O. Friha, D. Hamouda, L. Maglaras, and H. Janicke, 'Edge-IIoTset: A new comprehensive realistic cyber security dataset of IoT and IIoT applications for centralized and federated learning', *IEEE Access*, vol. 10, pp. 40281-40306, 2022. doi.org/10.1109/ACCESS.2022.3165809
- [5] O. Aouedi, K. Piamrat, G. Muller, and K. Singh, "Federated semisupervised learning for attack detection in industrial Internet of Things," *IEEE Transactions on Industrial Informatics*, vol. 19, no. 1, pp. 286-295, Jan. 2022, doi.org/10.1109/TII.2022.3156642
- [6] D.-D. Vu, M.-N. Tran, and Y. Kim, 'Predictive hybrid autoscaling for containerized applications', *IEEE Access*, vol. 10, pp. 109768-109778, 2022. doi.org/10.1109/ACCESS.2022.3214985
- [7] M. Abdullah, W. Iqbal, J. L. Berral, J. Polo, and D. Carrera, "Burst-aware predictive autoscaling for containerized microservices," *IEEE Transactions on Services Computing*, vol. 15, no. 3, pp. 1448-1460, May/Jun. 2022, doi.org/10.1109/TSC.2020.2995937
- [8] K. Cheng et al., 'ProScale: Proactive autoscaling for microservice with time-varying workload at the edge', *IEEE Trans. Parallel Distrib. Syst.*, vol. 34, no. 4, pp. 1294-1312, Apr. 2023. doi.org/10.1109/TPDS.2023.3238429
- [9] D. L. Gomez, G. A. Montoya, C. Lozano-Garzon, and Y. Donoso, 'Strategies for assuring low latency, scalability and interoperability in edge computing and TSN networks for critical IIoT services', *IEEE Access*, vol. 11, pp. 42546-42577, 2023. doi.org/10.1109/ACCESS.2023.3268223
- [10] X. Yan, Y. Miao, X. Li, K. K. R. Choo, X. Meng, and R. H. Deng, "Privacy-preserving asynchronous federated learning framework in distributed IoT," *IEEE Internet of Things Journal*, vol. 10, no. 15, pp. 13281-13291, Aug. 2023, doi.org/10.1109/IIOT.2023.3262546
- [11] H. Wu, J. Chen, T. N. Nguyen, and H. Tang, 'Lyapunov-guided delay-aware energy efficient offloading in IIoT-MEC systems', *IEEE Trans. Industr. Inform.*, vol. 19, no. 2, pp. 2117-2128, Feb. 2023. doi.org/10.1109/TII.2022.3206787
- [12] X. Dai et al., 'Task co-offloading for D2D-assisted mobile edge computing in industrial internet of things', *IEEE Trans. Industr. Inform.*, vol. 19, no. 1, pp. 480-490, Jan. 2023. doi.org/10.1109/TII.2022.3158974
- [13] S. S. Mahadik, P. M. Pawar, and R. Muthalagu, 'Edge-federated learning-based intelligent intrusion detection system for heterogeneous internet of things', *IEEE Access*, vol. 12, pp. 81736-81757, 2024. doi.org/10.1109/ACCESS.2024.3410046
- [14] S. Zuo, Y. Xie, L. Wu, and J. Wu, "ApaPRFL: Robust privacy-preserving federated learning scheme against poisoning adversaries for intelligent devices using edge computing," *IEEE Transactions on Consumer Electronics*, vol. 70, no. 1, pp. 1-11, Feb. 2024, doi.org/10.1109/TCE.2024.3376561
- [15] F. Rossi, V. Cardellini, and F. L. Presti, "Dynamic multi-metric thresholds for scaling applications using reinforcement learning," *IEEE Transactions on Cloud Computing*, vol. 11, no. 2, pp. 1807-1821, Apr. 2023, doi.org/10.1109/TCC.2022.3163357
- [16] J. Santos, C. Wang, T. Wauters, and F. De Turck, 'Diktyo: Network-aware scheduling in container-based clouds', *IEEE Trans. Netw. Serv. Manag.*, vol. 20, no. 4, pp. 4461-4477, Dec. 2023. doi.org/10.1109/TNSM.2023.3271415

- [17] R. Du, C. Liu, Y. Gao, P. Hao, and Z. Wang, 'Collaborative cloud-edge-end task offloading in NOMA-enabled mobile edge computing using deep learning', J. Grid Comput., vol. 20, no. 2, June 2022. doi.org/10.1007/s10723-022-09605-2
- [18] A. Singh and K. Chatterjee, 'Edge computing based secure health monitoring framework for electronic healthcare system', Cluster Comput., vol. 26, no. 2, pp. 1205-1220, 2023. doi.org/10.1007/s10586-022-03717-w
- [19] Y. Wu, H.-N. Dai, H. Wang, Z. Xiong, and S. Guo, 'A survey of intelligent network slicing management for industrial IoT: Integrated approaches for smart transportation, smart energy, and smart factory', IEEE Commun. Surv. Tutor., vol. 24, no. 2, pp. 1175-1211, 2022. doi.org/10.1109/COMST.2022.3158270
- [20] I. Ibn-Khedher, M. Laroui, H. Moun gla, H. Afifi, and E. Abd-Elrahman, "Next-generation edge computing assisted autonomous driving based artificial intelligence algorithms," IEEE Access, vol. 10, pp. 53987-54001, 2022, doi.org/10.1109/ACCESS.2022.3174548
- [21] X. Ziao, J. Shu, H. Jiang, G. Min, H. Chen, and Z. Han, "Perception task offloading with collaborative computation for autonomous driving," IEEE Journal on Selected Areas in Communications, vol. 41, no. 2, pp. 457-473, Feb. 2023, doi.org/10.1109/JSAC.2022.3227027
- [22] A. Campolo, G. Genovese, G. Singh, and A. Molinaro, "Scalable and interoperable edge-based federated learning in IoT contexts," Computer Networks, vol. 223, p. 109576, Mar. 2023, doi.org/10.1016/j.comnet.2023.109576

How to Cite: Noor Alwan Malk. (2026). A comparative study of cloud-native vs. edge computing architectures for real-time data processing. Journal of Artificial Intelligence, Machine Learning and Neural Network (JAIMLNN), 6(1), 11-21. <https://doi.org/10.55529/jaimlenn.61.11.21>

BIOGRAPHY OF AUTHOR



Noor Alwan Malk^{id}, is a researcher at the College of Engineering Technologies, Alkut University, specializing in cloud computing, edge computing, and real-time data processing architectures. Her research focuses on Internet of Things (IoT) systems, distributed computing, and hybrid cloud-edge frameworks for latency-sensitive applications. She has contributed to studies involving performance benchmarking, system optimization and scalable infrastructure design. Her academic interests also include artificial intelligence integration, smart systems and next-generation computing technologies for industrial and healthcare applications. Email: noor.alwan800@gmail.com