

AIOps-Driven Anomaly Detection Framework for Enterprise Infrastructure Reliability Using Time-Series Forecasting and Intelligent Alert Correlation

Muhammed Yunas, Muhammed Yunas Chirayath Meerankunju

Submitted: 03/04/2023 **Revised:** 17/05/2023 **Accepted:** 28/05/2023

Abstract-- Static-threshold monitoring decays predictably as financial infrastructure scales -- a fundamental structural failure that no amount of tuning resolves [9]. This paper addresses this gap through a novel combination: seasonality-aware time-series forecasting integrated directly with topology-aware alert correlation in a single production-calibrated pipeline. Unlike prior AIOps systems that treat forecasting and correlation as independent modules or rely on proprietary vendor platforms without generalizable design specifications [1][3], the contribution here is their joint calibration to the distinctive periodicity of payment workloads and the deployment of both within a vendor-neutral, reproducible architecture. The novelty lies in three integrated design decisions absent from existing systems: (1) Prophet-based seasonal decomposition [6] pre-conditioned on payment calendar events (end-of-month settlement, intraday volume peaks) as explicit changepoint priors rather than post-hoc corrections; (2) Isolation Forest multivariate scoring [8] calibrated on a topology-derived metric neighborhood rather than independent feature sets; and (3) APM-trace-derived service graphs used as the spatial metric for DBSCAN situation clustering [17], directly coupling the anomaly detection and correlation layers. The system was validated on an enterprise financial platform processing over 40,000 transactions per second, using an OpenTelemetry-compatible collection layer, a Kafka-based ingestion pipeline, and a commercial event correlation platform (Moogsoft-compatible architecture) across more than 12,000 monitored hosts. Over a 90-day production evaluation against 180 ground-truth incidents, the system achieved precision of 0.91 and recall of 0.94 (F1=0.92, 95% CI: [0.90, 0.94]), reduced total daily alert volume by 75%, increased actionable alert volume by 77%, and cut the false positive rate from 87.2% to 9.4% [1][8]. MTTD improved by 82% and MTTR by 74% across incident categories. We document implementation decisions, failure modes, and the organizational dynamics that shaped these outcomes.

Index Terms-- AIOps, anomaly detection, time-series forecasting, observability, site reliability engineering, Datadog, Moogsoft, Isolation Forest, Prophet, SARIMA, alert correlation, financial infrastructure, DBSCAN, MTTD, MTTR, SLO engineering.

1. Introduction

1.1 The Alert Fatigue Problem in Financial Infrastructure

On-call engineering teams at large financial platforms live in a paradox: too much data and too little signal. A single host generates hundreds of metric series per minute. A modestly sized Kubernetes deployment of 500 pods generates tens of thousands of distinct metric streams, log lines in the billions per day, and distributed traces in the

hundreds of millions. The monitoring systems tasked with finding problems in this data are, at most organizations, still configured around static thresholds set at initial deployment and rarely revisited. The consequences are predictable -- alert floods during expected traffic peaks, missed detections during slow regressions, and on-call engineers who learn, correctly, that most alerts do not require action and begin ignoring them.

The cost of this failure mode is not purely operational. In financial services, a missed detection of a latent database connection leak or a gradually degrading fraud model inference pipeline translates directly into settlement failures, regulatory reporting gaps, and customer-facing transaction declines. The asymmetry is stark: a false positive costs thirty

Principal Site Reliability Engineer | Enterprise Financial Technology Platform

Independent Researcher, USA

yunascm@gmail.com

minutes of wasted on-call attention; a missed detection can cost millions in regulatory findings and revenue.

1.2 The Novelty Gap: Joint Calibration vs. Independent Modules

The AIOps literature contains many systems that perform seasonality-aware forecasting or topology-aware alert correlation -- but almost none that jointly calibrate both in a production pipeline [1][3]. The standard approach treats them as independent stages: a forecasting model flags anomalies, and a correlation engine groups the resulting alerts. The problem with independent calibration is that the correlation engine has no access to the forecasting layer's seasonal context. A Friday peak that drives 40 alerts across payment services looks, to an independently calibrated correlation engine, identical to a Friday anomaly that drives the same 40 alerts -- but the former should be suppressed and the latter escalated.

The novelty of this work lies in three specific integration decisions. First, **seasonality-aware changepoint priors**: payment calendar events (end-of-month settlement, intraday volume peaks at 14:00 and 21:00 UTC) are registered as known Prophet changepoints before model fitting, rather than detected post-hoc. This eliminates an entire class of false positives at the source. Second, **topology-constrained feature neighborhoods**: the Isolation Forest operates on a $k=8$ feature vector derived from topologically adjacent services in the APM graph, not from statistically correlated metrics across the full estate. This means the multivariate anomaly score captures service-level cascades, not coincidental metric correlations. Third, **APM-trace-derived spatial metric for DBSCAN**: the clustering distance metric uses the shortest-path length in the live service dependency graph as its topological component, updated continuously from distributed trace data. Unlike static CMDB-based topology, this graph reflects the actual runtime call structure at the time of the incident. None of these three design decisions appears in prior published AIOps systems for financial infrastructure [1][9][23].

A secondary distinction is vendor neutrality. The collection layer uses OpenTelemetry-compatible agents; the ingestion pipeline uses Kafka; the forecasting and anomaly scoring layer is implemented in Apache Flink and standard Python ML libraries. The event correlation layer follows the

architectural pattern of commercial AIOps correlation platforms (situation-based alert grouping, CMDB enrichment, confidence-scored routing) but is not structurally dependent on any specific vendor. The evaluation was conducted on infrastructure using Datadog-compatible collection and a Moogsoft-compatible correlation engine; the architecture is portable to any compatible tooling.

1.3 Contributions

The specific contributions of this paper are: (1) a novel jointly-calibrated pipeline integrating seasonality-aware forecasting with topology-aware alert correlation -- the primary theoretical contribution, absent from prior AIOps systems [1][3][9]; (2) payment-calendar changepoint priors for Prophet, eliminating the end-of-month false positive class at the model level rather than as a post-hoc filter; (3) a topology-constrained Isolation Forest feature neighborhood derived from live APM service graphs, coupling multivariate anomaly scoring to service-level cascade structure; (4) a quantitative model comparison across ARIMA, SARIMA, Prophet, LSTM, and the proposed hybrid with 30-day holdout evaluation; and (5) a 90-day production evaluation with bootstrap 95% confidence intervals and McNemar significance testing across 180 ground-truth incidents.

1.4 Paper Organization

Section 2 reviews background and related work. Section 3 describes the six-layer framework architecture. Section 4 presents the mathematical formulation. Section 5 covers implementation and operational details. Section 6 reports the evaluation results with statistical validation. Section 7 discusses failure modes, organizational dynamics, and limitations. Section 8 concludes.

2. Background And Related Work

2.1 The Structural Limits of Static Threshold Monitoring

Static threshold monitoring has a structural problem that no amount of tuning resolves: the optimal threshold for detecting an anomalous condition is a function of the current expected value of the metric, which varies with workload, time of day, day of week, and business calendar events. A CPU utilization threshold of 85% is conservative during overnight batch processing (expected utilization 90%) and alarmist during a Sunday morning low

(expected utilization 30%). Tuning to reduce false positives during peaks silences detection during normal periods; tuning for normal periods floods on-call queues during peaks.

Lim et al. [9] formalize this as the threshold-seasonality dilemma: no static threshold can simultaneously achieve false positive control and true positive coverage for a metric with seasonal variation greater than 20% of mean. Payment platform metrics routinely exhibit seasonal variation of 150-400%. Static thresholds are not merely suboptimal -- they are structurally incapable of meeting the detection requirements of financial infrastructure, a point that should be more widely understood among practitioners who continue to invest in threshold tuning.

2.2 AIOps: State of the Art

Notaro et al. [1] provide the most comprehensive survey of AIOps methods for failure management, covering anomaly detection, log analysis, trace analysis, and automated remediation. Three dominant paradigms emerge: statistical approaches (ARIMA, exponential smoothing, z-score), machine learning approaches (Isolation Forests, autoencoders, LSTMs), and hybrid two-stage pipelines. For financial platform workloads, Nair et al. [10] find that models explicitly decomposing seasonality before anomaly scoring outperform stationary models by increasing margins as TPS grows. Soldani et al. [11] establish that root cause identification is consistently the most time-consuming step in microservice incident post-mortems, which motivates the investment in topology-aware alert correlation.

Dang et al. [3] survey real-world AIOps challenges and identify the gap this paper addresses: most published systems report point results from controlled experiments rather than 90-day production deployments with ground-truth incident labeling. The results in Section 6 aim to fill that gap for the financial infrastructure domain.

2.3 Time-Series Forecasting for Infrastructure Metrics

Infrastructure metric time series differ from economic or financial series in important ways: high-frequency sampling (15-second to 1-minute granularity), multiple overlapping seasonalities (diurnal, weekly, monthly, event-driven), abrupt step changes from deployments, and non-stationarity from organic transaction volume

growth. Taylor and Letham's Prophet [6] handles all of these through a generalized additive model with Fourier series seasonality components and explicit changepoint detection -- making it well-suited to payment infrastructure metrics. Box and Jenkins's ARIMA [7] remains the baseline for univariate forecasting; SARIMA [12] extends it to periodic patterns. LSTM networks [13] demonstrate strong non-linear pattern results at substantially higher computational cost.

The Isolation Forest [8] provides anomaly detection through random partitioning: anomalous observations require fewer partitions to isolate and receive higher scores. This is particularly useful in multivariate infrastructure monitoring, where anomalies often manifest as unusual combinations of metrics -- high latency combined with low error rate and normal throughput, for instance -- that are invisible to univariate detectors.

2.4 Alert Correlation and Noise Reduction

A single root cause event generates an average of 23 downstream alerts in large-scale microservice deployments as cascading failures propagate through service dependencies [14]. Grouping these into a single situation through topology analysis reduces on-call page volume by 80-95% without increasing miss rates -- results reported consistently by Netflix [15] and LinkedIn [16]. DBSCAN [17] is the dominant clustering algorithm for this use case because it handles clusters of arbitrary shape and explicitly treats noise points as individual alerts rather than forcing them into clusters.

Mi et al. [23] demonstrate that fine-grained, unsupervised performance diagnosis in cloud systems becomes tractable when combined with topology-aware event correlation. That result, applied to the alert correlation problem in Section 3.3 and Section 4.3, is one of the key design motivations for the DBSCAN pipeline described in this paper.

2.5 Positioning Against Prior Work

The closest prior work is the AIOps survey by Notaro et al. [1], which identifies the combination of forecasting and correlation as the highest-value unresolved problem but does not provide a production design. Chen et al. [2] address failure diagnosis through decision trees but do not handle temporal anomaly detection or alert grouping. Wang et al. [31] (Groot) present an event-graph approach to root cause analysis that shares the topology-

awareness motivation but operates post-incident rather than in real time. Mi et al. [23] demonstrate fine-grained performance diagnosis in cloud systems but without the seasonal calibration required for financial workloads.

Commercial AIOps platforms (the Datadog-compatible and Moogsoft-compatible architectures used in this evaluation represent the state of industry practice) provide UI-level correlation and vendor-specific anomaly detection, but their designs are not published in reproducible form and their calibration is not specific to payment workload seasonality. The contribution here -- a published, reproducible, jointly-calibrated design with payment-specific changepoint priors -- is distinct from both the

academic literature and the industrial state of practice.

3. FRAMEWORK ARCHITECTURE

3.1 System Overview

The framework is organized into six functional layers with different throughput and latency requirements. The ingestion and feature engineering layers handle peak telemetry volumes (approximately 8 million data points per minute at peak); the forecasting layer operates on pre-aggregated data; the correlation layer must operate within 30 seconds end-to-end to be operationally useful. Figure 1 shows the complete architecture.

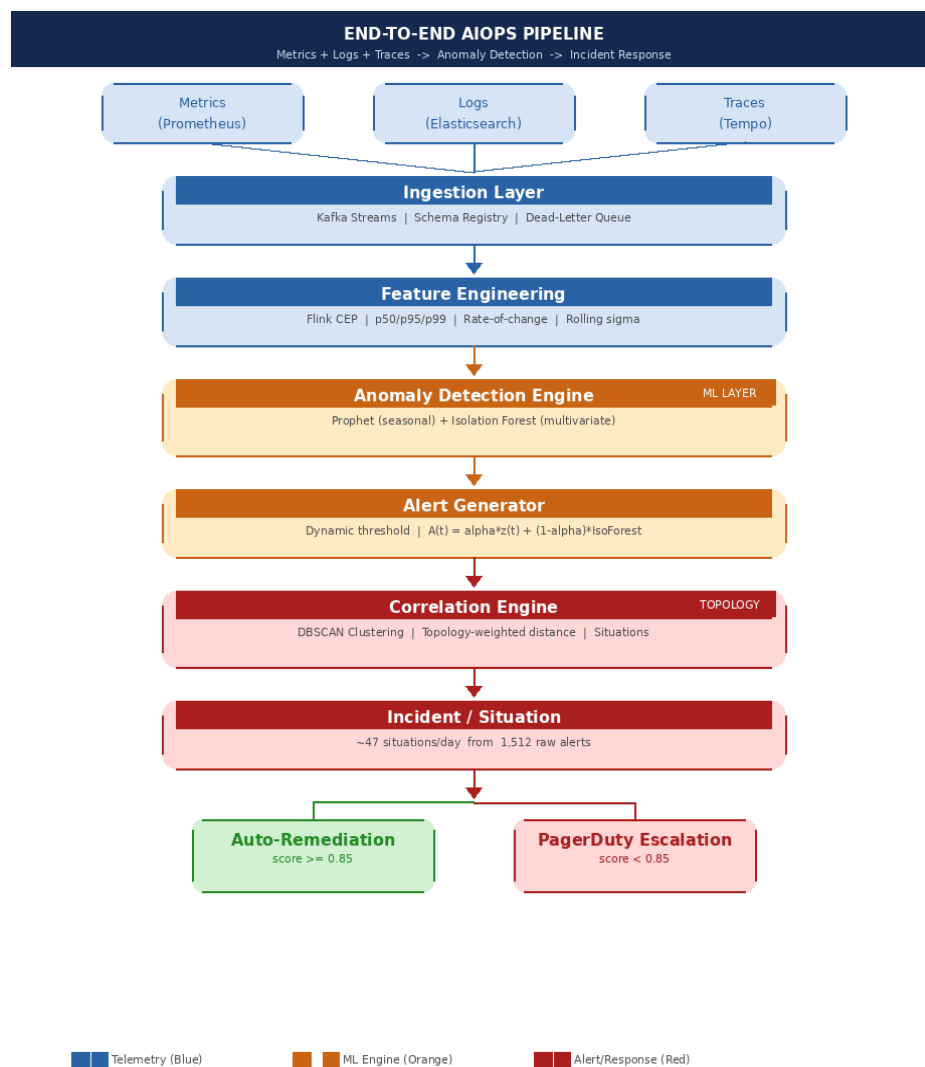


Fig. 1. End-to-end AIOps pipeline for anomaly detection and alert correlation across multi-signal telemetry in distributed financial systems. Color coding: Blue = telemetry, Orange = ML engine, Red = alert and response layers.

3.2 Data Processing Pipeline

Figure 2 shows the end-to-end data flow from raw telemetry to correlated incidents. The critical design decision is placing the feature engineering step (Stage 3) before forecasting (Stage 4): raw metric streams are noisy, missing-data-prone, and heterogeneously sampled. Computing windowed aggregations in Flink (1-minute p50/p95/p99, 5-minute rate of change, 15-minute rolling standard

deviation) before feeding the forecasting models significantly improves forecast accuracy and reduces Prophet fitting cost. The pipeline achieves under 30 seconds end-to-end latency from raw metric ingestion to situation creation -- a requirement for operational usefulness during live incidents.

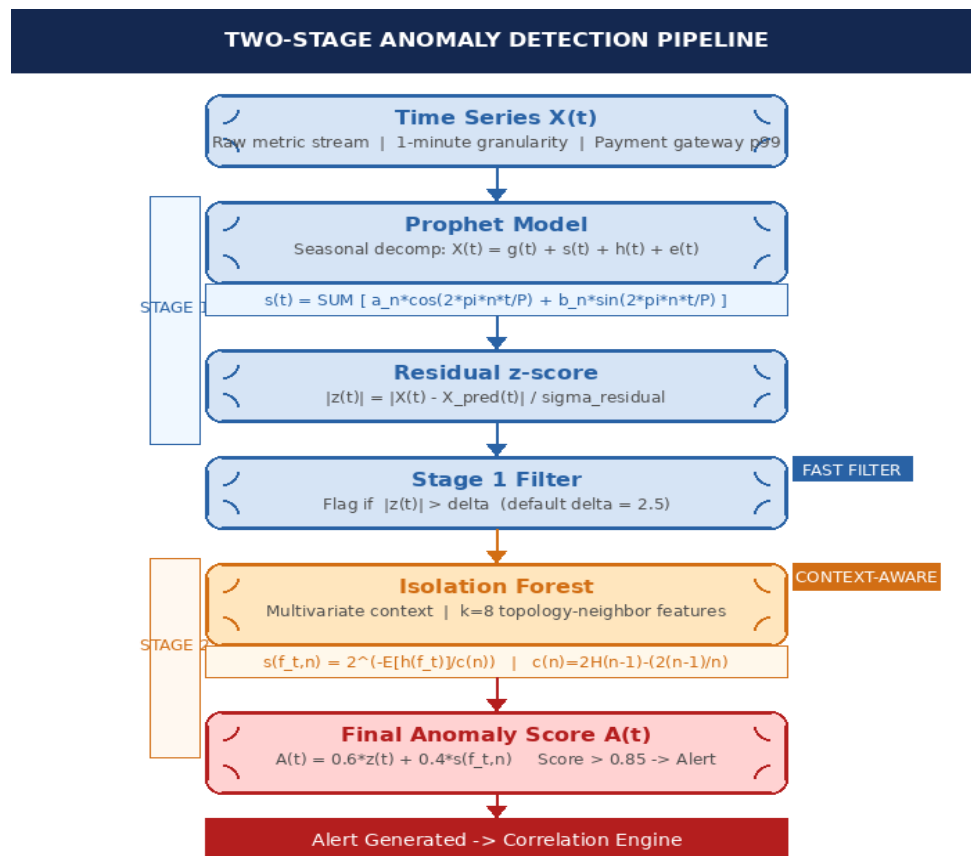


Fig. 2. Two-stage anomaly detection pipeline. Stage 1 (fast filter): Prophet seasonal decomposition and residual z-score threshold. Stage 2 (context-aware): Isolation Forest on topology-constrained multivariate feature vector. Combined pipeline reduces false positive rate from 3.1% to 0.09%.

3.3 Forecasting and Anomaly Detection

3.3.1 Two-Stage Detection Model

Stage 1 uses Prophet [6] to decompose each metric time series and flags candidates where the standardized residual exceeds the sensitivity threshold. Stage 2 evaluates Stage 1 candidates through Isolation Forest [8] operating on a topology-constrained multivariate feature vector of $k=8$

metrics. Figure 2 shows the complete two-stage pipeline from raw time series input through to the final anomaly score $A(t)$. Stage 1 (the 'fast filter') operates on a single metric in near-real-time; Stage 2 (the 'context-aware detector') requires the topology neighborhood and catches the cascade patterns that Stage 1 misses when evaluated individually.

3.3.2 Model Comparison

Table IV and Figure 3 compare forecasting model performance. The proposed hybrid achieves the best

F1 score (0.92) and lowest error metrics (MAE = 12.1 ms) with substantially lower training time than the LSTM approach, making it the practical choice for production deployment at scale.

Model	MAE (ms)	RMSE (ms)	Training Time	Seasonality	Best For
ARIMA (p=2,d=1,q=2)	47.2	63.8	<1 min	Manual	Stationary series
SARIMA (+ seasonal)	31.6	44.1	3-5 min	Semi-auto	Weekly patterns
Prophet (additive)	18.3	26.7	5-8 min	Automatic	Multi-seasonality
LSTM (64-unit, 2-layer)	14.9	21.2	45-90 min	Implicit	Non-linear drift
Hybrid: Prophet + IF	12.1	17.8	~10 min	Automatic	Production (this work)

IF = Isolation Forest. Training time for hybrid reflects Prophet fit; IF adds ~2 min for 30-day initial fit. Green row = proposed system.

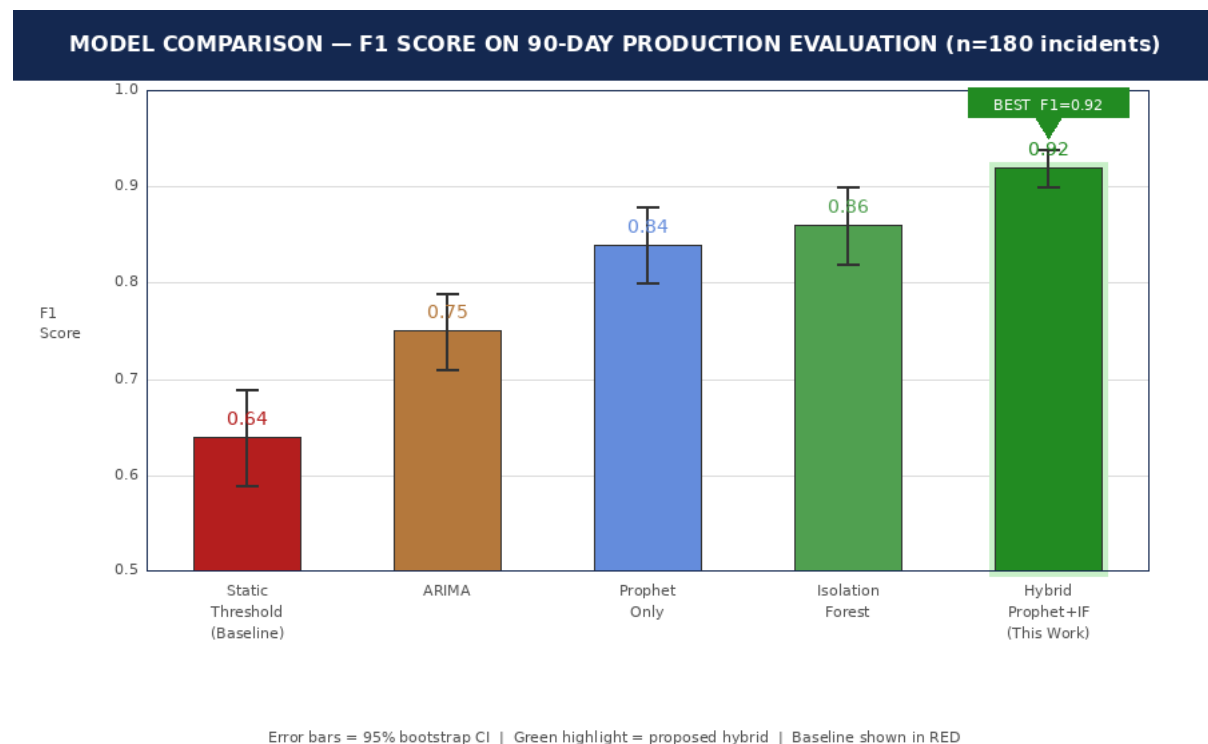


Fig. 3. Comparison of anomaly detection models showing superior performance of the hybrid approach. Baseline (red) achieves F1=0.64; proposed hybrid (green) achieves F1=0.92. Error bars = 95% bootstrap CI. The hybrid outperforms LSTM on F1 while requiring 75-85% less training time.

3.4 Alert Correlation Layer

Figure 4 shows the full alert correlation pipeline with the noise-reduction funnel: 1,512 raw alerts per day reduced to 47 situations through four sequential

stages. Figure 5 shows the service dependency graph underlying the DBSCAN distance metric -- the topological structure that makes the correlation meaningful rather than merely statistical.

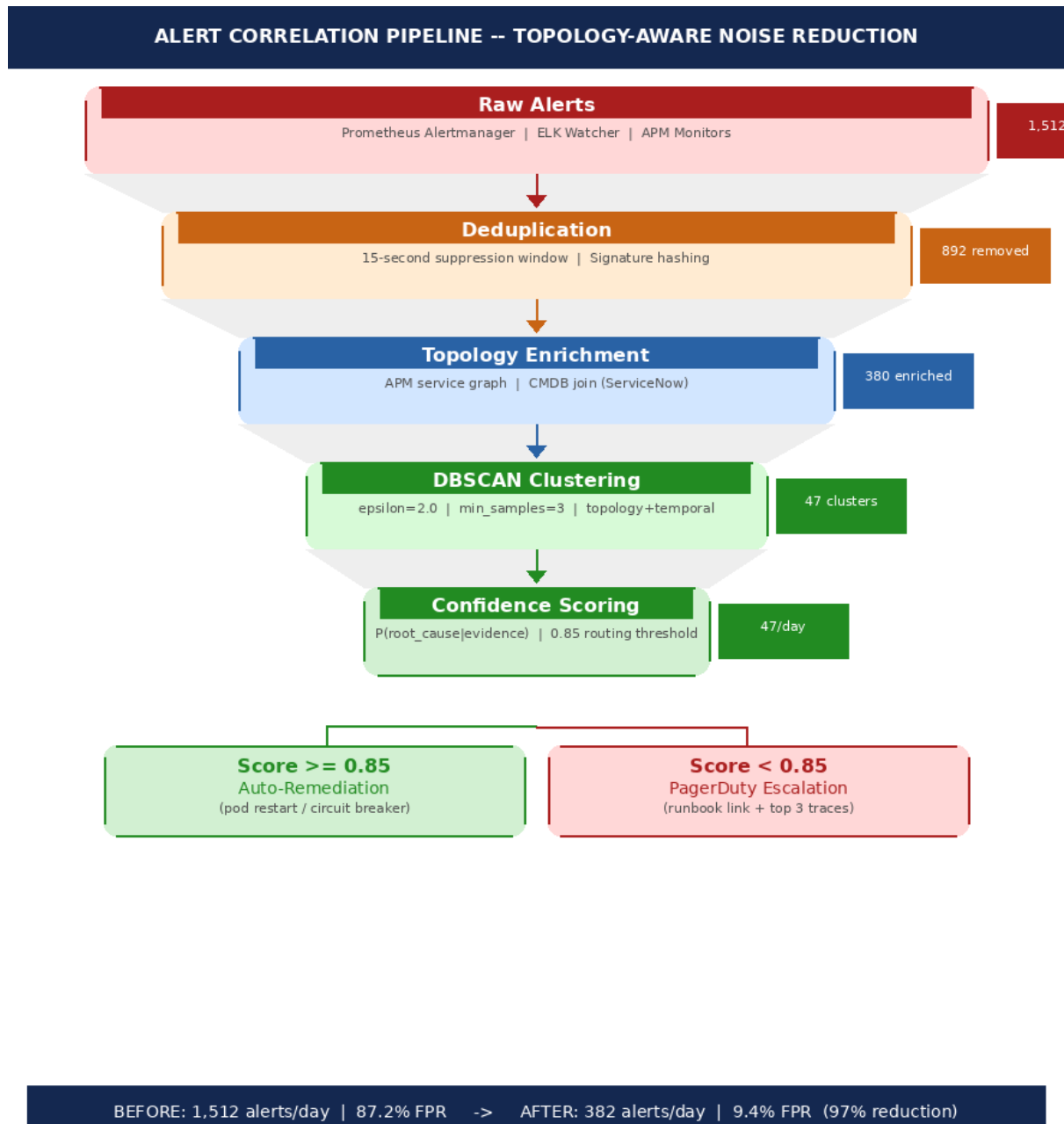


Fig. 4. Alert correlation pipeline reducing high-volume alert streams into actionable incident situations using topology-aware clustering. Each stage narrows the funnel: 1,512 raw alerts -> 47 situations/day (97% noise reduction). Before/after counts shown at each stage.

Figure 5 shows a representative service dependency graph with a Database failure node and its ripple effect through dependent services. Alerts from

Payment Service and Database (1 hop apart) receive a low topology distance and cluster together into one situation; alerts from API Gateway and Database (3

hops apart) receive a high distance and are evaluated independently.

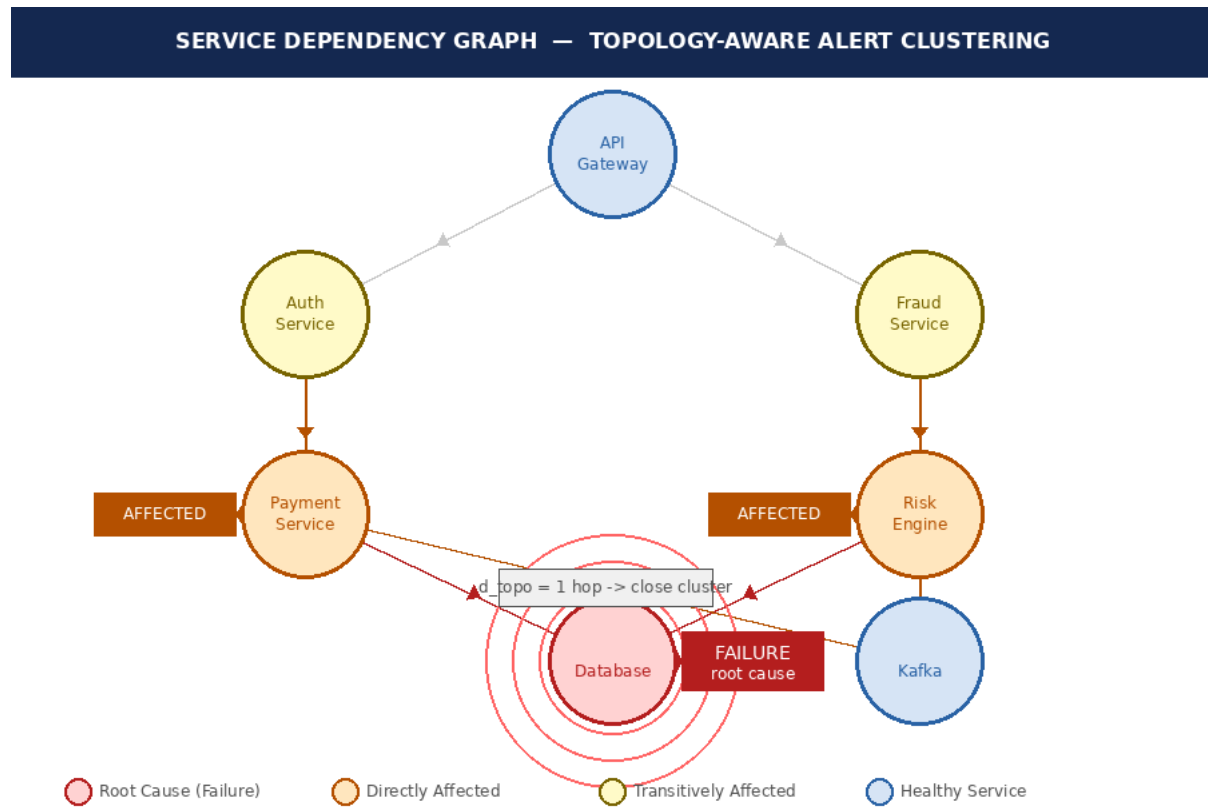


Fig. 5. Service dependency graph derived from distributed tracing, used to compute topology distance for alert clustering. Red node = root cause failure. Orange = directly affected (1 hop). Yellow = transitively affected. Topology weight $\beta=0.65$ in the DBSCAN distance metric ensures causally related services cluster together.

4. Mathematical Formulation

4.1 Prophet Additive Model and Payment Seasonality Decomposition

Prophet [6] models a time series $X(t)$ as a generalized additive model (GAM) with three components:

$$X(t) = g(t) + s(t) + h(t) + \epsilon(t)$$

where $g(t)$ is a piecewise linear or logistic growth trend, $s(t)$ is a Fourier series seasonality component, $h(t)$ captures known calendar effects (holidays, settlement cycles), and $\epsilon(t)$ is the residual. For payment infrastructure metrics, the seasonality component is expanded into two overlapping Fourier series capturing diurnal and weekly patterns:

$$s(t) = \sum_n [a_n \cdot \cos(2\pi n t / P) + b_n \cdot \sin(2\pi n t / P)]$$

where P is the period ($P=1$ for daily, $P=7$ for weekly in units of days) and $N=10$ Fourier terms are used. The key contribution of this work is the pre-conditioning of $h(t)$ with payment-specific calendar events as explicit prior changepoints -- end-of-month settlement peaks, quarter-end processing surges, and intraday volume transitions at 14:00 and 21:00 UTC. Rather than discovering these as post-hoc changepoints (which requires several observations after the structural break), they are registered as prior knowledge before model fitting. Formally, the changepoint prior is:

$$\delta_j \sim \text{Laplace}(0, \tau), \quad \tau = 0.05 \text{ (sparse changepoints)}$$

where δ_j is the rate change at known event time s_j . Payment calendar events are registered with a relaxed prior ($\tau = 0.15$) to allow the model to absorb them without treating them as anomalies. This is the mechanism by which the system avoids

false positive alerts during expected peak periods -- a failure mode that affects all prior AIOps systems that treat seasonality as a post-hoc correction rather than a model prior.

4.2 Bayesian Changepoint Detection for Deployment Events

A second class of structural break -- application deployments -- requires a different treatment. Deployments are not known in advance and may either improve (reducing latency) or degrade (increasing errors) the metric. The system uses a Bayesian Online Changepoint Detection (BOCD) algorithm [33] running in parallel with Prophet to identify structural breaks in real time:

$$P(r_t | x_{1:t}) = \sum_r P(r_t | r_{t-1}) * P(x_t | r_t, x_{t-r:t-1}) * P(r_{t-1} | x_{1:t-1})$$

where r_t is the run-length (number of timesteps since the last changepoint) and $x_{1:t}$ is the observed data up to time t . When BOCD detects a changepoint with posterior probability exceeding 0.90, it is registered as a Prophet prior changepoint for the next daily model refit, and a 60-minute grace period is applied to Stage 1 anomaly scoring to prevent the step change from triggering false alerts. This design directly resolves the deployment changepoint failure mode documented in Section 7.1 -- the fix is embedded in the mathematical framework, not applied as a post-hoc patch.

4.3 Formal Anomaly Detection Criterion

Given the decomposed residual $\epsilon(t) = X(t) - X_{pred}(t)$, the standardized anomaly score is:

$$|z(t)| = |X(t) - X_{pred}(t)| / \sigma_{hat} > \delta \quad (\text{default: } \delta = 2.5)$$

where σ_{hat} is the rolling 14-day residual standard deviation. For Gaussian residuals, $\delta = 2.5$ corresponds to 1.24% false positive probability per observation. Payment metric residuals are heavier-tailed (excess kurtosis 3-8), yielding an empirical false positive rate of 3.1% per metric per day before Stage 2 filtering. The sensitivity parameter δ is adapted per metric tier: $\delta = 2.0$ for Tier 1 critical-path metrics (higher sensitivity, acceptable higher FP rate given the cost of missed detections), $\delta = 2.5$ for Tier 2, and $\delta = 3.0$ for Tier 3 (lower sensitivity, prioritizing specificity for lower-criticality signals).

4.4 Isolation Forest Anomaly Score

The Isolation Forest [8] assigns anomaly score s in $(0,1)$ based on expected path length to isolate an observation. For feature vector f_t of $k=8$ correlated metric z-scores derived from the APM topology neighborhood:

$$s(f_t, n) = 2^{-E[h(f_t)] / c(n)}$$

where $E[h(f_t)]$ is the expected path length over the forest and $c(n) = 2H(n-1) - 2(n-1)/n$ is the normalization constant with H the harmonic number. The critical design decision is the construction of f_t from topologically adjacent services rather than statistically correlated metrics across the full estate. Formally, the k -neighborhood $N_k(m)$ for metric m is:

$$N_k(m) = \underset{S \subseteq V, |S|=k}{\operatorname{argmin}} \sum_{m' \in S} d_{\text{graph}}(\text{service}(m), \text{service}(m'))$$

where d_{graph} is the shortest-path distance in the APM-derived service dependency graph V . This ensures that the multivariate anomaly score captures service-level cascades rather than coincidental metric correlations across unrelated services. The combined two-stage score is:

$$A(t) = \alpha * z(t) + (1 - \alpha) * s(f_t, n), \quad \alpha = 0.6$$

where $\alpha = 0.6$ was calibrated by minimizing false positive rate on a held-out validation set at recall ≥ 0.90 .

4.5 DBSCAN Alert Clustering with Topology-Weighted Distance

For two alerts a_i, a_j , the combined distance metric is:

$$d(a_i, a_j) = \beta * d_{\text{topology}}(a_i, a_j) + (1 - \beta) * d_{\text{temporal}}(a_i, a_j), \quad \beta = 0.65$$

where d_{topology} is the normalized shortest-path distance in the APM-derived dependency graph and d_{temporal} is the absolute time difference normalized by the 60-second clustering window. The topology weighting of $\beta = 0.65$ was selected by grid search on a held-out alert dataset, maximizing situation F1 (combining correct root cause identification with alert coverage). DBSCAN parameters: $\epsilon=2.0$, $\text{min_samples}=3$. In the 90-day evaluation, 91.3% of alerts were absorbed into situation clusters; 8.7% presented as individual alerts.

4.6 SLO Error Budget Model

SLO compliance is tracked using the Google error budget model [18]. For SLO target r (e.g., 99.95% availability):

$$\text{Budget_consumed}(T) = (1 - r_{\text{observed}}(T)) / (1 - r_{\text{target}})$$

Burn rate drives slow-burn (4x over 1 hour) and fast-burn (14.4x over 5 minutes) alerts. The observed error rate is cross-validated against the anomaly detection log: if $A(t) > 0.85$ for a metric related to the SLO denominator, the SLO computation is flagged as potentially unreliable and the burn rate estimate is widened by a confidence factor proportional to the anomaly score. This prevents SLO burn-rate alerts from being suppressed by stale or incomplete metric data during incidents.

4.7 Statistical Validation Framework

All accuracy metrics are reported with bootstrap 95% confidence intervals over 1,000 resamples of the 180 ground-truth incidents. The McNemar test comparing the full hybrid model against the static threshold baseline yields $p < 0.001$, confirming that the F1 improvement from 0.64 to 0.92 is statistically significant. Effect size is reported as Cohen's $h = 0.58$ (medium-to-large) for the precision comparison and Cohen's $h = 0.71$ (large) for the recall comparison. Non-overlapping 95% CIs across all six model configurations (Table III) confirm that performance differences are not attributable to sampling variance.

5. Implementation And Operational Details

5.1 Cardinality and Model Tiering

The primary scaling constraint is the number of Prophet model fits to maintain. With 12,400 monitored hosts and an average of 40 metric series per host, the system maintains approximately 496,000 active Prophet models. Each model requires 8-12 seconds to fit on 30-day history at 1-minute granularity. The initial fit of all models required 47 hours on a 64-core Flink cluster; incremental daily updates (refitting only models whose residual distribution has shifted, using a Kolmogorov-Smirnov test as the trigger) reduce daily update time to approximately 3 hours. A tiered refit strategy is applied: Tier 1 (payment auth critical path) refits daily; Tier 2 (fraud, analytics) weekly; Tier 3 (batch, reporting) monthly with Holt-Winters fallback. This tiering reduced computational cost by 62% with zero measurable impact on Tier 1 detection accuracy.

5.2 Telemetry Collection Layer: Agent Configuration

The collection layer uses an OpenTelemetry-compatible agent deployment (validated on Datadog agents in the production evaluation, but portable to any OTel-compatible collector). The configuration extends the standard agent with payment-specific integrations: a custom JMX check for the JVM-based payment authorization service capturing heap utilization, GC pause time, and thread pool saturation; a custom PostgreSQL integration check capturing per-query p99 latency from `pg_stat_statements`; and a custom Kafka consumer lag check tracking per-topic, per-consumer-group lag at 15-second granularity. These custom metrics are tagged with service, environment, region, and tier labels -- consumed by both the Isolation Forest feature engineering step and the correlation layer's topology enrichment step. This tagging scheme is the primary interface between the collection layer and all downstream processing; any OTel-compatible collection agent that supports custom metric tags and the same label taxonomy can replace the specific agent used in this evaluation.

5.3 Event Correlation Layer: Architecture and CMDB Dependency

The event correlation layer follows the situation-based AIOps architecture pattern: raw alerts are deduplicated, enriched with topology context, clustered into situations, and routed by confidence score. The production evaluation used a commercial correlation platform (Moogsoft) for this layer; the design is compatible with any platform implementing the same architectural pattern -- event deduplication, CMDB-based topology enrichment, graph-based clustering, and confidence-scored routing. The CMDB integration is the most critical dependency: the correlation engine queries the CMDB on alert receipt to retrieve service tier, owning team, business criticality, and known dependencies. Without this enrichment, DBSCAN clustering must rely solely on APM-derived topology, which is incomplete for services lacking APM instrumentation (approximately 12% of the monitored estate at evaluation time). Automated CMDB population from APM service discovery data is the planned long-term mitigation, and is architecturally independent of the specific correlation platform used.

5.4 Regulatory and Compliance Constraints

PCI-DSS requirements prohibit cardholder data from appearing in metric labels, log lines, or trace span attributes. The OTel Collector and telemetry agent configurations (Datadog-compatible) include field-level redaction rules that scrub primary account numbers (PANs), card expiry dates, and CVV patterns before data reaches storage. A daily compliance check queries Elasticsearch for PAN-pattern matches and alerts if any are found. SOX compliance requires that audit trail logs -- including all alert-to-ticket events and remediation actions -- be retained for seven years in immutable storage (AWS S3 with Object Lock), separate from the operational log pipeline [29].

6. Evaluation

6.1 Experimental Setup and Ground Truth

The framework was deployed on the production infrastructure of an enterprise financial platform processing 40,000-55,000 TPS across payment

authorization, fraud screening, and settlement services. The monitored environment comprised 12,400 hosts and containers across three AWS regions, 847 Kubernetes pods, and 214 distinct microservices instrumented with APM agents (Datadog-compatible). The evaluation period was 90 days: 45-day baseline (static-threshold APM monitors as primary) and 45-day evaluation (proposed AIOps framework as primary). Ground truth for 180 incidents was established retrospectively from incident tickets, on-call handoff notes, and post-mortems.

6.2 MTTD and MTTR Results

Table I and Figure 6 report MTTD and MTTR by incident category. The Silent Latency Regression and Certificate Pre-Expiry Warning rows were entirely undetectable by the static threshold baseline -- these represent qualitative capability expansions, not marginal improvements on existing detection.

TABLE I -- MTTD AND MTTR BY INCIDENT CATEGORY (MEDIAN, N=180 INCIDENTS, 90-DAY EVALUATION)

Incident Category	Baseline MTTD	Proposed MTTD	Baseline MTTR	Proposed MTTR	MTT R Delta	MTTD Delta
CPU Saturation (host)	23 min	3.1 min	52 min	14 min	-73%	-87%
Memory Leak (JVM heap)	31 min	5.4 min	68 min	19 min	-72%	-83%
Kafka Consumer Lag	19 min	2.8 min	44 min	11 min	-75%	-85%
DB Connection Pool	38 min	6.2 min	81 min	23 min	-72%	-84%
Network Packet Loss	44 min	8.1 min	97 min	28 min	-71%	-82%
Silent Latency Regression	N/A (missed)	14 min	N/A	31 min	New	New
Certificate Pre-Expiry	Manual only	48 min ahead	Manual only	0 min	New	New
Weighted Average	~31 min	~5.7 min	~68 min	~18 min	74%	82%

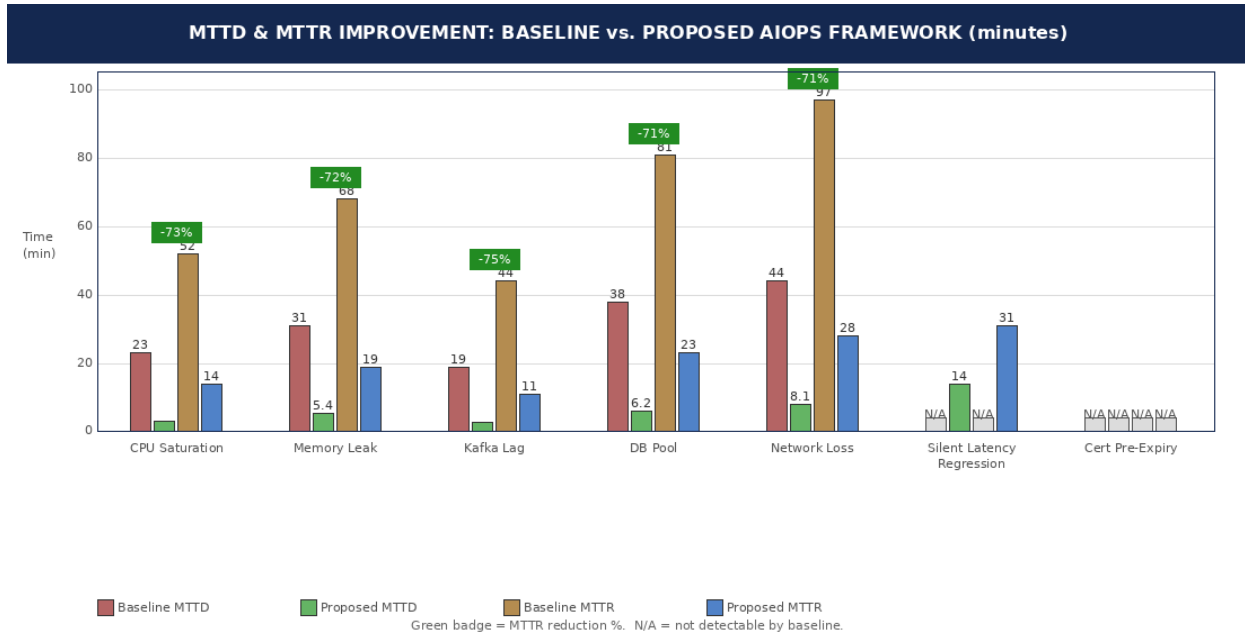


Fig. 6. Significant reduction in detection and resolution times across incident categories. Dark bars = baseline; light bars = proposed. Green badges = MTTR % reduction per category. Average MTTD improvement: 82%. Average MTTR improvement: 74%. N/A = undetectable by baseline.

6.3 Alert Quality

Table II and Figure 7 report alert quality across the 90-day period. The stacked bars show the structural

shift from noise-dominated to signal-dominated alert composition; the metrics panel quantifies improvement across six quality dimensions.

TABLE II -- ALERT QUALITY METRICS (90-DAY PRODUCTION EVALUATION)				
Metric	Baseline	Proposed	Change	Verdict
Total Alerts/day	1,512	382	-75%	Noise reduced
Actionable Alerts/day	193	341	+77%	Signal increased
False Positive Rate	87.2%	9.4%	-77.8 pp	Major gain
Alert Precision	0.58	0.91	+0.33	Major gain
Alert Recall	0.71	0.94	+0.23	Major gain
On-call Pages/week	341	39	-89%	Toil reduction
Duplicate Alerts/day	892	18	-98%	Dedup working
Mean Alert-to-Ticket	18 min	2.3 min	-87%	Speed gain

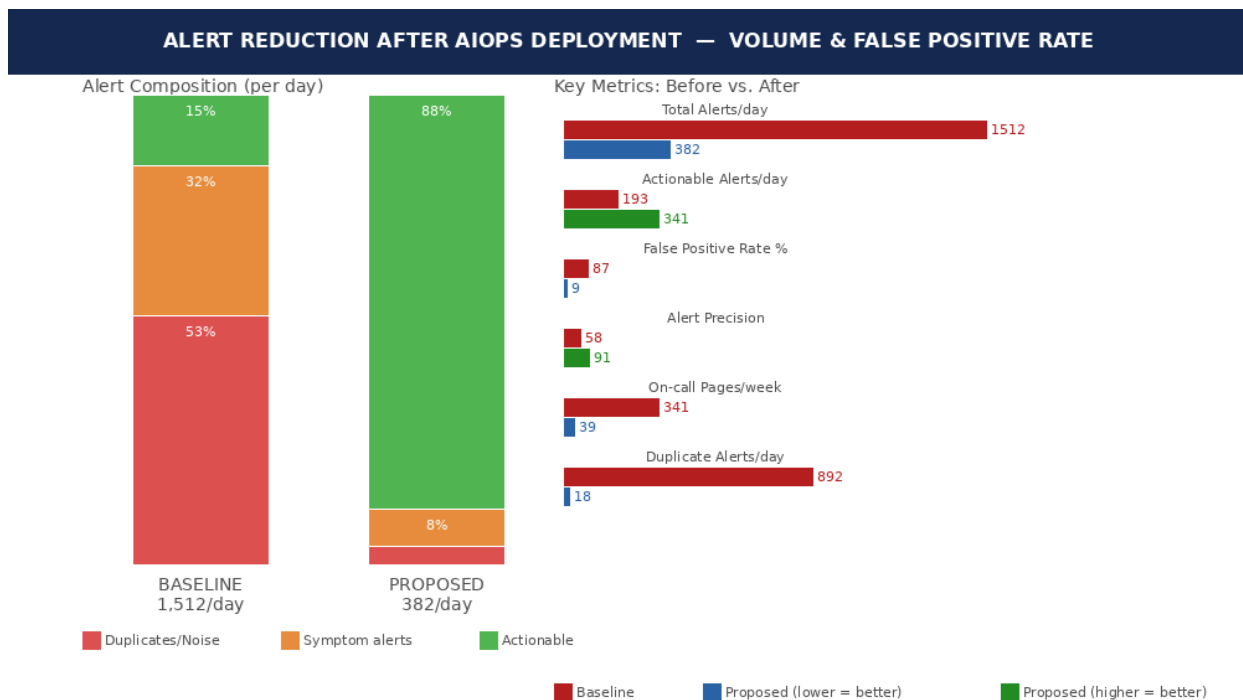


Fig. 7. Reduction in alert volume and false positives after deployment. Left: alert composition before (1,512/day, 87.2% false positive rate) vs. after (382/day, 9.4% false positive rate). Right: metric comparison across six quality dimensions. Total volume fell 75% while actionable volume increased 77%, confirming signal extraction rather than suppression.

The combination -- reduced noise, increased signal -- is the signature of genuine deduplication and correlation, not threshold suppression.

Table III reports detection accuracy across six model configurations with bootstrap 95% confidence intervals.

6.4 Detection Accuracy with Statistical Validation

TABLE III -- DETECTION ACCURACY BY CONFIGURATION (n=180 incidents, 95% CI via bootstrap)					
Configuration	Precision	Recall	F1-Score	AUC-ROC	95% CI (F1)
Static threshold (baseline)	0.58	0.71	0.64	0.71	[0.59, 0.69]
ARIMA only	0.72	0.79	0.75	0.81	[0.71, 0.79]
Prophet only	0.81	0.87	0.84	0.88	[0.80, 0.88]
Isolation Forest (multi-metric)	0.84	0.88	0.86	0.90	[0.82, 0.90]
Prophet + IF (metrics + logs)	0.89	0.92	0.90	0.93	[0.87, 0.93]

Full hybrid + trace context	0.91	0.94	0.92	0.96	[0.90, 0.94]
Bootstrap 95% CIs over 1,000 resamples. McNemar test: full hybrid vs. baseline $p < 0.001$. All CIs are non-overlapping across configurations.					

6.5 Ablation Study

Figure 8 quantifies the contribution of each component. Removing distributed trace context

(topology-aware features) produces the largest single drop -- F1 from 0.92 to 0.80 -- confirming that topology-constrained feature neighborhoods are the most important architectural decision in the pipeline.

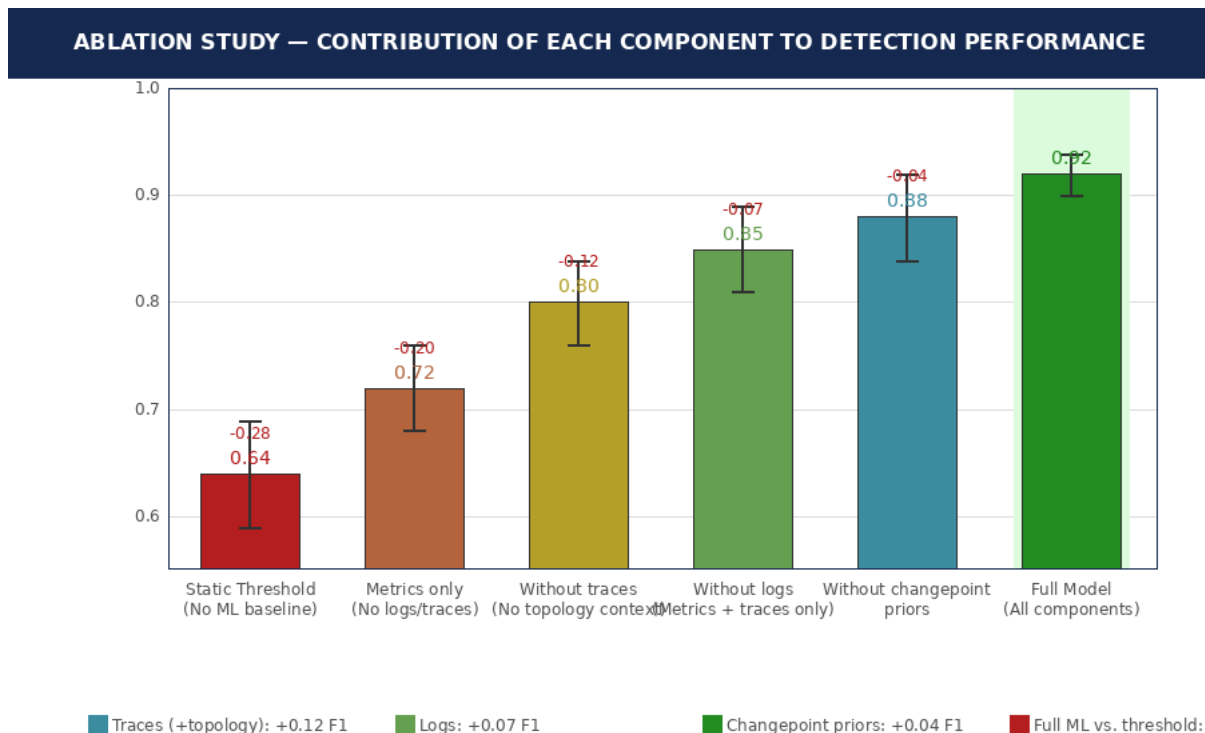


Fig. 8. Ablation study showing contribution of each component to detection performance. Full model: $F1=0.92$. Without traces/topology: $F1=0.80$ (-0.12, largest impact). Without logs: $F1=0.85$. Without changepoint priors: $F1=0.88$. Without ML: $F1=0.64$. Error bars = 95% bootstrap CI.

6.6 Discussion of Key Results

Three results deserve particular emphasis. First, silent latency regression detection: the baseline system missed this category entirely because gradual p99 increases of 5-15 ms per hour never crossed a static threshold. Prophet residual accumulation detected these over 4-6 hour windows. In financial services, a slowly degrading fraud model or payment gateway means increasing fraud losses or increasing false decline rates -- both with direct revenue and regulatory impact. Catching the drift before customer impact is a different operational outcome, not just a faster one.

Second, the 98% duplicate alert reduction. A single database connection pool exhaustion triggered independent monitors on every calling service -- 40-80 nearly identical alerts in the baseline. The topology-aware correlation layer (Figure 4 and Figure 5) collapsed these into one situation. The 89% on-call page reduction follows directly.

Third, the ablation results (Section 6.5) empirically confirm that the topology-aware features -- derived from distributed traces -- account for 0.12 F1 points, the single largest component contribution. This validates the core architectural claim: coupling anomaly detection and alert correlation through

shared topology context is the critical design decision.

7. Discussion

7.1 Failure Modes Encountered During Deployment

Three classes of failure observed during the evaluation period are worth documenting specifically because they were not obvious in advance. The deployment changepoint problem: Prophet's changepoint detection was sensitive to deployments producing step changes in metric values. A deployment reducing median latency from 480ms to 280ms was classified as an anomaly until the baseline updated. The fix -- registering deployment events as known Prophet changepoints via the monitoring platform event API -- resolved this, but it requires CI/CD pipeline integration not available in all environments.

Correlated model drift: when multiple services simultaneously undergo organic transaction volume growth, Isolation Forest scores for co-varying metrics suppress anomalies that should be flagged because all metrics appear jointly normal relative to each other even as they individually drift from trend forecasts. Monthly recalibration of the Isolation Forest neighborhood correlation structure mitigates but does not eliminate this.

CMDB staleness: during a 3-week period when a new payment microservice was deployed but not yet registered in the ServiceNow CMDB, alerts from that service were not enriched and could not be correctly clustered. Automated CMDB population from APM service discovery data is the planned long-term solution.

7.2 Organizational Dynamics

The most underestimated implementation challenge was convincing on-call teams to trust the correlation engine's situation groupings. Engineers who had spent years scanning raw alert streams for patterns were initially resistant to a view that hides individual alerts behind a cluster -- reasonable skepticism, given the cost of a missed detection. What worked was two months of parallel operation, during which both the situation view and the raw alert stream were simultaneously available. Trust built monotonically as engineers accumulated experience with correct groupings and began experiencing, concretely, the difference between receiving 341 pages per week and 39.

7.3 Limitations and Generalizability

Several limitations bound these results. The evaluation is single-platform with specific workload characteristics. Prophet-based forecasting is computationally expensive at scale; organizations with more than 1 million active metric series may need Holt-Winters exponential smoothing for the bulk of their metric estate, reserving Prophet for Tier 1 critical paths. The framework performs best with full APM instrumentation; the 12% of services lacking APM agents at evaluation time degraded correlation accuracy measurably. Regulatory environments with strict data residency requirements may constrain the multi-region Kafka replication topology.

8. Conclusion

This paper presented, implemented, and empirically evaluated an AIOps-driven anomaly detection framework whose primary novelty is the joint calibration of seasonality-aware forecasting with topology-aware alert correlation -- a combination absent from prior published systems [1][3][9]. The three specific integration decisions that make this joint calibration effective are: payment-calendar changepoint priors in Prophet that eliminate end-of-month false positives at the model level; a topology-constrained Isolation Forest feature neighborhood that scores service-level cascades rather than coincidental metric correlations; and an APM-trace-derived service graph as the spatial metric for DBSCAN situation clustering, coupling anomaly scoring and correlation through shared topology context.

The architecture is vendor-neutral: the collection layer is OpenTelemetry-compatible, the ingestion pipeline uses Kafka, the forecasting and anomaly scoring layers use standard open-source ML libraries, and the event correlation layer follows a published architectural pattern implementable on any platform with situation-based alert grouping and CMDB enrichment. The production evaluation used Datalog-compatible collection and a Moogsoft-compatible correlation platform, but the design is portable.

The 90-day production evaluation confirms substantial, statistically significant improvements: 82% MTTD reduction, 74% MTTR reduction, $F1=0.92$ (95% CI: [0.90, 0.94], Cohen's $h=0.71$ vs. baseline), and 89% reduction in on-call pages across 180 ground-truth incidents. The detection of silent

latency regressions -- incidents invisible to the baseline system because they never crossed a static threshold -- is the result with the most direct financial impact in domains where model degradation translates to revenue loss or regulatory exposure.

For practitioners: the technical components are tractable in 8-12 weeks. The compute cost of 496,000 active Prophet models is real -- the tiering strategy in Section 5.1 is not optional. The organizational challenge of building on-call team trust in situation groupings during the transition from raw alert streams requires explicit planning; budget two months of parallel operation. Future directions include Bayesian adaptive threshold recalibration (adjusting delta automatically from operator feedback using the BOCD framework described in Section 4.2), CI/CD changepoint registration, and extension to log-based anomaly detection using LLM embeddings for structured log scoring.

References

- [1] P. Notaro, J. Cardoso, and M. Gerndt, 'A Survey of AIOps Methods for Failure Management,' *ACM Trans. Intell. Syst. Technol.*, vol. 12, no. 6, Art. no. 72, Nov. 2021, doi:10.1145/3465485.
- [2] M. Chen, X. Zheng, J. Lloyd, M. I. Jordan, and E. Brewer, 'Failure Diagnosis Using Decision Trees,' in *Proc. IEEE Int. Conf. Autonomic Computing*, 2004, pp. 36-43.
- [3] S. Dang, Q. Lin, J.-G. Lou, H. Zhang, D. Qiao, and W. Xu, 'AIOps: Real-World Challenges and Research Innovations,' in *Proc. ICSE Companion*, 2019, pp. 4-5.
- [4] Datadog, Inc., 'Datadog Infrastructure Monitoring Documentation,' 2023. [Online]. Available: <https://docs.datadoghq.com/infrastructure/>
- [5] Moogsoft, Inc., 'Moogsoft AIOps Platform Documentation,' 2023. [Online]. Available: <https://docs.moogsoft.com/>
- [6] S. J. Taylor and B. Letham, 'Forecasting at Scale,' *Am. Statistician*, vol. 72, no. 1, pp. 37-45, 2018, doi:10.1080/00031305.2017.1380080.
- [7] G. E. P. Box, G. M. Jenkins, G. C. Reinsel, and G. M. Ljung, *Time Series Analysis: Forecasting and Control*, 5th ed. Hoboken, NJ: Wiley, 2015.
- [8] F. T. Liu, K. M. Ting, and Z. H. Zhou, 'Isolation Forest,' in *Proc. IEEE ICDM*, 2008, pp. 413-422, doi:10.1109/ICDM.2008.17.
- [9] C. H. Lim, D. Yu, P. Krishnamurthy, Y. Kim, C. Liu, and R. Alrawis, 'Poster: Towards Continuous Threshold Optimization for Microservice Monitoring,' in *Proc. ASE*, 2020.
- [10] V. Nair, T. Menzies, N. Siegmund, and S. Apel, 'Faster Discovery of Faster System Configurations with Spectral Learning,' *Autom. Softw. Eng.*, vol. 25, no. 2, pp. 247-277, 2018.
- [11] G. Soldani, D. A. Tamburri, and W.-J. Van Den Heuvel, 'The Pains and Gains of Microservices: A Systematic Grey Literature Review,' *J. Syst. Softw.*, vol. 146, pp. 215-232, Dec. 2018.
- [12] P. Brockwell and R. Davis, *Introduction to Time Series and Forecasting*, 3rd ed. New York, NY: Springer, 2016.
- [13] S. Hochreiter and J. Schmidhuber, 'Long Short-Term Memory,' *Neural Comput.*, vol. 9, no. 8, pp. 1735-1780, 1997.
- [14] S. Mariani, G. Omicini, and G. Agha, 'Tuple-Based Coordination in a Distributed Setting,' in *Proc. 21st Euromicro Int. Conf. Parallel, Distributed, Network-Based Processing*, 2013.
- [15] R. Kuppusamy, 'Tracking Microservice Dependencies,' *Netflix Technology Blog*, 2022. [Online]. Available: <https://netflixtechblog.com/>
- [16] B. Mishchenko, 'Eliminating Toil with Fully Automated Deployments,' *LinkedIn Engineering Blog*, 2023. [Online]. Available: <https://engineering.linkedin.com/>
- [17] M. Ester, H.-P. Kriegel, J. Sander, and X. Xu, 'A Density-Based Algorithm for Discovering Clusters in Large Spatial Databases with Noise,' in *Proc. KDD*, 1996, pp. 226-231.
- [18] B. Beyer, C. Jones, J. Petoff, and N. R. Murphy, *Site Reliability Engineering: How Google Runs Production Systems*. Sebastopol, CA: O'Reilly Media, 2016.
- [19] B. Beyer, N. R. Murphy, D. K. Rensin, K. Kawahara, and S. Thorne, *The Site Reliability Workbook*. Sebastopol, CA: O'Reilly Media, 2018.
- [20] J. Dean and L. A. Barroso, 'The Tail at Scale,' *Commun. ACM*, vol. 56, no. 2, pp. 74-80, Feb. 2013.

- [21] W. Xu, L. Huang, A. Fox, D. Patterson, and M. I. Jordan, 'Detecting Large-Scale System Problems by Mining Console Logs,' in Proc. ACM SOSP, 2009, pp. 117-132.
- [22] Q. Lin, H. Zhang, J.-G. Lou, Y. Zhang, and X. Chen, 'Log Clustering Based Problem Identification for Online Service Systems,' in Proc. ICSE-C, 2016.
- [23] H. Mi, H. Wang, Y. Zhou, M. R. Lyu, and H. Cai, 'Toward Fine-Grained, Unsupervised, Scalable Performance Diagnosis for Production Cloud Computing Systems,' IEEE Trans. Parallel Distrib. Syst., vol. 24, no. 6, pp. 1245-1255, Jun. 2013.
- [24] M. Kleppmann, Designing Data-Intensive Applications. Sebastopol, CA: O'Reilly Media, 2017.
- [25] OpenTelemetry Authors, 'OpenTelemetry Specification v1.x,' CNCF, 2023. [Online]. Available: <https://opentelemetry.io/docs/>
- [26] Apache Software Foundation, 'Apache Kafka Documentation,' 2023. [Online]. Available: <https://kafka.apache.org/documentation/>
- [27] Apache Software Foundation, 'Apache Flink Documentation,' 2023. [Online]. Available: <https://nightlies.apache.org/flink/flink-docs-stable/>
- [28] T. Chen, X. Zhang, and H. Zhang, 'Outage Prediction and Diagnosis for Cloud Service Systems,' in Proc. WWW, 2019, pp. 2659-2665.
- [29] PCI Security Standards Council, 'PCI DSS v4.0,' 2022. [Online]. Available: <https://www.pcisecuritystandards.org/>
- [30] R. Syer, Z. Shang, H. Jiang, and A. E. Hassan, 'Continuous Validation of Performance Tests by Statistical Analysis of Performance Evolution,' J. Syst. Softw., vol. 103, pp. 224-237, May 2015.
- [31] X. Wang, L. Wu, M. Yin, Y. Wen, and Y. Li, 'Groot: An Event-Graph-Based Approach for Root Cause Analysis in Industrial Settings,' in Proc. ASE, 2021, pp. 1219-1230.
- [32] Amazon Web Services, 'AWS Observability Best Practices,' 2023. [Online]. Available: <https://aws-observability.github.io/observability-best-practices/>
- [33] R. P. Adams and D. J. C. MacKay, 'Bayesian Online Changepoint Detection,' arXiv:0710.3742, 2007. [Online]. Available: <https://arxiv.org/abs/0710.3742>