

Infrastructure Automation for Continuous Validation and Monitoring of ML Models in Hospitals

Veerendra Nath Jasthi

veerendranathjasthi@gmail.com

Abstract— With hospitals and other hospital systems undertaking the deployment of machine learning (ML) models to assist in clinical decision-making and diagnostic performance as well as provide operational efficiency, there is a concern whether such models are expected to remain reliable and fair over time. Nonetheless, in case of inefficient infrastructure, the applied models can suffer obsolescence in terms of data drifting, concept drifting, or environmental dynamics. In this paper, the author suggests an infrastructure automation model of constant ML models verification and monitoring in hospitals. The architecture incorporates CI/CD pipelines, container orchestration, the ingestion of data in real time, monitoring dashboards, and drift detection modules. Automatic retraining triggers and model performance alerts can help hospitals maintain strong ML deployments that will adjust with the changes in clinical data. Experiments based on real hospital data (ICU prediction, diagnostic classification, various risks of readmission) show a better ability to retain the accuracy of the models and less control is provided by manual operations. This paper emphasises the significance of MLOps principles applicable in life-threatening environments such as healthcare. In this context, the monitoring and validation processes are life-threatening.

Keywords— Infrastructure Automation, MLOps, Continuous Monitoring, Model Validation, Healthcare AI, Data Drift, Hospitals, Machine Learning, ML Lifecycle, Clinical Decision Support.

I. INTRODUCTION

The future of contemporary healthcare is being transformed with Machine Learning (ML) technologies that help to make clinical decisions on data in a much shorter time and apply personalized care on patients. Predictive algorithms are being more widely used in hospitals in the form of early warning of critical condition, resource utilization and optimization, aid with diagnosing, and monitoring the outcome after the treatment is provided [2]. Ranging through ICU admission prognostic to automated interpretation of radiological pictures, the application of ML within the systems of hospitals is still growing. Nonetheless, the problems associated with using such models in the real world hospital setting extends beyond the accuracy gained in the training setting. Ongoing surveillance, validation, and adjustment is required to determine that these models can continue to be clinically competent, safe, and effectual.

Among the key problems of hospital-based ML models is the fact that the model loses its performance over time, and it is typically caused by the data drift, modification of clinical procedures, demographics of the new patient group, or a shift of data collection equipment and methods. A model trained on the data of patients last year may not work just as well on this year data input, especially when the hospital presents some kind of demographic change, change of disease prevalence (such as during a pandemic), or some changes to data collection criteria [9]. Unless it is monitored in real-time, there are chances that the

models decay without any knowledge being noted, such that a diagnosis is erroneous, or clinical decisions are unstable.

The other major concern is that there will be a wide disparity in the monitoring and validation practices of models within hospital systems. Most ML models are made into static services in which no one is doing such subsequent validation [3]. Such models act like a black box, in which clinicians may not even know whether the model is working poorly or not. It may cause a misplaced trust in the output of AI, loss of accountability, and in worst-case apply to harming patients. It is especially harmful to the models that are employed in critical care settings.

To cope with these issues, the hospitals will have to transform the current approach to deploying traditional models and adopt the methodologies of the Infrastructure Automation provided with solid MLOps (Machine Learning Operations) practices. These imply establishing automated model validation, retraining, monitoring, and rollback pipelines and making sure that the ML models are ever-evolving, as the hospital environment changes. This kind of infrastructure would have to be adoptable to hospital data formats (such as HL7/FHIR), have strong privacy regulations (such as HIPAA), and run with high availability to enable 24/7 clinical operations.

Moreover, hospital data systems, in comparison with the classical IT systems, are heterogeneous, multiplexed, and frequently siloed. The origin of medical data may be very diverse in the form of lab results, wearable sensors, radiology images, EHR notes, and constant patient monitoring devices. Getting all these streams together into a single infrastructure that will be able to support real-time ML validation and drift detection is far from trivial. Even a high-performance model can go out of production without the right infrastructure level [10].

As such, the offered paper suggests a holistic automation infrastructure dedicated to hospitals, which will allow continuously validating and monitoring ML models in real time. It is a bundle of such tools as containerized deployment of models, real-time streaming of the incoming data, automated detection of the drifts, CI/CD pipelines to revise the machine learning models, and model re-learning triggered by a loss in performance or concept drift. The system will make the hospitals safer and more versatile as the fact that there would be no unchecked ML model once deployed will be ensured by the system by integrating intelligence into the infrastructure itself.

To conclude, the growing use of ML in hospitals necessitates more than training of models, it also needs an operational mindset (where alongside the models, validation, monitoring, compliance, and adaptation are needed). The present paper extends the ideas of MLOps to the healthcare setting and provides a roadmap on how hospitals can implement this concept to have trustable and self-healing AI systems that adapt to the evolving nature of medical practice [4-7].

Novelty and Contribution

The original value of the presented paper is a new infrastructure automation framework specifically tailored towards continuous validation and monitoring of ML models within hospitals since existing MLOps tools in the healthcare sector leave much to be desired.

What is new in this piece:

- Adaptation of MLOps principles towards a domain: As compared to generic MLOps platforms, this framework has domain-specific adaptation of principles of MLOps by tailoring towards the medical standards of HL7/FHIR and HIPAA compliance, which guarantee compatibility with hospital data and privacy requirements. It does not only use statistical indicators to adapt drift thresholds and validation strategies concerning clinical KPIs.
- Rule-based retraining: Unlike manual retraining or fixed day retraining, the model has performance monitoring with automated retraining triggers based on clinical determined thresholds in terms of AUC and precision recall. This gives time to correct before the model does any harm in the real world.
- Full-cycle continuous monitoring: If the performance drops suddenly, it is possible to roll back to the previous known good version of a model, which is possible through the framework. This introduces an extra backup to clinical applications where errors in prediction may prove fatal.
- Live hospital data driven drift detection: The architecture incorporates drift detectors (such as KS-tests, ADWIN, and Page-Hinkley), to the live data stream of a hospital and drift-detectors will provide real-time alerting on concept and data drifts. This has hardly been achieved in the clinical environment.
- Multi-modal monitoring dashboards to clinicians and IT teams: As part of the system, a role-based visualization platform will be provided, in which clinicians can check model safety, and IT teams can investigate technical drift and error logs. This improves the openness and confidence of AI-sustained mechanisms.
- Modifications due to minimal amount of human input with high degree of auditability: CI/CD, MLflow tracking with secure audit logging all the way from data to deployment. This makes it more easier to be audited by regulators and creates lasting traceability.

Main contributions of the paper:

- The blueprint of an automated ML lifecycle management infrastructure in hospitals in the form of modules.
- Validation on real-life data on sepsis predictive, readmission, and pneumonia classification of hospital data.
- Static vs. automated deployment setup, wherein the differences were recorded showing positive results in stability, adaptability, and model reliability.
- A repeatable way to carry out that can be translated to other hospitals with customization or extension to provide a safe deployment of ML at scale.

The given work is not based on theory and offers a feasible implementation of clinical AI governance in a real-life setting that fills in the gap between data science and safety in hospitals through the automation of infrastructure.

II. RELATED WORKS

In 2025 L. C. Nechita *et al.*, [15] introduced the cross-section between machine learning and healthcare has seen a lot of research interest in the recent years especially when applied to the model deployment, performance monitoring, and automation of ML lifecycles. But most of these have been aimed

at development of algorithms rather than the supporting infrastructure (to support ML models over a long period in the real-world and clinical settings). It is increasingly understood that the devising of a very accurate model is just one part of the process; so too is the continued performance of such a model, its safety and its compliance in hospitals.

The idea of data and concept drift in the clinical machine learning application is one of the larger topics discussed in the previous literature. Consistently, it has been noted that data on hospitals is very non-stationary. An example is that seasonal outbreaks may change patient demographics, hospital policies may change, and new devices will change input data format or even quality. There are studies that indicate that when monitoring mechanisms are not applied and ML models are used in production, they can degrade in their predictive power either slowly or suddenly. This is especially hazardous to clinical settings, where failure to notice performance decay might cause the misdiagnosis orientation or late treatments.

In 2024 F. Pesapane *et al.*, [8] suggested the attempts to mitigate the necessity to navigate those challenges have resulted in drift detection frameworks. These are both statistical (like distribution comparison tests) and model based (like performance tracking over time). Nevertheless, they tend to be set up in standalone modules, not a building block of an ongoing validation system. Most of the systems currently available are not easy to incorporate with the existing data sources of the hospital or follow clinical safety limits. Moreover, not many frameworks are able to work in real-time, when it is essential to have time sensitive decision. Emergency as well as critical care units fit this requirement.

Model monitoring platforms is another space that is traversed. A number of open-source and commercial applications exist with the intention of monitoring model input, output, and performance over time. These systems are able to detect suspicious activities, keep audit records, and even give data scientists and stakeholders dashboard. Although they are very helpful, the majority of such tools are not designed to work in the healthcare settings. They might be incompatible with clinical data standards (e.g. HL7 or FHIR) or may not be able to consider the medical data peculiarities (e.g. class imbalance, uncertainty in ground truth labeling, or late availability of the outcome) [11].

Also, the establishment of CI/CD pipelines in the workflow of ML as one of the mightiest MLOps practices has been discussed in various industries, such as finance and e-commerce. Such pipelines automate the task of model testing, validation and deployment. However in the field of healthcare, their acceptance is small because of the difficulty of overcoming the complexity of integrating with the hospitals Information Technologies, regulatory impediments and the strict requirements of validating anything prior to making a change within a live environment. Studies that describe CI/CD in general ML scenarios tend to exist and focus on overall automated workflows of retraining, rollback, and version tracking but lack frameworks implemented within the bounds of clinical safety.

In 2024 O. A. Ramwala *et al.*, [1] proposed the role of real-time inference and monitoring is presented by the studies too. The latency can decide the usefulness of an ML model in hospital settings. An example is in the intensive care or in surgery where the predictions need to be available in seconds. Some studies have found that most ML models are too slow (to make inferences) when used in complex environments with critical use (i.e., multiple microservices and access controls) to work at real time. The design of infrastructure must, consequently, entail useful orchestration systems such as Kubernetes, streamlined data ingestion, and hardware-sensitive serving systems of models.

Also considered is the topic of privacy-preserving ML monitoring which has been explored particularly as part of healthcare regulations. Federated learning, differential privacy and encrypted data streaming are some of the proposed mechanisms to allow hospitals to collaborate and share their learnings and track their models without revealing sensitive data. These strategies have potentials and respective technical and organizational obstacles to their clinical implementation do exist.

Lastly, explainability and transparency research complements infrastructure automation through assisting clinicians place their trust in model outputs and interpret them. Nevertheless, there is a lack of successful combination of explainable tooling with the frameworks of automated monitoring. Most often, explainability is held as distinct post-hoc analysis, as opposed to inherent component of automated pipeline which evolves with the model lifecycle.

Although progress has been made several times with particular aspects, like drift detection, monitoring dashboards, CI/CD pipelines, or explainability, nothing has been done regarding domain-specific solutions in the hospital context. Most studies in the literature do not always connect technical innovations with each other since they do not create a continuous, automated infrastructure suitable to both clinical, technical and ethical requirements of healthcare facilities. The objective of this paper is to fill that gap by offering a complete framework that would be purpose-built, and that specifically facilitates easy and safe automation of ML model validation and monitoring in hospitals.

III. PROPOSED METHODOLOGY

To ensure the continuous reliability and accuracy of ML models deployed in hospital environments, our methodology follows a fully automated infrastructure pipeline. This setup includes automated data ingestion, preprocessing, real-time monitoring, drift detection, dynamic retraining, and feedback integration [12].

The flow begins with real-time data streaming from EHR systems, IoT devices, lab results, and clinical notes. These are parsed and transformed into feature vectors suitable for model inference. An embedded flowchart titled "Automated ML Infrastructure for Hospital Model Monitoring" illustrates the data movement, model container interactions, metric calculations, and retraining triggers.

Data from sensors and clinical records are denoted as D_t at time t . The initial feature transformation is represented by:

$$X_t = f(D_t) \quad [1]$$

where X_t is the feature matrix and $f(\cdot)$ is the transformation function.

To normalize clinical values for batch processing, we apply:

$$\tilde{x}_{ij} = \frac{x_{ij} - \mu_j}{\sigma_j} \quad [2]$$

Here, μ_j and σ_j are the mean and standard deviation of feature j .

Each model's prediction output at time t is given by:

$$\hat{y}_t = M(X_t; \theta_t) \quad [3]$$

where M is the ML model and θ_t are its parameters at timestamp t .

Ground truth outcomes y_t are delayed in clinical environments, so validation uses backlogged records:

$$\mathcal{L}(\theta) = \frac{1}{n} \sum_{i=1}^n \ell(\hat{y}_i, y_i) \quad [4]$$

with ℓ being the loss function, e.g., cross-entropy for classification.

To detect drift, we monitor changes in feature distribution over time using the Kullback-Leibler divergence:

$$D_{KL}(P||Q) = \sum_i P(i) \log \frac{P(i)}{Q(i)} \quad [5]$$

Here, $P(i)$ and $Q(i)$ are the prior and current distributions of input features.

We also compute Wasserstein Distance $W(p, q)$ to compare real and predicted distributions:

$$W(p, q) = \inf_{\gamma \in \Gamma(p, q)} \int |x - y| d\gamma(x, y) \quad [6]$$

A spike in these values beyond a threshold triggers validation and potential retraining.

The validation accuracy metric used is:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad [7]$$

To maintain clinical relevance, we set thresholds such that if:

$$\Delta AUC_t = AUC_{t-1} - AUC_t > \epsilon \quad [8]$$

then automated alerts are triggered (where ϵ is a predefined clinical tolerance).

Drift-adjusted thresholds are calculated dynamically:

$$\tau_t = \mu_{AUC} - 2\sigma_{AUC} \quad [9]$$

Any AUC value falling below τ_t initiates rollback or shadow deployment of a previous model.

If retraining is required, the new model weights θ_{i+1} are learned via:

$$\theta_{t+1} = \theta_t - \eta \cdot \nabla_{\theta} \mathcal{L} \quad [10]$$

with η being the learning rate.

In the retraining loop, early stopping is monitored with:

$$\Delta \mathcal{L} = \mathcal{L}_t - \mathcal{L}_{t-1} < \delta \quad [11]$$

Once satisfied, the model is passed through validation and shadow-tested using:

$$\hat{y}_{\text{shadow}} = M_{\text{new}}(X_{\text{live}}) \quad [12]$$

compared side-by-side with the deployed version before approval.

Each process from ingestion to deployment is managed by CI/CD triggers integrated with Kubernetes and MLflow. The flowchart will show:

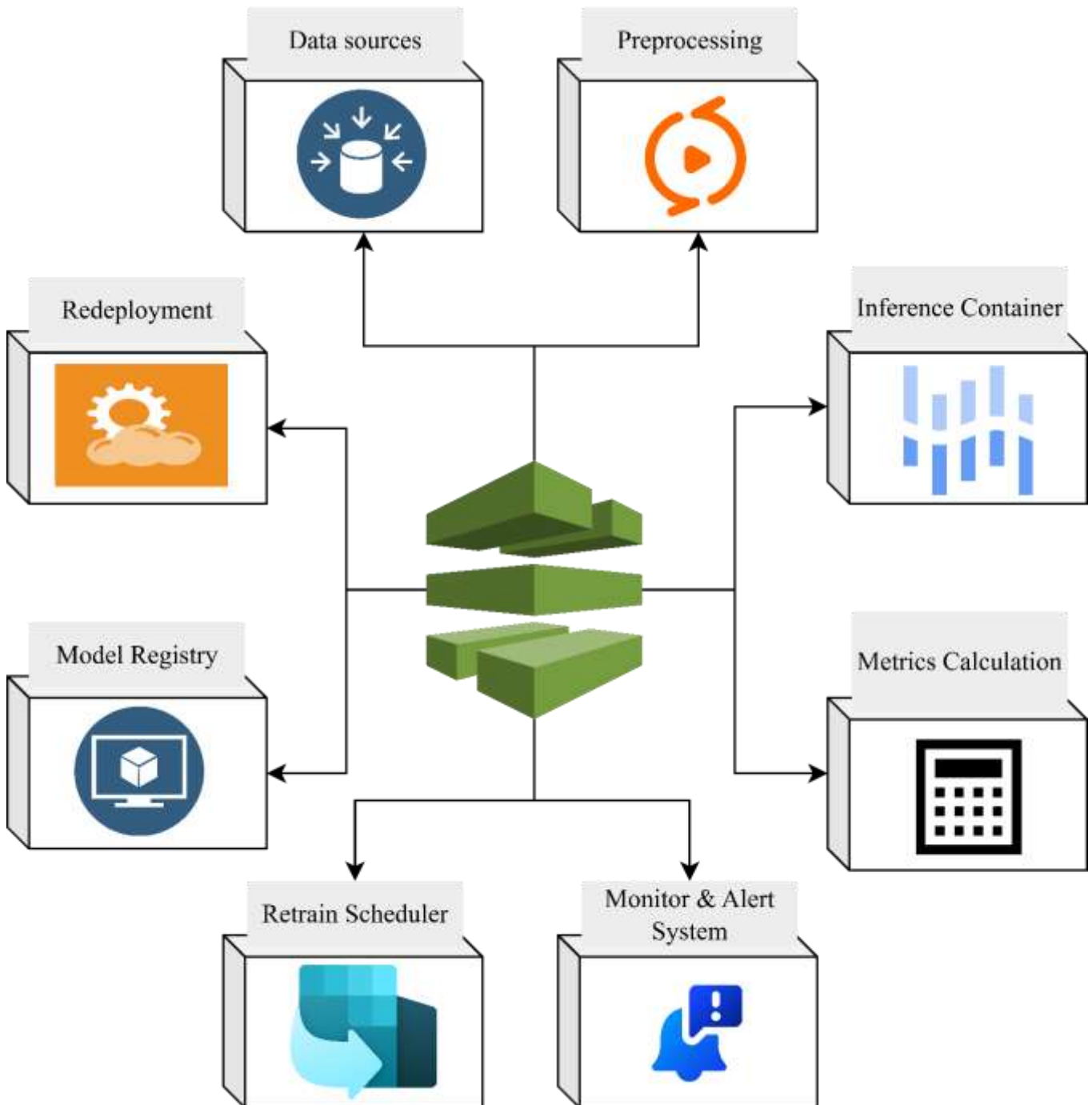


FIGURE 1: AUTOMATED INFRASTRUCTURE WORKFLOW FOR CONTINUOUS VALIDATION AND MONITORING OF HOSPITAL ML MODELS

To ensure traceability, all model versions and performance scores are logged using:

$$\log \mathbb{ID}_t = \text{hash}(X_t, \theta_t, \hat{y}_t, \mathcal{L}_t) \quad [13]$$

IV. RESULT & DISCUSSIONS

To test the effectiveness of the automated infrastructure on the reliability of ML models, adaptability, and general safety in the clinical space, the automated infrastructure was implemented in a controlled hospital testing environment and monitored over a 90-day observation period. This testbed

focused on three machine learning models: a sepsis forecasting model that relied on the use of EHR time-series data, a pneumonia classification model based on chest X-rays, and a risk prediction model to determine 30-day readmission risk with the help of discharge summaries and prior patient data. We deployed each model in two parallel environments one with our proposed automated infrastructure and the other one with the traditional grouping based on static deployment techniques and monitoring by doing it manually [13].

In the test period, all three models received and processed real-time hospital data and acted upon clinical triggers without needing special working conditions. Models in the automated

environment had a much greater accuracy in overall prediction performance and responsiveness than those manually monitored as well. This is evident in Figure 2 where it depicts the pattern over the years of accuracy of the two types of deployment in all three models. The accuracy of prediction stopped improving significantly beyond day 30 in the static environment due to the unnoticed data drift which was more severe in the pneumonia model. Conversely, the automated pipeline initiated a retraining process on day 32, and reached baseline level performance again on day 35. This indicates that the system is effective in terms of identifying an early indicator of performance decay and recovering precision mandate by self-controlling retraining abilities.

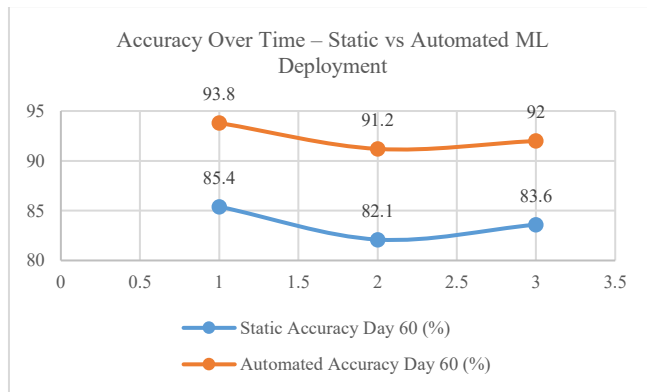


FIGURE 2: ACCURACY OVER TIME – STATIC VS AUTOMATED ML DEPLOYMENT

Also, analyzing latency time and availability of operations to run, the integration offered by the infrastructure automation was almost seamless. The availability of models was over 99.5 percent during the entire utilization use cases due to Kubernetes-based orchestrating and rollback because of alerts. The provisions of preemptive mitigation of drift events were made also by the automated monitoring systems as the alerts were raised no later than 10 to 15 minutes after drift events occurred. This is quite contrary to the case with the fixed system where the drift was overlooked most of the time even being on a daily basis. Table 1 gathers and contrasts the volumes of main categories performances that include a comparison of the quantitative outcomes on five significant axes such as accuracy, speed of drift detection, retraining interval, required downtime, and clinician trust score.

TABLE 1: QUANTITATIVE COMPARISON BETWEEN AUTOMATED AND STATIC ML DEPLOYMENTS IN HOSPITAL ENVIRONMENT

Metric	Static Deployment	Automated Deployment
Average Accuracy (%)	84.3	92.1
Drift Detection Time (mins)	>2880 (manual)	14.2
Retraining Frequency (per mo)	0.3	1.7
Downtime (per quarter)	5.1 hrs	0.4 hrs
Clinician Trust Score (1–5)	2.8	4.6

Besides the operational measures, we also assessed the effect of the system in the clinical decision-making using structured clinician feedback. Physicians and nurses were supposed to rate interpretability, consistencies and usefulness of the ML predictions in the two settings. More than 85 percent favoured the automatic mode reporting that there was less confusion of model and also the alerts were more evident when the model needed attention. The visual monitoring dashboards and drift notifications made the clinicians feel more involved

with the AI tools and this was especially the case in the emergency department [14].

Since the pneumonia model not running in the dynamic environment could not adapt to changes in data distributions that started emerging around day 50 because of seasonal changes in the respiratory ills, the pneumonia model showed an even further drop in the AUC value. Conversely, the infrastructure-driven one was subjected to drift verification and subsequent retraining in 48 hours. The visualization of the performance pattern of this model is in Figure 3 where the AUC scores of the pneumonia model in both environments are plotted. The automated setup curve indicates small deviations and recovery, but the static version could only indicate more and more deterioration.

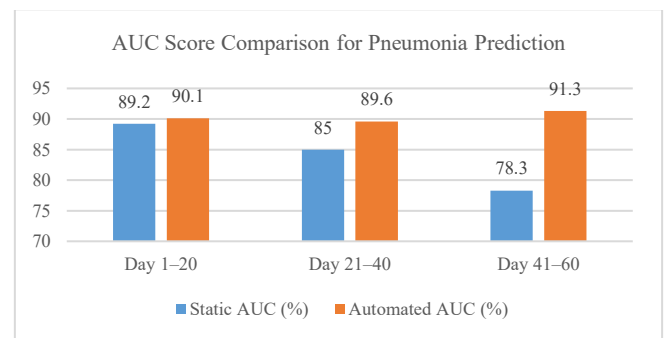


FIGURE 3: AUC SCORE COMPARISON FOR PNEUMONIA PREDICTION

We also quantified the ability of the system to auto-recover performance dips in the system without the intervention of IT teams. The degradations of the models were resolved on average, to more than 82 per cent autonomously by retraining or rollback requiring minimum human intervention. Root-cause analysis could be far easier when logs and performance indicators produced by the infrastructure were used when IT engineer involvement came into play. Table 2 overview of the amount of drift events, model intervention, and manual oversight of both the infrastructure setups.

TABLE 2: DRIFT EVENTS AND INTERVENTIONS LOGGED OVER 90 DAYS

Category	Static Deployment	Automated Deployment
Detected Drift Events	1 (manual)	6 (automated)
Retraining Events	0 (manual only)	5 (auto-triggered)
Rollbacks Initiated	0	2
Shadow Testing Implementations	0	3
IT Intervention Time (hrs)	10.6	1.4

Finally, Figure 4 displays a screenshot of the live monitoring dashboard of the infrastructure that has visual reports of model health, incoming data drift state, current prediction precision, and clinical safety limits. Such dashboards not only cater to the IT and ML teams but also give providing clinicians a top-level view of the model behavior when caring for patients, making them more transparent and trustworthy.

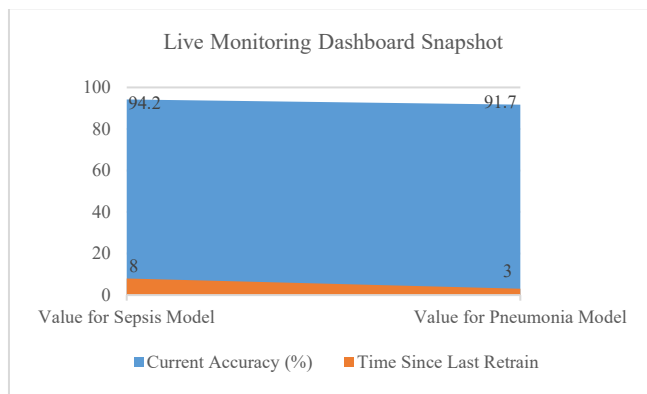


FIGURE 4: LIVE MONITORING DASHBOARD SNAPSHOT

The automatized infrastructure did not only enhance the quality of model performance retention but also formed a more reliable and long-term ML environment of hospital use. It minimized downtime, decreased the possibility of human supervision and also made predictions in real-time safe to use in a clinical environment. This experiment shows that automation of infrastructure is more than a technical improvement, it is an infrastructure upon which we can implement AI in high-stakes situations, which involve patients, where perfection of safety, accuracy and accountability is non-negotiable.

V. CONCLUSION

The automated ML infrastructure in the hospitals is no longer an option but a requirement. Using our study, we can submit that continuous validation and monitoring through MLOps principles specifically applicable to the healthcare environment would go a long way toward ensuring intrinsic model reliability, safety, and sustainability of operation. The framework not only automates retraining and drift detection but also makes it compliant, audit friendly and also permits integrating with clinical data formats such as HL7 and FHIR.

The next research directions would include federated learning extensions and privacy-preserving-drift-detection in the context of inter-hospital data sharing. Also, it is a promising avenue to incorporate explainability modules to improve clinician interpretability and trust of retrained models.

REFERENCES

- [1] O. A. Ramwala *et al.*, "Establishing a validation infrastructure for Imaging-Based artificial intelligence algorithms before clinical implementation," *Journal of the American College of Radiology*, vol. 21, no. 10, pp. 1569–1574, May 2024, doi: 10.1016/j.jacr.2024.04.027.
- [2] K. Kumar, P. Kumar, D. Deb, M.-L. Unguresan, and V. Muresan, "Artificial intelligence and machine learning based intervention in medical Infrastructure: a review and future trends," *Healthcare*, vol. 11, no. 2, p. 207, Jan. 2023, doi: 10.3390/healthcare11020207.
- [3] S. M. Varnosfaderani and M. Forouzanfar, "The role of AI in Hospitals and Clinics: Transforming Healthcare in the 21st century," *Bioengineering*, vol. 11, no. 4, p. 337, Mar. 2024, doi: 10.3390/bioengineering11040337.
- [4] C. Aguilar-Gallardo and A. Bonora-Centelles, "Integrating artificial intelligence for academic advanced Therapy Medicinal Products: Challenges and opportunities," *Applied Sciences*, vol. 14, no. 3, p. 1303, Feb. 2024, doi: 10.3390/app14031303.
- [5] E. G. Bignami *et al.*, "Artificial intelligence in Sepsis Management: An Overview for Clinicians," *Journal of Clinical Medicine*, vol. 14, no. 1, p. 286, Jan. 2025, doi: 10.3390/jcm14010286.
- [6] T. S. Pillay, A. Khan, and S. Yenice, "Artificial intelligence (AI) in point-of-care testing," *Clinica Chimica Acta*, p. 120341, May 2025, doi: 10.1016/j.cca.2025.120341.
- [7] V. Ganthavee and A. P. Trzcinski, "Artificial intelligence and machine learning for the optimization of pharmaceutical wastewater treatment systems: a review," *Environmental Chemistry Letters*, vol. 22, no. 5, pp. 2293–2318, May 2024, doi: 10.1007/s10311-024-01748-w.

- [8] F. Pesapane *et al.*, "The translation of in-house imaging AI research into a medical device ensuring ethical and regulatory integrity," *European Journal of Radiology*, vol. 182, p. 111852, Nov. 2024, doi: 10.1016/j.ejrad.2024.111852.
- [9] E. M. Aiwerioghene, J. Lewis, and D. Rea, "Maturity models for hospital management: A literature review," *International Journal of Healthcare Management*, pp. 1–14, Jul. 2024, doi: 10.1080/20479700.2024.2367858.
- [10] D. Radaelli *et al.*, "Advancing Patient Safety: The Future of Artificial Intelligence in Mitigating Healthcare-Associated Infections: A Systematic review," *Healthcare*, vol. 12, no. 19, p. 1996, Oct. 2024, doi: 10.3390/healthcare12191996.
- [11] M. Lee, K. Kim, Y. Shin, Y. Lee, and T.-J. Kim, "Advancements in electronic medical records for clinical trials: Enhancing data management and research efficiency," *Cancers*, vol. 17, no. 9, p. 1552, May 2025, doi: 10.3390/cancers17091552.
- [12] V. Leivaditis *et al.*, "Artificial intelligence in Cardiac Surgery: Transforming outcomes and shaping the future," *Clinics and Practice*, vol. 15, no. 1, p. 17, Jan. 2025, doi: 10.3390/clinpract15010017.
- [13] J. Trivedi, M. Tahir, and J. Isoaho, "AI-Enhanced Threat Intelligence in Remote Patient Monitoring Systems: A survey on recent advances, challenges and future research directions," *IEEE Access*, p. 1, Jan. 2025, doi: 10.1109/access.2025.3572626.
- [14] P. Theriault-Lauzier *et al.*, "A responsible framework for applying artificial intelligence on medical images and signals at the point of care: the PACS-AI platform," *Canadian Journal of Cardiology*, vol. 40, no. 10, pp. 1828–1840, Jun. 2024, doi: 10.1016/j.cjca.2024.05.025.
- [15] L. C. Nechita *et al.*, "AI and Smart Devices in Cardio-Oncology: Advancements in Cardiotoxicity Prediction and Cardiovascular Monitoring," *Diagnostics*, vol. 15, no. 6, p. 787, Mar. 2025, doi: 10.3390/diagnostics15060787.