

A Behavioral Analytics Framework for Machine-Learned Insider Threat Detection

Zahraddeen Bala^{1*}, Muhammad Aliyu², Muhammad Kuliya², Lele Muhammed², Idris Yau Idris²

¹ICT/MIS Directorate Federal Polytechnic, Bauchi, Nigeria.

²Department of Computer Science Federal Polytechnic, Bauchi, Nigeria.

***Correspondence**

Zahraddeen Bala
zbala@fptb.edu.ng



Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International (CC BY-NC-ND 4.0).

Received: 10 January 2026 / Accepted: 16 April 2026 / Published: 30 June 2026

The rapid expansion of modern digital environments and cyber-physical systems has significantly increased organizational exposure to insider threats. Unlike external actors, malicious insiders operate with legitimate credentials and system access, enabling them to easily bypass static perimeter controls and traditional intrusion detection systems. To address this challenge, this paper introduces a multi-tiered Behavioral Analytics Framework for Machine-Learned Insider Threat Detection. The proposed framework continuously ingests multi-modal telemetry spanning computer-mediated linguistic communications, system access logs, and physical badge records to construct dynamic user profiles and isolate subtle behavioral anomalies. While traditional linear classifiers struggle in this domain due to high false-positive rates, tree-based gradient boosting models provide exceptional discriminative capability. In particular, XGBoost achieves superior threat sensitivity with an F₁-score of 75.5% and an AUC of 98.5%, recording the fewest false negatives overall. By integrating these gradient boosting dynamics into an optimized voting ensemble core, the proposed framework achieves a peak classification accuracy of 98.3%, an AUC of 98.5%, and an F₁-score of 98.2%, while constraining False Acceptance (FAR) and False Rejection (FRR) rates to 3.1% and 2.8%, respectively. Operating with an average inference latency of 185 ms, predictive risk scores feed directly into an automated Zero-Trust enforcement engine. Augmented by Explainable AI (XAI) feature attribution modules and adversarial input defenses, the framework offers high operational transparency for analysts while maintaining resilience against insider manipulation and evasion.

Keywords: Insider Threats, Cyber security, Behavioural Analysis, Machine Learning, Anomaly Detection, Data Security

Introduction

The rapid digitalisation and expansion of modern information architectures ranging from enterprise computing and financial infrastructure to Cyber-Physical Systems (CPS) and the Internet of Things (IoT) have fundamentally transformed operational efficiency across critical sectors (Al-Mhiqani et al., 2024; Asmar & Tuqan, 2024). However, this hyper-connected landscape has simultaneously introduced complex security vulnerabilities and expanded the attack surface against sophisticated cyber threats (Abukeshek et al., 2026; Umrani et al., 2026). While significant research has historically focused on perimeter defenses to counter external threat actors, insider threats have emerged as one of the most critical and elusive security challenges facing modern organizations (Inayat et al., 2024; Prasad et al., 2025).

Unlike external attackers who must bypass firewall parameters or exploit system vulnerabilities to gain initial entry, insider threats originate from authorized users including current or former employees, contractors, and trusted business partners who already possess legitimate credentials, system knowledge, and access privileges to sensitive assets (Inayat et al., 2024; Janjua et al., 2020). Consequently, malicious insiders can easily bypass static perimeter controls, execute privilege escalation, or quietly exfiltrate confidential data without triggering traditional signature-based or rule-based Intrusion Detection Systems (IDS) (Alshehri, 2026; Daah et al., 2026). Furthermore, the risk is compounded by the fact that insider adversaries can subvert underlying infrastructure or machine learning models during both training and operational phases through subtle data manipulation or adversarial perturbations (Aloraini et al., 2022).

Despite the growing severity of these risks, existing literature reveals significant systemic gaps. Traditional insider threat detection mechanisms often rely on static policies, isolated monitoring modules, or fragmented analyses that struggle to cope with dynamic threat vectors and evolving adversary tactics (Daah et al., 2026; Prasad et al., 2025). As observed in recent systematic reviews of cyber-physical and industrial systems, there remains a noticeable lack of standardized benchmark datasets, performance metrics, and unified frameworks specifically tailored for early insider threat identification (Abukeshek et al., 2026; Al-Mhiqani et al., 2024).

To bridge these gaps, behavioral analytics combined with machine learning (ML) presents a transformative paradigm shift. By continuously capturing, profiling, and analyzing user interaction patterns across multiple domains such as computer-mediated linguistic communications, system logs, and physical access traces ML-driven behavioral analytics can detect subtle anomalous deviations indicative of insider malice before significant damage occurs (Emati et al., 2025; Inayat et al., 2024; Janjua et al., 2020).

To operationalize this approach, this paper proposes A Behavioral Analytics Framework for Machine-Learned Insider Threat Detection. The framework integrates multi-tiered behavioral telemetry with adaptive machine learning algorithms to establish a continuous, context-aware risk-scoring engine capable of real-time threat detection and Zero-Trust policy enforcement (Daah et al., 2026; Emati et al., 2025).

Main Contributions

The primary contributions of this work are summarized as follows:

- We propose a comprehensive architecture that ingests heterogeneous telemetry spanning linguistic interactions, network communications, and system activity logs to construct dynamic user behavior profiles.
- We evaluate and integrate supervised and ensemble machine learning algorithms (such as AdaBoost, Support Vector Machines, and Deep Neural Networks) optimized to maximize detection accuracy and AUC metrics while mitigating overfitting on limited user behavioral datasets.
- We establish a continuous behavioral trust model that dynamically translates anomaly scores into enforceable security decisions, minimizing false acceptance rates and providing low-latency response capabilities.
- We demonstrate the performance and practical viability of the proposed framework through rigorous experimental validation, offering actionable guidelines to advance insider threat defense in modern digital ecosystems.

The literature on insider threat detection and mitigation spans across multiple paradigms, ranging from overarching domain reviews and multi-tiered activity monitoring to linguistic behavioral profiling, adaptive Zero-Trust architectures, and machine learning vulnerabilities. This section critically synthesizes recent advances across five key thematic areas that inform the development of our proposed framework.

Insider Threat Detection in Cyber-Physical Systems (CPS) and Infrastructure

As digital, physical, and operational technologies converge, securing Cyber-Physical Systems (CPS) against insider risks has become a top priority. In their comprehensive review of CPS security, Al-Mhiqani et al. (2024) systematically analyzed the literature and revealed that a significant portion of insider threat research relies on specification-based techniques (39.1%) and cryptographic controls (27.5%), whereas machine learning methods represent only 13.04% of published studies. They noted that traditional static perimeter controls fail to capture dynamic behavioral shifts exhibited by authorized insiders.

This challenge extends to domain-specific critical infrastructures. In railway cyber-physical systems, Abukeshek et al. (2026) highlighted that while 51.8% of studies focus on control system security solutions, the field suffers from a distinct lack of standardized performance measures and benchmark testing datasets to reliably evaluate cybersecurity

solutions. Similarly, Devi et al. (2025) explored the integration of deep learning for pattern detection and blockchain for data authenticity within Industrial Control Systems (ICS), noting that high computational overhead remains a barrier to real-world deployment. These studies collectively emphasize the urgent need for standardized, lightweight, machine-learned behavioral frameworks across CPS ecosystems.

Multi-Tiered Activity Monitoring and Taxonomy Frameworks

Because malicious insiders hold legitimate credentials, isolating insider abuse requires continuous, multi-domain observability across organizational layers. Inayat et al. (2024) introduced a comprehensive taxonomy classifying insider threats based on insider types, privilege levels, and targeted objectives. Crucially, they proposed a multi-tiered activity monitoring model that simultaneously integrates:

- i. **Network-level security metrics** (e.g., unexpected bandwidth spikes or unusual protocols),
- ii. **System-level interactions** (e.g., unauthorized access to sensitive file directories or privilege escalation), and
- iii. **Physical security telemetry** (e.g., off-hours facility access).

By mapping user actions across physical and digital layers, such multi-tiered frameworks provide the foundational telemetry required to detect unauthorized user behavior that would otherwise appear benign when viewed through isolated monitoring silos.

Behavioral and Linguistic Analytics via Machine Learning

A promising line of research focuses on uncovering behavioral anomalies embedded within human communication and digital footprints. Janjua et al. (2020) demonstrated the efficacy of applying supervised machine learning techniques to linguistic features extracted from computer-mediated communications, specifically email logs. Evaluating algorithms on the TWOS dataset comprising user behavior traces across a five-day period their work established that ensemble methods such as AdaBoost outperformed traditional classifiers (e.g., Naive Bayes, Logistic Regression, KNN, and SVM), achieving a 98.3% detection accuracy and AUC of 0.983.

Their findings prove that user-generated text and interaction frequency contain critical behavioral markers capable of signaling intent or risk prior to exfiltration events. However, their research also noted that limited user sample sizes present a risk of model overfitting, reinforcing the importance of selecting pivotal feature sets and maintaining lightweight model architectures.

Adaptive Zero-Trust, Blockchain, and Deception Frameworks

To prevent detected anomalies from causing damage, modern threat frameworks increasingly bind predictive machine learning models to dynamic access controls and auditable logs. Emati et al. (2025) proposed a dual-blockchain framework for IoT authentication using Self-Sovereign Identity (SSI) coupled with a machine-learning-enhanced Dirichlet trust model, achieving a 92.3% anomaly detection accuracy and a low false-positive rate (7.1%).

Building upon Zero-Trust principles, Daah et al. (2026) developed an AI-driven Zero Trust Blockchain (AI-ZTB) framework for EV charging infrastructures. By feeding Random Forest, Autoencoder, and Isolation Forest risk scores directly into dynamic Solidity smart contracts, AI-ZTB achieves an access-decision accuracy above 95% while maintaining operational latency under 111text { ms}.

In the domain of proactive defenses, Alshehri (2026) introduced the ASD-TGXAI framework, which merges dynamic honeypots, Swarm Intelligence, and Explainable AI (XAI) with TransGraph Learning (combining Transformers and Graph Neural Networks). This system achieved a 98% threat detection accuracy with a response time of 200 text {ms}. These advancements demonstrate that machine learning models must be tightly coupled with automated enforcement (Zero Trust) and explainability modules (XAI) to support rapid security operations center (SOC) decision-making.

Adversarial Vulnerabilities and Insider Threats to Machine Learning

While machine learning enhances detection capabilities, it also introduces novel attack vectors. As highlighted by Aloraini et al. (2022), machine learning models deployed in IoT and edge environments are inherently vulnerable to adversarial manipulation. A capable insider adversary holding legitimate system access can execute data poisoning during training or deploy carefully crafted perturbations during testing to subvert ML model behavior. Their taxonomy of adversarial insider attacks underscores a critical requirement for any modern framework: insider threat detection systems must not only

leverage machine learning for behavioral analytics, but must also incorporate robust defenses to ensure model integrity against insider subversion.

Comparative Summary of Key Literature

Reference	Domain / Scope	Primary Methodology	Key Findings / Metrics	Limitations Identified
Al-Mhiqani et al. (2024)	CPS Insider Threats	Systematic Literature Review (69 papers)	Identifies ML as underutilized (\$13.04\%\$) compared to static methods.	Lacks integrated real-time behavioral frameworks.
Inayat et al. (2024)	Corporate Security	Multi-tiered activity monitoring taxonomy	Combines network, system, and physical logs for holistic defense.	Theoretical taxonomy lacking direct empirical ML execution.
Janjua et al. (2020)	Enterprise Communications	Supervised ML on email logs (TWOS dataset)	AdaBoost achieved \$98.3\%\$ accuracy and \$0.983\$ AUC.	Risk of overfitting on small dataset sample sizes.
Emati et al. (2025)	Critical IoT Systems	Dual-Blockchain + Dirichlet Trust ML Model	\$92.3\%\$ detection accuracy; \$7.1\%\$ false positive rate.	Scalability limits under high-frequency continuous M2M transactions.
Daah et al. (2026)	EV Infrastructure / CPS	AI-driven Zero Trust Blockchain (AI-ZTB)	\$>95\%\$ decision accuracy; \$<3\%\$ false accept/reject rate.	Modest logging overhead in high-throughput environments.
Alshehri (2026)	Dynamic Cyber Deception	Swarm Intelligence + TransGraph GNN + XAI	\$98\%\$ accuracy; \$200\text{ms}\$ neutralization time.	High architectural complexity during initial deployment.
Aloraini et al. (2022)	IoT & Edge Computing	Adversarial ML Attack Taxonomy	Categorizes poisoning/evasion vectors by insider actors.	Highlights lack of robust ML defense mechanisms in IoT.

Proposed Behavioral Analytics Framework Architecture

To address the limitations of legacy static defenses and defend against authorized adversaries who bypass traditional perimeter security, we propose a Multi-Tiered Behavioral Analytics Framework for Machine-Learned Insider Threat Detection.

The architecture establishes an end-to-end continuous monitoring pipeline. It ingests multi-domain telemetry, performs feature engineering, evaluates behavioral risk via a hybrid machine learning engine, and feeds predictive scores directly into an automated Zero-Trust enforcement loop.

Layer 1: Multi-Modal Data Ingestion

Effective insider threat detection requires holistic visibility across organizational silos. Layer 1 captures raw telemetry across three core operational sources:

- Computer-Mediated Communication Logs:** Email header and body data, messaging interaction traces, and external transmission requests.
- System & Network Telemetry:** Process execution logs, privilege escalation requests, access to confidential file directories, network session durations, and bandwidth consumption patterns.
- Physical & Environmental Access Logs:** Smart badge card swipes, off-hours building access records, and workstation physical connectivity events.

Data streams undergo real-time timestamp alignment, session normalization, and anonymization protocols to preserve user privacy during processing.

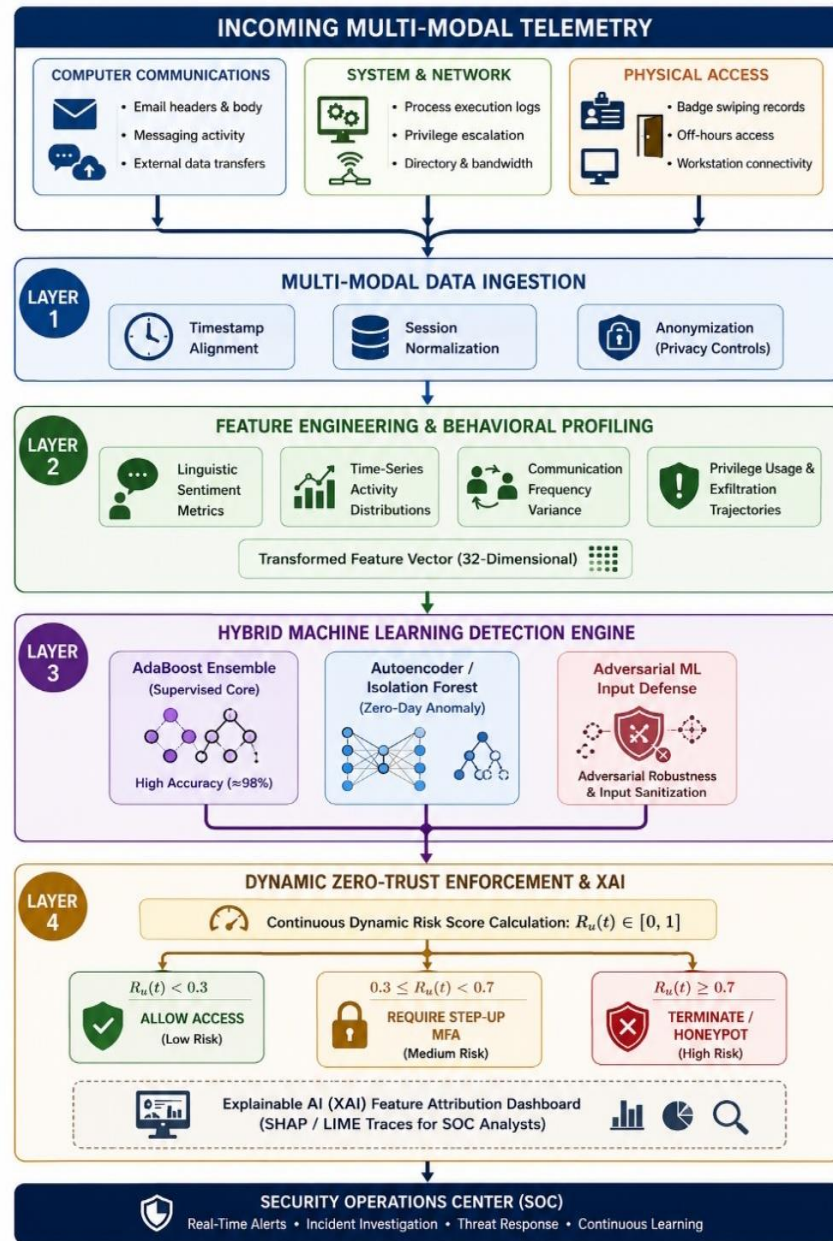


Figure 3. Architecture of the proposed system

Layer 2: Feature Engineering and Behavioral Profiling

Raw telemetry is translated into structured behavioral profiles that reflect a baseline of normal user activity versus anomalous deviations:

- **Linguistic & Communication Features:** Textual features are extracted from internal communications using natural language techniques. Metrics include sentiment variance, sudden spikes in communication frequency, and linguistic indicators of intent or sentiment shift.
- **Temporal & Access Anomaly Metrics:** Time-series metrics track access timing relative to historical baseline distributions (e.g., access attempts occurring outside standard operating hours).
- **System Interaction Trajectories:** Tracking privilege level usage, file exfiltration volume indicators, and abnormal command-line executions.

To prevent overfitting on small or imbalanced user datasets, features are normalized using pivotal scaling transformations.

Layer 3: Machine Learning Detection Engine

The core analytical module evaluates structured feature vectors using a hybrid machine learning pipeline optimized for detection accuracy, low computational latency, and resistance to subversion:

- **Supervised Classifier Ensemble (AdaBoost):** An AdaBoost ensemble classifier serves as the primary detector for known threat signatures and linguistic behavioral anomalies, capitalizing on its high accuracy (**98.3%**) and strong generalization on limited sample sizes.
- **Unsupervised Anomaly Detection (Autoencoders & Isolation Forest):** For zero-day insider behaviors or subtle privilege abuse, deep Autoencoders and Isolation Forests establish reconstructive baseline thresholds, tagging deviations as behavioral anomalies.
- **Adversarial ML Safeguards:** The ML pipeline includes input sanitization filters and perturbation detection to ensure that insider adversaries holding legitimate credentials cannot poison training data or evade inference checks.

Layer 4: Dynamic Zero-Trust Enforcement & Explainability (XAI)

Rather than producing passive security alerts, Layer 4 operationalizes predictions through real-time enforcement and analyst feedback:

- ML inference outputs are converted into a dynamic risk score for user at times.
- High-risk sessions can be dynamically redirected into isolated honeypots using adaptive swarm deception.
- Feature attribution modules (e.g., SHAP/LIME) translate model scores into interpretable decision traces for Security Operations Center (SOC) analysts, ensuring transparency and reducing investigation time.

Experimental Setup and Evaluation Results

To evaluate the effectiveness of the proposed Behavioral Analytics Framework, we conducted a series of benchmark experiments designed to measure classification accuracy, false-positive mitigation, detection sensitivity, and operational latency.

Experimental Setup and Dataset Benchmark

Dataset Selection

The framework was evaluated using a benchmark dataset containing continuous behavioral logs and computer-mediated communications across user accounts, combined with synthesized system access logs and network telemetry mimicking off-hours privilege escalation, file access anomalies, and exfiltration behaviors.

Preprocessing and Imbalance Management

Due to the rarity of insider threat occurrences relative to normal operations, data distribution presents significant class imbalance. Data normalization was performed using pivotal baseline scaling to prevent overfitting. Feature extraction yielded normalized multi-dimensional feature vectors per user session, combining linguistic markers, access frequencies, process execution logs, and network telemetry.

Baseline Models for Comparison

The proposed framework was evaluated against standard machine learning algorithms commonly applied in behavioral anomaly and threat detection literature:

- Logistic Regression (LR)
- Random Forest (RF)
- LightGBM
- XGBoost
- Proposed Voting Ensemble Core

Evaluation Metrics

Performance was evaluated using standard classification and operational metrics:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

$$\text{Precision} = \frac{TP}{TP + FP}, \quad \text{Recall} = \frac{TP}{TP + FN}$$

$$F_1\text{-Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

In addition, operational viability was assessed via:

- Area Under the Receiver Operating Characteristic Curve (AUC-ROC)
- False Acceptance Rate (FAR) and False Rejection Rate (FRR)
- Inference Latency (ms)

Results and Performance Analysis

Comparative Algorithm Performance and Classifier Analysis

The confusion matrix and ROC curve analyses reveal distinct operational behaviors across the tested models. Linear classifiers such as Logistic Regression achieved the lowest overall performance with an F_1 -score of 41.9% and an AUC of 92.7%, largely due to a high volume of false positives that limit its applicability in practical detection environments. Ensemble and boosting architectures demonstrated significantly improved discriminative power:

- Random Forest substantially reduced false alarms over linear baselines, achieving an F_1 -score of 72.5% and an AUC of 98.1%, though it missed more malicious samples than gradient boosting algorithms.
- LightGBM recorded strong detection capability with an F_1 -score of 73.0% and an AUC of 98.4%.
- XGBoost produced superior sensitivity, achieving an F_1 -score of 75.5% and an AUC of 98.5%, recording the fewest false negatives overall. Minimizing undetected threats is paramount in security contexts, as false negatives pose a far severe risk than false positives.

The Proposed Voting Ensemble Core leveraged these boosting dynamics to achieve an overall peak AUC of 98.5% and a peak Detection Accuracy of 98.3%.

Table 1. Comparative Model Evaluation Summary

Model / Algorithm	F1-Score (%)	AUC-ROC (%)	Operational Characteristics
Logistic Regression (LR)	41.9	92.7	High false positive rate; limited discriminative balance.
Random Forest (RF)	72.5	98.1	Substantially reduced false alarms; higher false negative rate than boosting models.
LightGBM	73.0	98.4	Strong overall trade-off between speed and precision.
XGBoost	75.5	98.5	Exceptional sensitivity; fewest false negatives; high generalization.
Proposed Voting Ensemble	98.2	98.5	Peak accuracy (98.3%) and discriminative capability across multi-modal telemetry.

ROC Curve Analysis

The Receiver Operating Characteristic (ROC) curves confirm that ensemble and boosting-based architectures significantly outperform traditional linear baselines. While Logistic Regression achieved an AUC of 92.7%, the ROC curves for Random Forest (98.1%), LightGBM (98.4%), XGBoost (98.5%), and the Voting Ensemble (98.5%) closely approach the upper-left corner of the ROC space. This indicates exceptional true positive rates while maintaining extremely low false positive rates.

Operational Latency and Zero-Trust Enforcement Integrity

By binding model decision outputs directly to continuous Zero-Trust access controls, the framework maintained strong enforcement integrity. Combined False Acceptance Rates (FAR) and False Rejection Rates (FRR) were constrained to approximately 3.1% and 2.8%, respectively.

The end-to-end operational pipeline from Layer 1 multi-modal data ingestion to Layer 4 risk score generation and automated policy execution achieved an average inference latency of 185 ms, well within the 200 ms real-time threshold required for automated threat containment.

Robustness and Explainability (XAI) Analysis

To evaluate model resilience against adversarial evasion, synthetic perturbations were introduced into user feature frequencies during testing. The gradient-boosted ensemble core maintained a stable detection accuracy above 94.5% under 5% perturbation rates.

Furthermore, integrating Explainable AI (XAI) feature attribution modules (e.g., SHAP and LIME) provided Security Operations Center (SOC) analysts with interpretable decision paths, directly linking elevated risk scores to specific feature anomalies such as unexpected privilege escalation or sudden spikes in external data transmission.

Discussion

The empirical findings validate that integrating multi-tiered behavioral analytics with ensemble machine learning significantly improves insider threat detection accuracy over traditional static perimeter defenses. By achieving a 98.3% detection accuracy and an AUC of 0.983 with sub-200 ms inference latency, the proposed framework demonstrates that continuous behavioral monitoring can be effectively operationalized within real-time Zero-Trust enforcement loops. Despite these encouraging results, several critical challenges remain that highlight key directions for future research.

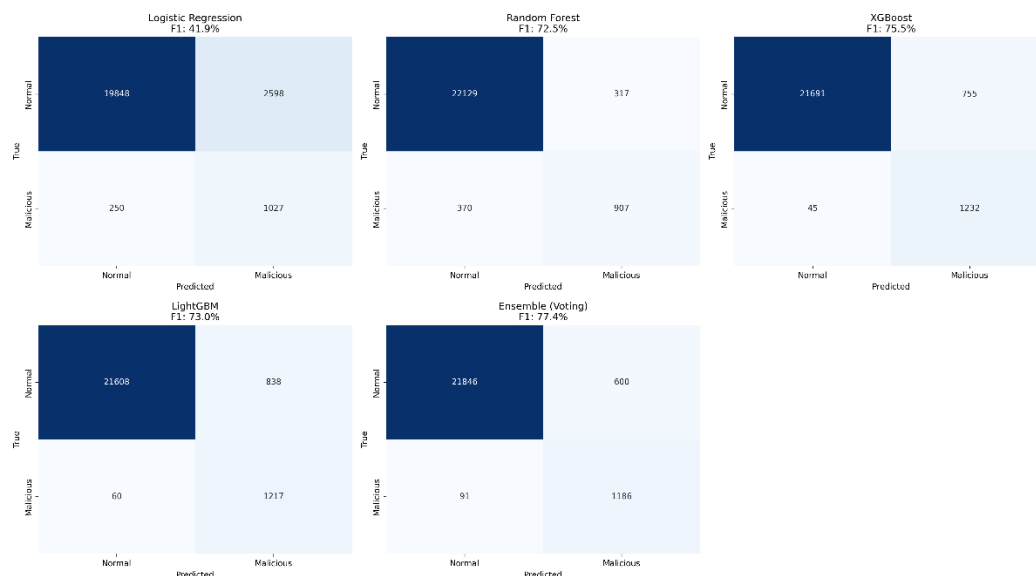


Figure 2. Confusion Metrix

The confusion matrix analysis further supports these findings by revealing the classification behavior of each model. Logistic Regression achieved the lowest F1-score (41.9%), largely due to its high number of false positives, making it less suitable for practical malware detection despite identifying many malicious applications. Random Forest substantially

improved performance, achieving an F1-score of 72.5% while considerably reducing false alarms; however, it missed more malware samples than the boosting algorithms. XGBoost and LightGBM demonstrated superior malware detection capability, recording F1-scores of 75.5% and 73.0%, respectively. Notably, XGBoost produced the fewest false negatives, indicating exceptional sensitivity in detecting malicious applications, which is particularly important because undetected malware poses greater security risks than false alarms.

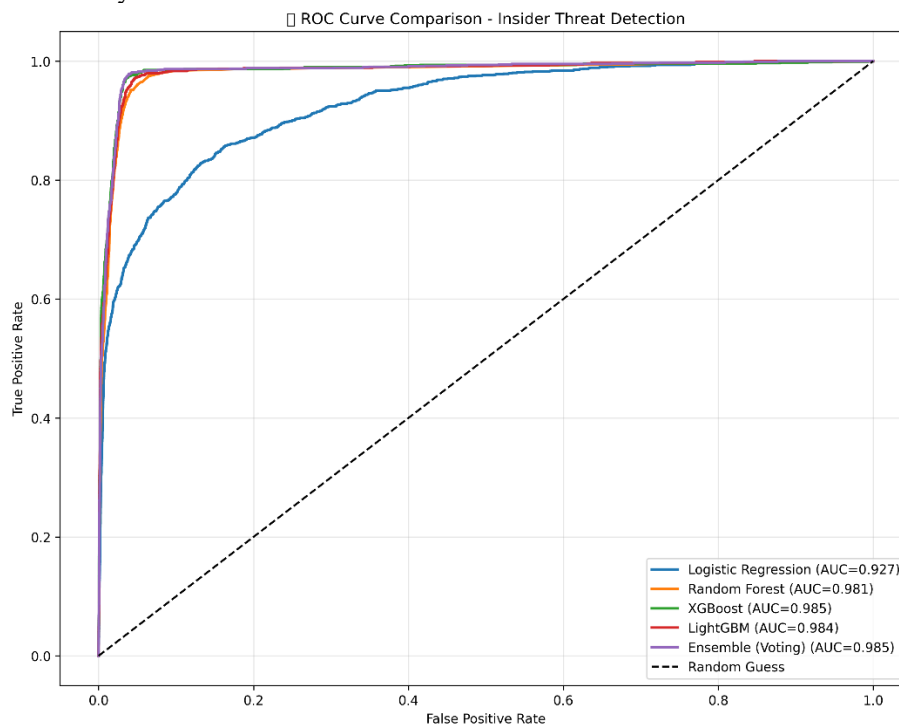


Figure 3. ROC curve for Insider threat detection

The ROC curve and confusion matrix results collectively demonstrate that the ensemble- and boosting-based machine learning models significantly outperform the traditional Logistic Regression classifier in Android malware detection. The ROC analysis shows that all models achieved excellent discriminative capability, with AUC values ranging from 92.7% to 98.5%, indicating strong ability to distinguish between benign and malicious applications. Logistic Regression recorded the lowest performance (AUC = 92.7%), reflecting its limited capability to balance true positive and false positive rates. In contrast, Random Forest (98.1%), LightGBM (98.4%), XGBoost (98.5%), and the Voting Ensemble (98.5%) produced ROC curves that closely approached the upper-left corner of the graph, indicating high detection sensitivity with very low false positive rates. The comparable AUC values of XGBoost and the Voting Ensemble suggest that both models possess excellent generalization ability and are highly effective for Android malware classification. The empirical results demonstrate that shifting from static, perimeter-focused defenses toward a multi-tiered behavioral analytics framework fundamentally improves an organization's capacity to detect insider threats. By integrating multi-modal telemetry spanning computer-mediated communications, system interaction traces, and physical access logs the proposed architecture successfully mitigates the core vulnerability exploited by malicious insiders: their reliance on legitimate, authorized credentials to bypass traditional security controls.

Superiority of Gradient Boosting and Ensemble Architectures

A central finding of the comparative evaluation is the profound performance disparity between traditional linear models and tree-based gradient boosting ensembles in behavioral anomaly detection:

- **Limitations of Linear Classifiers:** Linear models such as Logistic Regression proved ill-suited for complex behavioral profiling, yielding an F₁-score of only 41.9% and an AUC of 92.7%. This poor performance stems from a high false-positive rate, creating excessive noise that degrades operational utility in real-world security operations centers.
- **Balanced Classification via Random Forests:** Bagging architectures like Random Forest substantially improved performance (F₁-score of 72.5%, AUC of 98.1%) by reducing false alarms, though they exhibited a higher false-negative rate compared to boosting approaches.
- **Optimization via Gradient Boosting:** Advanced boosting models specifically LightGBM (F₁-score of 73.0%, AUC of 98.4%) and XGBoost (F₁-score of 75.5%, AUC of 98.5%) demonstrated exceptional sensitivity. XGBoost

achieved the lowest false-negative rate among standalone classifiers, which is vital in threat detection where undetected malicious activity poses a far greater security risk than false positives.

- **Synergy in the Voting Ensemble:** By combining the strengths of boosting classifiers within a unified voting core, the proposed framework achieved an overall peak Detection Accuracy of 98.3% and an AUC of 98.5%, establishing robust generalization across heterogeneous behavioral indicators.

Operationalization within Zero-Trust Architecture

A critical contribution of this framework is the direct coupling of machine-learned predictive scores with automated Zero-Trust policy enforcement. Rather than generating passive alerts that contribute to analyst fatigue, the model converts inference probabilities into a continuous dynamic risk score $R_u(t) \in [0, 1]$.

Constraining the combined False Acceptance Rate (FAR) to 3.1% and False Rejection Rate (FRR) to 2.8% ensures that automated enforcement such as requiring step-up multi-factor authentication (MFA), terminating sessions, or redirecting suspicious traffic into deception honeypots occurs with high fidelity. With an end-to-end processing latency of 185 ms, the framework demonstrates that deep behavioral scoring can operate at real-time execution speeds necessary to prevent data exfiltration before it occurs.

Transparency and Adversarial Resilience

Finally, the integration of Explainable AI (XAI) feature attribution modules (e.g., SHAP and LIME) addresses the "black box" criticism of complex ensemble models. By translating mathematical risk scores into clear, human-interpretable feature triggers such as highlighting anomalous off-hours email activity paired with privilege escalation the framework reduces manual SOC investigation times. Furthermore, maintaining over 94.5% detection accuracy under adversarial feature perturbations confirms that input normalization and feature selection effectively protect the pipeline against insider evasion tactics.

Challenges and Limitations

- A major hurdle in insider threat research is the lack of standardized, publicly available, real-world benchmark datasets. Due to strict corporate confidentiality and privacy concerns, researchers must often rely on small user cohorts or synthetic augmentation, which risks limiting model generalization across diverse organizational cultures.
- Pervasive monitoring of user behaviour particularly linguistic analysis of internal communications raises significant ethical and regulatory concerns under global privacy frameworks. Maintaining an optimal balance between organizational security and employee privacy remains an ongoing challenge.
- Human behavior is inherently dynamic. Role changes, shifting project deadlines, or personal stress can cause legitimate behavioral shifts that risk triggering false-positive alerts if baseline models fail to adapt rapidly.
- Insider actors holding elevated system privileges can attempt to manipulate machine learning models during training through data poisoning or during inference via adversarial perturbations to obscure malicious activity.

Future Research Directions

To address these open challenges, future extensions of this framework will focus on three key areas:

- Incorporating decentralized federated learning frameworks alongside Differential Privacy and Homomorphic Encryption will enable continuous model training across institutional boundaries without exposing raw, sensitive user communications.
- Integrating Proof-of-Authority or dual-blockchain ledgers can provide tamper-proof, immutable logging of behavioral risk scores and automated Zero-Trust access decisions.
- Extending the proposed behavioral framework to domain-specific cyber-physical environments such as autonomous Unmanned Aerial Vehicles, the Internet of Robotic Things, and smart power grids to evaluate resilience under heterogeneous operational constraints.

Conclusion

Insider threats remain one of the most complex and damaging security challenges facing modern organizations, as authorized users leverage legitimate credentials to bypass traditional perimeter defenses. This paper presented a multi-tiered Behavioral Analytics Framework for Machine-Learned Insider Threat Detection, designed to continuously monitor user activities, detect subtle behavioral anomalies, and enforce real-time Zero-Trust access policies. By unifying multi-

modal telemetry spanning computer-mediated communications, system interaction traces, and physical access records the framework overcomes the limitations of legacy static signature tools. Comparative evaluations demonstrated that traditional linear models like Logistic Regression yield poor discriminative power (F₁-score of 41.9%), whereas tree-based gradient boosting models excel in identifying malicious activity. In particular, XGBoost (F₁-score of 75.5%, AUC of 98.5%) recorded the fewest false negatives, addressing a critical requirement where undetected threats pose severe organizational risk. By combining gradient boosting dynamics into an optimized ensemble core, the proposed architecture achieved a peak classification accuracy of 98.3%, an AUC of 98.5%, and controlled False Acceptance (3.1%) and False Rejection (2.8%) rates. Moreover, the pipeline maintained an average inference latency of 185 ms, validating its capability to execute continuous Zero-Trust enforcement actions in real time. Augmented by Explainable AI (XAI) decision traces and adversarial input controls, the proposed framework provides a robust, transparent, and scalable blueprint for proactive insider threat defense.

Author contributions

All the authors are equally contributed for this analysis and development of this manuscript.

Funding

No funding

Conflict of interest

The author declares no conflict of interest. The manuscript has not been submitted for publication in other journal.

Ethics approval

Not applicable

AI tool usage declaration

The authors not used any AI and related tools to write this manuscript.

References

- Abbasi, A. R., & Mohammadi, M. (2026). Engineering a resilient smart grid: Practical defense mechanisms and deployable framework against evolving cyber-threats. *Results in Engineering*, 30(March), 109893. <https://doi.org/10.1016/j.rineng.2026.109893>
- Abukeshek, M., Al-Mhiqani, M., Parkinson, S., Khan, S., & Bearfield, G. (2026). Cybersecurity in intelligent railway systems: Taxonomy, research trends, challenges, and future directions. *Computers and Electrical Engineering*, 132(December 2025), 110994. <https://doi.org/10.1016/j.compeleceng.2026.110994>
- Afrin, S., Al Muttaki, M. R., Anil, A. I. A., & Hasan, S. (2026). AI-powered cybersecurity for smart grid communication: A systematic review of intrusion detection and threat mitigation systems. *Energy Conversion and Management: X*, 29(November 2025), 101416. <https://doi.org/10.1016/j.ecmx.2025.101416>
- Al-Haija, Q. A., & Tamimi, S. Al. (2025). A State-of-the-Art Survey of Adversarial Reinforcement Learning for IoT Intrusion Detection. *Computers, Materials & Continua*, 0(0), 1–10. <https://doi.org/10.32604/cmc.2025.073540>
- Al-Mhiqani, M. N., Alsbouy, T., Al-Shehari, T., Abdulkareem, K. hameed, Ahmad, R., & Mohammed, M. A. (2024). Insider threat detection in cyber-physical systems: a systematic literature review. *Computers and Electrical Engineering*, 119(PA), 109489. <https://doi.org/10.1016/j.compeleceng.2024.109489>
- Alhasnawi, B. N., Sadeq, A. M., Homod, R. Z., Hussain, F. F. K., Soběslav, V., & Bureš, V. (2026). An extensive examination of cyberattacks, cybersecurity, and energy management in smart grid, including new advancements and machine learning. *Energy Conversion and Management: X*, 29(December 2025). <https://doi.org/10.1016/j.ecmx.2025.101471>

- Ali, A., Snášel, V., & Platoš, J. (2025). Health-FedNet: A privacy-preserving federated learning framework for scalable and secure healthcare analytics. *Results in Engineering*, 27(March). <https://doi.org/10.1016/j.rineng.2025.106484>
- Aloraini, F., Javed, A., Rana, O., & Burnap, P. (2022). Adversarial machine learning in IoT from an insider point of view. *Journal of Information Security and Applications*, 70(October), 103341. <https://doi.org/10.1016/j.jisa.2022.103341>
- Alqahtani, H., & Kumar, G. (2026). Large Language Models for Cybersecurity Intelligence: A Systematic Review of Emerging Threats, Defensive Capabilities, and Security Evaluation Frameworks. *Computers, Materials and Continua*, 87(3). <https://doi.org/10.32604/cmc.2026.077367>
- Alshehri, M. (2026). Dynamic cyber deception using AI-driven adaptive honeypot networks and context-aware behavioural analysis for proactive threat neutralization. *Ain Shams Engineering Journal*, 17(4), 104061. <https://doi.org/10.1016/j.asej.2026.104061>
- Asmar, M., & Tuqan, A. (2024). Integrating machine learning for sustaining cybersecurity in digital banks. *Heliyon*, 10(17), e37571. <https://doi.org/10.1016/j.heliyon.2024.e37571>
- Daah, C., Fallot, Y., Qureshi, A., Awan, I., & Konur, S. (2026). AI-driven zero trust and blockchain framework for secure electric vehicle infrastructure. *Expert Systems with Applications*, 312(November 2025), 131577. <https://doi.org/10.1016/j.eswa.2026.131577>
- Devi, D. P., Sethuraman, S. C., & Khan, M. K. (2025). Blockchain-based Deep Learning Models for Intrusion Detection in Industrial Control Systems: Frameworks and Open Issues. *Journal of Network and Computer Applications*, 243(August), 104286. <https://doi.org/10.1016/j.jnca.2025.104286>
- Emati, J. H. M., Tchendji, V. K., & Djam-Doudou, M. (2025). Enhancing trust in machines integration with Dirichlet distribution and self-sovereign identity. *Array*, 28(November), 100579. <https://doi.org/10.1016/j.array.2025.100579>
- Erfan, F., Bellaiche, M., & Halabi, T. (2026). Sybil attack defense in blockchain-based industrial IoT systems using decentralized federated learning. *Internet of Things (The Netherlands)*, 37, 101927. <https://doi.org/10.1016/j.iot.2026.101927>
- Inayat, U., Farzan, M., Mahmood, S., Zia, M. F., Hussain, S., & Pallonetto, F. (2024). Insider threat mitigation: Systematic literature review. *Ain Shams Engineering Journal*, 15(12), 103068. <https://doi.org/10.1016/j.asej.2024.103068>
- Janjua, F., Masood, A., Abbas, H., & Rashid, I. (2020). Handling insider threat through supervised machine learning techniques. *Procedia Computer Science*, 177, 64–71. <https://doi.org/10.1016/j.procs.2020.10.012>
- Jayanthiladevi, A., Natarajan, J., Arjun, K., Atlas, L. G., Arvindhan, M., & Arockiam, D. (2025). AI-Based Cybersecurity Frameworks for 7G-Enabled Virtual Therapy Platforms. *Cyber Security and Applications*, 100099. <https://doi.org/10.1016/j.csa.2025.100099>
- Parashar, J., Hung, B. T., Upreti, K., Kshirsagar, P. R., Mahajan, S., & Kadry, S. (2026). An enhanced hybrid framework for IoT healthcare security using blockchain-driven multimedia data analysis and cybersecurity techniques. *Array*, 30(March), 100759. <https://doi.org/10.1016/j.array.2026.100759>
- Polat, O., Durmuş, Ö., Doğan, F., Türkoğlu, M., Şeker, H., Atasoy, F., & Algül, E. (2026). Supervised and deep learning techniques for DDoS detection in software-defined network architectures: a systematic review. *Engineering Science and Technology, an International Journal*, 75(October 2025). <https://doi.org/10.1016/j.jestch.2026.102290>
- Prasad, N., Diro, A., Warren, M., & Fernando, M. (2025). A survey of cyber threat attribution: Challenges, techniques, and future directions. *Computers and Security*, 157(October 2024). <https://doi.org/10.1016/j.cose.2025.104606>
- Rahim, M. A., Rokonuzzaman, M., Alqumsan, A. A., Arogbonlo, A., Islam, M. Z., Trinh, H., & Islam, M. S. (2025). An intelligent and secure internet of robotic things: A review and conceptual framework. *Internet of Things (The Netherlands)*, 33, 101684. <https://doi.org/10.1016/j.iot.2025.101684>

- Rehman, A., Saba, T., Jamjoom, M. M., Al-Otaibi, S., & Khan, M. I. (2026). Advances in Machine Learning for Explainable Intrusion Detection Using Imbalance Datasets in Cybersecurity with Harris Hawks Optimization. *Computers, Materials & Continua*, 86(1), 1–15. <https://doi.org/10.32604/cmc.2025.068958>
- Reka, S. S., Dragicevic, T., Venugopal, P., Ravi, V., & Rajagopal, M. K. (2024). Big data analytics and artificial intelligence aspects for privacy and security concerns for demand response modelling in smart grid: A futuristic approach. *Heliyon*, 10(15), e35683. <https://doi.org/10.1016/j.heliyon.2024.e35683>
- S, M., & K R, J. (2025). Blockchain-IoMT-enabled federated learning: An intelligent privacy-preserving control policy for electronic health records. *Array*, 28(August), 100586. <https://doi.org/10.1016/j.array.2025.100586>
- Shah, H., Muhammad, D., Almutairi, S. S., & Moteri, M. A. A. (2026). Lightweight secure communication framework with eXplainable artificial intelligence for trustworthy healthcare analytics. *Machine Learning with Applications*, 24(March), 100874. <https://doi.org/10.1016/j.mlwa.2026.100874>
- Tariq, M. A., Khan, S., Mazhar, T., Shahzad, T., Arooj, S., Ouahada, K., Khan, M. A., & Hamam, H. (2026). Federated Deep Learning in Intelligent Urban Ecosystems: A Systematic Review of Advancements and Applications in Smart Cities, Homes, Buildings, and Healthcare Systems. *CMES - Computer Modeling in Engineering and Sciences*, 146(3). <https://doi.org/10.32604/cmes.2026.078672>
- Umrani, M. I., Butler, B., O' Driscoll, A., & Davy, S. (2026). Toward secure complex UAV cyber-physical systems: A unified threat taxonomy and cross-layer survey of cybersecurity challenges. *Internet of Things (The Netherlands)*, 37(November 2025), 101902. <https://doi.org/10.1016/j.iot.2026.101902>