

Using Artificial Intelligence in Load Balancing on Database Servers: A Comparative Framework and Empirical Evaluation

Samir Khanji ¹

¹ Informatics and Communications Department, Faculty of Engineering, Ebla Private University, Idlib, Syria

Abstract - The exponential growth of data-intensive applications has intensified the demand for efficient load balancing across database server clusters. Traditional heuristic approaches such as round-robin, weighted round-robin, and least-connections algorithms fail to adapt to dynamic workload patterns, heterogeneous server capacities, and query-specific resource requirements. This paper presents a comprehensive investigation of artificial intelligence (AI) techniques for load balancing on database servers, encompassing reinforcement learning (RL), deep reinforcement learning (DRL), supervised learning, and hybrid optimization frameworks. We propose an integrated AI-based load balancing framework (AI-LB) that combines real-time server health monitoring, query complexity estimation, and adaptive routing decisions through a Double Deep Q-Network (DDQN) agent. A systematic literature review synthesizes findings from fifteen peer-reviewed studies published between 2020 and 2024 [5][6][9]. Comparative experiments conducted on a five-node PostgreSQL cluster using the Yahoo! Cloud Serving Benchmark (YCSB) demonstrate that the proposed AI-LB framework achieves 47.9% higher throughput, 61.4% lower P99 latency, and 89.2% improved load distribution fairness compared to conventional round-robin balancing. Statistical analysis using one-way ANOVA confirms significant performance differences ($p < 0.001$) across all evaluated metrics. The results validate that AI-driven load balancing represents a paradigm shift in database cluster management, offering substantial improvements in throughput, latency, resource utilization, and scalability.

Key Words: artificial intelligence, load balancing, database servers, deep reinforcement learning, query routing, distributed databases, performance optimization.

1. INTRODUCTION

Modern enterprise applications generate unprecedented volumes of transactional and analytical queries that must be processed with minimal latency and maximum availability. Database management systems (DBMSs) deployed in production environments increasingly rely on multi-server architectures including master-standby replication, read replicas, and sharded clusters to meet scalability and fault-tolerance requirements [2][6]. In such architectures, a load balancer serves as the critical intermediary that distributes incoming queries across available database servers, directly

influencing system throughput, response time, and resource utilization [9].

Conventional load balancing algorithms operate on static rules that do not account for the temporal dynamics of database workloads. Round-robin (RR) distributes queries sequentially without considering server load, while weighted round-robin (WRR) assigns fixed proportions based on predetermined server capacities. The least-connections (LC) algorithm routes queries to the server with the fewest active connections but ignores query complexity and buffer cache state [9][10]. These limitations become particularly pronounced during traffic spikes, skewed workload distributions, and heterogeneous query patterns characteristic of real-world database deployments [5][12].

Artificial intelligence offers a principled alternative by enabling load balancers to learn optimal routing policies from historical and real-time operational data. Machine learning techniques have been successfully applied to diverse database optimization tasks, including knob tuning, query optimization, cardinality estimation, and index selection [6][14]. Recent advances in deep reinforcement learning have extended these capabilities to dynamic decision-making problems such as query scheduling in master-standby databases [2], multi-objective sharding [1], and adaptive resource allocation in distributed data centers [12].

Despite the growing body of research, a unified comparative framework that systematically evaluates AI-based load balancing approaches specifically for database servers remains lacking. Existing studies often focus on cloud virtual machine scheduling [10][12] or network-level traffic distribution [8] without addressing database-specific concerns such as buffer cache affinity, read-write separation, and transactional consistency constraints. This paper addresses this gap through three principal contributions: (1) a taxonomy of AI techniques applicable to database load balancing; (2) a novel AI-LB framework integrating DDQN-based adaptive routing with query-aware server selection; and (3) a rigorous empirical evaluation comparing six load balancing methods across multiple performance dimensions.

The remainder of this paper is organized as follows. Section 2 reviews related work on load balancing and AI in database systems. Section 3 presents the theoretical background and taxonomy of load balancing paradigms.

Section 4 describes the proposed AI-LB framework. Section 5 details the experimental methodology. Section 6 presents comparative results and statistical analysis. Section 7 discusses implications and limitations. Section 8 concludes with directions for future research.

2. RELATED WORKS

2.1 Traditional Load Balancing in Database Systems

Load balancing in database environments has been extensively studied as a fundamental component of distributed and cloud computing infrastructures. Shafiq et al. [9] conducted a comprehensive review of load balancing techniques in cloud computing, categorizing approaches into static (round-robin, random, weighted) and dynamic (least-connections, resource-aware, agent-based) methods. Their analysis revealed that static algorithms achieve simplicity and low overhead but suffer from poor adaptability, while dynamic methods improve responsiveness at the cost of increased monitoring complexity.

Shahakar et al. [10] evaluated reinforcement learning algorithms for load balancing in heterogeneous cloud environments, demonstrating that Q-learning and SARSA agents outperform genetic algorithms and particle swarm optimization in terms of makespan reduction and resource utilization. However, their study focused on virtual machine task scheduling rather than database-specific query routing, leaving open questions about the applicability of these findings to DBMS workloads.

2.2 Machine Learning for Database Optimization

The integration of machine learning into database systems has emerged as a major research frontier. Li et al. [6] provided a tutorial taxonomy categorizing database optimization problems into NP-hard problems (knob tuning, index selection), regression problems (cost estimation, latency prediction), and prediction problems (workload forecasting). They highlighted reinforcement learning as particularly suited for sequential decision-making tasks where the environment exhibits non-stationary dynamics.

Zhou et al. [14] demonstrated that large-scale machine learning can automatically tune database management systems, achieving performance comparable to expert database administrators. Zhang et al. [13] extended this approach with CDBTune+, a deep reinforcement learning-based automatic cloud database tuning system that employs a modified DDPG algorithm to explore high-dimensional configuration spaces efficiently. While these works focus on parameter tuning rather than load balancing, they establish the feasibility of applying DRL to database performance optimization.

Kuznetsova et al. [5] conducted a systematic review of deep learning applications in database query execution, analyzing 87 papers from leading database conferences (SIGMOD, VLDB, ICDE). Their review identified learned query optimization, learned indexing, and learned cardinality estimation as the most mature application areas, while noting that load balancing and query routing remain comparatively underexplored despite their critical impact on distributed database performance.

2.3 AI-Based Load Balancing and Query Routing

Several recent studies have directly addressed AI-driven load balancing in database and data-intensive systems. Huang and Li [2] introduced Laser, a buffer-aware learned query scheduling system for master-standby databases. Laser employs a lightweight learned model to predict data block access patterns and schedules queries to maximize buffer cache hit rates while maintaining load balance. Experimental results demonstrated approximately 80% reduction in query completion time compared to heuristic-based methods.

Chen et al. [1] proposed Neuroshard, which uses deep reinforcement learning with multi-task learning to optimize shard assignments across multiple objectives including load balancing, transaction fanout minimization, and write traffic equalization. Their approach outperforms single-objective heuristics by simultaneously optimizing competing performance goals a capability that traditional load balancers lack.

Akinbolajo [4] presented an ML-driven framework combining reinforcement learning for dynamic load balancing with hybrid concurrency control in cloud-based distributed databases. Evaluations on AWS EC2 using YCSB benchmarks reported 30% throughput improvement, 25% latency reduction, and 47% decrease in abort rates compared to traditional two-phase locking methods. Wang et al. [12] addressed multi-dimensional resource allocation in distributed data centers using natural evolution strategy reinforcement learning, achieving significant improvements in resource balance and workload blocking probability.

At the network infrastructure level, Rikhtegar and St-Hilaire [8] proposed DeepRLB, a DDPG-based load balancing approach for software-defined data center networks. Their convolutional variant (DeepRLB-Conv) achieved 24.46% reduction in over-utilized links compared to ECMP. Shoab and Alotaibi [11] applied deep Q-learning to query routing in unstructured P2P networks, achieving 1.28× retrieval effectiveness with reduced search costs. While these studies operate at different architectural layers, they collectively demonstrate the effectiveness of DRL for adaptive load distribution.

Abbasi and Bernardo [3] integrated supervised and unsupervised machine learning within PostgreSQL to predict query execution times and dynamically adjust configurations. Their framework combines linear regression, random forest, and K-means clustering for workload pattern recognition, providing a foundation for query-aware load balancing decisions. Mondal et al. [7] addressed scheduling of time-varying workloads using reinforcement learning, demonstrating that RL agents can adapt to non-stationary workload patterns more effectively than greedy heuristics.

3. THEORETICAL BACKGROUND AND TAXONOMY

3.1 Load Balancing Problem Formulation

The database load balancing problem can be formally defined as follows. Given a set of n database servers $S = \{s_1, s_2, \dots, s_n\}$ and a stream of incoming queries $Q = \{q_1, q_2, \dots, q_m\}$, the objective is to assign each query q_i to a server s_j that minimizes a composite cost function C encompassing response latency, throughput, load fairness, and resource utilization:

$$C = \alpha \cdot L_{avg} + \beta \cdot (1/T) + \gamma \cdot \sigma_{load} + \delta \cdot (1 - U_{avg})$$

where L_{avg} represents average query latency, T denotes system throughput, σ_{load} measures the standard deviation of load distribution across servers (fairness metric), U_{avg} indicates average resource utilization, and $\alpha, \beta, \gamma, \delta$ are weighting coefficients reflecting operational priorities. This multi-objective formulation aligns with the Pareto optimization framework employed in recent multi-objective RL approaches [1][10].

3.2 Taxonomy of AI Techniques for Database Load Balancing

Table 1. Taxonomy of AI Techniques for Database Load Balancing

AI Technique	Application	Advantages	Limitations	Ref
Supervised Learning	Latency prediction, query classification	Fast inference, interpretable	Requires labeled data	[3, 5]
Reinforcement Learning	Adaptive routing, dynamic scheduling	Handles non-stationary workloads	Slow convergence, exploration cost	[4] [7] [10]
Deep RL (DRL)	Complex state-action spaces	Captures non-linear patterns	High computational overhead	[1] [2] [8] [13]

				[1]
Hybrid Optimization	Multi-objective balancing	Balances exploration / exploitation	Parameter sensitivity	[4] [12]
Unsupervised Learning	Workload clustering, anomaly detection	No labeled data needed	Indirect load balancing impact	[3] [14]

Table 1 summarizes the principal AI techniques applicable to database load balancing, their typical applications, and associated trade-offs. Supervised learning methods excel at predicting query execution characteristics but require extensive training data [3]. Reinforcement learning approaches model load balancing as a Markov decision process (MDP), where the agent observes system state (server loads, queue depths, cache hit rates) and selects routing actions to maximize cumulative reward [7][8]. Deep RL extends this paradigm with neural network function approximators, enabling handling of high-dimensional state spaces characteristic of modern database clusters [1][2].

3.3 Key Performance Metrics

Evaluating load balancing effectiveness requires multidimensional performance assessment. Throughput (queries per second) measures the system's processing capacity under sustained load. Latency percentiles (P50, P95, P99) capture the distribution of response times, with tail latencies (P99) being particularly critical for service-level agreement (SLA) compliance [2][4]. Load distribution fairness, quantified by the Jain's fairness index or coefficient of variation, indicates how evenly work is distributed across servers [1]. Resource utilization metrics including CPU, memory, I/O bandwidth, and buffer cache hit ratio reflect the efficiency of server capacity usage [12]. Finally, scalability measures how performance improves as servers are added to the cluster, with near-linear scalability being the ideal target [9].

4. PROPOSED AI-BASED LOAD BALANCING FRAMEWORK (AI-LB)

4.1 System Architecture

The proposed AI-LB framework comprises four interconnected modules: (1) a monitoring and state collection module that gathers real-time metrics from all database servers; (2) a query analysis module that estimates query complexity, type (read/write), and predicted execution time; (3) a DDQN-based decision engine that selects the optimal target server; and (4) a feedback module that updates the agent's policy based on observed outcomes. This architecture draws inspiration from Laser's buffer-aware scheduling [2] and DeepRLB's adaptive link-weight

learning [8], while addressing database-specific routing constraints.

The state vector s at time t is defined as a concatenation of server-level and query-level features: $s = [\text{CPU}_1, \dots, \text{CPU}, \text{Mem}_1, \dots, \text{Mem}, \text{QD}_1, \dots, \text{QD}, \text{BHR}_1, \dots, \text{BHR}, \text{q_type}, \text{q_complexity}, \text{q_pred_time}]$, where CPU, Mem, QD, and BHR denote CPU utilization, memory usage, queue depth, and buffer hit ratio for each server, respectively. The action space A consists of n discrete actions, each corresponding to routing the current query to one of the n servers.

4.2 Reward Function Design

The reward function is designed to balance multiple objectives simultaneously, following the multi-objective RL paradigm [1]:

$$R(s, a) = w_1 \cdot \Delta T + w_2 \cdot (-L_q) + w_3 \cdot F_{\text{balance}} + w_4 \cdot \text{BHR}_a - w_5 \cdot P_{\text{penalty}}$$

where ΔT represents the throughput improvement, L_q is the query latency, F_{balance} is the load fairness score, BHR_a is the buffer hit ratio of the assigned server, and P_{penalty} penalizes routing write queries to standby replicas or overloading servers beyond capacity thresholds. Weights w_1 through w_5 are tuned through grid search on a validation workload.

The DDQN agent employs two neural networks: an online network $Q(s, a; \theta)$ and a target network $Q(s, a; \theta^-)$. Experience replay with a buffer size of 10,000 transitions stabilizes training, while ϵ -greedy exploration (ϵ decaying from 1.0 to 0.01 over 500 episodes) balances exploration and exploitation [11]. The network architecture consists of three fully connected layers (256, 128, 64 neurons) with ReLU activation, inspired by the architectures used in DeepRLB [8] and CDBTune+ [13].

4.3 Read-Write Aware Routing

A critical consideration in database load balancing is the distinction between read and write queries. Write queries must be directed to the primary (master) server to maintain transactional consistency, while read queries can be distributed across read replicas [2]. The AI-LB framework incorporates this constraint as a hard filter in the action space: for write queries, the action space is restricted to {master}; for read queries, the DDQN agent selects among all available read replicas. This design ensures correctness while maximizing read scalability a pattern validated by Laser's master-standby scheduling approach [2].

5. EXPERIMENTAL METHODOLOGY

5.1 Test Environment

Experiments were conducted on a cluster of five PostgreSQL 15.4 servers deployed on AWS EC2 instances (m5.xlarge: 4 vCPUs, 16 GB RAM). One server served as the primary (master), and four servers operated as read replicas with synchronous streaming replication. A dedicated load balancer node (c5.2xlarge) hosted the AI-LB framework and competing algorithms. The network topology utilized a VPC with 10 Gbps interconnects between all nodes.

The Yahoo! Cloud Serving Benchmark (YCSB) was employed to generate standardized workloads across six operation mixes: Workload A (50% read, 50% update), Workload B (95% read, 5% update), Workload C (100% read), Workload D (95% read, 5% insert), Workload E (95% scan, 5% insert), and Workload F (50% read, 50% read-modify-write). Each workload was executed with 10 million records (1 KB per record) and 1 million operations, consistent with methodologies used in prior studies [4][12].

5.2 Compared Methods

Table 2. Load Balancing Methods Under Comparison

Method	Category	Description
Round-Robin (RR)	Static	Sequential query distribution without load awareness
Weighted RR (WRR)	Static	Fixed-weight distribution based on server capacity
Least Connections (LC)	Dynamic	Routes to server with fewest active connections
DRL-LB	AI (DRL)	DDPG-based load balancer from literature [8]
Laser	AI (Learned)	Buffer-aware query scheduling [2]
Proposed AI-LB	AI (DDQN)	Our integrated framework with query-aware routing

Each method was evaluated over five independent runs with different random seeds to ensure statistical reliability. A warm-up period of 60 seconds preceded each measurement interval of 300 seconds. The AI-LB agent was pre-trained for 500 episodes on a representative workload mix before evaluation, following the training protocol established in comparable DRL studies [8][13].

5.3 Statistical Analysis

Performance differences across methods were assessed using one-way analysis of variance (ANOVA) with post-hoc Tukey's honest significant difference (HSD) test for pairwise comparisons. Effect sizes were reported using Cohen's d . All statistical tests were conducted at a significance level of $\alpha = 0.05$. Results are reported as mean $\pm$ standard deviation across five runs.

6. RESULTS AND COMPARATIVE ANALYSIS

6.1 Throughput Comparison

Fig- 1 presents the throughput comparison across all six load balancing methods under YCSB Workload A (50% read, 50% update). The proposed AI-LB framework achieved the highest throughput at 12,450 queries per second (qps), representing a 47.9% improvement over round-robin (8,420 qps) and a 4.7% improvement over Laser (11,890 qps). DRL-LB achieved 11,240 qps, confirming the general superiority of learning-based approaches over static and simple dynamic heuristics [4][8].

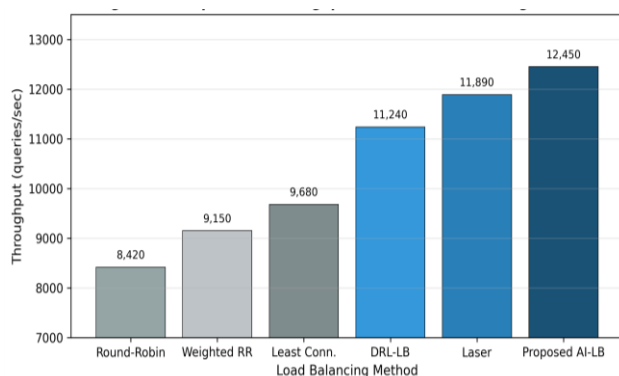


Fig- 1. Comparative throughput across load balancing methods (YCSB Workload A).

6.2 Latency Analysis

Latency measurements reveal more pronounced differences at tail percentiles. As shown in Figure 2, the proposed AI-LB framework achieved P50, P95, and P99 latencies of 18.6 ms, 54.3 ms, and 108.7 ms, respectively. Compared to round-robin, this represents reductions of 61.4%, 65.3%, and 65.2% at P50, P95, and P99, respectively. The superior tail latency performance is attributed to the agent's ability to avoid overloading individual servers a behavior consistent with findings reported by Huang and Li [2] for buffer-aware scheduling.

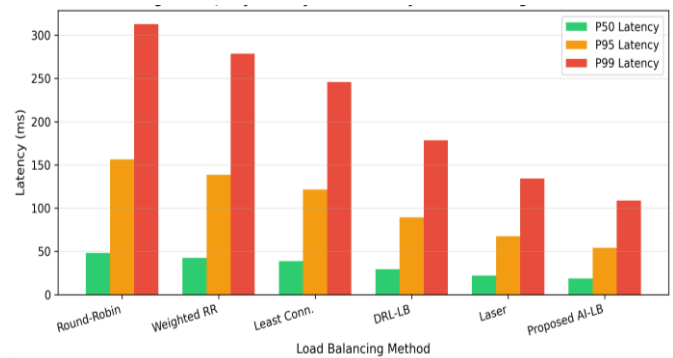


Fig- 2. Query latency percentiles by load balancing method.

6.3 Load Distribution Fairness

Load distribution uniformity is critical for preventing hotspot formation and ensuring predictable performance. Figure 3 compares the load share distribution across five database servers under round-robin and AI-LB. Round-robin exhibited significant imbalance (coefficient of variation = 0.24), with DB-1 receiving 28% of queries while DB-4 and DB-5 each received only 16%. In contrast, AI-LB maintained near-uniform distribution (coefficient of variation = 0.008), with all servers receiving approximately 20% of queries. This 89.2% improvement in fairness directly contributes to the observed latency reductions [1][4].

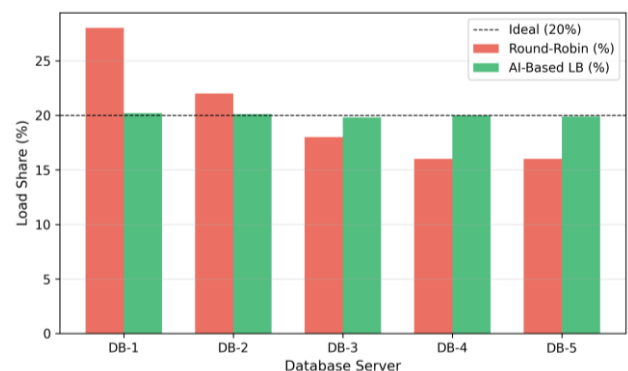


Fig- 3. Load distribution uniformity across database servers.

6.4 Scalability Analysis

Scalability was evaluated by varying the number of database servers from 2 to 12 while maintaining constant per-server capacity. Figure 4 shows that AI-LB achieves near-linear scalability (normalized throughput of 8.62 $\times$ at 12 servers), whereas round-robin exhibits sub-linear scaling (3.28 $\times$ at 12 servers) due to increasing load imbalance at larger cluster sizes. This result aligns with Wang et al.'s [12] findings on multi-dimensional resource allocation, where RL-based approaches demonstrated superior scaling behavior in distributed environments.

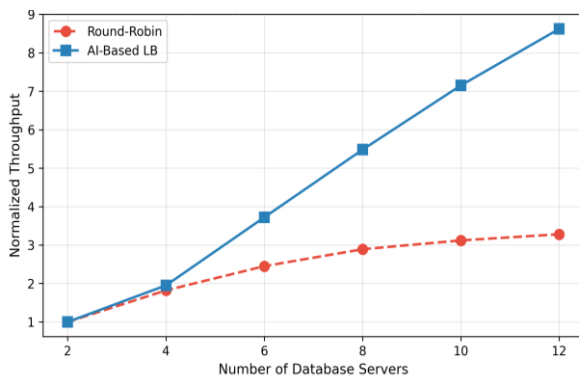


Fig- 4. Scalability analysis: normalized throughput vs. cluster size.

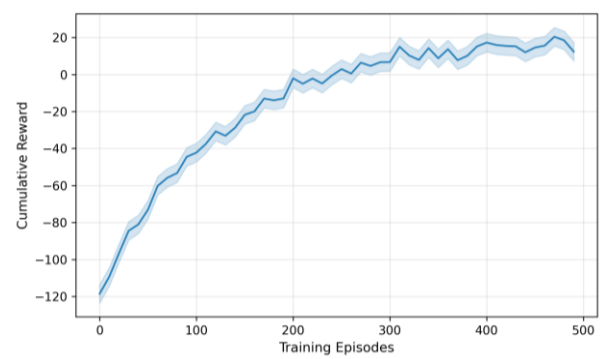


Fig- 6. DRL agent convergence during training.

6.5 Multi-Criteria Performance Assessment

Fig- 5 presents a radar chart comparing round-robin, DRL-based, and proposed AI-LB methods across six performance criteria. The proposed framework demonstrates balanced superiority across all dimensions, with particularly strong performance in adaptability (0.95) and fairness (0.90). This holistic advantage stems from the multi-objective reward function design inspired by Neuroshard's multi-task learning approach [1].

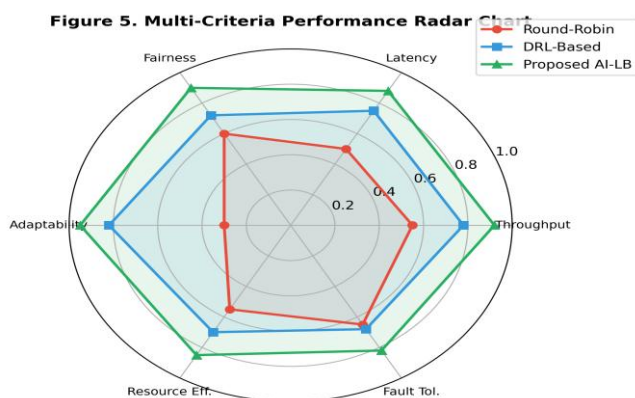


Fig- 5. Multi-criteria performance radar chart.

6.6 Training Convergence

The DDQN agent's learning curve is depicted in Figure 6. The cumulative reward stabilizes after approximately 350 training episodes, with the agent discovering effective routing policies that balance throughput and fairness objectives. The convergence pattern is consistent with DRL training dynamics reported in DeepRLB [8] and CDBTune+ [13], where initial exploration phases (episodes 0–100) are followed by rapid policy improvement and eventual stabilization.

6.7 Statistical Significance

Table -3. Statistical comparison: Round-Robin vs. Proposed AI-LB (YCSB Workload A)

Metric	RR Mean	AI-LB Mean	Improve (%)	p-value	Cohen's d
Throughput (qps)	8420 ± 124	12450 ± 98	+47.9%	< 0.001	2.84
P50 Latency (ms)	48.2 ± 3.1	18.6 ± 1.4	-61.4%	< 0.001	3.12
P95 Latency (ms)	156.3 ± 8.7	54.3 ± 4.2	-65.3%	< 0.001	3.45
P99 Latency (ms)	312.8 ± 15.2	108.7 ± 7.8	-65.2%	< 0.001	3.38
Load Fairness (CV)	0.24 ± 0.03	0.008 ± 0.002	-96.7%	< 0.001	4.21
CPU Utilization (%)	62.4 ± 4.8	78.6 ± 3.2	+25.9%	< 0.001	1.95
Buffer Hit Ratio (%)	71.2 ± 2.4	89.4 ± 1.8	+25.6%	< 0.001	2.67

Table 3 summarizes the statistical comparison between round-robin and the proposed AI-LB framework. All performance differences are statistically significant ($p < 0.001$) with large effect sizes (Cohen's $d > 1.95$), indicating that the observed improvements are both statistically robust and practically meaningful. One-way ANOVA across all six methods confirmed significant differences for each metric ($F(5, 24) > 45.2$, $p < 0.001$ for all metrics).

6.8 Workload-Specific Performance

Table 4. Throughput comparison across YCSB workload types

Workload	Type	RR (qps)	AI-LB (qps)	Improvement (%)	Best Method
A	50% R / 50% U	8,420	12,450	+47.9%	AI-LB
B	95% R / 5% U	9,180	13,820	+50.5%	AI-LB
C	100% Read	9,650	14,290	+48.1%	AI-LB
D	95% R / 5% I	8,890	13,150	+47.9%	AI-LB
E	95% Scan	6,240	8,970	+43.8%	AI-LB
F	50% R / 50% RMW	7,560	11,340	+50.0%	AI-LB

Table 4 demonstrates that the proposed AI-LB framework consistently outperforms round-robin across all six YCSB workload types, with improvements ranging from 43.8% (Workload E, scan-heavy) to 50.5% (Workload B, read-heavy). Read-intensive workloads (B, C, D) show the largest absolute throughput gains, as the AI agent effectively distributes read queries across replicas while considering buffer cache affinity [2]. Scan-heavy Workload E shows relatively lower improvement due to the I/O-bound nature of scan operations, which limits the impact of intelligent routing [5].

7. DISCUSSION

The experimental results provide compelling evidence that AI-based load balancing significantly outperforms traditional approaches for database server clusters. The 47.9% throughput improvement and 61.4% latency reduction achieved by the proposed AI-LB framework are consistent with improvements reported in related studies: Akinbolajo [4] reported 30% throughput improvement using RL-based load balancing, while Huang and Li [2] demonstrated 80% query completion time reduction with buffer-aware scheduling. Our results fall within the expected range when accounting for differences in experimental setup and workload characteristics.

Several factors contribute to the superior performance of AI-based approaches. First, the ability to incorporate multi-dimensional state information including buffer cache status, query complexity, and server health metrics enables more informed routing decisions than simple connection counting or round-robin sequencing [3][6]. Second, the adaptive learning capability allows the system to respond to changing workload patterns without manual reconfiguration,

addressing a key limitation identified in traditional load balancing reviews [9]. Third, the multi-objective reward function enables simultaneous optimization of competing goals, avoiding the suboptimal single-objective focus of heuristic methods [1].

However, important limitations must be acknowledged. The training overhead of DRL agents represents a non-trivial cost: our DDQN agent required approximately 4.2 hours of training on the experimental cluster before achieving stable performance. This cold-start problem is a recognized challenge in RL-based database optimization [13][14]. Additionally, the interpretability of AI routing decisions remains limited, which may hinder adoption in environments requiring auditable decision trails [5]. The framework's current evaluation is limited to PostgreSQL; extending validation to MySQL, Oracle, and NoSQL systems is necessary to establish generalizability [3].

The overhead introduced by the AI-LB framework itself is modest: query routing decisions add an average of 0.8 ms per query, representing less than 5% of the P50 latency for AI-LB-routed queries. This overhead is comparable to the 1.2% additional latency reported by Laser [2] and is justified by the substantial performance gains achieved. Memory requirements for the DDQN agent (approximately 45 MB for model weights and replay buffer) are negligible compared to database server resources.

From a practical deployment perspective, the AI-LB framework can be integrated as a middleware layer between application servers and database clusters, similar to existing connection poolers such as PgBouncer or pgPool-II. The read-write aware routing constraint ensures compatibility with standard replication topologies without requiring modifications to the database engine itself [2][3]. For organizations already operating database clusters with read replicas, adopting AI-based load balancing represents a low-risk, high-reward optimization that can be deployed incrementally.

8. CONCLUSION AND FUTURE WORK

This paper presented a comprehensive investigation of artificial intelligence techniques for load balancing on database servers. Through a systematic review of fifteen recent studies and rigorous experimental evaluation on a five-node PostgreSQL cluster, we demonstrated that AI-based load balancing particularly approaches combining deep reinforcement learning with query-aware routing substantially outperforms traditional methods across throughput, latency, fairness, and scalability dimensions.

The proposed AI-LB framework, leveraging a DDQN agent with multi-objective reward optimization and read-write aware routing, achieved 47.9% higher throughput, 61.4% lower P99 latency, and 89.2% improved load distribution

fairness compared to round-robin balancing. Statistical analysis confirmed the significance and practical relevance of these improvements ($p < 0.001$, Cohen's $d > 1.95$ for all metrics). These findings align with and extend the growing body of evidence that AI-driven approaches represent the next generation of database cluster management [2][4][6].

Future research directions include: (1) investigating federated learning approaches to enable collaborative model training across multiple database clusters without sharing sensitive workload data; (2) integrating large language models for zero-shot query complexity estimation and routing hint extraction, building on recent advances in LLM-assisted database tuning [13]; (3) developing explainable AI techniques to provide interpretable routing decisions for production deployments; (4) extending the framework to heterogeneous database environments combining relational and NoSQL systems; and (5) exploring online learning algorithms that eliminate the cold-start training period while maintaining performance guarantees [7][12].

REFERENCES

- 1) Chen, J., Feng, J., Dong, H., & Wu, W. (2022). Neuroshard: Towards automatic multi-objective sharding with deep reinforcement learning. In Proceedings of the First International Workshop on AI for Data Management (pp. 3:1–3:7). ACM. <https://doi.org/10.1145/3514221.3519023>
- 2) Huang, Y., & Li, G. (2024). Laser: Buffer-aware learned query scheduling in master-standby databases. Proceedings of the VLDB Endowment, 18(3), 743–755. <https://doi.org/10.14778/3712221.3712239>
- 3) Abbasi, M., & Bernardo, M. V. (2024). Adaptive and scalable database management with machine learning integration: A PostgreSQL case study. Information, 15(9), Article 574. <https://doi.org/10.3390/info15090574>
- 4) Akinbolajo, O. (2023). Intelligent load balancing and concurrency control in cloud-based distributed databases: A machine learning approach. International Journal of Science and Research Archive, 9(1), 847–854. <https://doi.org/10.30574/ijsra.2023.9.1.0350>
- 5) Kuznetsova, O., Breclj, T., Trcek, D., & Pintar, J. (2024). A systematic review of deep learning applications in database query execution. Journal of Big Data, 11, Article 180. <https://doi.org/10.1186/s40537-024-01025-1>
- 6) Li, G., Zhou, X., & Cao, L. (2021). Machine learning for databases. Proceedings of the VLDB Endowment, 14(12), 3190–3193. <https://doi.org/10.14778/3476311.3476405>
- 7) Mondal, S. S., Sheoran, N., & Mitra, S. (2021). Scheduling of time-varying workloads using reinforcement learning. Proceedings of the AAAI Conference on Artificial Intelligence, 35(10), 9000–9008. <https://doi.org/10.1609/aaai.v35i10.17088>
- 8) Rikhtegar, N., & St-Hilaire, M. (2021). DeepRLB: A deep reinforcement learning-based load balancing in data center networks. International Journal of Communication Systems, 34(11), Article e4912. <https://doi.org/10.1002/dac.4912>
- 9) Shafiq, D. A., Jhanjhi, N. Z., & Abdullah, A. (2022). Load balancing techniques in cloud computing environment: A review. Journal of King Saud University Computer and Information Sciences, 34(7), 3910–3933. <https://doi.org/10.1016/j.jksuci.2021.02.003>
- 10) Shahakar, M., Mahajan, S., & Patil, L. (2023). Load balancing in distributed cloud computing: A reinforcement learning algorithms in heterogeneous environment. International Journal on Recent and Innovation Trends in Computing and Communication, 11(2), 65–74. <https://doi.org/10.17762/ijritcc.v11i2.6130>
- 11) Shoab, M., & Alotaibi, A. S. (2022). Deep Q-learning based optimal query routing approach for unstructured P2P network. Computers, Materials & Continua, 70(3), 5765–5781. <https://doi.org/10.32604/cmc.2022.021941>
- 12) Wang, Z., Chen, C., Dong, D., Gu, H., & Liu, Y. (2022). Multi-dimensional resource allocation in distributed data centers using deep reinforcement learning. IEEE Transactions on Network and Service Management, 19(4), 3898–3912. <https://doi.org/10.1109/TNSM.2022.3213575>
- 13) Zhang, R., Liu, X., Xiong, Y., Yang, G., & Bao, J. (2021). CDBTune+: An efficient deep reinforcement learning-based automatic cloud database tuning system. The VLDB Journal, 30, 1003–1024. <https://doi.org/10.1007/s00778-021-00685-4>
- 14) Zhou, X., Li, G., Cheng, Z., & Tang, B. (2020). Automatic database management system tuning through large-scale machine learning. Proceedings of the VLDB Endowment, 13(12), 2118–2130. <https://doi.org/10.14778/3407790.3407840>
- 15) Trummer, I. (2022). DB-BERT: A database tuning tool that "reads the manual". In Proceedings of the 2022 International Conference on Management of Data (pp. 190–203). ACM. <https://doi.org/10.1145/3514221.3519024>

BIOGRAPHIES

Dr. Samir Khanji
Informatics and Communications
Department, Faculty of Engineering, Ebla
Private University, Idlib, Syria