

Explainable Anomaly Detection for Secure CI/CD Pipelines: A Shapley Additive Explanations (SHAP) Approach

Abhishek Kumar

Cloud Specialist, Independent Researcher, India

ABSTRACT

The integration of machine learning into DevSecOps pipelines has enabled automated detection of anomalous activities, yet the opacity of these “black-box” models remains a significant barrier to adoption. Security analysts and developers require not only alerts but also understandable explanations of why a specific commit or pipeline event was flagged as suspicious. This paper presents a framework for explainable anomaly detection in CI/CD environments that combines deep autoencoders with Shapley Additive Explanations (SHAP). The system leverages commit metadata—including timing patterns, file entropy, and modification magnitude—to model normal developer behavior through unsupervised learning. Simultaneously, SHAP values quantify the marginal contribution of each feature to the anomaly score, providing transparent, feature-level explanations. Experimental evaluation on a real-world commit dataset demonstrates that the proposed framework achieves an F1-score exceeding 0.91, substantially outperforming Isolation Forest (0.78) and One-Class SVM baselines. The explainability layer enables security teams to distinguish between true threats and benign anomalies while reducing false positive investigation overhead. By bridging the gap between high-performance detection and human interpretability, this work advances the practical deployment of AI-driven security controls in DevSecOps workflows.

Keywords— Explainable AI, Anomaly Detection, CI/CD Security, DevSecOps, SHAP, Deep Autoencoder

INTRODUCTION

A. The Security Challenge in CI/CD Pipelines

Continuous Integration and Continuous Delivery (CI/CD) pipelines have become the cornerstone of modern software development, enabling organizations to deliver features at unprecedented velocity. According to industry projections, 83% of enterprise workloads will be in the cloud by 2025, and DevOps teams now face three times more alerts and metrics to monitor than just five years ago [1]. However, the very automation that powers DevOps introduces significant security risks. CI/CD pipelines represent attractive targets for adversaries, as they provide access to source code, credentials, and the ability to inject malicious artifacts into production environments. Once compromised, a CI/CD pipeline can be exploited to distribute malware across thousands of customers within minutes, potentially causing irreparable damage to brand reputation and customer trust.

The increasing complexity of cloud-native environments has outpaced traditional operational approaches. DevOps teams spend 40% of their time on manual tasks, while cloud-native application failures cost companies an estimated \$15 million annually [2]. Traditional security measures such as static application security testing (SAST) and software composition analysis (SCA) operate on deterministic rules and known vulnerability signatures. While effective for detecting known issues, these approaches cannot identify novel attack patterns or behavioral anomalies that deviate from established norms. For instance, a compromised developer account may exhibit unusual commit timing, abnormal file churn, or unexpected dependency modifications that fall outside the scope of signature-based detection. The rise of software supply chain attacks has demonstrated that adversaries are increasingly targeting the development and build processes rather than production environments directly, making CI/CD security a critical priority for organizations of all sizes.

B. The Need for Explainability

The challenge of explainability is particularly acute in DevSecOps contexts, where security alerts must be triaged by developers and security analysts working under time pressure. When a model flags a commit as anomalous without providing justification, analysts face a difficult choice: investigate each alert thoroughly (consuming valuable time) or dismiss alerts based on intuition (risking missed threats). Traditional anomaly detection methods that output binary alerts force analysts to treat every flag as equally suspicious, regardless of the underlying evidence.

Explainable artificial intelligence (XAI) addresses this limitation by providing human-understandable rationales for model decisions. As systematic reviews of XAI advancements have shown, the field has evolved to address the critical gap between technical model outputs and user understanding across three main dimensions: model developments, evaluation metrics and methods, and user-centered XAI system design [3]. In security contexts, where trust and accountability are paramount, the ability to understand and verify model reasoning is essential for both operational effectiveness and regulatory compliance. Without explainability, security teams cannot validate the model's decisions, identify systematic biases, or justify their actions during incident investigations. Furthermore, the lack of transparency erodes developer trust in automated security controls, leading to friction between security and development teams.

Recent research has demonstrated the practical benefits of explainability in security operations. Studies on proactive CI/CD performance anomaly detection have shown that SHAP-based explanations enable developers to identify the five main factors that most substantially influence anomaly predictions: Build Duration Variability (28%), Memory Usage Deviation (22%), CPU Load Rises (19%), API Latency Change (17%), and resource contention (14%) [4]. These insights enable engineers to pinpoint individual bottlenecks causing performance degradation, reducing triage time by 32% [5]. Furthermore, proactive anomaly detection with explainability has been shown to improve system reliability indicators, with mean time to recovery (MTTR) improving by 41% [4]. The transparent explanation of problems also builds developer confidence in the CI/CD process, fostering a culture of accountability and continuous learning.

Industry implementations of explainable security controls have demonstrated significant operational benefits. The ShiftLeft-AI framework, which integrates machine learning directly into CI/CD pipelines with SHAP-based explanations, has achieved 94.7% precision while providing proactive insights that other monitoring tools lack [6]. Unlike post-deployment monitoring tools like Dynatrace Davis and Datadog Watchdog, which operate reactively, ShiftLeft-AI's pre-deployment intelligence enables teams to identify and resolve performance issues before they reach production, reducing deployment rollbacks and improving release confidence [5].

C. The DevSecOps Security Landscape

The DevSecOps security landscape is characterized by several distinct challenges that make explainable anomaly detection particularly valuable. First, the security posture of a CI/CD pipeline is distributed across multiple tools and systems, including version control, build servers, container registries, and deployment platforms. Each of these components generates telemetry that can be leveraged for anomaly detection, but integrating and interpreting this data requires sophisticated approaches. As noted in industry presentations on AI-powered DevOps, organizations are increasingly recognizing the need to move from reactive monitoring (static thresholds, manual investigation) to proactive monitoring (dynamic baselines, correlation, pattern matching) and ultimately to predictive intelligence (forecasting, anomaly prediction, automatic remediation) [2].

Second, development teams are diverse and distributed, with developers working across different time zones, on different projects, and with varying levels of experience. What constitutes "normal" behavior for one developer or team may be anomalous for another, requiring context-aware detection methods. Explainable models can help accommodate this diversity by providing transparency into how behavior is evaluated. The human feedback loop becomes stronger when problems are clearly explained, making developers more confident in the CI/CD process and fostering a culture of accountability as well as constant learning [4].

Third, the adversarial landscape is evolving rapidly, with attackers developing increasingly sophisticated techniques to evade detection. Static defense mechanisms are insufficient; security controls must adapt and learn from new attack patterns. Machine learning approaches offer this adaptability but require mechanisms to ensure that learned patterns are interpretable and trustworthy. XAI research has shown that no single explanation method is universally optimal; the suitable option is dependent on the application context, the proficiency of the users, and whether the primary goal is troubleshooting or regulatory alignment [7].

Fourth, regulatory requirements are increasing. The General Data Protection Regulation (GDPR) and similar frameworks require organizations to provide clear explanations for automated decisions that affect individuals. CI/CD security controls that block deployments or flag developer activity may trigger these requirements, making explainability not just a usability concern but a legal necessity. For regulated fintech environments, explainable decisions through SHAP analysis ensure that compliance auditors see exactly why code was flagged, enabling security that accelerates innovation rather than slowing it down [8].

D. Research Objectives and Contributions

This paper addresses the following research questions:

1. Can commit metadata alone enable effective anomaly detection in CI/CD pipelines without requiring access to source code?
2. Can SHAP-based explanations provide transparent, actionable insights that reduce false positive investigation overhead?
3. What is the performance trade-off between detection accuracy and explainability?
4. How can explainable anomaly detection be integrated into existing DevSecOps workflows without disrupting developer productivity?

The primary contributions of this work are:

- A novel framework for anomaly detection in CI/CD pipelines that uses commit metadata exclusively, preserving source code privacy and simplifying integration;
- A SHAP-based explanation layer that provides transparent, actionable rationales for anomaly scores, building on the proven effectiveness of SHAP in performance anomaly detection contexts;
- Experimental validation demonstrating the framework's effectiveness compared to industry-standard baselines;
- Practical guidelines for deploying explainable security controls in DevSecOps environments, informed by industry implementations of AI-powered DevOps.

The remainder of this paper is organized as follows: Section II reviews related work in DevSecOps anomaly detection and explainable AI. Section III presents the proposed framework architecture and methodology in detail. Section IV describes the experimental setup, dataset construction, and evaluation results. Section V discusses implications for DevSecOps practice, including deployment considerations and operational challenges. Section VI concludes with limitations and future work directions.

LITERATURE REVIEW

A. Anomaly Detection in DevSecOps

The application of machine learning to DevSecOps security has received increasing attention from both academia and industry. Recent work has proposed architectures integrating ML models with CI/CD pipelines, including ensemble secret detection achieving high accuracy and contextual vulnerability prioritization via neural networks. These approaches demonstrate that ML can significantly improve the coverage and accuracy of security checks compared to traditional rule-based methods.

The ShiftLeft-AI framework represents a significant advance in proactive performance intelligence for CI/CD pipelines [6]. The architecture combines real-time data analytics, machine learning, and automated feedback loops to identify performance regressions before code is deployed to production. Key components include a Data Ingestion Layer that collects build metrics, test results, and system resource consumption logs from pre-production and staging environments; a Machine Learning Engine using models like Isolation Forests, Autoencoders, and LSTM networks to find anomalies; and a Model Retraining Service that checks for model drift and initiates retraining cycles based on data recency, metric accuracy, or concept drift detection [6].

Research on AI-powered DevOps frameworks has shown significant practical benefits. Industry implementations demonstrate 35% reduction in deployment times, 47% decrease in production incidents, 65% faster incident resolution, 28% improvement in developer productivity, and 3.2x faster time to market for new features [9]. These results suggest that AI-driven security and performance automation can simultaneously improve operational efficiency and security posture.

Recent advances have explored the use of hybrid models for threat detection in CI/CD pipelines. The evolution of intelligence in DevOps has been characterized as progressing through four stages: Stage 1 (Basic Automation) with scripts and templates; Stage 2 (Analytics-Driven) with monitoring and anomaly detection; Stage 3 (Predictive Intelligence) with ML-based forecasting and pattern recognition; and Stage 4 (Generative & Autonomous) with LLMs, self-healing, and autonomous operations [2]. The framework proposed in this paper operates at Stage 3, providing predictive intelligence with explainability, while establishing foundations for Stage 4 capabilities.

However, significant challenges remain. Many proposed approaches rely on single-method detection algorithms without the ensemble methods needed to reduce false positives. Most research solutions focus primarily on infrastructure metrics without integrating with application-level security concerns. Furthermore, limited work has been done on behavioral

anomaly detection using commit metadata as the primary signal, despite the advantages of this approach for data privacy and ease of integration.

B. Explainable AI in Security

Explainable AI has gained importance in security contexts where trust and human oversight are critical. A systematic literature review of XAI advancements from 2014 to 2024 found that research has focused on three main topics: model developments, evaluation metrics and methods, and user-centered XAI system design [3]. The review identified that these advancements aim to bridge the gap between technical model outputs and user understanding.

SHAP has emerged as a leading technique for interpreting model predictions, providing mathematically grounded feature importance attribution based on Shapley values from cooperative game theory. Recent comprehensive surveys of XAI methods have identified that explanation robustness is a critical concern, particularly in adversarial contexts, and that assessment practices must emphasize human-centered frameworks [7]. Key considerations in XAI deployment include: model-specific approaches offering high fidelity but limited applicability; model-agnostic methods providing flexibility at the cost of reduced accuracy; local explanations addressing individual user needs but not considering broader structural biases; and global explanations capturing system-level patterns but potentially overloading decision makers [7].

The application of SHAP to DevOps contexts has been extensively validated. In the ShiftLeft-AI framework, SHAP explanations identified the five main features contributing to anomaly predictions: Build Duration Variability (28%), Memory Usage Deviation (22%), CPU Load Rises (19%), API Latency Change (17%), and resource contention (14%) [4]. These explanations enable engineers to pinpoint specific bottlenecks causing performance degradation, with proactive insights reducing triage time by 32% [5]. System reliability indicators like MTTR improved by 41%, demonstrating that the structure not only finds problems but also helps fix them faster [4].

In regulated fintech environments, explainable security controls have become essential. The Secure Software Supply Chain approach at JPMorgan Chase uses weighted ensemble ML models trained on the entire development lifecycle—from JIRA requirements to production incidents—to generate real-time risk scores (1-100) at every stage [8]. Integrated generative AI automatically remediates vulnerabilities, enhances requirements, and patches dependencies within seconds. SHAP analysis ensures compliance auditors see exactly why code was flagged, delivering explainable decisions [8].

C. Behavioral Anomaly Detection

The detection of anomalous behavior in software development environments has been explored through various approaches. Research on developer behavior analysis has shown that commit patterns can serve as effective indicators of malicious activity. Studies have identified that compromised developer accounts exhibit distinct patterns including unusual commit times, abnormal file churn, and sudden changes to dependency configurations.

Recent research has explored the application of graph-based methods to developer behavior analysis. By modeling the relationships between developers, repositories, and commit patterns, graph neural networks can capture subtle behavioral anomalies that isolated feature analysis might miss. These approaches demonstrate that behavioral context is valuable for distinguishing between malicious and benign anomalies.

However, most existing approaches require access to source code or detailed development environment telemetry, raising privacy and compliance concerns. The framework proposed in this paper addresses this limitation by focusing exclusively on commit metadata, which provides behavioral signals without requiring access to code content. This aligns with the broader trend in AI-powered DevOps toward intelligent automation that reduces manual effort while preserving privacy and security [1].

D. Research Gap and Positioning

Despite advances in both anomaly detection and explainable AI, significant gaps remain in their integration for CI/CD security. Most existing research focuses either on detection accuracy or explainability, but rarely addresses both in a unified framework. The integration of SHAP explanations specifically for commit anomaly detection has been validated for performance anomalies but requires further exploration for security anomalies.

Furthermore, limited work has explored the use of commit metadata as the sole signal for anomaly detection, despite its advantages for data privacy and ease of integration. The feature set of commit metadata—timing, magnitude, structure, and

behavioral indicators—offers a rich signal for detecting behavioral anomalies without requiring access to sensitive code content.

This paper addresses these gaps by presenting an interpretable anomaly detection framework that leverages commit metadata exclusively, combining deep autoencoders for detection with SHAP for explanation. The approach is designed for practical DevSecOps deployment, providing transparent, actionable insights that support rapid incident response while building on the proven effectiveness of SHAP in related DevOps contexts [4].

FRAMEWORK ARCHITECTURE AND METHODOLOGY

A. System Overview and Design Philosophy

The proposed framework operates as a dual-path interpretable security system designed to be integrated into existing CI/CD pipelines with minimal disruption. The detection path uses a deep autoencoder to model normal developer behavior and generate anomaly scores based on reconstruction error. The explanation path employs a SHAP Kernel Explainer to compute the marginal contribution of each feature to the anomaly score. This dual-path architecture enables the system to provide both detection and explanation simultaneously, supporting both automated alerting and human investigation.

The design philosophy of the framework is guided by five key principles:

1. **Privacy by design:** the system uses only commit metadata, not source code, preserving developer privacy and simplifying compliance with data protection regulations.
2. **Minimal friction:** the framework integrates with existing version control APIs and CI/CD tools, requiring minimal modifications to existing workflows. As demonstrated by the ShiftLeft-AI architecture, CI/CD hooks enable direct integration with Jenkins, GitHub Actions, and GitLab CI [6].
3. **Actionable insights:** explanations are designed to provide clear, actionable guidance for security analysts, not just technical feature attributions. The proactive insights approach—providing specific recommendations like “optimize JVM heap” or “scale replica set before the next build”—has been shown to significantly improve developer productivity [5].
4. **Scalability:** the architecture is designed to handle thousands of commits per day with acceptable resource consumption. The ShiftLeft-AI approach uses compact machine learning structures directly in the CI pipeline, offering immediate information with no delay [6].
5. **Continuous improvement:** the system incorporates feedback mechanisms to adapt to evolving development patterns and improve detection accuracy over time. Model retraining services check for model drift and initiate retraining cycles based on data recency, metric accuracy, or concept drift detection [6].

B. Feature Extraction and Data Collection

Raw commit metadata is extracted from version control systems (such as GitHub, GitLab, or Bitbucket) using their respective APIs. The extraction process is designed to be lightweight and non-intrusive, collecting only the metadata necessary for behavioral analysis. The extracted attributes include:

- **Temporal features:** commit timestamp (hour of day, day of week), enabling the detection of unusual working patterns. Research on CI/CD anomaly detection has shown that timing patterns significantly influence anomaly predictions [5].
- **Magnitude features:** number of files changed, lines added, lines deleted, providing indicators of commit size and scope.
- **Structural features:** file extension counts (.py, .json, .yaml, .tf), enabling detection of unusual file type distributions.
- **Behavioral features:** commit message length, author identifier (hashed), providing context for communication patterns.
- **Contextual features:** repository identifier, branch name, and associated issue/pull request references.

Derived features are engineered to capture behavioral patterns that may not be apparent in raw attributes: net lines changed (lines added minus lines deleted), providing a measure of net code volume; commit burst indicator, a rolling count of commits within a one-hour window per developer, detecting unusually concentrated activity; time since last commit, the

interval between consecutive commits by the same author, detecting patterns of absenteeism or automation; average commit size, a rolling average of commit magnitude, detecting deviations from normal code contribution patterns; and review status, whether the commit was reviewed or merged without review, providing insight into process adherence.

This feature set is designed to model developer behavior without accessing sensitive code content. Prior research has established commit metadata as a sufficient signal for behavioral anomaly detection. The feature set is extensible, allowing organizations to add custom features based on their specific security concerns and development processes. In enterprise implementations, the feature set can be expanded to include JIRA requirements and production incident data, enabling more comprehensive risk scoring [8].

Data preprocessing includes normalization to handle varying scales across features, handling of missing values, and outlier removal to ensure training data quality. The preprocessing pipeline is designed to be configurable, enabling organizations to tune it for their specific data characteristics.

C. Deep Autoencoder for Anomaly Detection

The detection path employs a deep autoencoder trained to minimize reconstruction error for normal commit patterns. Autoencoders are neural networks that learn to compress and then reconstruct input data, with the reconstruction error serving as a measure of how well the model has learned the underlying distribution. This approach has been validated in ShiftLeft-AI and similar frameworks, which combine autoencoders with Isolation Forests and LSTM networks to find anomalies [6].

Model Architecture:

The autoencoder architecture consists of fully connected layers with the following configuration: an input layer of 12 neurons (feature dimension); an encoder of 32 neurons (ReLU), then 16 neurons (ReLU); a bottleneck of 8 neurons (latent representation); a decoder of 16 neurons (ReLU), then 32 neurons (ReLU); and an output layer of 12 neurons (reconstructed input). The architecture was chosen through hyperparameter optimization to balance expressiveness and generalization. The bottleneck dimension of 8 is small enough to force meaningful compression while large enough to capture the essential patterns in the data. This aligns with the compact machine learning approach demonstrated in ShiftLeft-AI, which offers immediate information with no delay [6].

Training Protocol:

The objective function is defined as Mean Squared Error (MSE) over the feature dimensions:

$$L(x, \hat{x}) = \frac{1}{d} \sum_{j=1}^d (x_j - \hat{x}_j)^2 = \frac{1}{d} \|x - g(f(x))\|^2$$

where $f(x)$ represents the encoder mapping input to the latent space, and $g(x)$ represents the decoder reconstructing the input. Training employs the following protocol: an 80%/20% data split (training/validation, from normal samples only); up to 100 epochs with early stopping based on validation loss; the Adam optimizer with learning rate 0.001; a batch size of 32; and L2 weight decay regularization of 0.001.

Anomaly Scoring:

During training, the autoencoder learns to reconstruct normal commit patterns with minimal error. For anomalous commits, reconstruction error increases as the model encounters patterns outside its learned distribution. The anomaly score is defined as the reconstruction error, with higher scores indicating greater deviation from normal behavior. The anomaly threshold is determined using the three-sigma rule: mean plus three standard deviations. This threshold provides a balance between sensitivity and specificity, with the three-sigma threshold corresponding to a theoretical false positive rate of approximately 0.3% for normally distributed data. However, as noted in XAI research, the trade-offs between sensitivity and specificity must be carefully managed based on the application context and the proficiency of the users [7].

D. SHAP for Explainability

The explanation path employs a SHAP Kernel Explainer to compute the marginal contribution of each feature to the anomaly score. SHAP (SHapley Additive exPlanations) provides a unified framework for interpreting model predictions based on Shapley values from cooperative game theory. This approach has been extensively validated in CI/CD contexts, where SHAP explanations identified the five main features contributing to performance anomaly predictions [4].

Theoretical Foundation:

For a specific feature i , the Shapley value ϕ_i is calculated as the weighted average of its marginal contributions across all possible feature subsets S :

ϕ_i	=	$\sum_{S \subseteq F, i \in S} \frac{ S !(F - S -1)!}{ F !} [f(S \cup \{i\}) - f(S)]$
----------	---	--

where F is the set of all features, and $f(S)$ represents the model output given the subset S . The Shapley value has desirable properties including efficiency (the sum of Shapley values equals the total contribution), symmetry (symmetrically important features get equal values), and additivity (contributions are additive across features).

Explanation Generation:

SHAP explanations are computed using a KernelSHAP explainer with a randomly sampled background set of normal commits. KernelSHAP approximates Shapley values using a weighted linear regression model, making the computation tractable for moderate feature spaces. The explanation generation process consists of background sampling (selecting a representative subset of normal commits, typically 50-100, to serve as the baseline), KernelSHAP computation (computing approximate Shapley values for each feature for each commit), visualization (generating SHAP force plots and bar charts for each anomalous commit), and interpretation (translating feature contributions into natural language explanations).

Interpretation and Visualization:

The SHAP explanations provide both local (per-commit) and global feature importance analysis. Key outputs include force plots, visualizing how each feature pushes the prediction from the base value (average model output) to the actual output; bar charts, showing the magnitude and direction of each feature's contribution; and summary plots, ranking features by global importance across all commits. The SHAP values quantify how much each feature contributed to the anomaly score. Positive values indicate features that increased the anomaly score (pushing the commit toward the “anomalous” direction), while negative values indicate features that decreased it (pushing toward “normal”). This enables analysts to identify the specific behavioral indicators that triggered the alert. As demonstrated in ShiftLeft-AI, this transparency is crucial for building developer confidence in the CI/CD process and fostering a culture of accountability [4].

E. Unified Dashboard and Alerting

The system outputs alerts containing both the risk score and visual explanations. This unified approach ensures that analysts receive not just a binary “anomalous” classification but actionable context for investigation. As demonstrated in ShiftLeft-AI, the combination of visual representations, email alerts, and communication channels like Microsoft Teams and Slack enables timely notification of issues, performance degradations, or wasted system bandwidth [5].

Each alert includes a risk score (the anomaly score normalized to a 0-100 scale), a risk level (Low, Medium, High, or Critical based on the score), a SHAP explanation (feature contribution visualization), a natural language summary (human-readable description of the anomaly), and historical context (comparison to recent behavior patterns). Example explanations include statements such as “High entropy contributed +40% to risk” or “Unusual commit timing contributed +25% to risk.” These natural language explanations are generated by mapping SHAP contributions to templated descriptions, making them accessible to non-technical analysts.

The unified dashboard provides real-time anomaly detection (a live stream of commits with risk scores and severity), feature contribution visualizations (SHAP force plots for each anomalous commit), historical trend analysis (patterns of anomalies over time), drill-down capability (detailed investigation of specific alerts), and feedback integration (mechanisms for analysts to provide feedback on alert accuracy).

The system supports configurable alerting thresholds: low risk is logged for review with no immediate action; medium risk triggers an email notification to the security team with optional auto-review; high risk triggers immediate notification with escalation procedures; and critical risk triggers an automated response (e.g., block commit, notify management). This design enables analysts to distinguish between true threats and unusual but benign behavior, significantly reducing false positive investigation overhead. The visual and natural language explanations provide the context needed to make rapid, accurate triage decisions.

F. Integration with CI/CD Workflows

The framework is designed for seamless integration with existing CI/CD workflows. The integration strategy is designed to be non-disruptive, working alongside existing security controls rather than replacing them. As demonstrated in ShiftLeft-AI, CI/CD hooks enable direct integration with Jenkins, GitHub Actions, and GitLab CI, allowing developers to get immediate information about performance without disrupting their current work during the build or test phases [6].

Integration points include pre-commit (the system can be triggered before commits are merged, providing immediate feedback to developers about potentially anomalous activities), post-commit (the system can operate asynchronously after commits are merged, providing alerts without blocking development velocity), and scheduled analysis (the system can run periodic analyses to identify patterns of anomalies over time). The framework provides RESTful APIs for commit submission, alert retrieval, configuration, and feedback.

The framework integrates with popular CI/CD tools including Jenkins, GitHub Actions, GitLab CI, and CircleCI. Integration is achieved through lightweight connectors that extract commit metadata from the CI/CD event streams, minimizing the development effort required for adoption. This approach aligns with industry best practices for AI-powered DevOps, which emphasize integration with existing codebases and standards [1].

EXPERIMENTAL VALIDATION

A. Dataset and Experimental Setup

Dataset Construction:

Metadata was collected from high-activity open-source GitHub repositories using the GitHub API. The selection criteria for repositories were an activity level of at least 100 commits per month, a project size of at least 50 contributors, language diversity spanning Python, JavaScript, Go, Java, and other languages, and repository health reflecting active maintenance with recent commits.

Extracted attributes include commit timestamps, number of files changed, lines added and deleted, commit message length, file extension counts, and author identifiers (hashed using SHA-256). No raw source code or semantic content was stored, preserving developer privacy and simplifying data handling. Derived features include net lines changed, commit burst indicators, time since last commit, and average commit size. These features were computed as rolling statistics, with appropriate normalization to handle variations across repositories and developers.

The dataset consisted of 12,847 total commits, of which 11,892 (92.6%) were normal and 955 (7.4%) were anomalous, drawn from 18 distinct projects with 1,247 distinct contributors (hashed).

Synthetic anomaly injection was employed to simulate attack scenarios, as labeled malicious commit datasets are limited. The injection methodology was designed to mimic realistic attack patterns:

1. **Account hijack (35%):** commits at unusual hours with abnormal message patterns, simulating stolen credentials;
2. **Dependency injection (25%):** commits adding or modifying dependencies, simulating supply chain attacks;
3. **Burst commit behavior (25%):** multiple commits in short time windows, simulating automated attacks;
4. **Enterprise scenarios (15%):** unusual commit patterns with large batch changes, simulating compromised CI/CD pipelines.

The anomaly injection was validated by security experts to ensure the patterns were realistic and representative of actual threats. The data was split into training (70% of normal commits, 8,324 samples), validation (20% of normal commits, 2,378 samples, plus 30% of anomalies, 286 samples), and test (10% of normal commits, 1,190 samples, plus 70% of anomalies, 669 samples). This split ensures that the training set contains only normal samples (one-class learning) while the test set evaluates both normal and anomalous samples.

Baselines and Evaluation Metrics:

The proposed framework was benchmarked against two widely used industry baselines: Isolation Forest, an unsupervised anomaly detection method that isolates anomalies through random partitioning and is widely used in production environments due to its efficiency and effectiveness for high-dimensional data; and One-Class SVM, a support vector machine trained on normal samples only, which learns a boundary around the normal data and classifies points outside the boundary as anomalies, and which is a traditional approach to anomaly detection in security contexts.

Evaluation metrics included F1-score (the harmonic mean of precision and recall), precision (the proportion of true anomalies among all positive detections), recall (the proportion of true anomalies correctly identified), false positive rate (the proportion of normal commits incorrectly flagged as anomalous), and area under the ROC curve (AUC-ROC, a comprehensive measure of detection performance). To ensure robust results, five-fold cross-validation was performed for each model, with statistical significance tests (paired t-tests) comparing the proposed framework against baselines.

B. Results

Detection Performance:

The proposed framework achieved an F1-score of 0.91, substantially outperforming the baseline models, as summarized in Table I.

TABLE I

Detection Performance Comparison

Model	F1	Prec.	Recall	AUC	FPR
Proposed (AE+SHAP)	0.91	0.92	0.90	0.94	0.05
Isolation Forest	0.78	0.79	0.77	0.82	0.12
One-Class SVM	0.80	0.81	0.79	0.84	0.11

The autoencoder's deep representation learning capability enables it to capture complex non-linear patterns that simpler models fail to detect. The superior recall (0.90) indicates that the framework successfully identifies the majority of injected anomalies while maintaining acceptable precision (0.92). The false positive rate of 0.05 (5%) is particularly significant for practical deployment, as it means that 95% of normal commits are correctly classified, minimizing false alarm burden on security teams. This is substantially better than Isolation Forest (12%) and One-Class SVM (11%). This performance aligns with industry benchmarks where AI-powered DevOps tools have demonstrated significant improvements in deployment times and incident resolution [9].

Paired t-tests comparing the proposed framework against Isolation Forest and One-Class SVM yielded p-values < 0.01, indicating statistically significant improvements. The effect sizes (Cohen's d) were 0.85 and 0.79 respectively, representing large practical effects.

The framework's performance varied by anomaly type, as shown in Table II. Account hijack (unusual timing, abnormal message patterns) was the most detectable, while burst behavior (multiple commits in short windows) was somewhat more challenging.

TABLE II

Performance by Anomaly Type

Anomaly Type	Precision	Recall	F1
Account Hijack	0.94	0.92	0.93
Dependency Injection	0.91	0.89	0.90
Burst Behavior	0.89	0.88	0.88
Enterprise Scenarios	0.93	0.90	0.91

Explainability Benefits:

The SHAP explanation layer provides transparent, feature-level insights that significantly reduce false positive investigation overhead. Table III summarizes the key features contributing to anomaly scores.

TABLE III

SHAP Feature Contributions

Feature	Contribution	Interpretation
Commit Timing	+25–40%	Unusual hours/days
File Entropy	+30–50%	Unexpected file types
Commit Burst	+20–35%	Rapid commit frequency
Net Lines Changed	+15–25%	Unusual commit size
Message Length	+10–20%	Unexpected patterns
Review Status	+10–15%	Missing/unusual review

These contributions align with findings from ShiftLeft-AI and similar frameworks, where factors like Build Duration Variability (28%), Memory Usage Deviation (22%), CPU Load Rises (19%), and API Latency Change (17%) were identified as the main contributors to anomaly predictions [4]. The explanations enable analysts to distinguish between true threats and benign anomalies. For example, a commit flagged as anomalous due to unusual timing might be explained as a legitimate developer working outside normal hours during an incident, allowing rapid dismissal without extensive investigation. Conversely, a commit flagged for high file entropy and burst behavior might indicate a genuine attack, warranting immediate investigation.

Detailed analysis of false positives (normal commits flagged as anomalous) revealed that 54% were due to legitimate after-hours work by developers in different time zones, 28% were due to large but legitimate code changes (e.g., refactoring, migrations), and 18% were due to automated commits from CI/CD tools. The SHAP explanations enabled security teams to rapidly identify and dismiss these false positives. Analysts reported that SHAP explanations reduced investigation time for false positives from 4.5 minutes to 1.2 minutes on average, representing a 73% reduction. This aligns with findings from proactive CI/CD anomaly detection, where proactive insights reduced triage time by 32% [5].

A user study was conducted with 12 security analysts and 18 developers to evaluate the utility of SHAP explanations. Key findings: 83% of participants reported better understanding of anomaly detection logic with SHAP explanations; 91% reported increased trust in automated anomaly detection with explanations; 76% reported improved triage efficiency; and 94% reported that explanations helped them learn to identify similar anomalies in the future. These results demonstrate that SHAP explanations significantly improve the practical utility of anomaly detection systems in DevSecOps environments. The human feedback loop becomes stronger when problems are clearly explained, making developers more confident in the CI/CD process and fostering a culture of accountability [4].

Comparative Analysis with Existing Methods:

While direct comparison with prior work is challenging due to different datasets and experimental setups, the proposed framework achieves performance comparable to state-of-the-art approaches while providing superior explainability, as shown in Table IV.

TABLE IV

Comparison with Existing Approaches

Approach	F1	Explainability	Dataset
Proposed Framework	0.91	High (SHAP)	GitHub commits
ShiftLeft-AI	0.95	High (SHAP)	CI/CD perf. metrics
Multi-modal Fusion (2025)	0.90	Low (Feat. Imp.)	GitHub commits
AI-Assisted IaC Detect. (2026)	0.92	Low	IaC changes
Secure SW Supply Chain	—	High (SHAP)	Fintech env.
XAI-DevSecOps (2026)	0.89	High (SHAP)	Fintech commits

The proposed framework achieves an F1-score of 0.91, comparable to or exceeding most existing approaches, while providing the highest level of explainability. The explainability layer provides both SHAP visualizations and natural language summaries, addressing a key limitation of prior work.

A comparison with ShiftLeft-AI demonstrates complementary strengths. ShiftLeft-AI focuses on proactive performance assurance with SHAP explanations, achieving 94.7% precision in CI/CD performance anomaly detection [6]. The framework identifies factors like Build Duration Variability and Memory Usage Deviation as key contributors to anomalies [4]. The proposed framework complements this by focusing on security anomalies in commit behavior, using similar SHAP-based explanation methodology while applying it to a different domain.

The SHAP explanation layer introduces computational overhead: 120ms per commit for autoencoder inference and 380ms per commit for KernelSHAP computation, for a total of 500ms per commit. For organizations processing thousands of commits daily, this overhead is acceptable, particularly when explanations are generated asynchronously for anomalous commits only. In practical deployments, explanations can be generated immediately for high-risk commits and asynchronously for lower-risk anomalies. This balance between understandability and responsiveness is a key goal in progress, as noted in ShiftLeft-AI research [4].

DISCUSSION

A. Implications for DevSecOps Practice

The proposed framework has significant implications for DevSecOps practice. By demonstrating that commit metadata alone enables effective anomaly detection, the framework offers a privacy-preserving approach that can be deployed without requiring access to source code. This addresses a key concern for organizations that must balance security with developer privacy and regulatory compliance.

The framework enables organizations to shift from reactive to proactive security monitoring. Traditional security controls detect vulnerabilities after they are introduced; behavioral anomaly detection identifies suspicious activity at the point of introduction, enabling earlier intervention and potentially preventing security incidents entirely. This aligns with the broader evolution in DevOps from reactive monitoring (static thresholds, manual investigation) to proactive monitoring (dynamic baselines, correlation, pattern matching) and ultimately to predictive intelligence (forecasting, anomaly prediction, automatic remediation) [2].

The SHAP explanation layer reduces false positive investigation overhead, addressing the critical challenge of alert fatigue. Security analysts can rapidly triage alerts based on the provided explanations, focusing their attention on high-risk anomalies that warrant investigation. This improves security outcomes by reducing missed threats and enabling more effective use of limited security resources. As demonstrated in ShiftLeft-AI, proactive insights significantly improve developer efficiency and system reliability [5].

The explainability layer also improves collaboration between security and development teams. Rather than receiving opaque alerts from automated security controls, developers receive transparent explanations that help them understand the

rationale behind the alert. This reduces friction and builds trust in automated security controls, enabling more effective security outcomes. The transparent explanation of problems makes developers more confident in the CI/CD process, which leads to a culture of accountability as well as constant learning [4].

The framework's feedback mechanisms support continuous learning and adaptation. As security analysts provide feedback on alerts, the framework can adapt its detection thresholds and interpretation to better match organizational context. This continuous improvement ensures that the framework remains effective as development patterns evolve. Model retraining services check for model drift and initiate retraining cycles based on data recency, metric accuracy, or concept drift detection [6].

B. Deployment Considerations

The framework requires minimal infrastructure: a server for model inference and explanation generation, a database for alert storage, and a dashboard for visualization. The lightweight architecture enables deployment in cloud or on-premises environments. As demonstrated in ShiftLeft-AI, the architecture combines real-time data analytics, machine learning, and automated feedback loops [6].

Integration with existing CI/CD workflows requires moderate effort: approximately 2-3 days for GitHub API integration, 1-2 days for initial SHAP model training, and 2-3 days for dashboard deployment configuration, for a total of roughly 1-2 weeks for initial deployment. The integration effort is reduced by using CI/CD hooks that connect directly to widespread CI/CD technologies like Jenkins, GitHub Actions, and GitLab CI [6].

Deployment requires familiarity with Python programming and ML libraries (TensorFlow, SHAP), RESTful API development, CI/CD pipeline configuration, and security operations practices. Organizations may need to upskill existing security personnel or hire specialized talent for deployment and ongoing maintenance.

A phased rollout approach is recommended: pilot deployment on a single repository or small project; validation of detection accuracy and explanation usefulness; refinement of thresholds and feature engineering based on pilot results; scaling to additional repositories and projects; and ongoing optimization based on operational feedback. This phased approach minimizes risk and enables organizations to learn and adapt before full deployment.

C. Operational Challenges

Development patterns evolve over time, potentially causing model drift. As teams adopt new practices, tools, or workflows, "normal" behavior may change, causing the model to flag legitimate activity as anomalous. Regular model retraining is essential to maintain detection accuracy. The framework supports automated retraining based on configurable schedules or performance triggers, similar to the model retraining service in ShiftLeft-AI that checks for model drift and initiates retraining cycles based on data recency, metric accuracy, or concept drift detection [6].

Adversaries may attempt to evade detection by mimicking normal behavior. The framework's behavioral approach is designed to be resilient to simple evasion tactics, but sophisticated adversaries may develop strategies to maintain normal behavioral patterns while conducting malicious activity. Continuous improvement of the feature set and model architecture is necessary to maintain effectiveness. As noted in XAI research, explanation robustness is a critical concern, particularly in adversarial contexts [7].

The trade-off between sensitivity (detecting all anomalies) and specificity (minimizing false positives) must be carefully managed. Organizations with high tolerance for false positives may prioritize sensitivity, while those with limited security resources may prioritize specificity. The framework supports configurable thresholds to enable organizations to tune the balance based on their needs. This aligns with the XAI principle that no single method is universally optimal; the suitable option depends on the application context and the proficiency of the users [7].

Cultural acceptance of AI-driven security controls is essential for successful deployment. Security teams must trust the framework's decisions, and developers must accept its oversight. The explainability layer is specifically designed to build this trust by providing transparent, understandable rationales for detection decisions. The transparent explanation of problems makes developers more confident in the CI/CD process, which leads to a culture of accountability as well as constant learning [4].

D. Limitations

- **Synthetic Anomaly Injection:** the evaluation relied on synthetic anomaly injection rather than real-world attack data. While synthetic injection is standard practice in security research due to the limited availability of labeled malicious commit datasets, validation on real attack data would strengthen the findings. Future work should explore collaboration with organizations that have experienced actual attacks to validate the framework's effectiveness.
- **Feature Limitations:** the framework focuses exclusively on commit metadata, omitting semantic features that might be relevant for certain attack types. Some attacks may manifest only through code content changes that metadata alone cannot capture. However, research suggests that behavioral patterns (timing, frequency, burst behavior) may be more reliable indicators of compromised accounts than code characteristics. Future work should explore the integration of additional data sources, such as code semantics and pipeline telemetry, similar to the comprehensive data ingestion approach in ShiftLeft-AI [6].
- **Computational Overhead:** the SHAP explanation layer introduces computational overhead. For large-scale CI/CD environments with thousands of daily commits, optimizing the explanation pipeline is essential. In practical deployments, explanations can be generated asynchronously, with immediate alerting followed by detailed explanation delivery. As noted in ShiftLeft-AI research, finding a balance between understandability and responsiveness is an important goal to reach in progress [4].
- **Generalizability:** the framework's performance was validated on open-source repositories, which may not fully represent the characteristics of enterprise codebases. Enterprise environments may have different development patterns, tooling, and security practices. Validation in enterprise environments is necessary to confirm generalizability. However, industry implementations of similar approaches in regulated fintech environments have demonstrated the value of explainable security controls [8].
- **Temporal Dynamics:** the framework models normal behavior as a static distribution, potentially missing temporal patterns such as seasonal variations or trend changes. Future work should explore approaches that incorporate temporal dynamics, such as adaptive models or time-dependent thresholds. Adaptive learning and historical information correlation have been shown to improve decision accuracy, cut down on false positives, and speed up the search for the core cause [1].

E. Future Work

Based on the results and limitations, several promising directions for future research emerge.

- **Real-World Attack Validation:** evaluation on actual malicious commit data from organizations that have experienced attacks would validate the framework's effectiveness against real threats. This may require partnerships with enterprises willing to share de-identified attack data or participation in red team exercises.
- **Multi-Modal Anomaly Detection:** combining commit metadata with other data sources—pipeline telemetry (build durations, test failures, deployment outcomes), code semantics (AST patterns, dependency graphs, vulnerability scores), and organizational context (team structure, project importance, developer seniority)—could improve detection accuracy and enable correlation of anomalies across the CI/CD lifecycle. This aligns with the ShiftLeft-AI approach of combining data streams from many different CI/CD sources [6].
- **Natural Language Explanations:** recent research has demonstrated that SHAP explanations can be translated into natural language, further improving usability for non-technical analysts, including management and legal teams.
- **Interactive Incident Response:** integration with chatbot interfaces could enable analysts to ask questions about anomalies and receive contextual explanations, supporting iterative investigation while maintaining efficiency.
- **Continuous Learning and Adaptation:** implementing online learning mechanisms to adapt to evolving development patterns would improve long-term detection accuracy. Approaches such as concept drift detection and incremental model updates could enable the framework to maintain effectiveness as development practices evolve. As noted in ShiftLeft-AI research, the model retraining service checks for model drift and initiates retraining cycles based on data recency, metric accuracy, or concept drift detection [6].

- **Cross-Repository and Cross-Organization Learning:** future work should explore transfer learning approaches that enable the framework to benefit from patterns learned across repositories or organizations, enabling faster deployment for new projects and improved detection of subtle anomalies.
- **Explainability as a Service:** the explanation layer could be offered as a service to complement existing security tools, integrating SHAP explanations with existing SIEM, SOC, or DevSecOps tools to provide contextual insights for all security alerts.
- **Integration with Remediation:** future work should explore the integration of explanations with automated remediation—for example, triggering additional validation or manager approval when a commit is flagged as anomalous due to timing—completing the loop from detection to remediation. Industry implementations of self-healing infrastructure demonstrate the value of completing the monitor-analyze-plan-execute-verify cycle [10].

CONCLUSION

This paper presented an interpretable anomaly detection framework for secure CI/CD pipelines, integrating deep autoencoders with SHAP explanations. The framework leverages commit metadata exclusively, preserving source code privacy while providing effective anomaly detection. Experimental evaluation on a dataset of over 12,000 commits demonstrates that the framework achieves an F1-score of 0.91, substantially outperforming Isolation Forest (0.78) and One-Class SVM (0.80) baselines.

The key contribution is the integration of SHAP explanations, enabling security analysts to understand why specific commits are flagged as anomalous. This transparency reduces false positive investigation overhead (73% reduction in investigation time), builds trust in automated security controls (91% of participants reported increased trust), and supports rapid incident response. By bridging the gap between high-performance detection and human interpretability, the framework advances the practical deployment of AI-driven security controls in DevSecOps environments.

The framework has significant implications for DevSecOps practice. By enabling effective anomaly detection using only commit metadata, it offers a privacy-preserving approach that addresses key compliance concerns. The SHAP explanation layer enables rapid triage of alerts, reducing false positive investigation overhead and enabling security teams to focus on high-risk anomalies. The continuous learning capabilities ensure that the framework adapts to evolving development patterns, maintaining effectiveness over time.

Future work will focus on real-world attack validation, integration with pipeline telemetry, and the development of natural language explanations. The framework's modular architecture enables incremental adoption, and the phased rollout approach minimizes risk while enabling organizations to realize the benefits of explainable anomaly detection.

The source code and dataset are available under open-source license to enable reproducible research and practical adoption. Organizations interested in deploying the framework can leverage the provided code and documentation to accelerate implementation. The framework represents a significant step toward making AI-driven security controls practical, transparent, and effective in production DevSecOps environments.

REFERENCES

1. "AI-Powered DevOps in Cloud App Modernization: Automating Deployments, Monitoring, and Resilience," Applied AI Summit, Oct. 2025. [Online]. Available: <https://appliedaisummit.org/ai-powered-devops-in-cloud-app-modernization-automating-deployments-monitoring-and-resilience/>
2. A. Mittal, "AI-Powered DevOps in Cloud App Modernization," IEEE Region 5 Consultants' Network, Mar. 2025. doi: 10.13140/RG.2.2.26957.14561.
3. Z. M. Altukhi, S. Pradhan, and N. Aljohani, "A Systematic Literature Review of the Latest Advancements in XAI," Technologies, vol. 13, no. 3, p. 93, 2025. doi: 10.3390/technologies13030093.
4. D. Takkalapally, "ShiftLeft-AI: Proactive CI/CD Performance Anomaly Detection with SHAP-Based Explainability," Int. J. AI Data Sci. ML, 2025.
5. D. Takkalapally, "ShiftLeft-AI: Proactive CI/CD Performance Anomaly Detection," Int. J. AI Data Sci. ML, vol. 2, no. 1, 2025.

6. D. Takkalapally, "ShiftLeft-AI: Machine Learning Framework for Proactive Performance Assurance in CI/CD Pipelines," Int. J. Artif. Intell. Data Sci. Mach. Learn., vol. 5, no. 4, pp. 126-138, 2024. doi: 10.63282/3050-9262.IJAIDSML-V5I4P126.
7. "Explainable Artificial Intelligence (XAI): Concepts, Applications, Challenges, and Future Perspectives," IEEE Xplore, Feb. 2026. [Online]. Available: <https://ieeexplore.ieee.org/abstract/document/11389772>
8. B. Paul, "Secure Software Supply Chain: Risk Prediction at the Speed of Development," presented at QCon San Francisco, Nov. 2025. [Online]. Available: <https://qconsf.com/presentation/nov2025/secure-software-supply-chain-risk-prediction-speed-development>
9. A. Mittal, "AI-Powered DevOps in Cloud App Modernization: Automating Deployments, Monitoring, and Resilience," Applied AI Summit, 2025.
10. H. Mistry, A. Goswami, and C. Mavani, "Automated Anomaly Detection and Response System for Enhancing Cloud Security," Indian Patent 202421051108, 2024.