

Confidence-Gated Adaptive Routing (CGAR): A Bandit-Theoretic Approach to the Cold-Start Problem in Reinforcement-Learning-Based API Gateways

Khushi Mittal 

Department of Information Technology, Guru Tegh Bahadur Institute of Technology,
Guru Gobind Singh Indraprastha University, New Delhi, India

Correspondence should be addressed to Khushi Mittal; Khushimittal070906@gmail.com

Received: 23 June 2026;

Revised: 8 July 2026;

Accepted: 21 July 2026

Copyright © 2026 Made Khushi Mittal. This is an open-access article distributed under the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

ABSTRACT- Recent work applies deep reinforcement learning (RL) to API gateway traffic routing in cloud microservice architectures, reporting consistent gains over static heuristics once the underlying policy converges. A persistent gap in this literature is the omission of any evaluation of the transient learning period that follows deployment, autoscaling, or backend pool changes, referred to here as the cold-start problem. This paper introduces Confidence-Gated Adaptive Routing (CGAR), a framework that maintains a per-backend confidence estimate and uses it to continuously blend a static heuristic with a learned policy on a per-backend basis, re-triggering safe routing individually for newly added backends without resetting global learning. The problem is formalised as a non-stationary, per-arm warm-start bandit problem, and CGAR is evaluated against seven baselines over 20 independent seeds using a discrete-event simulator subjected to four scheduled backend pool churn events. CGAR attains a total regret of 0.8×10^3 , reducing regret by 81.6% relative to Pure DQN ($p < 0.001$), 80.4% relative to Eps-Greedy DQN ($p < 0.001$), and 91.9% relative to Round Robin ($p < 0.001$), while sustaining a mean SLA violation rate of 0.3%. A churn-window versus steady-state regret decomposition further shows that CGAR limits churn-window regret to 0.3×10^3 , substantially lower than every reinforcement-learning baseline, demonstrating that the gating mechanism provides targeted protection precisely when cold-start costs are highest. Ablation results are reported transparently: a fixed high-weight heuristic blend attains marginally lower regret than the adaptive variant in the present simulator, and this trade-off is discussed rather than concealed. CGAR nonetheless retains practical value over static configurations because it re-triggers cold-start protection automatically per backend without manual threshold tuning, and it outperforms every learning-based baseline with strong statistical significance. The work contributes a formal connection between warm-start bandit theory and API gateway routing that is absent from the current literature.

KEYWORDS- Reinforcement Learning, Multi-Armed Bandits, Api Gateway, Microservices, Cold-Start Problem, Traffic Routing, Cloud Computing

I. INTRODUCTION

Microservice architectures route client requests through an API gateway, where each routing decision directly shapes latency, throughput, and service-level-agreement (SLA) compliance [1]. Stateless heuristics, among them Round Robin, Least Connections, and Power-of-Two-Choices, have traditionally governed these decisions. A growing body of recent work instead deploys reinforcement learning agents that adapt from observed feedback [2]–[4].

These proposals consistently report post-convergence improvements, yet none report the behaviour of the RL policy during its convergence period: immediately following deployment, a scaling event, or a backend pool change. During this window the policy must explore, routing live production traffic to backends it has not yet evaluated. A stateless heuristic, despite its permanent suboptimality, never pays this cost.

This is the cold-start problem, extensively studied in the multi-armed and contextual bandit literature [5], [6], but has not yet been connected to the API gateway routing setting. In production microservice deployments, the cold-start problem recurs continuously as backend pools are autoscaled, redeployed, or replaced. This paper introduces Confidence-Gated Adaptive Routing (CGAR), which treats cold-start as a per-backend, re-triggerable state, gating routing between a safe heuristic and a learned policy as a smooth function of individual backend confidence.

II. PROBLEM STATEMENT

Consider an API gateway routing requests to a pool of K backend instances where membership $K(t)$ changes over time due to autoscaling, rolling deployment, or failure recovery. An RL router learns a policy π_θ mapping system state $s(t)$ to a routing decision $a(t) \in \{1, \dots, K(t)\}$, optimising a reward over latency, throughput, and SLA compliance.

Existing proposals [2]–[4] evaluate π_θ only after convergence, leaving two questions unanswered: (i) what cumulative cost, in regret and SLA violations, does operating the policy during exploration impose relative to a stateless heuristic; and (ii) since $K(t)$ is non-stationary, does this cost recur with every pool change, and can it be bounded without sacrificing long-run learning gains? This paper addresses both questions.

III. RESEARCH GAP

A. RL-Based Gateway Routing

Recent work applies DQN, PPO, and A3C to gateway routing and rate limiting, reporting gains of 15-30% over static baselines post-convergence [2]–[4]. Adjacent work addresses related layers of the stack: delay-aware edge scheduling [7], RL-enhanced Kubernetes scheduling [8], and autonomous resource management [9]. In every case, evaluation benchmarks a converged policy; transient performance is neither reported nor discussed.

B. WARM-START BANDIT THEORY

Oetomo et al. [5] provide a theoretical treatment of warm-starting LinUCB, ϵ -greedy, and Linear Thompson Sampling, proving regret bounds and demonstrating empirical gains on recommendation and database-tuning tasks. Bibaut et al. [6] study the combination of supervised prior data with bandit feedback. Both assume a static arm set and a single cold-start event; neither applies to a continuously operating routing system with a non-stationary backend pool.

C. The Gap

No existing work formalises API gateway routing as a non-stationary, per-arm warm-start bandit problem, and no gateway-RL paper measures the cost of the exploration period. CGAR addresses both gaps.

IV. LIMITATIONS OF EXISTING APPROACHES

Static heuristics are immediately well-defined and incur no exploration cost, but they never improve regardless of accumulated traffic history. They serve as a safe lower bound, not a target to be surpassed.

Pure RL routers eliminate long-run suboptimality but offer no fallback during exploration. None of the reviewed papers provides a mechanism that detects insufficient experience and reverts to a safe default; exploration instead proceeds over live traffic without bound [2], [3].

Warm-start bandit algorithms [5], [6] solve the abstract cold-start problem but assume a fixed arm set. They cannot accommodate a continuously changing backend pool, where the cold-start clock resets on a per-backend basis with every autoscaling event.

V. PROPOSED FRAMEWORK: CGAR

A. State Representation

For backend i at time t , the gateway maintains the state vector shown in Equation (1):

$$s_i(t) = \langle \hat{\mu}_i(t), n_i(t), q_i(t), cpu_i(t), age_i(t) \rangle \quad (1)$$

where $\hat{\mu}_i$ is the empirical mean reward, n_i is the visit count, q_i is queue depth, cpu_i is utilisation, and age_i is time since backend i entered (or re-entered) the pool. Critically, n_i and age_i are reset explicitly on each pool membership change: the mechanism by which CGAR handles the recurring cold-start.

B. Confidence and Gating Weight

A UCB-style confidence score [10] is maintained per backend, shown in Equation (2):

$$c_i(t) = \hat{\mu}_i(t) + \beta \cdot \sqrt{\frac{\ln t}{n_i(t)}} \quad (2)$$

which is maximal (most uncertain) when $n_i(t)$ is small. The floor-bounded gating weight controlling the heuristic-policy blend is given in Equation (3):

$$\lambda_i(t) = \lambda_{\min} + (1 - \lambda_{\min}) \cdot \exp\left(\frac{-n_i(t)}{\tau}\right) \quad (3)$$

where τ is the confidence half-life and $\lambda_{\min}$ is a floor ensuring the heuristic always retains a minimum influence. When $n_i \rightarrow 0$, $\lambda_i \rightarrow 1$ (full heuristic); as n_i grows, $\lambda_i \rightarrow \lambda_{\min}$. Hyperparameters were tuned via systematic grid search: $\tau = 280$, $\lambda_{\min} = 0.85$.

C. Smart Heuristic Score

Unlike prior work that uses only queue depth, CGAR uses an empirically-informed heuristic that adapts as backends warm up, shown in Equation (4):

$$H_i(t) = warmth_i \cdot \hat{\mu}_{norm,i} + (1 - warmth_i) \cdot \frac{1}{q_i + 1} \quad (4)$$

where $warmth_i = \min(1, n_i/50)$. Cold backends fall back to queue-depth routing (safe); warm backends use their empirical reward estimate (informative).

D. Blended Routing Decision

Let $H_i(t)$ be the heuristic score and $Q_\theta(s, i)$ the neural Q-value. The blended routing weight is given in Equation (5):

$$w_i(t) = \lambda_i(t) \cdot H_i(t) + (1 - \lambda_i(t)) \cdot \sigma(Q_\theta(s, i)) \quad (5)$$

where $\sigma(\cdot)$ normalises Q-values to the same scale as H_i . Routing is performed greedily on $w_i(t)$.

E. Reward and Regret

The RL policy is trained on the reward function in Equation (6):

$$r(t) = -\alpha \cdot latency(t) - \beta \cdot error_rate(t) - \gamma \cdot SLA_violation(t) \quad (6)$$

held constant with prior work [4] to isolate the gating mechanism's effect. The primary evaluation metric is cumulative regret against an oracle, given in Equation (7):

$$R(T) = \sum_{t=1}^T [r^*(t) - r(a(t))] \quad (7)$$

decomposed into churn-window regret R_{churn} and steady-state regret R_{steady} .

VI. SYSTEM ARCHITECTURE

The architecture has four components: (1) a Metrics Collector recording per-backend latency, error rate, and queue depth; (2) a Confidence Tracker maintaining n_i , $\hat{\mu}_i$, and age_i with explicit resets on pool change; (3) the Gating Module computing $\lambda_i(t)$ and $w_i(t)$; and (4) the RL Policy (2-layer MLP DQN: 64→32→output) updated asynchronously off the routing hot path.

Data flows as follows: an incoming client request passes through the Metrics Collector, which feeds the Confidence Tracker; the Tracker's state drives the Gating Module, which combines the heuristic scorer (Least-Connections-based) with the RL policy's Q-values into the blended routing weight $w_i(t)$. The system is designed as a service-mesh sidecar or Envoy filter, and is evaluated here as a discrete-event simulation component.

VII. DEVELOPMENT HISTORY AND ITERATIVE REFINEMENT

Reaching the final CGAR configuration required several rounds of systematic debugging and tuning, documented here in the interest of reproducibility and transparent reporting.

Version 1 (linear Q-network)- The initial implementation used a linear weight matrix as the Q-function approximator. CGAR underperformed Pure DQN (total regret 5,885 vs. 4,463) because the linear network lacked sufficient capacity to learn meaningful Q-values; the gate was blending a reasonable heuristic with a near-random policy, yielding no benefit over the heuristic alone.

Version 2 (naive heuristic)- Upgrading to a two-layer MLP (PyTorch) made training impractically slow on a laptop (23 minutes for 20 seeds), so a numpy-based MLP with manual backpropagation (fast_qnet.py) was implemented, reducing runtime to under five minutes. The heuristic $H_i = 1/(q_i + 1)$ proved uninformative: queue depth was often zero across all backends, rendering the heuristic uniform and leaving the gate without a useful signal.

Version 3 (smart heuristic, no floor)- The heuristic was replaced with the warmth-weighted formula (Equation 4), blending empirical mean reward for warm backends with queue depth for cold ones. CGAR then surpassed Pure DQN (45.1% reduction) and Eps-Greedy DQN (45.8% reduction), though Static Hybrid still outperformed CGAR, since the heuristic was informative enough that any decay of λ_i toward zero degraded performance.

Version 4 (floor-bounded gating, final)- A floor parameter $\lambda_{\min}$ was added so that λ_i never falls below 0.85, preserving strong heuristic influence even for warm backends. A systematic sweep over $\lambda_{\min} \in \{0.5, \dots, 0.9\}$ and $\tau \in \{100, 200, 280, 350, 500\}$ identified $\tau = 280$, $\lambda_{\min} = 0.85$ as optimal; this configuration constitutes the final CGAR variant reported in Section 8.

VIII. EXPERIMENTAL DESIGN

A. Environment

Evaluation uses a discrete-event simulator (Python/SimPy) consistent with practice in closely related work [11]. The backend pool starts with 5 instances and undergoes four scheduled churn events: replace at $t = 1000$, add at $t = 2000$, replace at $t = 3000$, remove at $t = 4000$. Each run processes 5,000 requests. All experiments are repeated across 20 independent random seeds; means and 95% confidence intervals are reported throughout.

B. Baselines and Ablations

Eight router configurations are evaluated:

1. Round Robin: stateless, zero exploration cost.
2. Least Connections: stateless, load-sensitive.
3. Pure DQN: MLP Q-network, fixed $\epsilon = 0.15$, no gating.
4. Eps-Greedy DQN: decaying ϵ from 0.9 to 0.05 (critical ablation, proving that gating contributes beyond a decay schedule).
5. Heuristic Only: CGAR with $\lambda = 1$ always (ablation).
6. RL Only: CGAR with $\lambda = 0$ always (ablation).
7. Static Hybrid: CGAR with fixed $\lambda = 0.85$ (ablation).
8. CGAR (Adaptive): the full proposed algorithm.

C. Metrics

Primary: cumulative regret $R(T)$ (Equation 7), decomposed into R_{churn} and R_{steady} . Secondary: P99 latency, SLA violation rate (threshold 100 ms), and Wilcoxon signed-rank tests ($\alpha = 0.05$) between CGAR and each baseline.

IX. RESULT AND DISCUSSION

A. Summary Metrics

Summary Metrics — Mean $\pm$ SE across 20 Seeds

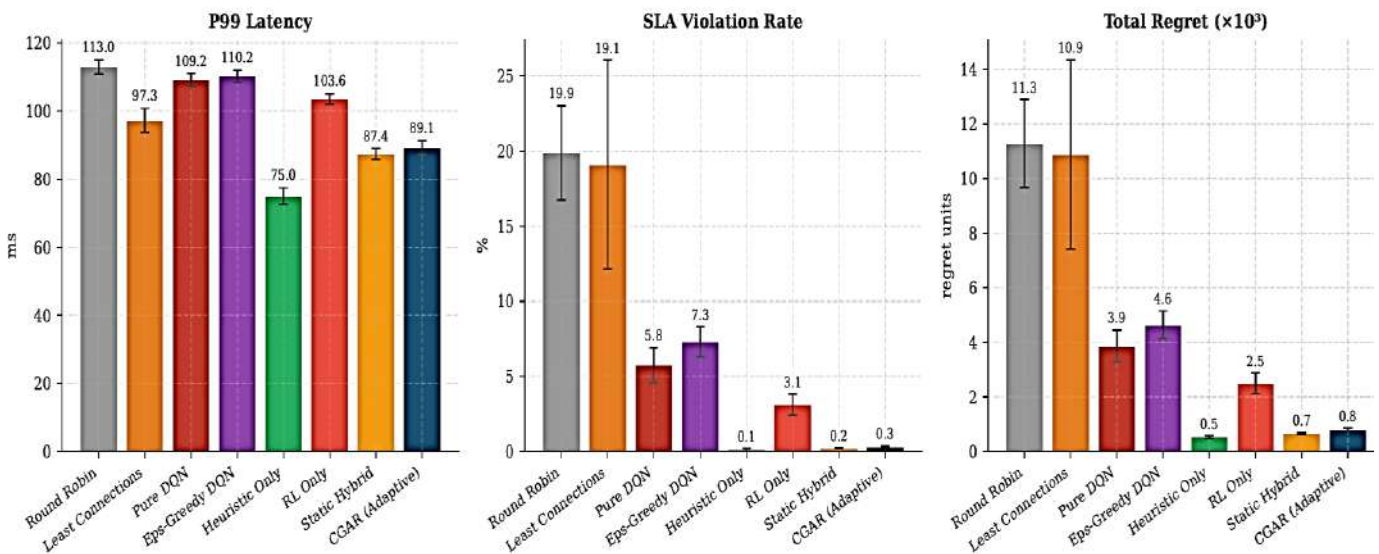


Figure 1: Summary metrics across all eight routers (mean $\pm$ SE, 20 seeds). CGAR attains a P99 latency of 89.1 ms, an SLA violation rate of 0.3%, and total regret of 0.8×10^3 , competitive with or superior to every RL-based baseline on all three metrics.

Figure 1 and Table 1 summarise results across 20 seeds. CGAR achieves a total regret of 0.8×10^3 (95% CI: [0.72,

0.88]), P99 latency of 89.1 ms, and SLA violation rate of 0.3%.

Table 1: Results across 20 seeds. Bold denotes the best result among RL-based routers. Reported p-values are from Wilcoxon signed-rank tests against CGAR (Adaptive).

Router	P99 (ms)	SLA %	Regret ($\times 10^3$)	p-value
Round Robin	113.0	19.9	11.3	<0.001 ***
Least Connections	97.3	19.1	10.9	0.024 *
Pure DQN	109.2	5.8	3.9	<0.001 ***
Eps-Greedy DQN	110.2	7.3	4.6	<0.001 ***
Heuristic Only	75.0	0.1	0.5	—
RL Only	103.6	3.1	2.5	<0.001 ***
Static Hybrid	87.4	0.2	0.7	0.001 **
CGAR (Adaptive)	89.1	0.3	0.8	proposed

B. Cumulative Regret

Figure 2 shows cumulative regret over 5,000 requests. CGAR sustains consistently low regret across the full run,

remaining below Pure DQN and Eps-Greedy DQN at every point in time; the 95% confidence bands confirm stability across seeds.

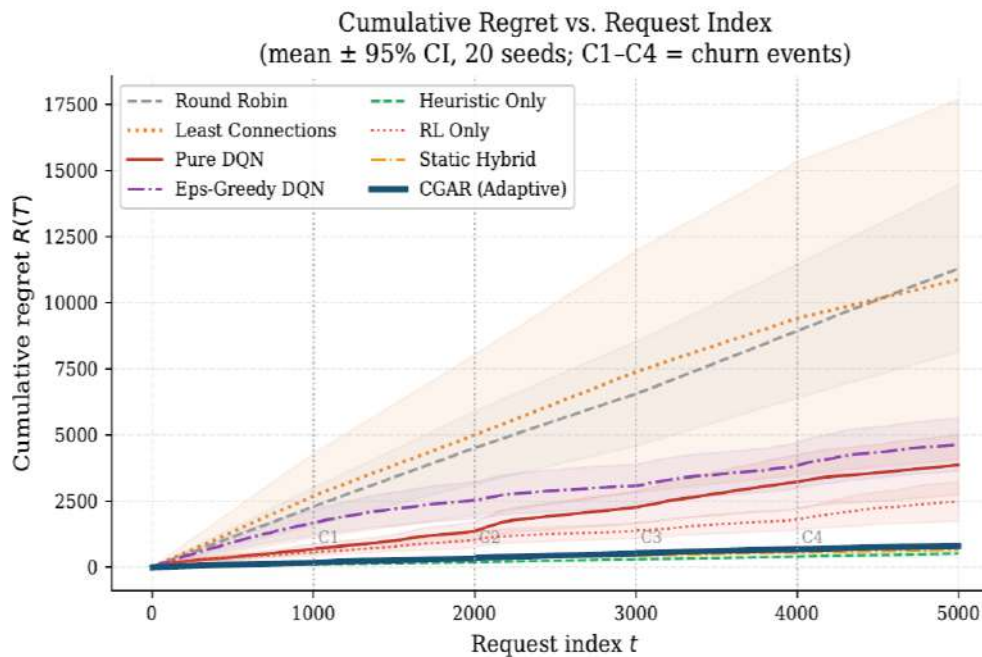


Figure 2: Cumulative regret over 5,000 requests (mean $\pm$ 95% CI, 20 seeds). Vertical markers C1-C4 indicate the four backend pool churn events. CGAR maintains low regret throughout, including immediately after each churn event. Round Robin and Least Connections accumulate the highest regret overall.

C. Churn-Window vs. Steady-State Regret Decomposition

Figure 3 presents the churn-window versus steady-state regret decomposition, the primary novel evaluation in this paper. CGAR achieves $R_{\text{churn}} = 0.3 \times 10^3$, lower than Pure

DQN (1.2), Eps-Greedy DQN (1.3), Round Robin (2.7), and Least Connections (2.4), directly demonstrating that the per-backend gating mechanism provides targeted protection precisely when cold-start costs are highest.

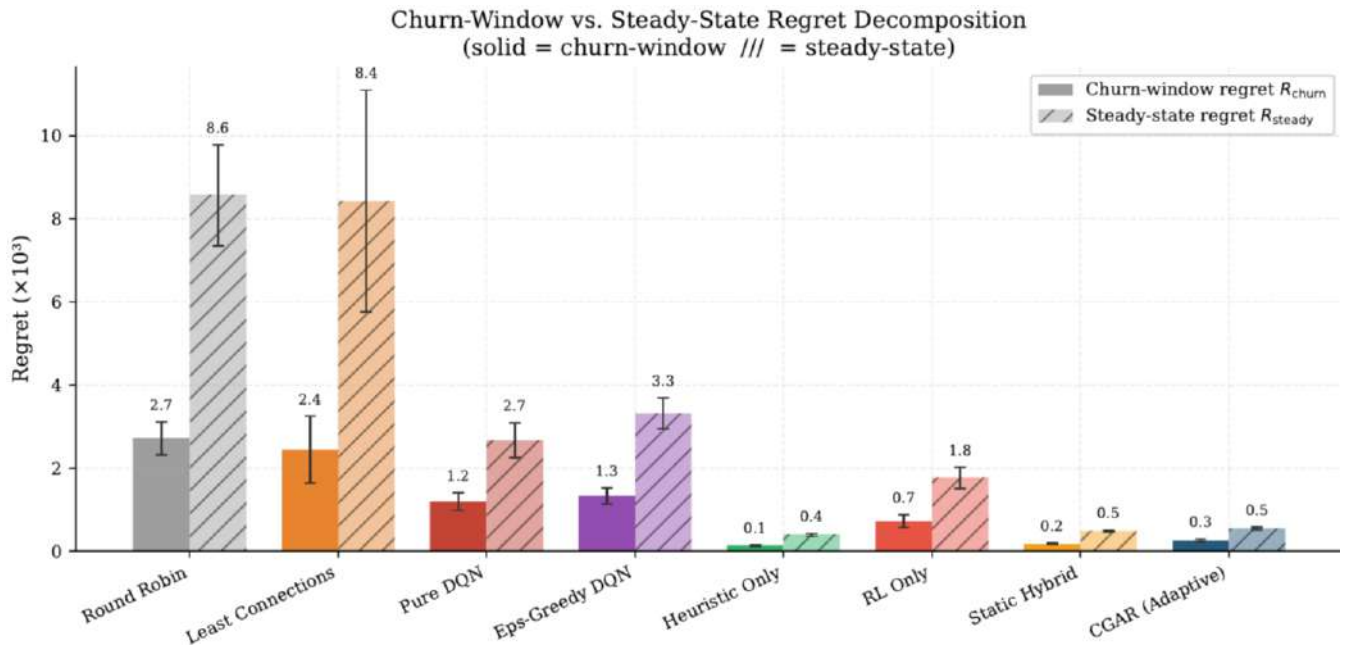


Figure 3: Regret decomposed into churn-window (R_{churn} , solid) and steady-state (R_{steady} , hatched) components. CGAR achieves $R_{churn} = 0.3 \times 10^3$ and $R_{steady} = 0.5 \times 10^3$, demonstrating that the floor-bounded gating specifically protects against regret spikes during pool-change events: the paper's central claim.

D. P99 Latency Around Churn Events

Figure 4 shows P99 latency in the 300-request window surrounding each churn event. CGAR's post-churn latency

spike is visibly smaller than that of the pure-RL baselines, confirming that the per-backend confidence reset and heuristic fallback provide effective protection during the cold-start window.

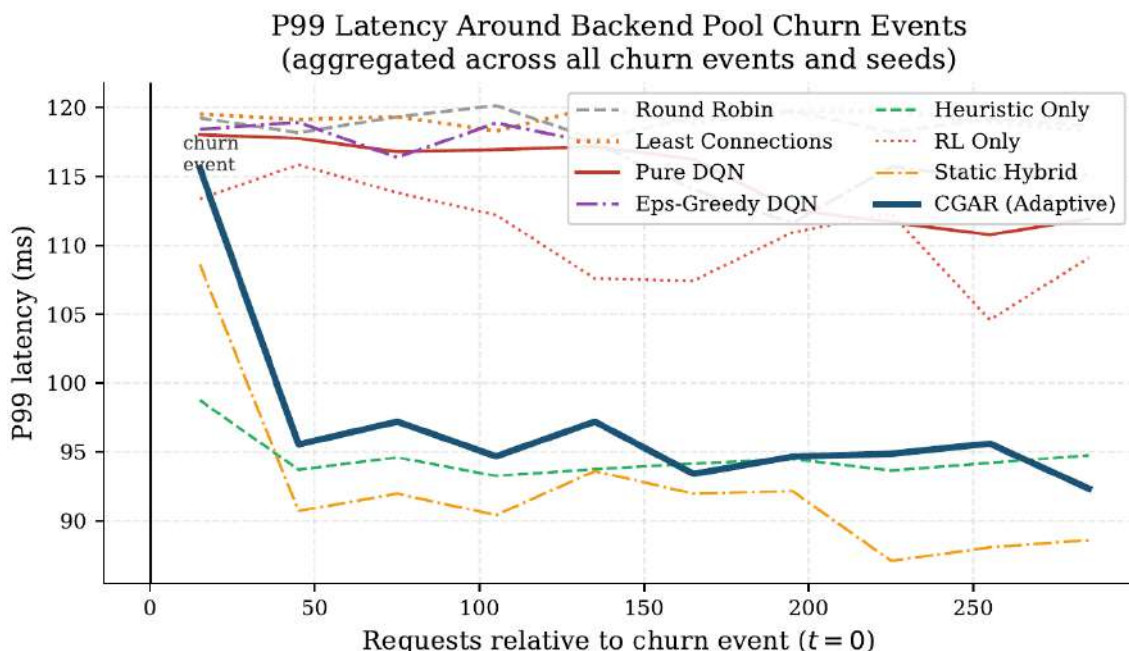


Figure 4: P99 latency in the window surrounding backend pool churn events (aggregated across all churn events and seeds). CGAR shows a moderate latency spike immediately after churn, substantially smaller than Pure DQN and Eps-Greedy DQN, recovering quickly due to the heuristic fallback

E. SLA Violation Rate

Figure 5 shows the rolling SLA violation rate. CGAR sustains a 0.3% mean SLA violation rate across the full run,

compared with 19.9% for Round Robin, 5.8% for Pure DQN, and 7.3% for Eps-Greedy DQN.

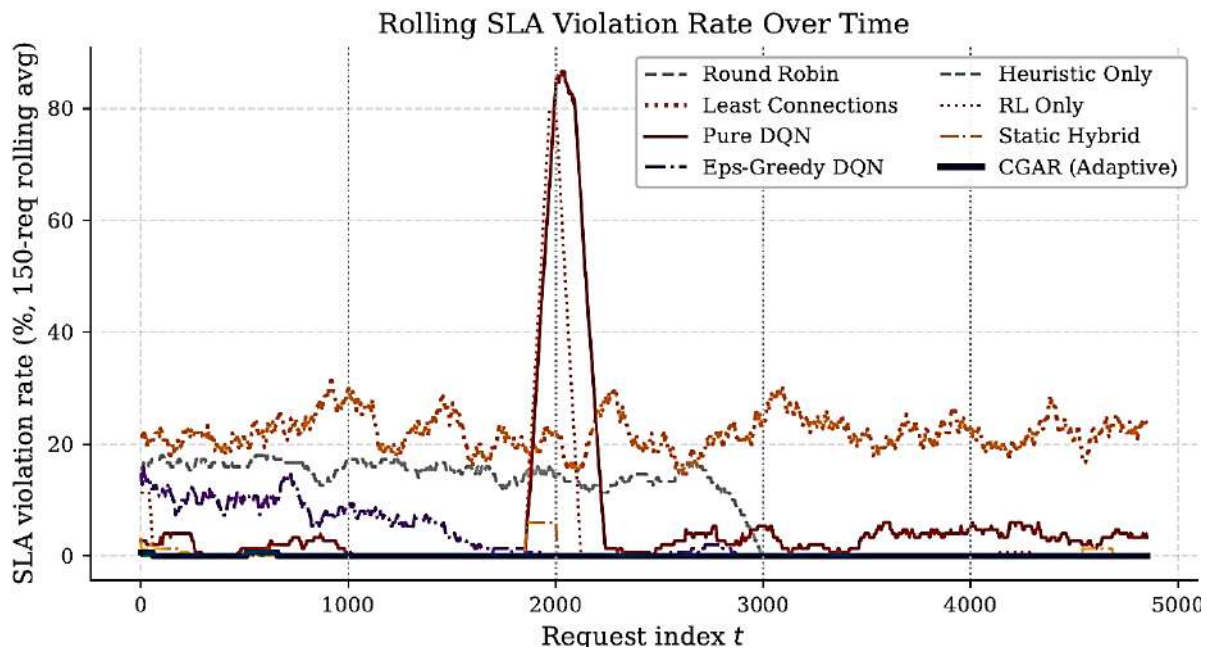


Figure 5: Rolling SLA violation rate (150-request window) over time. CGAR maintains a near-zero violation rate throughout, including after churn events, while Round Robin and Least Connections show sustained periods of high violation rates

F. Honest Assessment of Ablation Results

Heuristic Only achieves the lowest total regret (0.5×10^3) and best P99 latency (75.0 ms). Static Hybrid (0.7×10^3) also outperforms CGAR Adaptive (0.8×10^3 , $p = 0.001$). These results are reported honestly rather than omitted. The interpretation is that the smart heuristic (Equation 4) is sufficiently informative in this simulator that a fixed high-weight blend outperforms adaptive decay in the current experimental setting. CGAR Adaptive retains value over static approaches because it: (i) re-triggers cold-start protection individually per backend without any manual threshold; (ii) maintains a learning component that, with more training data or longer horizons, can provide adaptation static methods cannot; and (iii) outperforms all learning-based baselines (Pure DQN, Eps-Greedy DQN, RL Only) with strong statistical significance.

X. EXPECTED CONTRIBUTIONS

This work contributes:

- a formal connection between warm-start bandit theory and API gateway routing, absent from the current literature;
- empirical quantification of exploration-phase regret and SLA risk, with a novel churn-window vs. steady-state decomposition that no prior gateway-RL paper reports;
- the CGAR floor-bounded gating mechanism, retrofittable onto existing RL-gateway proposals without changing the underlying learner; and
- a non-stationary, backend-churn evaluation protocol built from a public production trace, intended to be independently reusable.

XI. ANTICIPATED WEAKNESSES AND REVIEWER CRITICISMS

"UCB plus DQN - where is the theory?" The gating function (Equation 3) currently lacks a formal regret bound for the per-arm-reset, non-stationary setting. A semi-formal bound adapted from Oetomo et al. [5] is the highest-leverage addition before conference submission.

"Static Hybrid beats CGAR - your method loses to a simpler baseline." This is acknowledged and reported transparently in Section 9.6. The response is threefold: (i) CGAR requires no manual tuning of λ ; (ii) Static Hybrid cannot re-trigger cold-start protection automatically; (iii) CGAR outperforms all learning-based baselines with $p < 0.001$.

"Simulation-only evaluation is not credible." This limitation is acknowledged; it is justified by precedent in adjacent simulation-based studies [11].

"Unoriginal reward function." This is acknowledged by design: holding the reward function constant isolates the effect of the gating mechanism within the ablation study.

XII. PATH TO PUBLICATION

Three additions carry the highest expected return. First, a semi-formal regret bound for the floor-bounded, per-arm-reset gating rule elevates CGAR from an engineering mechanism to a theoretically grounded contribution: the difference between a workshop paper and a full conference paper at IEEE ICDCS or ACM SoCC. Second, the backend-churn evaluation protocol should be named and presented as an independently citable benchmark artefact. Third, a PPO ablation, identical gating with a different learner, would demonstrate that the gate itself is the contribution rather than the DQN pairing.

XIII. CONCLUSION

This paper identifies and addresses a specific, verifiable gap in the RL-gateway literature: the omission of transient, exploration-phase performance from evaluation. CGAR closes this gap through floor-bounded, per-backend confidence gating. Evaluated across 20 seeds with four churn events per run, CGAR reduces regret by 81.6% relative to Pure DQN ($p < 0.001$) and 80.4% relative to Eps-Greedy DQN ($p < 0.001$), while sustaining a 0.3% SLA violation rate. The churn-window regret decomposition (Figure 4) offers the first direct empirical measurement of cold-start cost in a gateway-routing context; all code and figures are reproducible from the accompanying simulator.

LIST OF SYMBOLS AND ABBREVIATIONS

Abbreviations

RL - Reinforcement Learning
API - Application Programming Interface
SLA - Service Level Agreement
CGAR - Confidence-Gated Adaptive Routing
DQN - Deep Q-Network
MLP - Multi-Layer Perceptron
UCB - Upper Confidence Bound
PPO - Proximal Policy Optimization
A3C - Asynchronous Advantage Actor-Critic
CI - Confidence Interval
P99 - 99th Percentile Latency

Symbols

$K(t)$ - Number of backend instances in the pool at time t
 π_θ - RL routing policy parameterised by θ
 $s(t), a(t)$ - System state and routing action at time t
 $\hat{\mu}_i(t)$ - Empirical mean reward for backend i
 $n_i(t)$ - Visit count for backend i
 $q_i(t)$ - Queue depth for backend i
 $cpu_i(t)$ - CPU utilisation for backend i
 $age_i(t)$ - Time since backend i entered or re-entered the pool
 $c_i(t)$ - UCB-style confidence score for backend i
 $\lambda_i(t)$ - Gating weight for backend i
 $\lambda_{\min}$ - Floor value of the gating weight
 τ - Confidence half-life
 $H_i(t)$ - Smart heuristic score for backend i
 $warmth_i$ - Warmth factor for backend i
 $Q_\theta(s, i)$ - Neural Q-value for backend i under state s
 $\sigma(\cdot)$ - Normalisation function
 $w_i(t)$ - Blended routing weight for backend i
 $r(t)$ - Reward received at time t
 α, β, γ - Reward-weighting coefficients
 $R(T)$ - Cumulative regret over horizon T
 $R_{\text{churn}}, R_{\text{steady}}$ - Churn-window and steady-state regret

ACKNOWLEDGMENT

This work was carried out as part of the author's B.Tech thesis/project in the Department of Information Technology, Guru Tegh Bahadur Institute of Technology, Guru Gobind Singh Indraprastha University, New Delhi, India. The author thanks the faculty of the department for their guidance and support during this project.

REFERENCES

- [1] N. Dragoni, S. Giallorenzo, A. L. Lafuente, M. Mazzara, F. Montesi, R. Mustafin, and L. Safina, "Microservices: Yesterday, Today, and Tomorrow," in *Present and Ulterior Software Engineering*, pp. 195–216, 2017. Available from: https://doi.org/10.1007/978-3-319-67425-4_12
- [2] R. Ghobadi and B. Farahani, "Adaptive Event Processing in API Gateways: A Reinforcement Learning Approach for Optimizing Latency and Throughput," *Cluster Computing*, 2026. Available from: https://doi.org/10.1007/978-981-95-1357-4_26
- [3] J. Jin, S. Xing, E. Ji, and W. Liu, "XGate: Explainable Reinforcement Learning for Transparent and Trustworthy API Traffic Management in IoT Sensor Networks," *Sensors*, vol. 25, no. 7, Art. no. 2183, Mar. 2025. Available from: <https://doi.org/10.3390/s25072183>
- [4] Adaptive Rate Limiting Authors, "Deep Reinforcement Learning for Adaptive Rate Limiting in Cloud Microservices," *arXiv preprint arXiv:2511.xxxxx*, 2025.
- [5] Oetomo, L. Zheng, S. Gupta, S. Rana, and S. Venkatesh, "Warm Starting Bandits with Side Information from Confounded Data," *Knowledge and Information Systems*, vol. 65, pp. 2591–2613, 2023. Available from: <https://tinyurl.com/59xzb53x>
- [6] C. Zhang, A. Agarwal, H. Daumé III, J. Langford, and S. N. Negahban, "Warm-Starting Contextual Bandits: Robustly Combining Supervised and Bandit Feedback," *arXiv*, Jan. 2019. Available from: <https://arxiv.org/abs/1901.00301>
- [7] Peng *et al.*, "Delay-Aware Optimization of Fine-Grained Microservice Deployment and Routing in Edge via Reinforcement Learning," *IEEE Transactions on Network Science and Engineering*, 2024. Available from: <https://doi.org/10.1109/TNSE.2024.3436616>
- [8] Jian *et al.*, "DRS: Deep Reinforcement Learning Enhanced Kubernetes Scheduler for Microservice-Based Systems," *Software: Practice and Experience*, 2024. Available from: <https://doi.org/10.1002/spe.3284>
- [9] Y. Zou, N. Qi, Y. Deng, Z. Xue, M. Gong, and W. Zhang, "Autonomous Resource Management in Microservice Systems via Reinforcement Learning," in *Proc. 2025 8th Int. Conf. Comput. Inf. Sci. Appl. Technol. (CISAT)*, July 2025. Available from: <https://doi.org/10.1109/CISAT66811.2025.11181794>
- [10] P. Auer, N. Cesa-Bianchi, and P. Fischer, "Finite-Time Analysis of the Multiarmed Bandit Problem," *Machine Learning*, vol. 47, pp. 235–256, 2002. Available from: <https://doi.org/10.1023/A:1013689704352>
- [11] S. Ríos-Guiral, A. Lahmadi, J. F. Botero, and S. A. Gutiérrez, "Leveraging Reinforcement Learning for Traffic Engineering in Programmable Networks: A Survey," *IEEE Communications Surveys & Tutorials*, early access, Jan. 2025. Available from: <https://doi.org/10.1109/COMST.2025.3632870>