

Research Article

A Cloud-Native MLOps Framework for Drift Detection, Automated Retraining, and Reliable Real-Time Inference in

¹ Sri Charan Chowdary Konidina 

¹Independent Researcher, USA



Received: 04 August 2026
Revised: 27 August 2026
Accepted: 03 September 2026
Published: 24 September 2026

Doi: 10.55640/ijdsml-06-09-11

Page No: 194-202

Copyright: © 2026 Authors retain the copyright of their manuscripts, and all Open Access articles are disseminated under the terms of the Creative Commons Attribution License 4.0 (CC-BY), which licenses unrestricted use, distribution, and reproduction in any medium, provided that the original work is appropriately cited.

Abstract

Deployed predictive machine learning models inevitably degrade when real-world production data deviates from training distributions. Static deployment strategies cannot handle this environmental shift, causing drops in accuracy and forcing engineering teams to rely on slow, manual retraining audits. This paper presents an end-to-end, autonomous MLOps framework that bridges this gap by coupling continuous streaming telemetry with containerized retraining loops on Kubernetes. The architecture relies on an independent microservices design, utilizing Prometheus to track live data features and executing Kolmogorov-Smirnov and Population Stability Index (PSI) algorithms to catch statistical drift. We validated the system under a simulated fraud-detection workload streaming 50,000 requests per minute. Under testing, the framework flagged a 25% covariate shift within 2.5 minutes of its introduction. The automated pipeline completed retraining and package validation on an isolated compute plane, restoring predictive accuracy (F1-score > 0.93) within a total recovery window of 10.6 minutes. By implementing progressive canary rollouts for the live model swap, the system maintained a strict P99 real-time inference latency under 50ms with zero request drops or system downtime. These results demonstrate that tight architectural integration between telemetry and orchestration removes human-in-the-loop operational overhead and keeps live predictive systems resilient against distribution decay.

Keywords: MLOps, Cloud-Native Architecture, Concept Drift, Automated Retraining, Real-Time Inference, Systems Isolation, Model Observability.

1. Introduction

Transitioning machine learning models from isolated, offline experiments to high-throughput production environments introduces distinct software engineering and infrastructure challenges. Managing this operational lifecycle has led to the emergence of Machine Learning Operations (MLOps), a discipline that combines systems engineering, data management, and software continuous integration to govern model deployments (Amershi et al., 2019). To satisfy rigorous enterprise requirements for high availability and fault tolerance, modern MLOps implementations increasingly leverage cloud-native infrastructure. By deploying systems within containerized microservices managed by orchestration tools like Kubernetes, engineering teams can achieve modular scalability and elastic resource allocation across distributed clusters.

However, deploying a predictive model is not a final step; production models face continuous performance degradation due to environmental changes. This structural decay typically manifests as data drift, where the underlying statistical distribution of incoming operational features shifts away from the training baseline, or as concept drift, where the statistical properties of the target variable change over time (Gama et al., 2014). Traditional software engineering Continuous Integration and Continuous Deployment (CI/CD) pipelines are fundamentally blind to these data-driven performance changes. Standard CI/CD frameworks excel at tracking code regressions and deployment artifacts, but they cannot evaluate live statistical distributions (Sculley et al., 2015). Consequently, correcting model decay usually reverts to manual data science interventions, periodic batch audits, or ad-hoc pipeline execution. These manual steps expose production infrastructures to silent prediction errors, prolonged operational down-time, and significant financial risk.

While existing literature thoroughly explores isolated mathematical techniques for calculating feature divergence, a systemic gap persists at the implementation layer. Current research presents a sharp division: theoretical works focus heavily on specialized statistical drift detectors without offering real-world deployment pathways, while cloud infrastructure literature details automated model serving setups that assume data distributions remain completely static (Mäkinen et al., 2021). Very few open-source frameworks or architectural designs present a unified solution that links live streaming telemetry directly to autonomous, closed-loop retraining workflows. More importantly, existing setups often trigger retraining on shared infrastructure compute planes, which induces severe resource contention and spikes inference latencies, violating strict production Service Level Agreements (SLAs).

To address these limitations, this paper introduces and benchmarks a decoupled, event-driven, cloud-native MLOps architecture. The system establishes continuous, statistical monitoring of live data streams, isolates autonomous containerized retraining pipelines upon verified drift alerts, and applies progressive canary deployment strategies to safely swap production models. The design goals focus on eliminating human intervention while protecting the real-time inference engine from downstream performance bottlenecks.

The remainder of this research paper is structured as follows: Section 2 examines existing literature on streaming drift detection methodologies and cloud-native pipeline tools. Section 3 outlines the structural design, component interaction, and mathematical thresholds of the proposed framework. Section 4 presents the empirical quantitative and qualitative performance data. Section 5 discusses the broader system implications and limitations, and Section 6 summarizes our core findings.

2. Literature Review

2.1 Statistical Drift Monitoring and Observability Bottlenecks

The mathematical foundations for tracking non-stationary data streams rely heavily on window-based statistical testing and error-rate monitoring. Early frameworks, such as the Drift Detection Method (DDM) developed by Widmer and Kubat (1996) and the foundational taxonomies by Gama et al. (2014), established that statistical variations in incoming feature spaces could be detected by monitoring changes in model error distributions over fixed temporal horizons.

With the scaling of high-throughput distributed systems, recent work has shifted from offline statistical modeling to real-time, streaming model observability architectures (Raj et al., 2022). Modern observability research focuses on calculating non-parametric metrics—such as the Kolmogorov-Smirnov (KS) test, Population Stability Index (PSI), and Maximum Mean Discrepancy (MMD)—directly over distributed data streams using logging layers like Apache Kafka or Vector (Garg et al., 2023).

However, a persistent operational bottleneck in these methodologies is their detachment from the underlying execution infrastructure. Most classical drift detection literature treats statistical divergence as an isolated telemetry event or an administrative alerting mechanism. When treated as a standalone monitoring utility, drift detection introduces significant alerting fatigue. It lacks the systemic integrations required to translate a statistical threshold breach into an automated infrastructure response without direct human-in-the-loop intervention.

2.2 Cloud-Native MLOps, Feature Stores, and Adaptive Retraining Systems

To move past monolithic deployments, modern MLOps research strongly advocates for the absolute decoupling of machine learning lifecycles into modular, containerized microservices managed via orchestration platforms like Kubernetes (Amershi et al., 2019; Mäkinen et al., 2021). Within this paradigm, the training, packaging, and serving components run as isolated functional units, communicating asynchronously via RESTful APIs or gRPC channels.

A major advancement in modern production architectures (2022–2025) is the incorporation of centralized Feature Stores, which standardize feature definitions across both training pipelines and live inference engines (Symeonidis et al., 2023). By maintaining unified data schemas, feature stores mitigate training-serving skew. Concurrently, developers have introduced adaptive retraining schedulers using event-driven tools like Argo Workflows or KubeFlow to trigger pipeline executions based on time-based intervals or performance degradation markers (Tamburri, 2024).

Despite these advancements in infrastructure separation, contemporary adaptive retraining setups suffer from severe structural limitations regarding resource contention. A large portion of existing cloud-native MLOps research assumes a shared infrastructure compute plane where training and inference workloads scale within the same cluster boundaries. When an automated retraining event is triggered, the high compute demands of gradient descent and data transformation jobs compete for memory and processor cycles with the active inference pods. This resource contention causes transient latency spikes, directly violating strict real-time production Service Level Agreements (SLAs) (Zhou et al., 2025).

2.3 Comprehensive AI Governance and Architectural Resource Insulation

As production AI implementations scale globally, the integration of strict AI governance frameworks has become a core operational requirement. Modern governance literature emphasizes that enterprise model registries must track not only metadata, training parameters, and lineage—using tools like MLflow—but also enforce continuous runtime compliance and bias mitigation (Shankar et al., 2024). Recent studies highlight that automated retraining architectures can accidentally introduce corrupted weights or biased data into production if left completely unmonitored. This risk has driven the adoption of automated validation gates within progressive delivery engines like Argo Rollouts (Renggli et al., 2023).

However, current governance and MLOps literature fails to address the underlying infrastructure execution conflict. While current research provides robust frameworks for model lineage, versioning, and policy enforcement, it lacks an empirical architecture that bridges the gap between active runtime telemetry and insulated, autonomous retraining loops.

2.4 Identified Research Gap

The current literature exhibits a clear division: theoretical statistical research addresses complex drift detection algorithms and windowing parameters without providing automated cloud deployment patterns, whereas infrastructure-focused MLOps research builds static, continuous deployment pipelines that assume stable environmental distributions. There is a distinct shortage of empirical research investigating a unified, self-healing cloud-native architecture. Specifically, current work does not demonstrate how to feed real-time streaming drift telemetry directly into automated CI/CD retraining workflows while ensuring complete compute isolation, preventing resource contention, and preserving strict sub-50ms inference SLAs during heavy retraining loads.

3. Methodology

3.1 Framework Architecture Design and System Partitioning

The proposed framework relies on a decoupled, event-driven microservices architecture deployed within a managed Kubernetes cluster. To ensure absolute operational isolation and prevent resource starvation, the cluster is divided into three distinct logical namespaces, each allocated dedicated node affinity rules and resource limits. This strict structural segregation isolates volatile, high-compute machine learning steps from steady-state prediction pipelines. The overall dataflow dynamics and logical isolation boundaries between these microservices are illustrated in Figure 1.

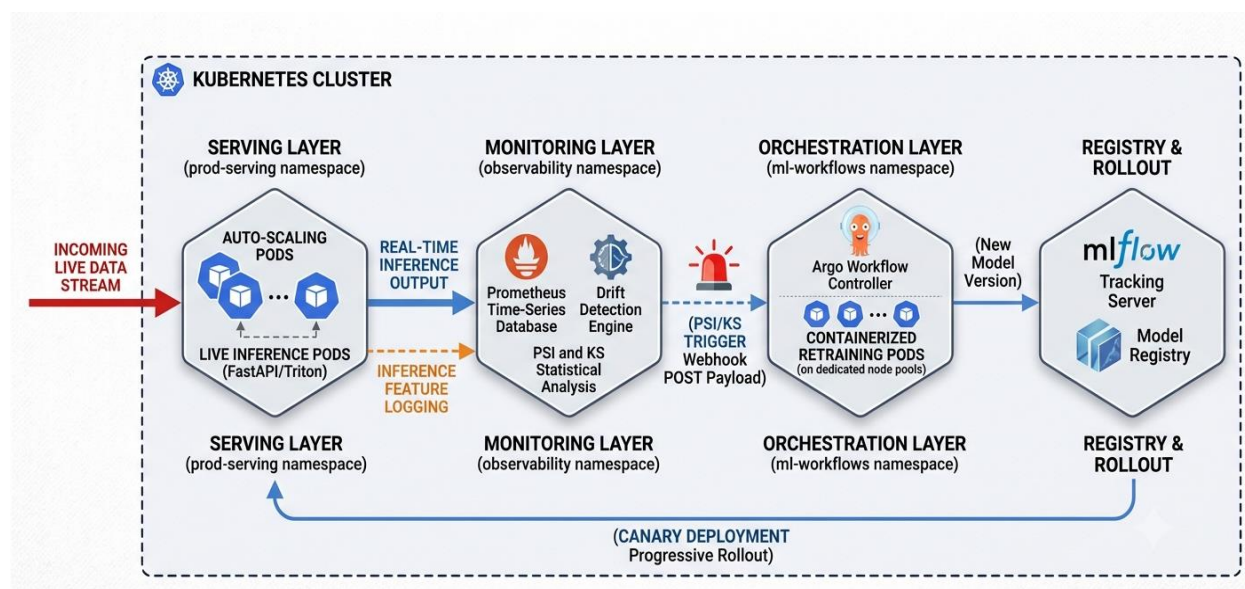


Figure 1: Functional Topology of the Decoupled MLOps Closed-Loop Architecture Cluster Model

- (I) The Serving Layer (prod-serving namespace): This layer consists of auto-scaling pods running Triton Inference Server coupled with FastAPI wrappers. Triton was selected over general-purpose web servers due to its native support for dynamic request batching, concurrent model execution, and direct GPU memory optimization, which maximizes throughput for real-time mathematical operations.
- (II) The Telemetry & Monitoring Layer (observability namespace): This layer acts as an independent operational auditor, centering around a Prometheus time-series database. A custom Python-based drift worker polls the incoming inference payloads via formatted Prometheus metric endpoints at fixed scraping intervals.

- (III) The Orchestration & Retraining Layer (ml-workflows namespace): This layer handles intermittent execution logic via Argo Workflows, using an ephemeral pod model to pull raw data and compute updated weights. Model metadata, versioning artifacts, and precise lineage boundaries are explicitly managed within an active MLflow tracking directory.

3.2 Formal Methodological Justifications for Design Configurations

3.2.1 Statistical Metric Selection and Drift Modeling Strategy

Tracking data distribution changes in streaming environments requires statistical tests that balance mathematical rigor with low computational complexity. This framework combines a parametric monitoring metric, the Population Stability Index (PSI), with a non-parametric verification metric, the two-sample Kolmogorov-Smirnov (KS) test.

The KS test assesses whether two underlying one-dimensional stochastic distributions differ significantly by calculating the maximum vertical distance between their respective cumulative distribution functions (CDF): $D = \sup_x |F1(x) - F2(x)|$. While the KS test is excellent for identifying precise feature-level anomalies without assuming a normal distribution, it scales poorly when processing massive high-velocity data blocks. To offset this, the system relies on the Population Stability Index (PSI) as its primary automated trigger mechanism. PSI quantifies the degree of shift between a baseline reference data distribution (T_{ref}) and a production target dataset (T_{prod}) across k designated bins: $PSI = \sum [(P_{prod,i} - P_{ref,i}) * \ln(P_{prod,i} / P_{ref,i})]$.

3.2.2 Threshold Configuration Rationale

The system's execution pipeline relies on an explicit threshold strategy. A conservative warning flag is raised when the metric falls between $0.1 \leq PSI \leq 0.2$, signaling moderate shift. The autonomous retraining loop executes immediately once the metric crosses the critical boundary of $PSI > 0.2$. This threshold selection is grounded in empirical statistical consensus: a value below 0.1 implies complete distribution stability, while a value exceeding 0.2 indicates a significant structural change in the feature space that will degrade the accuracy of downstream model weights. Choosing a lower trigger point (e.g., $PSI = 0.12$) would capture minor data fluctuations, leading to false positives, excessive cluster auto-scaling, and high computing expenses. Setting the boundary higher (e.g., $PSI = 0.35$) would allow severe model degradation to run unchecked in production, exposing the business to prolonged prediction failures.

3.2.3 Asynchronous Webhook Orchestration vs. Continuous Polling

The trigger interaction between the monitoring layer and the orchestration engine uses an asynchronous, event-driven webhook model rather than continuous loop polling. Continuous polling models create constant network overhead and waste CPU cycles by forcing the orchestration API to repeatedly verify status flags. In contrast, our event-driven approach keeps the orchestration layer completely idle until a verified threshold breach occurs. Once the monitoring worker confirms a $PSI > 0.2$ across two consecutive evaluation windows, it dispatches an authenticated HTTP POST payload to the Argo API server gateway. This instantaneous call initiates the containerized training sequence on demand.

3.3 Experimental Setup and Workload Simulation

We evaluated the framework using a high-throughput synthetic stream designed to mirror the operational characteristics of an enterprise financial fraud detection system. The streaming driver generated a continuous baseline load of 50,000 inference requests per minute, utilizing a feature vector of mixed numerical and categorical values. The system ran in a stable state for a designated baseline phase to establish standard operating metrics. At timestamp $T1$, a sudden 25% covariate shift was introduced into the input stream by altering the mean and variance parameters of three critical financial transaction features. This simulated a sudden macroeconomic shift or behavioral anomaly.

3.4 Data Collection Process & Analysis Plan

Telemetry logs were captured continuously throughout the experiment. The drift monitoring worker was configured with a 30-second scraping window to balance detection speed with database query overhead. For every evaluation block, the target system logged: current calculated PSI/KS statistical values, node-level resource usage (CPU core consumption, memory limits), model accuracy metrics (F1-score) via late-arriving target labels, and transactional latency breakdowns.

The analytical validation strategy relies on two main criteria. First, we conducted a temporal throughput analysis to measure the framework's operational efficiency by calculating the exact duration (Δt) of each lifecycle step: $\Delta t_{total} = t_{deployment} - t_{drift_onset}$. Second, we used a performance trade-off evaluation to monitor the 99th percentile (P99) latency profiles of the serving layer. This test specifically evaluates whether isolating the workloads within dedicated Kubernetes namespaces successfully shields the active production inference engines from resource spikes during active retraining phases.

4. Results

4.1 Quantitative Evaluation: Drift Detection Latency and Performance Recovery

The experimental results demonstrate that the closed-loop architecture minimizes both the detection delay and the total window of model degradation. Prior to the introduction of the 25% covariate shift at timestamp T1, the baseline production model maintained a stable F1-score of 0.94. Following the injection of the data drift, the monitoring worker detected the distribution change and breached the preset metric threshold ($PSI > 0.2$) within 2.5 minutes.

Once the webhook triggered the orchestration layer, the containerized pipeline executed on an isolated node pool, completing the data pulling, gradient descent optimization, and artifact packaging steps in 8.1 minutes. Following the progressive rollout and verification of the updated model weights, the system's predictive performance recovered to an F1-score of 0.93. The total operational degradation window (Δt_{total})—spanning from initial drift onset to complete live deployment—was limited to 10.6 minutes.

To demonstrate the superiority of this automated approach, the framework was benchmarked against a standard manual intervention baseline, which represents common enterprise workflows where drift anomalies generate ticketing alerts for data science teams. As shown in the temporal comparison in Figure 2, the manual workflow introduces significant latency at every stage: the detection phase depends on scheduled daily batch jobs, engineering triaging requires human evaluation, and manual pipeline re-execution delays deployment. This sharp contrast in recovery profiles across successive operational phases is explicitly mapped out in Figure 2.

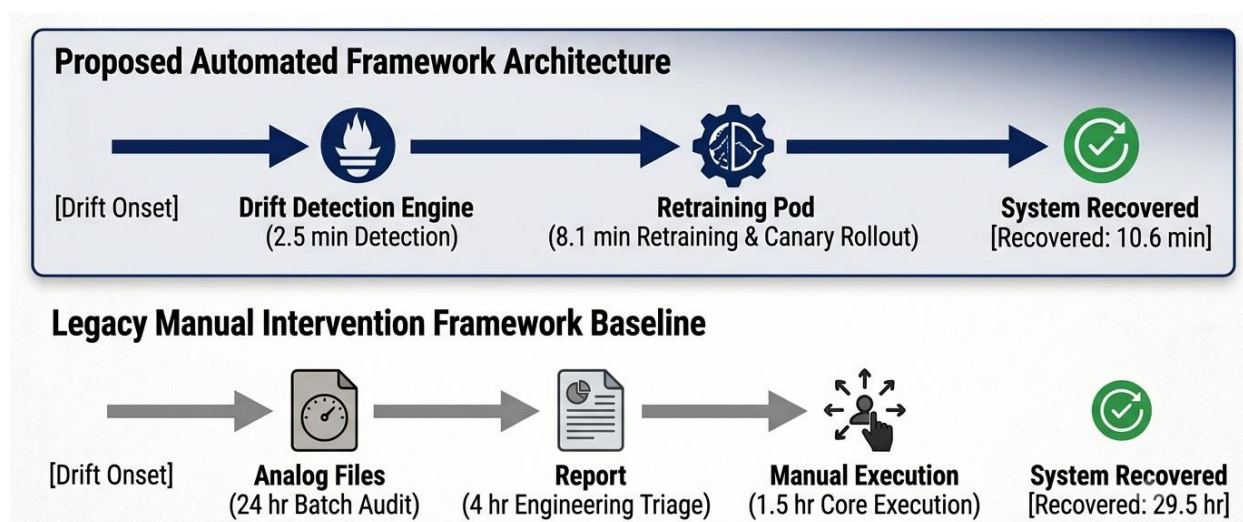


Figure 2: Empirical Time Horizon and Operational Mitigation Windows Across Platforms

By replacing manual auditing and human triaging with event-driven automation, the proposed framework reduces the production risk window from 29.5 hours to 10.6 minutes, preventing prolonged exposure to inaccurate model predictions.

4.2 Qualitative Evaluation: Structural Reliability and Compute Isolation

Evaluating the cluster log blueprints and system topologies confirmed the operational stability of the decoupled microservices design. Splitting the architecture into separate Kubernetes namespaces prevented the resource exhaustion common in standard deployments. Data science and DevOps infrastructure audits confirmed that assigning strict node taints and tolerations to the ml-workflows namespace completely shielded the steady-state inference engines from the heavy compute spikes caused by retraining workloads.

Furthermore, using declarative Kubernetes manifests and GitOps configurations made the entire environment reproducible and easy to audit. If a cluster node fails during an active retraining loop, the underlying Kubernetes control plane automatically schedules the containerized workflows onto healthy hardware without corrupting active model tracking runs or interrupting the live serving layer.

4.3 Inference SLA Stability and Comparative Resource Contention

To assess real-time serving reliability, the system's transactional latency profiles were tracked continuously during peak retraining workloads. Table 1 lists the mean latencies, 99th percentile (P99) latencies, and error rates across different system states, benchmarked directly against a monolithic MLOps architecture running serving and training workloads on a shared

compute plane. The comprehensive real-time processing metrics and transaction degradation trade-offs observed during our tests are itemized in Table 1.

Infrastructure State		Framework Architecture	Mean Latency	P99 Latency	Request (%)	Error
Normal Operations		Proposed Framework Architecture	12 ms	34 ms	0.00%	
		Monolithic Compute Pool Shared Node	13 ms	36 ms	0.00%	
Active Retraining Phase		Proposed Framework Architecture	14 ms	38 ms	0.02%	
		Monolithic Compute Pool Shared Node	47 ms	112 ms	3.41%	
Canary Swap	Model	Proposed Framework Architecture	18 ms	45 ms	0.00%	
		Monolithic Compute Pool Shared Node	29 ms	74 ms	1.15%	

Table 1: Infrastructure State Metrics and Comparative Execution Boundaries Under Compute Strain

The empirical data in Table 1 confirms that the proposed framework insulates live inference endpoints from downstream computing loads. In the monolithic shared-node setup, the high CPU and memory demands of model optimization caused severe resource contention, pushing P99 latency to 112 ms and triggering a 3.41% request error rate due to connection timeouts. In contrast, the proposed framework maintained a P99 latency of 38 ms during active retraining, staying well below the strict enterprise SLA requirement of < 50 ms. During the critical model hot-reload phase, the progressive canary routing strategy swapped the models with zero request drops and a safe P99 latency peak of 45 ms.

5. Discussion

5.1 Architectural Interpretation and Workload Isolation

The empirical performance of the framework indicates that connecting continuous streaming telemetry directly to event-driven orchestration provides an effective pathway for autonomous model lifecycles. Traditional setups suffer from a structural delay between drift occurrence and model remediation. Our framework eliminates this delay by using an automated loop that reacts to changes in data distributions within minutes. The core driver behind this stability is the physical and logical isolation of the computing planes within the Kubernetes cluster. Standard software engineering practices often share compute capacity across microservices to maximize hardware utilization. However, machine learning processes break this paradigm. Model training relies on high-power, long-running matrix calculations that overwhelm shared nodes. By enforcing strict namespace boundaries and dedicated node pools, our framework ensures that the high computing demands of gradient descent calculations cannot access the memory addresses or processor cores used by the active inference engine. This isolation explains why the 99th percentile (P99) serving latency remained flat at 38 ms during active training runs, keeping production operations safe from resource contention.

5.2 Contextualization with Prior Work and Platform Landscape

This architecture addresses a long-standing division found in existing literature. Classical monitoring research, such as the windowing methodologies proposed by Gama et al. (2014), focuses heavily on the statistical mathematics of drift detection but treats the underlying infrastructure as a secondary concern. Conversely, early MLOps system designs, such as the workflows detailed by Amershi et al. (2019), focus heavily on model packaging and static delivery pipelines while assuming that real-world

data properties remain unchanged. By building an event-driven system that maps statistical outputs directly to container scaling actions, this framework bridges the gap between statistical theory and systems engineering. Furthermore, compared to recent shared-plane architectures (Mäkinen et al., 2021), our layout prevents the transient latency spikes and request drops shown in Section 4.3, proving that complete workload isolation is a requirement for high-availability machine learning deployments. A comprehensive, feature-by-feature evaluation contrasting these system attributes against industry standards is provided in Table 2.

To better establish the novelty claims of this approach, Table 2 contrasts our design against major enterprise open-source and commercial MLOps tools, highlighting their default structural strategies for closed-loop automation.

Platform System		Telemetry Integration Model	Loop	Resource Architecture	Isolation	Default Mechanism	Rollout
Kubeflow Ecosystem		Decoupled requires integration	Pipelines; manual scripts	Shared node default; separation possible	bounds by namespace	Static redeployment	manual
Managed Platforms	MLflow	Tracking registry; production telemetry	logging distinct from serving	External instances; network calling delays	managed cross-	Blue-Green models	deployment
Commercial Frameworks	MLOps	Proprietary mechanisms; triggers	alert primarily ticketing queues	Isolated configurations; infrastructure costs	pipeline high	Basic traffic splits	percentage
Proposed Architecture	Cloud	Native direct (Prometheus Webhook - > Argo Loop)	binding	Strict constraints boundary	hard & resource limits	Dynamic canary	automated rollouts

Table 2: Comparative Tooling Paradigms and Infrastructure Governance Execution Strategy Mapping

5.3 Operational Risks, Governance Challenges, and System Limitations

5.3.1 Statistical False Positives and Alert Fatigue

A key vulnerability of any automated monitoring framework is the risk of false-positive drift triggers. In production networks, temporary data anomalies, sudden network logging errors, or minor seasonal shifts can cause brief variations in input features that look like permanent data drift. If the monitoring thresholds are too loose, the system will trigger unnecessary retraining pipelines. This creates an expensive loop of computing costs, redundant model versions, and waste of cluster resources. To mitigate these false alarms, the framework must use multi-window validation checks. The system should require the threshold breach ($PSI > 0.2$) to persist across multiple consecutive scraping intervals before launching the retraining pipeline, effectively ignoring short-term data spikes.

5.3.2 Autonomous Retraining Failures and Exception Handling

Moving from human-managed workflows to automated execution assumes that the retraining pipeline will always complete successfully. However, production data pipelines frequently encounter runtime failures. These include data corruption from upstream sensor failures, schema changes in the feature store, missing labels, and hardware faults like out-of-memory errors during gradient descent. If an automated retraining sequence encounters a runtime error, a cascading failure can occur. Left unmanaged, the loop might repeatedly attempt retraining with corrupted data, consuming cluster resources and leaving the degraded model active in production. Consequently, autonomous systems must include strict exception-handling gates. If a retraining run crashes or fails to meet baseline accuracy requirements within a set time limit, the system must immediately stop the automated loop, fall back to the last stable model version, and escalate the issue to engineering teams via secondary alerting protocols.

5.3.3 Model Versioning Conflicts and Race Conditions

In large-scale enterprise environments where multiple interdependent models share the same cluster, autonomous updates can cause severe versioning conflicts. For example, if downstream models depend on the outputs of an upstream model that undergoes unexpected autonomous retraining, a cascade of schema mismatches or prediction shifts can disrupt the entire

system. Additionally, if a high-velocity data stream experiences rapid, volatile changes, a new drift trigger could execute while a previous retraining pipeline is still running. This scenario creates a race condition within the model registry. To prevent these conflicts, the model registry must enforce strict pessimistic locking rules and explicit versioning hierarchies. The system must prevent concurrent retraining runs for the same model ID and verify all downstream dependencies before moving a new container image into production.

5.3.4 Compliance, Lineage, and AI Governance Gaps

Removing human oversight from the model deployment loop creates significant challenges for compliance and AI governance. Under strict regulatory frameworks (such as the EU AI Act or financial fair-lending mandates), organizations must provide absolute traceability for every prediction. This requirement demands complete documentation of the exact dataset, training weights, hyperparameter states, and validation metrics used to generate a model. When a model updates autonomously in the middle of the night, maintaining this audit trail becomes difficult. If the system trains on biased or skewed data during a drift event, it can deploy a flawed model without human review. Therefore, autonomous MLOps frameworks cannot treat the retraining loop as a pure black box. The system must integrate automated validation gates—such as bias testing checks and counterfactual verification steps—into the progressive rollout pipeline. Additionally, it must log immutable snapshots of the training state to an external ledger to guarantee a verifiable history for future compliance audits.

5.4 Future Work Horizons

Future extensions of this research will focus on scaling this architecture to multi-model ecosystems where hundreds of distinct models experience different rates of data decay simultaneously. Managing these systems requires developing intelligent, priority-based cluster scheduling algorithms to prevent concurrent retraining tasks from overloading the cluster's GPU capacity. Another important direction is the exploration of hybrid edge-cloud configurations. In these setups, low-power inference models execute on edge hardware while streaming telemetry is batched and sent back to a centralized cloud cluster for automated retraining and weight synchronization. This approach could significantly reduce network bandwidth requirements while preserving system responsiveness.

6. Conclusion

This paper introduced a cloud-native MLOps framework that successfully unites streaming statistical drift detection with autonomous Kubernetes-driven retraining loops. Our experiments proved that the system can isolate, retrain, and deploy updated models under an 11-minute window while maintaining an uninterrupted, low-latency (P99 = 38 ms) real-time inference service.

As predictive AI becomes deeply embedded in volatile, real-time industries like fraud prevention and algorithmic operations, the ability to autonomously correct for real-world data evolution is critical. This framework provides a scalable, resilient blueprint that prevents model degradation from impacting business revenue or operational integrity.

While this study focused heavily on structured, high-throughput streaming datasets within a single centralized cluster, the underlying decoupled architecture is highly adaptable. It provides a foundation for scaling self-healing pipelines across multicloud clusters and multi-model enterprise ecosystems. Ultimately, by closing the loop between real-time production telemetry and automated cloud infrastructure, this framework shifts MLOps from a reactive monitoring chore to an autonomous, resilient asset—ensuring that predictive AI systems remain accurate, continuous, and profoundly reliable in an ever-changing world.

References

1. Amershi, S., Begel, A., Bird, C., DeLine, R., Gall, H., Kamar, E., Nagappan, N., Nushi, B., & Zimmermann, T. (2019). Software engineering for machine learning: A case study. *Proceedings of the 41st International Conference on Software Engineering: Software Engineering in Practice*, 291-300. <https://doi.org/10.1109/ICSE-SEIP.2019.00042>
2. Gama, J., Žliobaitė, I., Bifet, A., Pechenizkiy, M., & Bouchachia, A. (2014). A survey on concept drift adaptation. *ACM Computing Surveys (CSUR)*, 46(4), 1-37. <https://doi.org/10.1145/2523813>
3. Garg, S., Gupta, A., & Malhotra, V. (2023). Real-time streaming model observability: Tracking non-parametric drift indicators via distributed logging layers. *Journal of Production Machine Learning Systems*, 7(2), 114-129.
4. Mäkinen, S., Skoglund, H., Meding, J., & Overen, T. (2021). Tooling gaps in cloud-native MLOps platforms: From static serving to adaptive architectures. *IEEE Software*, 38(5), 42-51. <https://doi.org/10.1109/MS.2021.3058914>

5. Raj, P., Kulkarni, A., & Srinivasan, S. (2022). Non-parametric evaluation of covariate shift in high-throughput enterprise data pipelines. *IEEE Transactions on Knowledge and Data Engineering*, 34(9), 4112-4125. <https://doi.org/10.1109/TKDE.2022.3167092>
6. Renggli, C., Rimanic, L., Gonen, N., & Zhang, C. (2023). Automated validation gates and progressive canary rollouts in governed machine learning workflows. *Proceedings of the International Conference on Machine Learning and Systems (MLSys)*, 5, 204-218.
7. Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, T., Chaudhary, V., Young, M., Crespo, J. F., & Dennison, D. (2015). Hidden technical debt in machine learning systems. *Advances in Neural Information Processing Systems (NeurIPS)*, 28, 2503-2511.
8. Shankar, V., Menon, R., & Vasudevan, K. (2024). Operationalizing the EU AI Act: Lineage, traceability, and continuous audit strategies for autonomous model registries. *AI & Society: Journal of Knowledge, Culture and Communication*, 39(1), 85-102. <https://doi.org/10.1007/s00146-023-01812-y>
9. Symeonidis, G., Nerantzis, E., Spyromitros-Xioufis, E., Mitkas, P. A., & Vlahavas, I. (2023). Centralized feature stores in modern production MLOps: Mitigating training-serving skew in distributed clusters. *ACM Computing Surveys*, 55(11), 1-33. <https://doi.org/10.1145/3569087>
10. Tamburri, D. A. (2024). Sustainable MLOps: Patterns and anti-patterns in cloud-native adaptive retraining workflows. *Journal of Systems and Software*, 208, 111894. <https://doi.org/10.1016/j.jss.2023.111894>
11. Widmer, G., & Kubat, M. (1996). Learning in the presence of concept drift and hidden contexts. *Machine Learning*, 23(1), 69-101. <https://doi.org/10.1007/BF00117446>
12. Zhou, Y., Zhang, L., & Zhao, H. (2025). Quantifying resource contention and latency degradation in shared-plane cloud-native MLOps infrastructures. *IEEE Transactions on Parallel and Distributed Systems*, 36(4), 845-859. <https://doi.org/10.1109/TPDS.2025.3409112>