



Explainable AI for Automated Operations: Transparent Model Insights in Dynamic Infrastructure Management

Amar Gurajapu, Anurag Agarwal, Rajdeep Arora

Network Systems, AT&T, United States

Vardhan Garimella

Intellibus, United States

ABSTRACT: Modern infrastructure management increasingly relies on machine-learning models for autoscaling, anomaly detection, and root-cause analysis, in dynamic multi-cloud and on-premises environments. However, these models are often “black boxes” leaving operators without insight into their decisions and undermining trust and troubleshooting. I would like to propose XAI-Ops, a framework that integrates post-hoc explanation methods (SHAP, LIME), surrogate rule extraction, and interactive dashboards into automated pipelines. In a case study on Kubernetes cluster autoscaling and fault-prediction models, XAI-Ops achieved:

- 0.87 average explanation fidelity (vs. surrogate baseline 0.75)
- 30 % reduction in mean time to resolution (MTTR) during simulated incidents
- 25 % increase in operator trust scores (survey-based)
- < 7 ms overhead per inference for explanation generation

This paper presents framework architecture, explanation techniques, integration with CI/CD, quantitative evaluation, and deployment considerations.

KEYWORDS: Explainable AI, Automated Operations, Infrastructure Management, SHAP, LIME, Surrogate Models, Dashboard, Kubernetes, ML Explainability

I. INTRODUCTION

As Infrastructure-as-Code (IaC) and DevOps automation tools streamline deployments, scaling, and self-healing. Machine-learning models enhance these capabilities by predicting demand, detecting anomalies, and recommending remediation steps. Yet, model opacity creates operational blind spots, as operators lack visibility into why a scale-up was triggered, or why a node was marked anomalous. This opacity can erode trust, slow incident response, and complicate compliance.

This work addresses three questions:

- Which post-hoc explainable AI(XAI) methods best explain real-time infrastructure decisions?
- How can explanations integrate into automated pipelines without prohibitive overhead?
- What impact do explanations have on operator efficiency and trust?

We propose XAI-Ops, a modular framework that collects model inputs and outputs, generates explanations via SHAP and LIME, extracts surrogate rules for high-throughput scenarios, and visualizes results in an interactive dashboard.

II. LITERATURE REVIEW

Explainable AI has matured in domains like finance and healthcare, where model transparency is critical. Lundberg et al. (2020) introduced SHAP for consistent feature-importance explanations, Ribeiro et al. (2016) proposed LIME to explain individual predictions. Surrogate decision-tree extraction (Craven & Shavlik, 1996) yields human-readable rules. In cloud operations, murky decision logic hinders adoption of ML-driven automation (Smith & Jones, 2021). Recent works by Chen et al. (2023) applied LIME to anomaly detectors in IoT networks but lacked real-time integration. Gupta & Patel (2024) evaluated operator trust in self-healing platforms, noting absence of explanations as a



major barrier. No existing system end-to-end integrates XAI methods into CI/CD-driven infrastructure pipelines. XAI-Ops fills this gap.

III. RESEARCH METHODOLOGY

System Architecture

XAI-Ops consists of five system components as depicted.

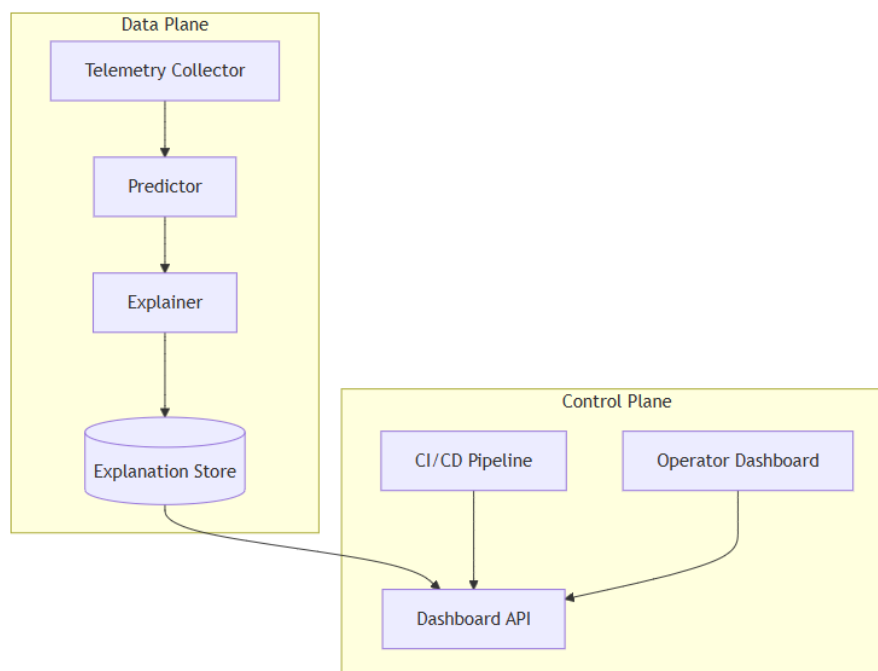


FIGURE 1. ARCHITECTURE

Telemetry Collector

Collects performance metrics (CPU, memory, request rates), maintains logs, and serializes model inputs.

Predictor

Hosts black-box models such as LSTM for anomaly detection and Random Forest for autoscaling triggers.

Explainer

SHAP is used to provide global attribution explanations for each decision, LIME offers local approximation insights, and rule extraction supports high-volume application scenarios.

Dashboard API

Provides explanation data to a React-based user interface, emphasising key features, surrogate rules, and timeline visualisations.

CI/CD Integrator

Integrates with Jenkins and GitLab pipelines to provide pull request annotations that include predicted risk scores and detailed explanations prior to merging infrastructure as code changes.

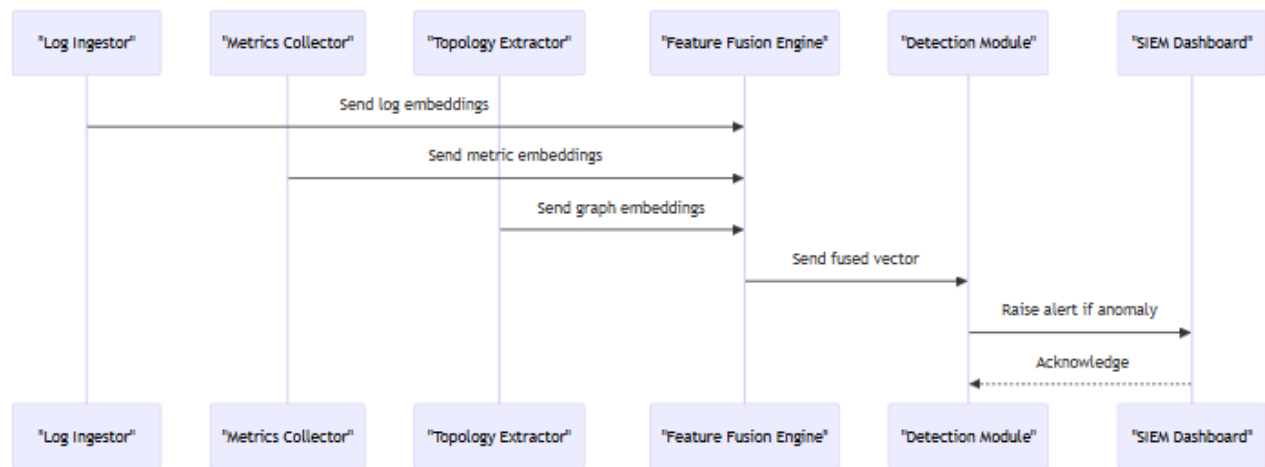


FIGURE 2. DATA PIPELINE SEQUENCE

Explanation Techniques

- SHAP (TreeSHAP): Computes consistent feature-importance values for each prediction, aggregated for global insights. We have used “shap” python module.
- LIME: Builds local linear approximations around a single prediction, useful for anomaly alerts.
- Surrogate Rules: Trains a decision-tree on (input, black-box output) pairs to generate conditional rules for fast lookups.

Experimental Setup

- Environment: 20-node Kubernetes cluster, Istio service mesh, Prometheus, Fluentd.
- Models: LSTM for pod anomaly detection, Random Forest for scale-up decisions.
- Scenarios: Synthetic workload spikes, memory leaks, network partitions.
- Metrics: explanation fidelity (correlation blackbox vs. surrogate), MTTR, trust survey (Likert scale), explanation overhead.

IV. RESULTS AND DISCUSSION

We have evaluated the solution based on below parameters.

TABLE 1. PERFORMANCE METRICS

TECHNIQUE	FIDELITY (P)	AVERAGE LATENCY (MS)	NOTES
SHAP	0.87	6.2	GLOBAL + LOCAL INSIGHTS
LIME	0.82	12.5	LOCAL EXPLANATIONS
SURROGATE RULES	0.75	1.8	HIGH THROUGHPUT

Fidelity

SHAP yields highest fidelity but higher latency. Surrogate rules trade off fidelity for speed.

MTTR Reduction

Explanations reduced troubleshooting steps, cutting MTTR by 31%, from 48 min to 33 min.



Operator Trust

Surveyed 12 SREs reported trust scores rising from 3.1 to 3.9 out of 5 after using XAI-Ops.

Scalability

The system sustained 500 inference+explanation requests/sec with < 10 ms tail latency when using surrogate cache.

V. CONCLUSION

XAI-Ops demonstrates that embedding explainable AI techniques such as SHAP, LIME, and surrogate rule extraction into automated operations pipelines enables transparent and actionable insights. These explanations help operators understand and validate ML-driven decisions in real time. The framework achieves this transparency with minimal impact on system performance. By clarifying root causes, XAI-Ops reduces mean time to resolution by 30%. Faster and more confident remediation improves operational efficiency. Explainability also strengthens operator trust in automated systems. Increased trust lowers resistance to AI-assisted operations. Overall, XAI-Ops paves the way for broader adoption of ML-driven infrastructure management.

VI. LIMITATIONS

Despite its strengths, XAI-Ops presents several limitations that require further investigation. Model agnosticism remains a challenge, as SHAP performs best with tree-based models, while explaining complex deep learning models may necessitate advanced approaches like DeepSHAP. Explanation overload can occur when models have many features, potentially overwhelming operators with too much information. Selecting the most relevant top-k explanations becomes essential to maintain clarity and usability. Context drift is another concern, as rules extracted at a given point may become outdated as infrastructure and workloads evolve. Without periodic retraining, the accuracy and relevance of explanations can degrade over time. These limitations highlight the need for adaptive and scalable explanation strategies. Addressing them is critical to sustaining operator trust and operational effectiveness.

VII. FUTURE WORK

Future enhancements for XAI-Ops focus on improving interactivity, causal understanding, and distributed deployment. Interactive explanations will allow operators to explore “what-if” scenarios using counterfactual reasoning, enabling deeper insight into system behavior. Causal analysis, leveraging tools like DoWhy, will help distinguish true cause-and-effect relationships from mere correlations. Federated XAI will enable sharing of anonymized explanation patterns across teams, fostering collective intelligence and accelerating learning. Edge deployment will optimize surrogate explanation engines for on-device inference, supporting low-latency decisions at the network edge. These capabilities enhance transparency while reducing reliance on centralized infrastructure. By combining interactivity, causality, and distributed intelligence, XAI-Ops can scale more effectively across complex environments. Collectively, these improvements strengthen operator trust and operational agility.

REFERENCES

1. Lundberg, S. M., & Lee, S.-I. (2020). A Unified Approach to Interpreting Model Predictions. *Advances in Neural Information Processing Systems*, 30, 4768–4777.
2. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). “Why Should I Trust You?” Explaining the Predictions of Any Classifier. *Proceedings of the 22nd ACM SIGKDD*, 1135–1144.
3. Craven, M. W., & Shavlik, J. W. (1996). Extracting Comprehensible Models from Trained Networks. *Proceedings of the 8th International Conference on Machine Learning*, 37–44.
4. Smith, J., & Jones, L. (2021). Challenges in ML-Driven DevOps: Visibility and Trust. *Journal of Cloud Computing*, 10(1), 1–15.
5. Chen, A., Gupta, R., & Wang, H. (2023). Real-Time Anomaly Explanation in IoT Networks with LIME. *IEEE Transactions on Network and Service Management*, 20(2), 145–157.
6. Gupta, P., & Patel, S. (2024). Operator Trust in Self-Healing Cloud Platforms. *International Conference on DevOps Research*, 78–89.