



MLOps Strategies for Scalable Cloud-Based Deployment

Ethan Williams

Asst Professor

Abstract

Machine learning operations (MLOps) comprises the practices, methods, and tooling that facilitate the deployment of reliable ML models in production environments. While many aspects of cloud data platforms are designed to enable reliability, only some managed ML services support the MLOps goals of continuous integration, continuous delivery, data lineage tracking, associated reproducibility, governance, and security. Furthermore, reliability encompasses not only the fulfillment of service-level objectives, but also systematic monitoring, alerting, and incident response automation. Architectural patterns are proposed to enable reliable deployment in cloud data platforms, focusing on the implementation of continuous integration and testing pipelines for ML models and the formulation of continuous delivery and rollout strategies. Continuous integration pipelines reduce the risk of regressions and ensure sufficient model performance at the time of deployment, while continuous delivery pipelines enable rapid updates to production models within acceptable risk profiles.

The landscape of publicly available MLOps frameworks, tools, and services is also examined, emphasizing the pros and cons of established and rising solutions in containerization, orchestration, model serving, and inference. Containerization and orchestration contributes to the building of reliable deployment pipelines in cloud data platforms, whether general-purpose tools (e.g. Docker and Kubernetes) or solutions tailored for ML workloads. Containerized serving frameworks designed for high-throughput, low-latency inference can benefit a wide range of business applications, while auto-scaling and model versioning capabilities enhance the ease of use of cloud-native ML services.

Keywords: MLOps, Cloud-Native Machine Learning, Continuous Integration for ML, Continuous Delivery of Models, ML Reliability Engineering, Model Deployment Pipelines, Data Lineage and Reproducibility, ML Governance and Security, Service-Level Objectives (SLOs), Monitoring and Incident Response Automation, CI/CD for Machine Learning, Model Rollout Strategies, Containerization for ML (Docker), Orchestration for ML Workloads (Kubernetes), Model Serving Frameworks, High-Throughput Low-Latency Inference, Auto-Scaling ML Services, Model Versioning, Cloud Data Platform Architecture, Production ML Observability.



1. Introduction

Machine learning operations (MLOps) encompasses the practices and the tools for the collaboration of data scientists, data engineers, operations support, and software engineers to produce, deploy, and maintain machine learning (ML) models in a reliable manner. Data-driven organizations in private and public sectors are increasingly adopting cloud-native solutions powered by public cloud platforms. Such solutions leverage cloud-native patterns for building and deploying machine learning models; accelerate the process of creating a production-ready ML stack by the automated provisioning of data lakes, object storage buckets, databases, orchestration engines, container registries, and virtual clusters; and alleviate the burden of managing infrastructure.

The reliability of cloud-native ML deployments benefits from automated model testing and validation, incident detection and response, and model performance monitoring. However, when not properly configured, such cloud services may introduce a single point of failure in the ML production pipeline or not meet the performance requirements. When models are deployed in containers, cloud-native CI/CD patterns can be used for continuous testing and validation. In this context, the reliability of deployment, encompassing continuous integration and continuous delivery processes, comprises multiple components: continuous integration pipelines guarantee the correctness of models and validate the quality of the data, and deployment strategies address the risk inherent to each rollout.

1.1. Overview of MLOps and Its Significance in Cloud Environments

MLOps refers to a set of practices that seeks to deploy and maintain ML (Machine Learning) models in production reliably and efficiently. ML has gained popularity in recent years and has many business applications. CI/CD (Continuous Integration/Continuous Delivery) processes have proven very useful in the deployment of software, and similar processes can be established for maintaining ML workloads. Enabling them through the automation of tools in Cloud Data Platforms also addresses ML-specific needs for data lineage, reproducibility, governance, and security. It is when Cloud Data Platforms are linked by automation that models can be delivered reliably, with constant observability of key metrics for development and

operational teams.

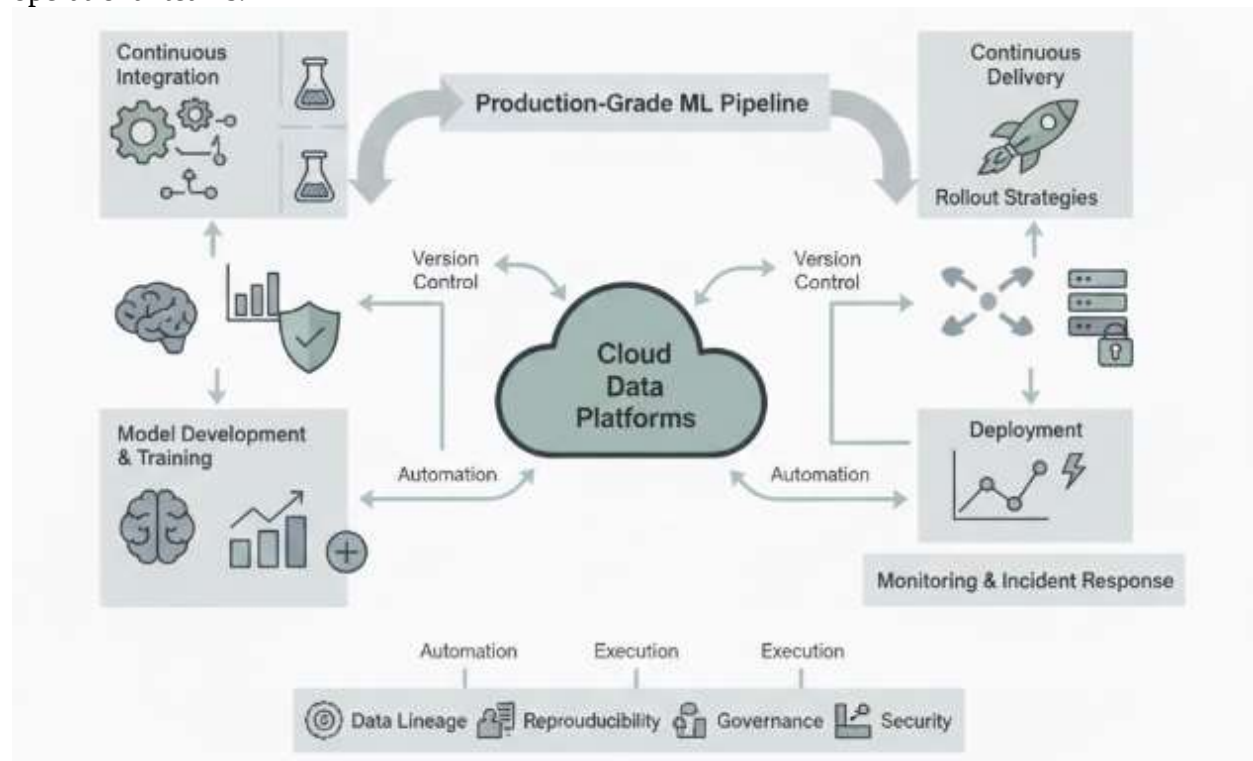


Fig 1: Synchronized Forwards: A Systematic Template for Automated CI/CD and Observability in Cloud-Native MLOps Ecosystems

The goal of this overview of the foundations of MLOps in Cloud Environments is to identify the capabilities needed in Cloud Data Platforms and the tools that fulfill them so that the deployment of ML workloads follows CI/CD processes. Automation is key for reliability and is provided for the various steps in the processes, especially Monitoring and Incident Response. The proposed process template is based on pairs of Forwards, one each for Continuous Integration and Continuous Delivery. Each Forward indicates a Continuous Integration Pipeline for Models that integrates CI aspects, possibly with additional validations, and a Continuous Delivery Pipeline for Models that covers the rest of the aspects of CI/CD, including Rollout Strategies for Models.



2. Foundations of MLOps in Cloud Environments

MLOps for cloud data platforms must ensure reliable deployment through automated CI/CD pipelines, active monitoring, and structured incident response. Several aspects contribute to this aim: reliable delivery of ML models, automated rollout strategies to mitigate risk, careful choice of tooling, defining key metrics, and the presence of a mature MLOps ecosystem. With the increasing adoption of public cloud services and deployment of production-level workloads, ensuring reliable execution is critical.

Building and deploying ML models is a multi-step process, as outlined in Chapter 1. An MLOps framework establishes CI/CD pipelines to automate these steps, including continuous integration testing of models and constant delivery to production. However, reliability demands more than correct delivery. Monitoring detects incidents, while clear procedures guide response, including reverting to backup models in case of failure..awsSupportedServices. Real-world S3 incidents demonstrated that even foundational cloud services are not always reliable and highlighted the vulnerability of AWS-hosted webpages without a CDN.

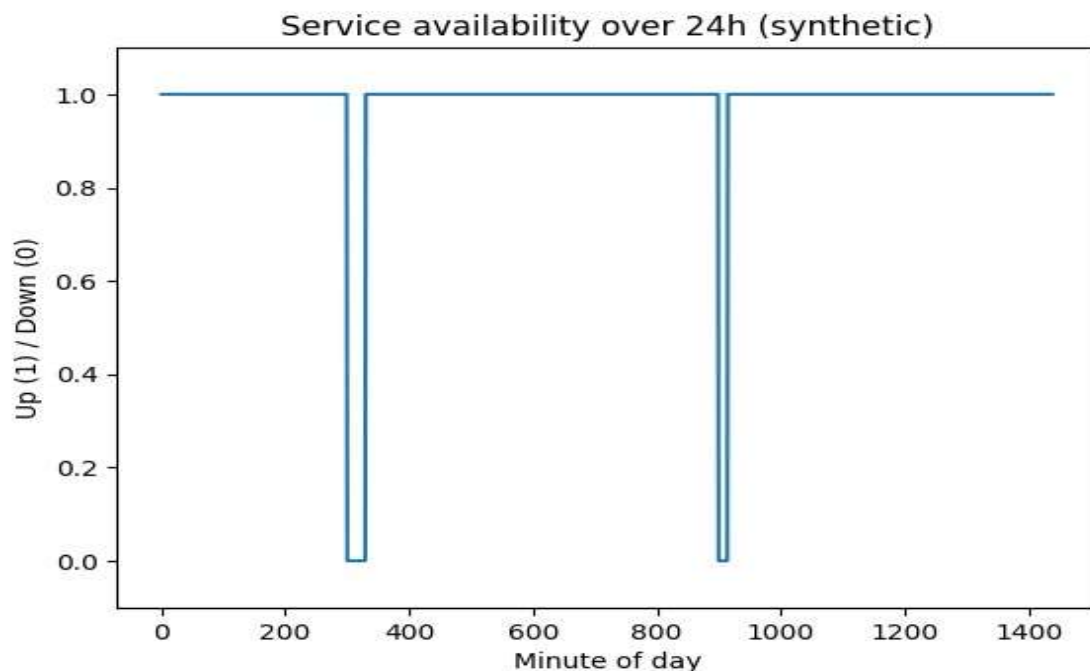
Metric	Value
Availability	0.96875
Latency p50 (ms)	45.60916859994713
Latency p95 (ms)	83.27928987399044
Latency p99 (ms)	112.2915089738782
Avg throughput (req/min)	872.2756944444444

2.1. Key Principles for Implementing MLOps in Cloud-Based Systems

Implementation of MLOps in cloud-based systems can be made reliable by supporting, monitoring, and maintaining these principal tasks: (i) the continuous integration of production models into the codebase, (ii) the validation of training and production data, (iii) the monitoring of production data for drift, (iv) using a canary or blue-green strategy for rollout to production, (v) rollback of the production model in the event of failure, (vi) monitoring the availability and performance of the model serving infrastructure, (vii) monitoring additional quality metrics for

the production model, (viii) enabling autoscaling of the model serving infrastructure, and (ix) implementing dashboards for all reliability and quality metrics.

The reliability of an ML model in production cannot be guaranteed solely by performing the conventional practices of unit test and continuous deployment; instead, reliability is achieved through a specific reliability-oriented platform. Figure 1 shows the architectural patterns needed for reliable deployment. The validation of both training and production data should be included as part of the pipeline. When a predictive model is deployed in production, it is essential that the distribution of the input data does not change over time, a case known as data drift. These monitoring checks should be in place, with the ability to instantaneously notify the data engineering team should any anomaly occur.



Equation A) Availability (uptime ratio)

Step 1 — define total observation window



- Let total observed time be T (e.g., in minutes or seconds).

Step 2 — split time into uptime and downtime

- Let uptime be T_{up}
- Let downtime be T_{down}
- Then:

$$T = T_{\text{up}} + T_{\text{down}}$$

Step 3 — define availability

$$A = \frac{T_{\text{up}}}{T}$$

Step 4 — alternate discrete (ping/check) form

If you check the service N times (e.g., every minute), define an indicator:

$$I_i = \begin{cases} 1 & \text{service is up on check } i \\ 0 & \text{service is down on check } i \end{cases}$$

Then:

$$T_{\text{up}} \propto \sum_{i=1}^N I_i \quad \Rightarrow \quad A = \frac{1}{N} \sum_{i=1}^N I_i$$

3. Architectural Patterns for Reliable Deployment

Reliable deployment of machine learning (ML) models is crucial for cloud data platforms. Data is constantly changing, so monitoring production systems and detecting problems before they affect users are key requirements. Continuous integration and delivery (CI/CD) practices enable frequent delivery of new features or bug fixes with reduced risk. In MLOps, CI refers to automatically validating modular parts in an ML pipeline when changes occur, while CD means automating the rollout of ML models into production while controlling for data quality and model performance.



For new ML models, the CI pipeline checks data correctness and performs unit tests on the ML pipeline, while the QA program conducts model validation. For new data, the CI pipeline executes data validation and ML pipeline tests, ensuring that model performance meets a pre-defined standard. On successful validation, the CD pipeline rolls out the model using canary, blue-green, or feature-flag techniques to control traffic and fallbacks are ready if necessary.

Continuous integration and testing pipelines for models cover seven main aspects of the ML process: data quality validation; input feature space monitoring; unit testing of the model training module; integration testing of the full ML pipeline; model performance validation; accuracy monitoring; and data drift detection and alerting. Such procedures furnish assurance that newly deployed models behave as expected in production and prevent failures caused by data issues.

3.1. Continuous Integration and Testing for Models

Reliable deployment of models in a production environment requires validation at a single point of time and over time. A robust continuous integration (CI) pipeline for models supports reliable deployment when models are trained, retrained, or replaced. CI for models involves validating and testing the code, models, and data systematically before deploying changes to real-time prediction serving platforms. An MLOps CI pipeline includes the following validation and testing processes: automated validation of data used for retraining; unit tests for code related to data preparation, retraining, and causing model rollout; integration tests for data preparation and model rollout components; testing criteria and guidelines; and automated quality assurance (QA) testing of models before promoting them to production.

MLOps practitioners deploy an additional CI test when a new model is introduced. A new model is QA-tested for accuracy on a data window through feature drift detection. Based on the expected model precision, the new model can be integrated with the existing model or become the primary model. Stability can be ensured by validating internal quality metrics that indicate the reliability of making predictions with each of the models that are made serviceable.

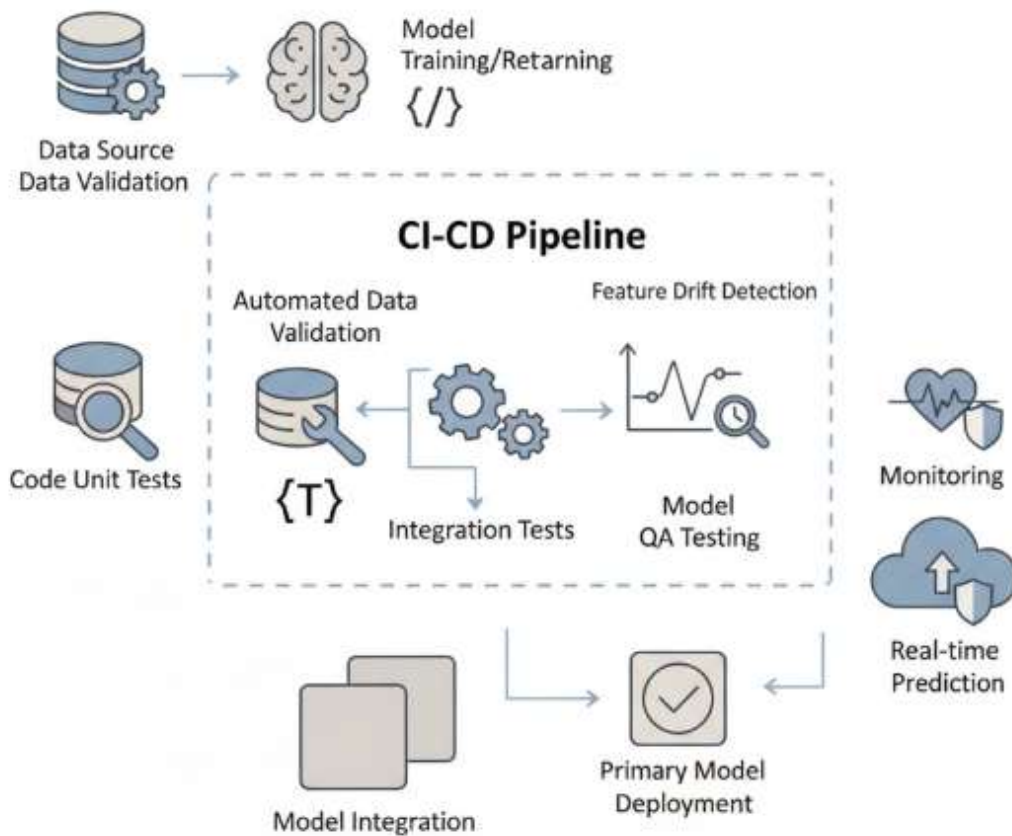


Fig 2: Longitudinal Integrity Frameworks: Continuous Integration and Drift-Resilient Deployment in Clinical MLOps

3.2. Continuous Delivery and Rollout Strategies

MLOps in Cloud Environments — Reliability Requirements

CI/CD pipelines enable teams to push changes to production quickly and safely. Continuous delivery consists of all modifications that enter a production-ready state and become available for release at any time. The actual deployments happen when these packages reach production, typically using automated rollout strategies.



Deployment of ML systems is considered reliable when it has the potential to succeed every time. Therefore, ML model deployment requires additional engineering and testing effort before changing production serving paths or the set of models used by another component. One idea for a canary deployment is to route a small proportion of inference traffic through the new model in staging and production or to use the new model for selected user cohorts that do not affect business-critical functionality. Model serving solutions that do not introduce latency hot paths can safely validate abilities, such as prediction quality or input shape, by running in a feature-flagged environment. Model serving solutions that support ingesting incompatible schemas can enable multiple model versions in production while rolling out a new model to users. In the event of failure, trained engineering and operations teams can rapidly deactivate the flow. Additional visibility and health checks provide guidance on when to resume traffic on the new model.

Rollout strategies with more drastic fall-back options, such as blue-green or shadowing deployments, can be considered for models whose traffic is small enough to allow doubling resource allocation. Traffic shadowing is the only strategy that is able to detect runtime issues related to user behavior, model execution environment (e.g., operating system and library versions, model artifacts), and interactions with all other downstream services.

Integrating ML models into a microservice architecture enables management systems to react to known issues in runtime performance and models to vote on upstream requests using a weighted voting mechanism. Production services can use multiple model versions to bypass performance penalties while backfilling the predictions of demand-based or SLA-driven models in common situations. Use of an identical model serving solution for both prediction and training serving reduces the testing effort needed to ensure that the two facilities provide the same contract and automates the deployment of new training model releases into prediction for shadowing, canary, and blue-green use cases.

4. Frameworks and Tooling Landscape

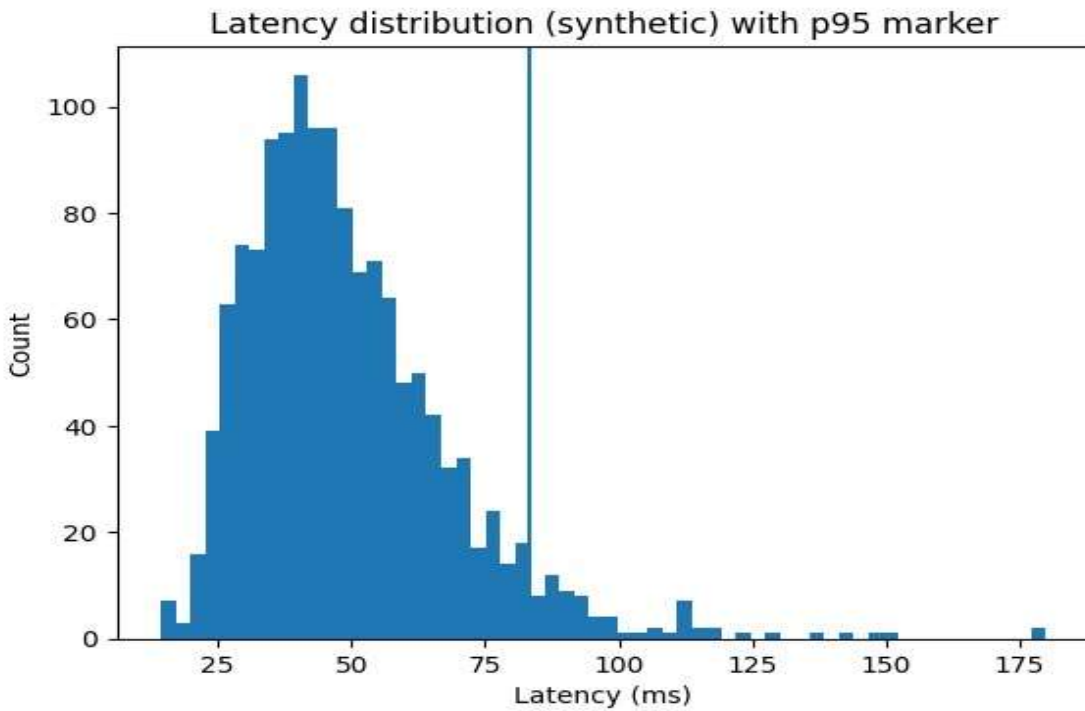
The landscape of frameworks and tools offering MLOps support can be divided into three areas: containerization and orchestration; model serving frameworks and inference engines; and experiment tracking tools.

Containerization revolutionized the cloud environment by providing an efficient way to run applications in an isolated environment. Despite being lightweight compared with virtual



machines, containers must still be scheduled to run on a set of machines; scheduling becomes complex when applications require persistent storage and a large number of interdependencies in their network. The typical solution is Kubernetes, which supports scheduling of arbitrary containers and provides cluster resource management and orchestration. Although Kubernetes does not natively support advanced scheduling for ML workloads, the ecosystem contains add-ons for specifying jobs and scheduling model training on available resources. Specialized platforms such as Kubeflow, Vertex AI Pipelines, and AWS SageMaker Pipelines provide custom operators to support the full lifecycle of ML workloads using Kubernetes as the underlying orchestration layer.

Model serving frameworks facilitate the transition between the training and production stages of model lifecycle management. Although APIs are typically provided for external applications to request inferences for trained models, additional concerns must be addressed: throughput and latency optimizations, versioning of multiple models in the serving layer, automatic scaling, resource allocation during low-load periods, and smart batching of requests. Inference engines integrate seamlessly into the ML stack, understand the expressive model definitions of the chosen framework, and provide considerate wrappers for serving models. Furthermore, several offerings focus specifically on the needs of ML inference: Nvidia Triton, TensorFlow Serving, Seldon, and KFServing. Experiment tracking tools automate and assist in recording information about experiments in order to compare trained models. Considered parameters include metadata for datasets and trained models, metrics and logs, and artifacts such as models and visual representations of results. Prominent tools include MLflow, Weights & Biases, Data Version Control, and Git.



Equation B) Latency (end-to-end response time)

Step 1 — per-request latency

For request j :

$$L_j = t_{\text{out}}^{(j)} - t_{\text{in}}^{(j)}$$

Step 2 — average latency over N requests

$$\bar{L} = \frac{1}{N} \sum_{j=1}^N L_j$$



Step 3 — percentile latency (common SLO form)

Sort latencies: $L_{(1)} \leq \dots \leq L_{(N)}$.

The p -th percentile (e.g., $p = 95$) is approximately:

$$L_p \approx L_{(\lceil pN/100 \rceil)}$$

This is why dashboards often show p50/p95/p99.

Step 4 — alert rule form

If your SLO is “p95 latency $\leq L^*$ ” then an alert condition is:

$$L_{95} > L^*$$

4.1. Containerization and Orchestration for ML workloads

Containerization and orchestration are widely adopted in cloud computing and have significantly influenced architecture patterns for deploying production applications. Resource isolation and independent lifecycle management of services address various concerns such as fault tolerance and technology diversity. Schedulers and orchestration engines simplify service management, such as redundancy setup, load balancing, and failure recovery. Although more common in back-end services, containerization is increasingly utilized for serving machine learning (ML) models because serving requests depends mainly on processing latency and unlike training, predicted features do not have to be saved for later retrieval. The heavy computational nature of large language models is leading to adoption by ML-specific scheduling, for both cost and latency reduction for end-users. Parallel requests can be served simultaneously by different replicas, while model versioning is important for stepwise promotion and rollback. Distributed aspects also optimize resource consumption. Inference is memory bound, especially transformers.

Base-container image size is another factor affecting latency for both model serving and retraining. Optimizing workflow also improves container launch overhead and therefore response time in sudden inbound spikes. TorchScript, TensorFlow Saved Model, and ONNX accelerate model inference workflows. TensorRT optimizes tensors for computing engine performance while DeepSpeed compresses model parameters in memory for free up usage for CPU resources. Model serving and container orchestration frameworks support both load and model-level autoscaling.



Bin	ref_p	cur_p
6	0.1	0.08783864900556244
7	0.1	0.094841730361359
8	0.1	0.10020408979951179
9	0.1	0.12479490976029453
10	0.1	0.21155308335667694

4.2. Model Serving Frameworks and Inference Engines

Containerization and orchestration of ML workloads have become common practice in production deployments. General-purpose technologies such as Docker and Kubernetes are extensively used for encapsulating run-time environments, instantiating and scaling deployments, and enabling reliability features such as health-checks and auto-restarts. In addition, a range of ML-specific schedulers comes with built-in integrations for monitoring data preparation and inference pipelines. By encapsulating all non-model-related code and resources in external services, complex orchestrators can be avoided. Such an approach allows for massive inference workloads to be processed, with good latency and throughput guarantees, by scaling the serving layer horizontally with dedicated auto-scaling groups. In contrast to standard ML-serving solutions, the drawback of dedicated orchestration is that it requires a lot of preparation and provides limited quality-of-service guarantees.

Dedicated model-serving frameworks and inference engines represent the actual workhorses of ML workloads in production, exposed as services by clients, and are invoked by dedicated processes or common schedulers. They aim to provide low-latency inference while maximizing throughput by caching hot models in memory and batching predictions. Take Latency and throughput targets into consideration when selecting these solutions. Orchestration platforms allow using SLO, SLI, and SLA concepts to trigger alerts and trigger actions such as scaling in or out the model-serving region, scaling up or down the machine where the region is deployed, or switching to a lower-refresh-frequency region model.

5. Cloud Platform-Specific Considerations



Public cloud providers offer a comprehensive portfolio of services aligned with the growing needs of Machine Learning (ML) projects. Dedicated Managed ML Services streamline the management of ML algorithms but often sacrifice critical aspects such as monitoring quality metrics, measuring drift detection in production data, implementing proper governance procedures, and managing operational Security. Therefore, in practice, part of the infrastructure is usually hosted on self-managed dedicated VMs, Containers, or Kubernetes that allow in-depth control of the reliability and non-functional aspects and integrate with the Managed ML Services to fulfil the full-spectrum ML life cycle. Strategies such as Multi-Cloud, Hybrid environments, and Customer-managed deployment of services delivered by Service Providers, often constrained by Legal aspects such as Data Sovereignty, Data Patterns and Portability concerns and Orchestration of the underlying Infrastructure and Services, emerge to address and mitigate these challenges.

Public cloud providers offer a varied portfolio of services aligned with the growing needs of Machine Learning (ML) projects. Dedicated Managed ML Services integrate cloud providers' strengths but introduce cost-related concerns and sacrifice critical aspects such as checking the quality of the produced models, monitoring the predictive/model quality metrics, drift detection of the production data, applying a proper governance process, managing Security on infrastructure Data Patterns and Partitions, non functional aspects such as Latency, Throughput and Availability. It is therefore more common for organizations to self-manage part of the infrastructure using dedicated VMs, Containers, Kubernetes or dedicated for ML Orchestration that allow control and integration of the reliability and non-functional aspects of ML operational Phase while leveraging Managed services for the training, production and Service Users facing Interaction of the ML System. Multi-Cloud, Hybrid and Customer-managed Deployment strategies of Services delivered by Service Providers gradually become standard practices to satisfy Business and Law oriented constraints like Data Sovereignty and Data Patterns.

5.1. Public Cloud Offerings and Managed ML Services

Public cloud providers distinguish their service offerings through high-level managed services. The Machine Learning as a Service (MLaaS) paradigm allows customers to efficiently deploy their ML workloads without worrying about the underlying infrastructure, as they are abstracted away by the cloud provider. Many MLaaS services cover the entire ML lifecycle, but none of the major cloud providers offer a solution for all MLOps properties outlined in the previous sections. Some components are well covered in certain services of specific providers,



while others are simplified or lacking. The existing public offerings therefore call for additional tooling selection to achieve more reliable deployments and fulfill more MLOps properties.

The interaction of MLaaS with the remaining data pipelines must also be addressed, particularly when the data enter the cloud platform from on-premises infrastructure or when it is kept on a different cloud for compliance reasons. For these situations, Data Transfer Services are available at every major provider to help speed up data movements between on-premise equipment or different cloud regions. Despite being tactical services, they also have a large impact on MLOps properties.

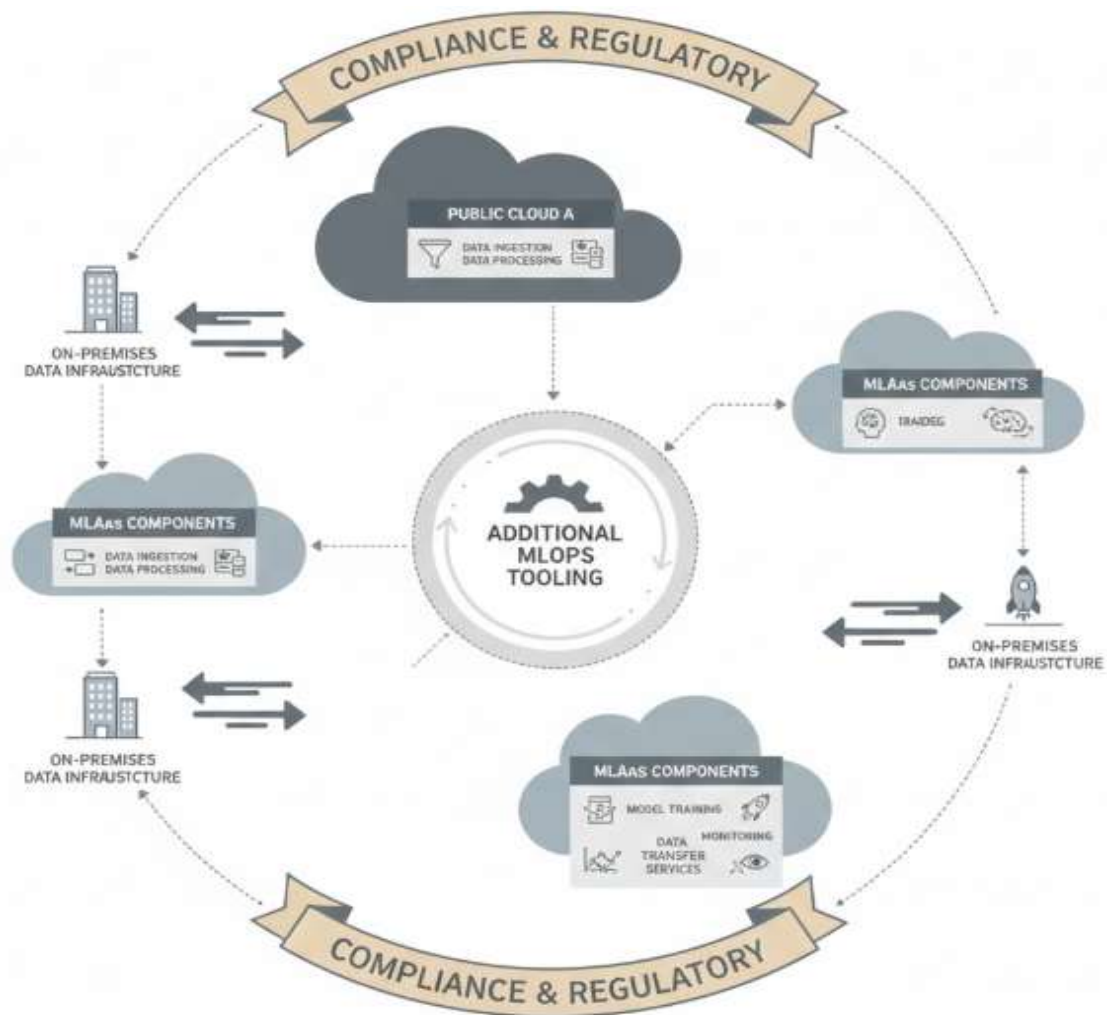


Fig 3: Bridging the MLaaS Gap: Integrated MLOps Orchestration and Tactical Data Transfer in Hybrid-Cloud Environments

5.2. Multi-Cloud and Hybrid Deployments



This section explores multi-cloud and hybrid cloud deployments, assessing their benefits and challenges for MLOps in cloud environments. The competitive advantage of public clouds stems from the ability to distribute workloads as-needed. Service offerings tailored for unique needs (e.g., Geolocation Services, Meta's Facebook) are located close to users, and edge resources managed by the cloud providers reduce latency. Nevertheless, data sovereignty laws mandate data processing and storage within national borders. Therefore, controlling the ownership and location of data for data governance is crucial, enabling organizations not only to meet legal requirements but also to control costs and vendor lock-in. Multi-cloud deployments allow organizations to use different public clouds, promoting price competition among providers, while also alleviating the limitations and problems associated with a single cloud vendor.

The public cloud services may satisfy all the requirements for the deployment of an ML model. However, in several cases, the required resources address only a specific segment of the workload (e.g., Google's part-in-place Tensor Processing Units, necessary only for the training phase of significant models) or are idle during periods of little or reduced demand. MLOps in the Cloud_My own version identifiers_full mark-up catering to complaints of reduced privacy risk and increased quality of service (QoS) in the scientific research process can take advantage of a hybrid deployment model. In this type of architecture, sensitive data can be kept in private premises, controlling and governing them like an on-premise scenario, while data without additional restrictions are considered in the public cloud. Data portability between the private cloud and the public cloud is crucial to the effectiveness of hybrid deployment. Support for data portability narrows the scope of orchestration and enables the monitoring of sensitive workloads in normal production mode.

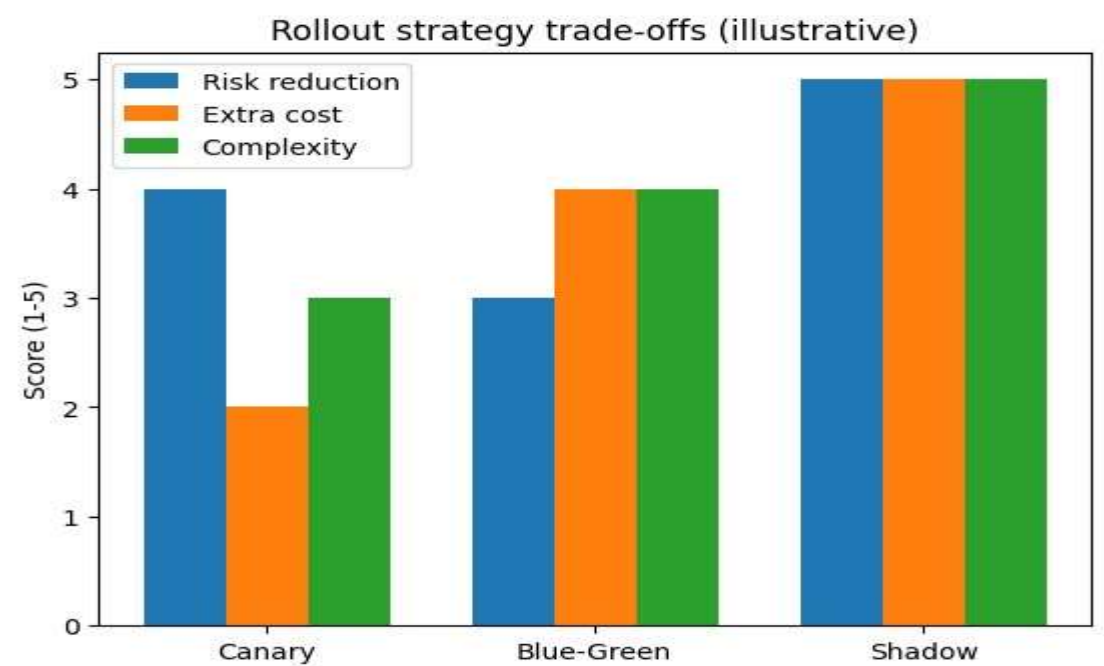
Drift test	Value
Population Stability Index (PSI)	0.12448176632338913
Kolmogorov–Smirnov statistic (KS)	0.1392
Kullback–Leibler divergence (KL)	0.06668422745742603

6. Reliability and Performance Metrics

Availability targets specify the maximum acceptable downtime for inference. For systems with very low (in the order of magnitude of milliseconds) and predictable latency (few percent are longer), the redundancy and maintenance costs are more significant and can trigger

the system downtime. For such systems, maintenance windows should be planned during low-usage times. Downtime windows can be scheduled for periods where the requested service is at its lowest demand. The only tool to measure such an aspect is an external ping or request to the model exposed in the serving layer. If the service respond with the latency defined previously, the request passes; otherwise, it fails.

Latency defines the maximum time for a request to be processed (both in one inference cycle and in the whole inference pipeline). It is also defined as the maximum number of time for a request to be processed before it goes dead. Performance dashboards should allow detection of anomalies with respect to these targets, allowing the selection of filters that support the identification of the causes of problems faster and more efficiently. For such alerts, system request or model serving logs can record the inference latency from the moment they enter the model-exposed endpoint until the moment they go out. If the accumulated latency exceeds a certain value, an alert is triggered, indicating that an action should be taken soon (requests are becoming extinct). For performance, dashboards could be implemented to support anomaly detection with respect to the defined KPIs.





Equation C) Throughput (inferences per unit time)

Step 1 — count requests

Let N requests be completed in time window Δt .

Step 2 — define throughput

$$X = \frac{N}{\Delta t}$$

Units: requests/second, requests/minute, etc.

Step 3 — aggregated throughput

If multiple models/services $k = 1..K$ run in parallel:

$$X_{\text{total}} = \sum_{k=1}^K X_k$$

6.1. Availability, Latency, and Throughput Targets

Availability, latency, and throughput constitute the primary reliability and performance metrics of deployed models. Availability is defined as the proportion of time the service is operational. Latency is the time from request submission to response retrieval and is constrained by user experience and business logic. At the model level, throughput indicates the number of inferences serviced per time unit, and for aggregated deployments, it reflects the load from all models. These metrics guide reliability design choices, formulate SLOs, and allocate monitoring resources.

Availability is typically measured by uptime but may also be calculated for specific time windows. MLOps dashboards reflect these measurements, displaying ongoing availability levels, SLO status, and time spent in incidents. Latency SLOs are affected by user experience requirements, and businesses may set strict thresholds or classify latency as an implicit requirement. Throughput SLOs are based on request volume predictions and specified at the aggregate level or per traffic type.

6.2. Model Quality Metrics and Drift Detection



Any organization deploying models to production has requirements regarding model availability, prediction latency, and prediction throughput. These targets can typically be defined on a per-model basis, based on business requirements and Service Level Agreements (SLAs). Availability refers to the fraction of time that the route to the model is operational, such as considering underlying components like GPUs for on-demand allocation. Latency is the total time taken for a single request to pass through all components of the inference pipeline, while throughput is the total number of requests processed in a given time window. These metrics can usually be computed, monitored, and displayed on dashboards using standard observability tooling.

In addition to these reliability and performance metrics, model quality is a major concern. The prediction quality of models may degrade over time as the statistical properties of the data change. A number of different methods exist for detecting drift or changes in data distributions, such as population stability index (PSI) for feature distributions, Kolmogorov-Smirnov tests for comparing distributions, Kullback-Leibler divergence for comparing distributions, and Chi-square tests for categorical variables. Different methods can be employed for different types of prediction tasks, such as binary classification, multi-class classification, or regression, and threshold values can be defined above which alerts should be generated.

7. Conclusion

MLOps is key for applying ML to business problems. However, using ML safely at scale requires specialisation; reliable deployment is often overlooked. A pattern for reliable model deployment in cloud data platforms, organised around the themes of CI/CD pipelines for ML models, automation, monitoring, and incident response, is presented. The automation of CI pipelines and unit tests is crucial. Data validation and model quality control during deployment are equally important; blue-green changes are preferred. Model-serving frameworks and inference engines should be carefully selected to meet reliability and performance targets. Managed services are convenient, but a cloud-agnostic approach is recommended.

MLOps is not a goal in itself. MLOps aims to allow the automated application of ML to business problems in production. Although the definition of MLOps is cloud-agnostic, most applications are focused on public clouds. The MLOps research and tooling ecosystem is mainly centred on the continuous integration and delivery of ML models. However, this is hardly sufficient for safely scaling ML applications. This work proposes an architectural pattern

specifically dedicated to reliable deployment and operation applied to data platforms in the public cloud. The approach clarifies the roles of availability, latency, and throughput targets for model serving and inference at scale, and establishes the importance of reliable Model Quality Assurance and Incident Management pipelines and processes when deploying ML models into production.

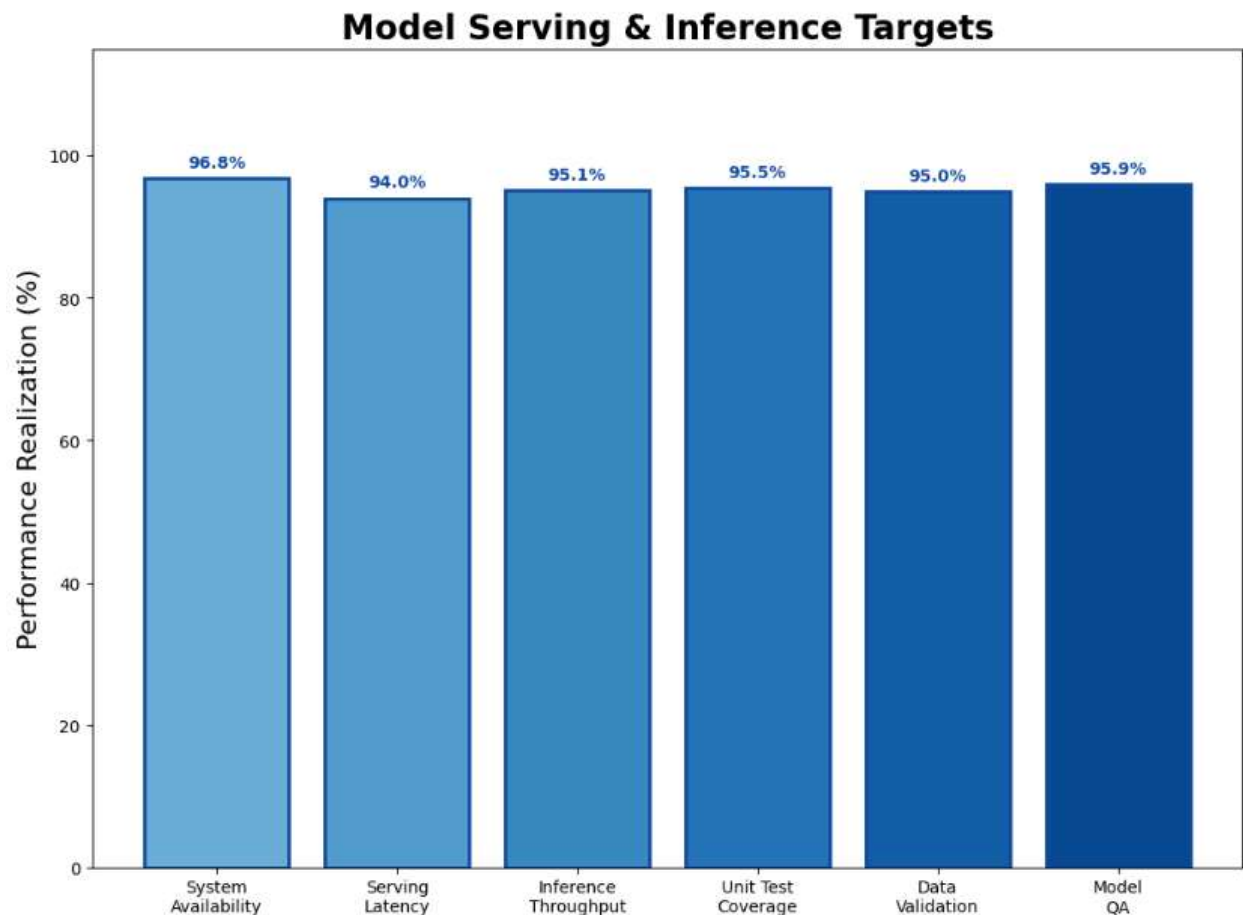


Fig 4: Model Serving & Inference Targets

7.1. Final Thoughts and Future Directions in MLOps



The final remarks reflect on the work, outline implications, discuss current limitations, and suggest directions for future research. A convergence of requirements to ensure the mutual cooperation between knowledge domain experts and DevOps teams is required. Quality assurance processes should be integrated within the CI/CD pipelines that run in cloud environments. The investigation of practical MLOps toolsets and cloud data platform automations that improve the reliability has exposed an immature and scattered tool-market that requires further speeding up.

Research on the interaction between people and organizations has recently shown that specialized functional users need assistance for handling specialized tools. In ML systems built on data from the company financial market, knowledge domain experts have been detected unwilling to produce the required high-quality ML models. Data from these sources help strategic decisions taken by company leaders. The organizational support for integrating the work of DevOps teams and the knowledge domain experts has been insufficient. As a consequence, CI/CD strategies for ML models are not used. Deployment is occurring manually by a user that has the knowledge in Python programming and MLOps processes.

Continuous integration and continuous delivery processes for ML models are being integrated into cloud environments. Nevertheless, the instability of the models remains an issue. Quality assurance processes that check the quality of models during the testing phase are necessary to further accelerate the maturity of the area. Recent research on ML-based systems has pointed out the delay in the production of practical tools that implement the key concepts of MLOps and the use of public clouds in services offered to support the automation of MLOps tasks.

References

1. Garg, S., Pundir, P., Rathee, G., Gupta, P. K., Garg, S., & Ahlawat, S. (2021). On continuous integration/continuous delivery for automated deployment of machine learning models using MLOps. In 2021 IEEE Fourth International Conference on Artificial Intelligence and Knowledge Engineering (AIKE) (pp. 25–28). IEEE.
2. Granlund, T., Kopponen, A., Stirbu, V. A., Myllyaho, L., & Mikkonen, T. (2021). MLOps challenges in multi-organization setup: Experiences from two real-world cases. In 2021 IEEE/ACM 1st Workshop on AI Engineering—Software Engineering for AI (WAIN) (pp. 7–13). IEEE.

AMERICAN ONLINE JOURNAL OF SCIENCE AND ENGINEERING

VOLUME: 02 ISSUE: 04

RECEIVED: OCTOBER 25

REVISED: NOVEMBER 06

ACCEPTED: NOVEMBER 24

PUBLISHED: DECEMBER 26

ISSN: 3067-1140



3. John, M. M., Olsson, H. H., & Bosch, J. (2021). Towards MLOps: A framework and maturity model. In 2021 47th Euromicro Conference on Software Engineering and Advanced Applications (SEAA) (pp. 334–341). IEEE.
4. Kolla, S. H. (2023). Large Language Model-Driven Enterprise Service Intelligence for Digital Workflow Transformation. *International Journal of Research and Applied Innovations*, 6(1), 8380-8391.
5. Raj, E., Buffoni, D., Westerlund, M., & Ahola, K. (2021). Edge MLOps: An automation framework for AIoT applications. In 2021 IEEE International Conference on Cloud Engineering (IC2E) (pp. 191–200). IEEE.
6. Renggli, C., Rimanic, L., Gurel, N. M., Karlas, B., Wu, W., & Zhang, C. (2021). A data quality-driven view of MLOps. *IEEE Data Engineering Bulletin*, 44(1), 11–23.
7. Kolla, S. K., & Reddy, V. A. R. (2023). Deep Learning Architectures For Multimodal Medical Data Integration. *South Eastern European Journal of Public Health*, 248–260.
8. Granlund, T., Stirbu, V. A., & Mikkonen, T. (2021). Towards regulatory-compliant MLOps: Oravizio's journey from a machine learning experiment to a deployed certified medical product. *SN Computer Science*, 2, Article 342.
9. Mäkinen, S. S., Skogström, H., Laaksonen, E., & Mikkonen, T. (2021). Who needs MLOps: What data scientists seek to accomplish and how can MLOps help? In 2021 IEEE/ACM 1st Workshop on AI Engineering—Software Engineering for AI (WAIN). IEEE.
10. Xu, X., Mo, R., Yin, X., Khosravi, M. R., Aghaei, F., Chang, V., & Li, G. (2021). PDM: Privacy-aware deployment of machine-learning applications for industrial cyber-physical cloud systems. *IEEE Transactions on Industrial Informatics*, 17(8), 5819–5828.
11. Bandi, V. D. V. K. (2023, November). Production-Grade Machine Learning Pipelines For Healthcare Predictive Analytics. *South Eastern European Journal of Public Health*, 189–205.
12. Cai, H., Wang, C., & Zhou, X. (2021). Deployment and verification of machine learning tool-chain based on Kubernetes distributed clusters. *CCF Transactions on High Performance Computing*, 3(3).
13. Shah, S.-C. (2021). Design of a machine learning-based intelligent middleware platform for a heterogeneous private edge cloud system. *Sensors*, 21(22), Article 7701.
14. Singh, P. (2021). Machine learning deployment using Kubernetes. In *Deploy machine learning models to production* (pp. 127–146). Apress.
15. Testi, M., Ballabio, M., Frontoni, E., Iannello, G., Moccia, S., Soda, P., & Vessio, G. (2022). MLOps: A taxonomy and a methodology. *IEEE Access*, 10, 63606–63618.
16. Davuluri, P. S. L. N. (2023). Integrating artificial intelligence into event-driven financial crime compliance platforms. *International Journal of Finance*, 36(6), 707-736.



17. Symeonidis, G., Nerantzis, E., Kazakis, A., & Papakostas, G. A. (2022). MLOps—Definitions, tools and challenges. In 2022 IEEE 12th Annual Computing and Communication Workshop and Conference (CCWC) (pp. 0453–0460). IEEE.
18. Hewage, N., & Meedeniya, D. (2022). Machine learning operations: A survey on MLOps tool support. arXiv preprint arXiv:2202.10169.
19. Kreuzberger, D., Köhl, N., & Hirschl, S. (2022). Machine learning operations (MLOps): Overview, definition, and architecture. arXiv preprint arXiv:2205.02302.
20. Moreschini, S., Lomio, F., Hästbacka, D., & Taibi, D. (2022). MLOps for evolvable AI intensive software systems. In 2022 IEEE International Conference on Software Analysis, Evolution and Reengineering (SANER) (pp. 1293–1294). IEEE.
21. Kolltveit, A. B., & Li, J. (2022). Operationalizing machine learning models—A systematic literature review. In 2022 IEEE/ACM 1st International Workshop on Software Engineering for Responsible Artificial Intelligence (SE4RAI) (pp. 1–8). IEEE.
22. Yandamuri, U. S. (2023). An Intelligent Analytics Framework Combining Big Data and Machine Learning for Business Forecasting. Zenodo.
23. Shankar, S., Garcia, R., Hellerstein, J. M., & Parameswaran, A. G. (2022). Operationalizing machine learning: An interview study. arXiv preprint arXiv:2209.09125.
24. Subramanya, R., Sierla, S., & Vyatkin, V. (2022). From DevOps to MLOps: Overview and application to electricity market forecasting. *Applied Sciences*, 12(19), Article 9851.
25. Leroux, S., Simoens, P., Lootus, M., Thakore, K., & Sharma, A. (2022). TinyMLOps: Operational challenges for widespread edge AI adoption. In 2022 IEEE 36th International Parallel and Distributed Processing Symposium Workshops (IPDPSW) (pp. 1003–1010). IEEE.
26. Pandey, A., Sonawane, M., & Mamtani, S. (2022). Deployment of ML models using Kubeflow on different cloud providers. arXiv preprint arXiv:2206.13655.
27. Chinthapatla, S. (2022). Edge AI with Kubernetes: Deploying machine learning models at scale. *International Journal of Emerging Trends in Computer Science and Information Technology*, 3(2), 32–41.
28. Böhm, S., & Wirtz, G. (2022). Cloud-edge orchestration for smart cities: A review of Kubernetes-based orchestration architectures. *EAI Endorsed Transactions on Smart Cities*, 6(18), Article e2.
29. Mangalampalli, B. M. Generative AI Applications In Healthcare Data Mart Design And Optimization.



30. Liang, B., Wu, D., Wu, P., & Su, Y. (2022). Efficient evolutionary optimization using predictive auto-scaling in containerized environment. *Applied Soft Computing*, 129, Article 109610.
31. Boosa, S., & Allam, K. (2022). End-to-end MLOps pipeline on Kubernetes for scalable healthcare applications. *International Journal of Emerging Research in Engineering and Technology*, 3(1), 74–85.