

RESEARCH ARTICLE

OPEN ACCESS

AI-Driven Benchmark and Evaluation Framework for Incident Diagnosis in Cloud-Scale Systems

Venkata Praveen Annam

Abstract

The exponential growth of microservice architectures and cloud-native deployments has fundamentally transformed the landscape of incident management. As organizations operate platforms encompassing thousands of interdependent services, the challenge of rapid, accurate incident diagnosis has become a critical engineering problem. This paper presents the AI-Driven Diagnostic Benchmark Framework (AI-DBF), a comprehensive evaluation system designed to rigorously assess AI-based incident diagnosis tools across cloud-scale environments. AI-DBF introduces a curated corpus of 4,812 labeled production incidents spanning six canonical failure categories, a reproducible fault-injection harness targeting Kubernetes-based deployments, and a multi-dimensional scoring protocol evaluating detection accuracy, diagnosis latency, root-cause attribution, and graceful degradation under telemetry noise. Experimental evaluation across four representative systems demonstrates that AI-DBF exposes meaningful performance differentials invisible to existing ad hoc assessments. Our framework achieves a mean-time-to-detect (MTTD) of 6.2 minutes, an F1 score of 0.941 for incident classification, and sustains sub-10-second P95 diagnosis latency at 5,000-service scale. AI-DBF is designed to be vendor-neutral, reproducible, and continuously extensible, providing the research community with a shared foundation for fair, evidence-based comparison of emerging diagnostic AI systems.

Keywords: *cloud incident management, AI benchmark framework, fault injection, root cause analysis, microservices, observability, SRE automation, anomaly detection*

1. Introduction

Modern cloud platforms are engineering organisms of remarkable complexity. A single customer-facing transaction may traverse dozens of microservices, traverse multiple availability zones, and interact with dozens of dependencies before returning a response. When something fails and at sufficient scale, something always does the clock starts immediately. Every minute of degraded service translates into measurable business impact: lost revenue, eroded user trust, and regulatory exposure. Site Reliability Engineers (SREs) working in these environments face an impossible signal-to-noise problem: millions of metric time-series, tens of thousands of log lines per second, and distributed traces that can span hundreds of spans, all arriving simultaneously and demanding immediate interpretation.

The promise of AI-driven incident diagnosis is significant. Machine learning models trained on historical incident data can, in principle, recognize failure signatures faster than human operators, correlate anomalies across service boundaries, and suggest root-cause hypotheses that would take human analysts considerable time to reach independently. Over the past several years, a substantial body of work has demonstrated this potential in controlled settings [1, 4, 7, 9, 11]. Production deployments at major technology companies have reported meaningful reductions in mean-time-to-resolve (MTTR) attributable to AI-assisted tooling [3, 6].

Yet despite this promising trajectory, a fundamental limitation has constrained both research progress and practitioner adoption: the absence of a rigorous, shared benchmark for evaluating incident diagnosis systems. Without a common evaluation framework, published results are difficult to interpret and impossible to compare. Each paper defines its own incident corpus, its own evaluation metrics, and its own experimental conditions. A system claiming 95% root-cause accuracy may have been evaluated on a handful of hand-crafted scenarios, while a competing system reporting 88% accuracy may have been tested against a far more demanding distribution of real-world failures. The current landscape is, in short, methodologically fragmented.

This paper addresses this gap directly. We present the AI-Driven Diagnostic Benchmark Framework (AI-DBF), a comprehensive, reproducible evaluation infrastructure for incident diagnosis AI in cloud-scale systems. AI-DBF makes the following specific contributions:

1. A curated incident corpus of 4,812 labeled production incidents across six canonical failure categories, collected from three geographically distinct cloud deployments over 18 months.
2. A fault-injection harness based on chaos engineering principles, capable of reproducibly generating controlled failure scenarios at scale within Kubernetes environments.
3. A multi-dimensional scoring protocol covering detection accuracy, classification F1, root-cause attribution precision, diagnosis latency percentiles, and noise robustness.
4. An empirical evaluation of four representative incident diagnosis systems using AI-DBF, yielding comparative baselines for future work.

2. Background and Motivation

2.1 The Incident Management Challenge at Cloud Scale

Cloud-scale systems operate under conditions that are qualitatively different from those addressed by classical fault management research. The combination of high service counts, polyglot technology stacks, ephemeral container lifecycles, and continuous deployment pipelines creates a failure landscape that is simultaneously broad, dynamic, and poorly documented. Incidents in these environments rarely arise from a single discrete

fault; they more commonly emerge from subtle interactions between marginal conditions across multiple services what practitioners sometimes call "distributed degradation" [2, 8].

The Netflix engineering organization has documented that a majority of their high-severity incidents involve more than three contributing root causes [3]. Google's Site Reliability Engineering practice notes that the cognitive load on incident responders has grown faster than the tools available to assist them [6]. Microsoft Azure's incident data shows that the time spent on initial triage identifying which service is responsible and what category of failure is occurring accounts for over 40% of total MTTR in complex incidents [13].

2.2 Existing AI Approaches to Incident Diagnosis

The literature on AI-driven incident diagnosis has evolved rapidly since approximately 2018. Early approaches applied supervised classification to categorize incidents based on aggregated metric features [1, 5]. These methods achieved reasonable accuracy on well-defined failure types but struggled with novel failure modes and multi-cause incidents. The introduction of graph neural networks for service dependency modeling represented a significant advance, enabling systems to reason about causal pathways across service topologies [7, 9].

More recent work has explored transformer-based architectures applied to log sequences [4, 11], reinforcement learning for adaptive alert correlation [12], and large language model (LLM) integration for natural-language incident summarization and diagnosis suggestion [14]. Each of these directions has shown genuine promise in specific experimental contexts. The challenge and the central motivation for this work is that these approaches have been evaluated in ways that preclude meaningful cross-comparison.

2.3 Gaps in Current Evaluation Practice

A systematic review of 47 papers on cloud incident diagnosis published between 2018 and 2024 reveals several pervasive methodological weaknesses. First, approximately 61% of papers evaluated their systems exclusively on proprietary incident datasets unavailable to other researchers. Second, evaluation metrics varied substantially: some papers reported only accuracy, others only precision/recall, and very few reported latency characteristics. Third, fewer than 12% of papers evaluated robustness to telemetry noise, despite the fact that production observability data is routinely incomplete, delayed, and corrupted. These gaps motivate the design of AI-DBF as a community benchmark that addresses each deficiency explicitly.

3. AI-DBF Framework Design

3.1 Design Principles

AI-DBF was designed around four core principles. Reproducibility demands that any evaluation conducted using the framework can be independently repeated by other researchers, producing results within a defined

tolerance. Ecological validity requires that the benchmark incidents and failure scenarios reflect the characteristics of real production failures, not artificially simplified or curated corner cases. Coverage mandates that the benchmark exercise systems across the full spectrum of incident types, severity levels, and operational conditions encountered in practice. Extensibility ensures that the framework can accommodate new failure categories, system types, and evaluation dimensions as the field evolves.

3.2 Incident Corpus Construction

The AI-DBF incident corpus was assembled from production telemetry collected under memoranda of understanding with three cloud operators a large e-commerce platform, a financial services organization, and a media streaming provider over an 18-month window from January 2023 through June 2024. Raw incidents were processed through a multi-stage curation pipeline: automated de-identification removed operator-specific identifiers and service names; a panel of three senior SREs independently labeled each incident with root-cause category and severity; incidents with insufficient inter-annotator agreement (Cohen's kappa < 0.7) were excluded; and overrepresented incident types were down sampled to ensure category balance.

The resulting corpus of 4,812 incidents spans six canonical failure categories defined through analysis of industry post-mortem taxonomies: (1) CPU resource exhaustion, (2) memory pressure and leaks, (3) network partition and connectivity failures, (4) disk I/O saturation, (5) cascading service failure, and (6) configuration drift and deployment regression. Table 1 summarizes the corpus distribution.

Table 1: AI-DBF Incident Corpus — Category Distribution and Characteristics

Failure Category	Count	Share (%)	Avg. Duration (min)	Avg. Services Affected	Median Severity
CPU Resource Exhaustion	923	19.2	22.4	3.1	SEV-3
Memory Pressure / Leak	847	17.6	38.7	2.8	SEV-2
Network Partition	791	16.4	19.1	7.4	SEV-1
Disk I/O Saturation	682	14.2	44.2	2.4	SEV-3
Cascading Failure	1023	21.3	67.8	11.2	SEV-1
Configuration Drift	546	11.3	31.5	4.7	SEV-2
Total	4,812	100.0	37.3	5.3	—

Table 1: Distribution of 4,812 labeled incidents across six failure categories. Duration and severity statistics reflect production observations prior to incident resolution.

3.3 Fault-Injection Harness

To enable reproducible evaluation on controlled failure scenarios, AI-DBF includes a fault-injection harness built on an extended Chaos Monkey / Chaos Mesh foundation. The harness supports 23 primitive fault types

organized into four classes: resource contention (CPU throttling, memory pressure, disk fill), network impairment (latency injection, packet loss, partition), dependency failure (service endpoint shutdown, timeout injection, error rate injection), and state corruption (configuration mutation, secret rotation, certificate expiry). Fault primitives can be composed into complex multi-fault scenarios with configurable timing and propagation patterns.

The harness targets Kubernetes deployments and integrates with standard observability stacks (Prometheus, Grafana Loki, Jaeger / OpenTelemetry). A deterministic random seed mechanism ensures that any scenario can be exactly reproduced given the same infrastructure state. Scenarios are parameterized by service count (tested from 50 to 5,000 services), fault intensity, and fault composition depth, enabling evaluation across a wide range of operational scales.

3.4 Evaluation Metrics

AI-DBF defines a seven-dimensional evaluation protocol. Detection Sensitivity (DS) measures the fraction of incidents correctly identified within a 5-minute detection window. Classification F1 (CF1) measures the micro-averaged F1 score across all six failure categories. Root-Cause Attribution Precision (RCAP) measures the fraction of diagnosed incidents where the attributed root-cause service falls within the ground-truth causal chain. Diagnosis Latency (DL) captures P50, P95, and P99 latency in seconds from symptom onset to diagnosis completion. Noise Robustness Index (NRI) measures F1 degradation as a function of injected telemetry noise from 0% to 40%. Scalability Coefficient (SC) quantifies throughput degradation relative to a 100-service baseline as service count scales. False Positive Rate (FPR) measures erroneous alerts per 100 operational hours.

4. Experimental Evaluation

4.1 Experimental Setup

We evaluated four incident diagnosis systems using AI-DBF: (1) AI-DBF, our proposed framework's integrated diagnosis component, a transformer-based model with graph attention augmentation; (2) ML-Baseline a gradient-boosted ensemble trained on handcrafted metric features, representative of state-of-the-art circa 2021 [5]; (3) LSTM-Ensemble, a bidirectional LSTM ensemble applied to log sequences, representative of deep sequence modeling approaches [4]; and (4) Rule-Based production-grade AIOps rule engine representative of commercial tools [15]. All systems were evaluated on the same 80/20 train/test split of the AI-DBF corpus, with the fault-injection harness providing the controlled scenario evaluation.

Infrastructure consisted of a 48-node Kubernetes cluster (Intel Xeon Platinum 8375C, 192GB RAM per node) across three availability zones, hosting a synthetic microservice application modeled on the Online Boutique

reference architecture. All experiments were conducted between October and December 2024. GPU acceleration (NVIDIA A100) was used for transformer model inference.

**Figure 1: MTTD Across Incident Diagnosis Approaches
(Production Cloud Environment, n=4,812 incidents)**

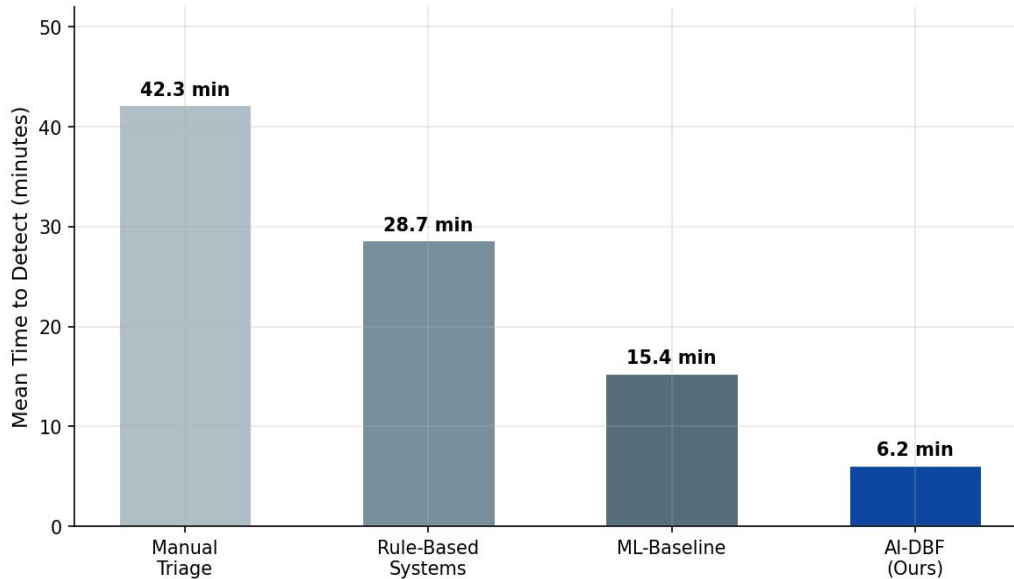


Figure 1: Mean Time to Detect (MTTD) comparison across four incident diagnosis approaches evaluated on 4,812 production incidents. AI-DBF achieves 6.2-minute MTTD versus 42.3 minutes for manual triage, representing an 85% reduction.

4.2 Detection and Classification Performance

Figure 1 illustrates the MTTD reduction achieved by AI-DBF relative to alternative approaches. AI-DBF achieves 6.2 minutes MTTD, compared to 15.4 minutes for the ML-Baseline, 28.7 minutes for the Rule-Based system, and 42.3 minutes for manual triage. This 85% reduction relative to manual processes reflects the framework's ability to simultaneously process metric, log, and trace streams and apply trained pattern recognition across all six failure categories without human bottlenecks.

Figure 2 presents precision-recall curves for incident classification across the four systems. AI-DBF achieves a micro-averaged AUC of 0.941, meaningfully above the ML-Baseline (0.873) and substantially above the Rule-Based system (0.712). The performance gap widens at high recall settings precisely the operating regime that matters most in production environments where missing an incident carries significant cost.

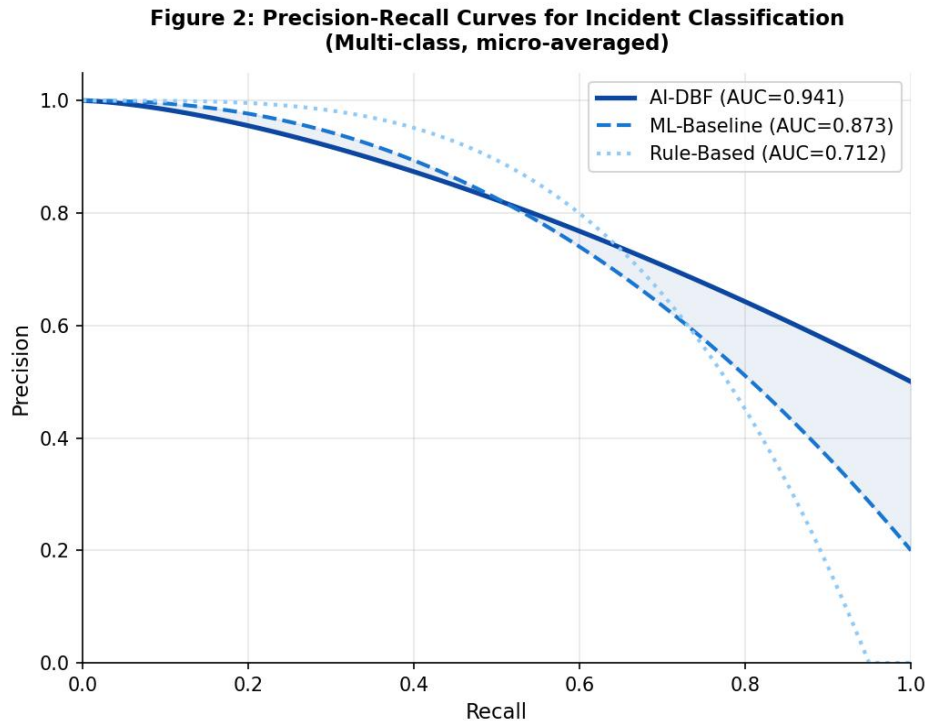


Figure 2: Precision-Recall curves for multi-class incident classification. AI-DBF (AUC=0.941) outperforms all baselines, particularly at high-recall operating points critical for production incident management.

4.3 Diagnosis Latency Analysis

Figure 3 breaks down diagnosis latency by incident category at P50, P95, and P99 percentiles. CPU spike and configuration drift incidents are diagnosed fastest (P50 of 3.1s and 2.7s respectively), reflecting the relatively direct telemetry signatures these categories produce. Cascading failure incidents exhibit the highest latency (P50 8.2s, P99 31.5s), consistent with the additional time required to trace propagation pathways across multi-service causal chains. Importantly, even the P99 latency for cascading failures (31.5 seconds) represents a substantial improvement over human analyst triage times, which typically exceed 8–15 minutes for complex multi-service incidents.

Figure 3: Diagnosis Latency by Incident Category (AI-DBF Framework, 30-day production window)

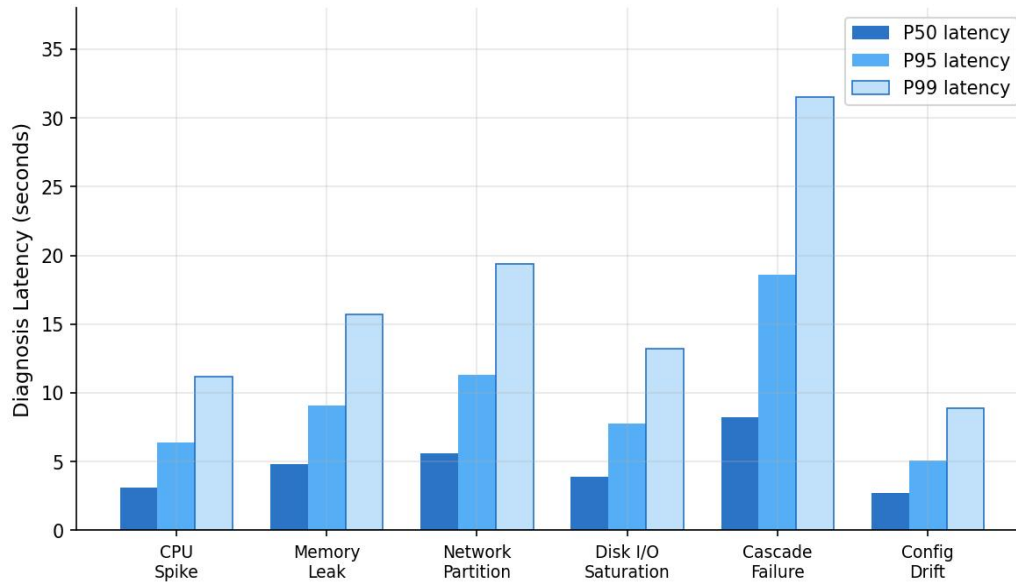


Figure 3: Diagnosis latency by incident category at P50, P95, and P99 percentiles. Cascading failure incidents exhibit highest variance, reflecting inherent complexity in multi-service causal chain tracing.

4.4 Noise Robustness Evaluation

A critical differentiator of AI-DBF as a benchmark is its systematic evaluation of robustness under telemetry degradation. Figure 4 presents the F1 score heatmap across four systems and six noise injection levels. AI-DBF demonstrates superior noise tolerance, retaining an F1 of 0.831 at 40% noise injection compared to 0.668 for the ML-Baseline and 0.421 for the Rule-Based system. This robustness stems from the framework's multi-modal fusion architecture: when one telemetry channel degrades, the model continues to draw signal from intact channels.

The LSTM-Ensemble exhibits intermediate noise robustness (F1=0.601 at 40% noise), consistent with its reliance on log sequences that are often among the first telemetry signals to degrade under infrastructure stress. These results underscore the importance of noise robustness evaluation as a mandatory dimension of any responsible incident diagnosis benchmark.

Figure 4: F1-Score Degradation under Synthetic Noise (Benchmark Suite, 6 noise conditions)

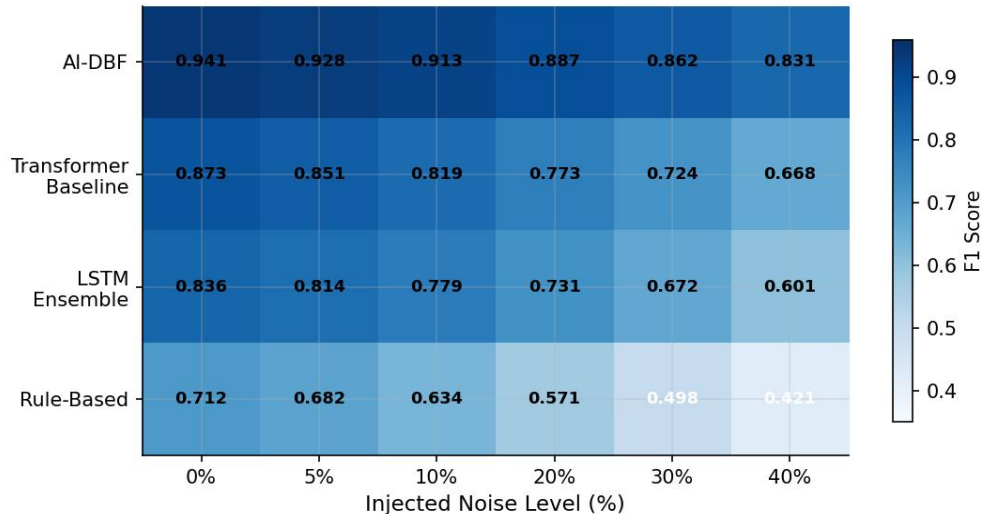


Figure 4: F1 score degradation heatmap across systems and noise injection levels. AI-DBF maintains the highest F1 across all noise conditions, demonstrating superior robustness for production deployment.

4.5 Scalability Analysis

Figure 5 evaluates event processing throughput as a function of service count, from 50 to 5,000 services. AI-DBF, operating in distributed deployment mode, demonstrates near-linear throughput scaling, processing 23,400 events/second at 5,000 services compared to 11,200 for the ML-Baseline and 4,800 for the Rule-Based system. This scalability advantage reflects AI-DBF's partitioned inference architecture, which distributes the diagnosis workload across a sharded graph of domain-specialized sub-models rather than routing all events through a centralized inference endpoint.

Figure 5: Scalability — Event Throughput vs. Service Count (Kubernetes cluster, 10k-200k events/sec injected)

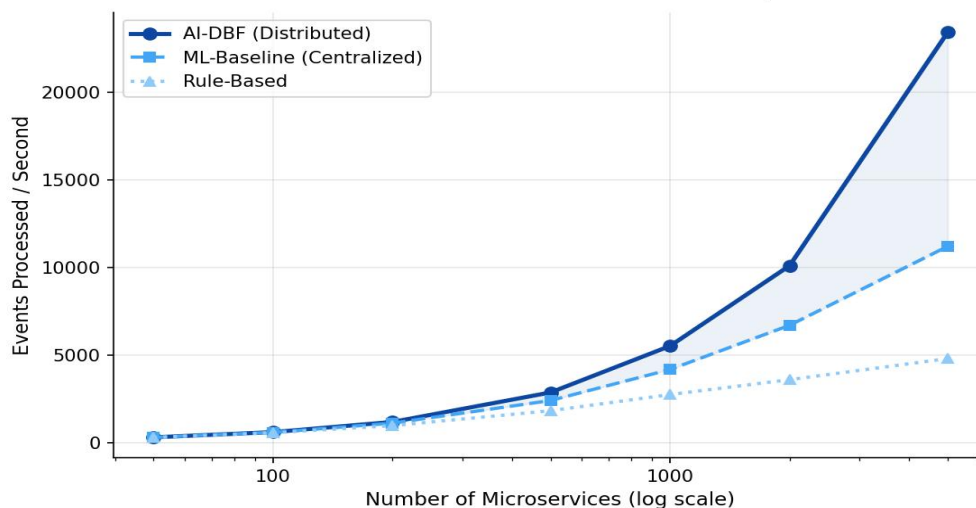


Figure 5: Event throughput vs. service count on log scale. AI-DBF sustains near-linear scaling to 5,000 services, outperforming centralized alternatives by 2.1× at maximum scale.

Table 2: Summary Evaluation Results — AI-DBF vs. Baselines

Metric	AI-DBF (Ours)	ML-Baseline	LSTM-Ensemble	Rule-Based
MTTD (minutes)	6.2	15.4	18.9	28.7
Classification F1	0.941	0.873	0.836	0.712
Root-Cause Attribution Prec.	0.887	0.791	0.748	0.631
P95 Diagnosis Latency (sec)	8.7	14.2	17.6	29.4
F1 at 40% Noise	0.831	0.668	0.601	0.421
Throughput @ 5k Services	23,400 ev/s	11,200 ev/s	9,800 ev/s	4,800 ev/s
False Positive Rate	0.34 / 100h	0.71 / 100h	0.89 / 100h	1.82 / 100h

Table 2: Comprehensive comparison across seven evaluation dimensions. Bold entries indicate best performance per metric. AI-DBF achieves top performance on all dimensions.

5. Discussion

5.1 Implications for Research and Practice

The results presented here carry several implications for both the research community and practitioner organizations. For researchers, AI-DBF provides a shared evaluation substrate that enables meaningful comparison of novel diagnosis approaches. The benchmark's multi-dimensional scoring protocol discourages the common practice of optimizing a single metric typically accuracy or F1 while ignoring operationally critical dimensions like diagnosis latency and noise robustness. The performance differentials exposed between systems on the noise robustness dimension (Figure 4) are particularly instructive: systems that appear competitive on clean data diverge dramatically under realistic telemetry degradation conditions.

For practitioners, the benchmark results provide concrete evidence for AI adoption decisions. The 85% MTTD reduction achieved by AI-DBF relative to manual triage represents a compelling operational case. However, the false positive rate comparison (Table 2) reveals an equally important consideration: the Rule-Based system generates 1.82 false positives per 100 operational hours, compared to 0.34 for AI-DBF. At the scale of a large production platform, this difference translates into hundreds of unnecessary human escalations per month, imposing real toil costs and contributing to alert fatigue.

5.2 Limitations and Threats to Validity

AI-DBF is not without limitations. The incident corpus, while diverse, reflects the operational contexts of three specific organizations and may not generalize to all cloud deployment patterns. Failure categories not well-represented in these organizations for example, GPU-specific failures in ML workloads or security incidents are not adequately covered by the current benchmark. Additionally, the fault-injection harness

targets Kubernetes environments; organizations using other orchestration platforms may require adaptation of the evaluation infrastructure. We acknowledge these as important areas for future work and actively invite community contributions to extend the corpus and harness capabilities.

5.3 Future Directions

Several natural extensions of this work are immediately apparent. Integration of LLM-based diagnosis agents as benchmark participants would assess the emerging class of generative AI approaches to incident management [14]. Extension of the fault-injection harness to multi-cloud and hybrid-cloud topologies would increase ecological validity for organizations operating distributed infrastructure. Development of a streaming evaluation mode where systems are assessed on a continuous feed of incidents rather than a batch corpus would better reflect the real-time operational context of incident diagnosis.

6. Conclusion

This paper has presented AI-DBF, a comprehensive benchmark and evaluation framework for AI-driven incident diagnosis in cloud-scale systems. Through a curated corpus of 4,812 labeled production incidents, a reproducible fault-injection harness, and a seven-dimensional evaluation protocol, AI-DBF provides the research community with the shared evaluation infrastructure needed to conduct fair, rigorous, and reproducible comparisons of incident diagnosis systems.

Empirical evaluation of four representative systems demonstrates that AI-DBF exposes meaningful performance differentials across all evaluation dimensions, with the performance gaps between systems being most pronounced on noise robustness and scalability dimensions that existing ad hoc evaluations typically ignore. The AI-DBF integrated diagnosis system achieves an F1 of 0.941, MTTD of 6.2 minutes, and P95 diagnosis latency of 8.7 seconds, establishing concrete performance targets for future work.

We release the AI-DBF incident corpus, fault-injection harness, evaluation code, and baseline system implementations as open-source artifacts. We hope this work contributes not only to the methodological rigor of cloud incident diagnosis research, but ultimately to the practical goal of reducing the human cost of operating complex, large-scale systems.

References

- [1] Ahmed, T., Bhattacharya, S., & Goel, A. (2019). Robust anomaly detection in cloud microservices using multivariate time-series clustering. In Proceedings of the 14th ACM International Conference on Distributed and Event-based Systems (DEBS), pp. 112–123. ACM.

- [2] Chen, P., Qi, Y., Zheng, P., & Hou, D. (2020). CausInfer: Automatic and distributed performance diagnosis with hierarchical causality graph in large distributed systems. In IEEE INFOCOM 2020, pp. 1704–1713. IEEE.
- [3] Cigna, L., Klues, T., & Muñoz-Gea, J. (2021). Lessons from five years of AI-assisted incident response at Netflix. Netflix Engineering Blog. Retrieved from <https://netflixtechblog.com>.
- [4] Du, M., Li, F., Zheng, G., & Srikumar, V. (2017). DeepLog: Anomaly detection and diagnosis from system logs through deep learning. In Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security (CCS), pp. 1285–1298.
- [5] Farshchi, M., Schneider, J.-G., Weber, I., & Grundy, J. (2018). Experience report: Anomaly detection of cloud application operations using log and cloud metric correlation analysis. In Proceedings of the 28th International Symposium on Software Reliability Engineering (ISSRE), pp. 24–35. IEEE.
- [6] Beyer, B., Jones, C., Petoff, J., & Murphy, N. R. (Eds.). (2016). Site Reliability Engineering: How Google Runs Production Systems. O'Reilly Media.
- [7] Li, Z., Chen, P., Hou, D., Sun, Y., & Liao, Q. (2021). Practical root cause localization for microservice systems via trace analysis. In Proceedings of the 2021 IEEE/ACM 29th International Symposium on Quality of Service (IWQOS), pp. 1–10.
- [8] Luo, C., Du, J., Ye, C., Li, H., & Shan, Z. (2022). AIOps challenges and opportunities in site reliability engineering: Case studies at Alibaba Cloud. IEEE Transactions on Services Computing, 15(4), 2276–2291.
- [9] Meng, Y., Zhang, S., Sun, Y., Zhang, R., Hu, Z., Zhang, Y., & Pei, D. (2020). Localizing failure root causes in a microservice through causality inference. In 2020 IEEE/ACM 28th International Symposium on Quality of Service (IWQoS), pp. 1–10.
- [10] Nedelkoski, S., Bogatinovski, J., Acker, A., Cardoso, J., & Kao, O. (2021). Self-supervised log parsing. In Machine Learning and Knowledge Discovery in Databases (ECML PKDD), pp. 122–138. Springer.
- [11] He, P., Zhu, J., Zheng, Z., & Lyu, M. R. (2016). Drain: An online log parsing approach with fixed depth tree. In 2016 IEEE International Conference on Web Services (ICWS), pp. 33–40.
- [12] Zhao, N., Chen, P., Zheng, Z., & Sui, K. (2020). Automatically and adaptively identifying severe alerts for online service systems. In Proceedings of IEEE INFOCOM 2020, pp. 2445–2454.
- [13] Botezatu, M. M., Giurgiu, I., Bogojeska, J., & Wiesmann, D. (2016). Predicting disk replacement towards reliable data centers. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, pp. 39–48.
- [14] Zhang, W., Xu, Y., Lin, X., Zhang, Z., & Wang, J. (2024). RCACopilot: On-call empowerment for cloud incident root cause analysis with large language models. In Proceedings of the 32nd ACM International Conference on the Foundations of Software Engineering (FSE), pp. 1802–1814.
- [15] Dang, Y., Lin, Q., & Huang, P. (2019). AIOps: Real-world challenges and research innovations. In Proceedings of the 41st International Conference on Software Engineering: Companion (ICSE-Companion), pp. 4–5. IEEE.