

Application of Machine Learning in Identifying Anomalous User Behavior Patterns for Insider Threat Forensics

Usman Ayobami
Department of Computer
Science
Western Michigan University
USA

Abstract: The rising complexity of digital infrastructures has significantly increased organizations' vulnerability to insider threats—malicious or negligent actions originating from within the organization. Traditional security mechanisms often fail to detect such threats, as insiders typically operate with legitimate credentials, making their malicious intent difficult to distinguish from routine activity. As cybersecurity paradigms shift toward proactive threat identification, machine learning (ML) has emerged as a pivotal tool in uncovering subtle, context-dependent behavioral anomalies indicative of insider compromise. This paper explores the application of machine learning in identifying anomalous user behavior patterns to enhance forensic capabilities in insider threat detection. It begins by examining the evolving nature of insider threats, including data exfiltration, privilege misuse, and sabotage, and highlights the limitations of signature-based detection systems. The study then categorizes ML techniques into supervised, unsupervised, and hybrid approaches, analyzing their efficacy in behavior profiling, anomaly scoring, and adaptive learning under real-world enterprise constraints. Particular emphasis is placed on unsupervised learning models—such as clustering and autoencoders—which can operate without labeled threat data and dynamically model user activity baselines. The paper also presents a forensic framework that integrates ML-based behavioral analytics with system logs, access records, and contextual metadata, enabling security teams to trace the timeline, method, and intent of insider incidents with higher precision. Through case studies and experimental simulations, the proposed ML-based approach is shown to improve detection accuracy, reduce false positives, and support incident response workflows. The paper concludes by offering design recommendations for secure, ethical, and interpretable deployment of ML in insider threat forensics.

Keywords: Insider threat forensics, anomalous user behavior, machine learning, behavioral analytics, unsupervised learning, cybersecurity.

1. INTRODUCTION

1.1 Rise of Insider Threats in Digital Enterprises

As digital transformation intensifies across industries, the complexity and breadth of enterprise networks have grown substantially, creating new vulnerabilities within internal ecosystems. Insider threats—whether driven by malicious intent or negligent behavior—pose a significant and often underestimated risk to organizational security [1]. Unlike external cyberattacks, insider threats originate from users with legitimate access to systems, making them difficult to detect through perimeter-focused defenses.

Incidents of data leaks, privilege misuse, and sabotage have shown a steady increase, particularly in sectors such as finance, healthcare, and critical infrastructure, where high-value data is routinely accessed [2]. Many of these breaches evade detection for weeks or even months, causing not only reputational damage but also regulatory repercussions and substantial financial losses.

Moreover, the advent of remote work and bring-your-own-device (BYOD) policies has expanded the range of access points, often without proportional advances in internal monitoring. Insiders now operate in more flexible, less observable digital environments, complicating conventional methods of surveillance and detection [3].

Understanding and preempting insider threats thus require deeper insights into behavioral patterns within networks—insights that go beyond static rules and involve contextual awareness of user actions, deviations, and emerging intent signals.

1.2 Limitations of Traditional Detection Methods

Traditional cybersecurity measures typically rely on predefined rules, signature databases, and static thresholds to detect threats. While effective against known attack vectors, these techniques struggle to identify novel behaviors or evolving insider strategies that do not trigger predefined alerts [4]. Static rule-based systems often produce excessive false positives or overlook subtle anomalies that may indicate emerging threats.

Furthermore, insider threats are inherently contextual. A user accessing sensitive files at unusual hours or transferring data to an external drive may not breach any explicit policy but could signify suspicious behavior when compared to their historical activity profile [5]. Signature-based intrusion detection systems (IDS) are particularly ineffective in these scenarios, as they are not designed to adapt to individual user baselines or detect deviations that are behaviorally anomalous but technically permissible.

Auditing and manual log reviews are time-intensive and retrospective in nature. They rarely offer real-time detection and typically require post-incident analysis, which limits opportunities for proactive response [6]. Additionally, fragmented data across systems and departments further hampers coherent threat assessment.

To address these gaps, organizations are increasingly turning to adaptive, learning-based solutions that can analyze vast behavioral datasets, model user baselines, and detect outliers with greater speed and precision.

1.3 Role of Machine Learning in Behavioral Analysis

Machine learning (ML) offers a promising paradigm shift in insider threat detection by enabling dynamic analysis of user behavior across time and systems. Unlike static rules, ML algorithms can learn baseline activity patterns for individual users and departments, allowing them to detect deviations indicative of insider threat activity [7].

Through techniques such as unsupervised clustering, anomaly detection, and behavioral profiling, ML systems can flag suspicious activities in real time. These include sudden spikes in data access, atypical login locations, or access to resources outside a user's role—without requiring prior knowledge of specific attack signatures [8].

1.4 Article Roadmap and Objectives

This article explores the application of machine learning in identifying anomalous user behavior patterns within enterprise environments to detect insider threats. It begins by reviewing key types of insider threats and the data sources most relevant to behavioral modeling. It then outlines the core machine learning techniques suitable for behavioral analytics, including supervised, unsupervised, and hybrid approaches.

Subsequent sections illustrate real-world case applications, discuss integration challenges in enterprise security infrastructures, and provide implementation guidelines. The article concludes with future research directions and policy considerations to institutionalize machine learning frameworks for forensic and real-time insider threat mitigation.

2. UNDERSTANDING INSIDER THREATS AND BEHAVIORAL PATTERNS

2.1 Definition and Taxonomy of Insider Threats

Insider threats are security risks that originate from individuals within an organization who have authorized access to systems and data. These individuals may intentionally or unintentionally compromise the confidentiality, integrity, or availability of sensitive information [5]. The taxonomy of insider threats typically

falls into three primary categories: malicious insiders, negligent insiders, and infiltrators.

Malicious insiders are employees or contractors who exploit their access privileges to harm the organization for personal, financial, or ideological reasons. These actors often operate with a high degree of planning and covertness, targeting valuable data or critical systems [6]. Their actions may include intellectual property theft, data leaks, or system sabotage.

Negligent insiders differ in that they do not act with harmful intent but inadvertently cause security breaches due to carelessness or poor security practices. Examples include mishandling sensitive files, using weak passwords, or falling victim to phishing attacks that provide external actors with access to internal systems [7].

A third and increasingly relevant category is **infiltrators**, who are external actors posing as insiders. This includes compromised accounts or adversaries hired under false pretenses to gain insider access. These actors benefit from the legitimacy of their credentials, making them harder to detect than traditional attackers [8].

Understanding these distinctions is crucial for designing machine learning models that effectively differentiate between harmless anomalies and genuine threats. Behavioral baselines must be adapted based on the user's role, intent profile, and access level, to avoid misclassification and ensure effective monitoring [9].

2.2 Common Behavioral Indicators of Insider Risk

Behavioral indicators of insider risk often manifest subtly and over time, requiring continuous monitoring and context-aware analysis. One of the most consistent signs is access misuse—when a user accesses files or systems outside the scope of their job role or during unusual hours [10]. This could involve downloading large volumes of data, opening sensitive documents repeatedly, or bypassing standard operational processes.

Another common signal is privilege escalation, which involves users requesting or gaining unauthorized higher-level access. While sometimes justified for task completion, unmonitored escalations may indicate intent to exfiltrate data or modify audit trails [11]. Machine learning systems can model expected privilege request patterns and flag deviations that exceed historical norms.

Temporal anomalies also raise red flags. These include logins during off-hours, sudden bursts in system activity, or erratic use of applications not typically associated with a user's role [12]. For instance, a finance employee initiating multiple secure shell (SSH) sessions or accessing a software repository may indicate risk.

Geolocation inconsistencies, such as simultaneous logins from different countries, can highlight compromised credentials or

malicious usage. Similarly, use of anonymizing technologies like VPNs or TOR nodes within an enterprise network may serve as early indicators of concealment behavior [13].

Behavioral drift, which refers to a gradual but significant change in digital behavior over time, is often detectable via trend analysis. This includes changing login times, new application usage, or shifts in communication channels. These patterns are especially significant when paired with non-technical indicators, such as recent HR issues or upcoming terminations [14].

Machine learning algorithms can aggregate and analyze these indicators over time to differentiate benign anomalies from truly suspicious activity, providing both scalability and accuracy in threat detection.

2.3 Challenges in Detecting Insider Threats

Detecting insider threats presents several challenges due to the legitimacy of credentials, the context-dependent nature of behavior, and the high rate of false positives produced by traditional systems. Unlike external attacks, insiders do not need to bypass firewalls or intrusion detection systems—they already have legitimate access, making their actions difficult to distinguish from routine operations [15].

One core difficulty is contextual ambiguity. The same action—such as accessing a confidential file—can be either completely acceptable or deeply suspicious, depending on the user's role, timing, and prior behavior. Without contextual interpretation, security systems are prone to false alarms or, conversely, may fail to flag dangerous activity [16].

Machine learning models can help by learning individualized behavior baselines; however, these models require extensive, clean training data. Behavioral variance across departments, job functions, and user personas can complicate efforts to generalize features or thresholds for anomaly detection [17]. Overfitting to historical patterns may also limit the model's adaptability to evolving behaviors or new employee roles.

Another complication is the dynamic evolution of user roles and workflows. As job responsibilities shift, so too do access patterns and system usage. Without mechanisms to update baselines dynamically, detection systems risk becoming outdated or misleading [18].

Data fragmentation is also a barrier. Behavioral signals are spread across access logs, email records, endpoint monitors, and HR databases. Integrating these into a unified model poses technical and organizational challenges, particularly in environments with siloed or legacy systems [19].

Finally, the interpretability of machine learning outputs is essential for adoption. Security teams must be able to understand and justify alerts. Complex or opaque models can hinder this process, reducing confidence in automation and slowing incident response.

To be effective, insider threat detection frameworks must balance sensitivity with contextual understanding and present results in a transparent and actionable format.

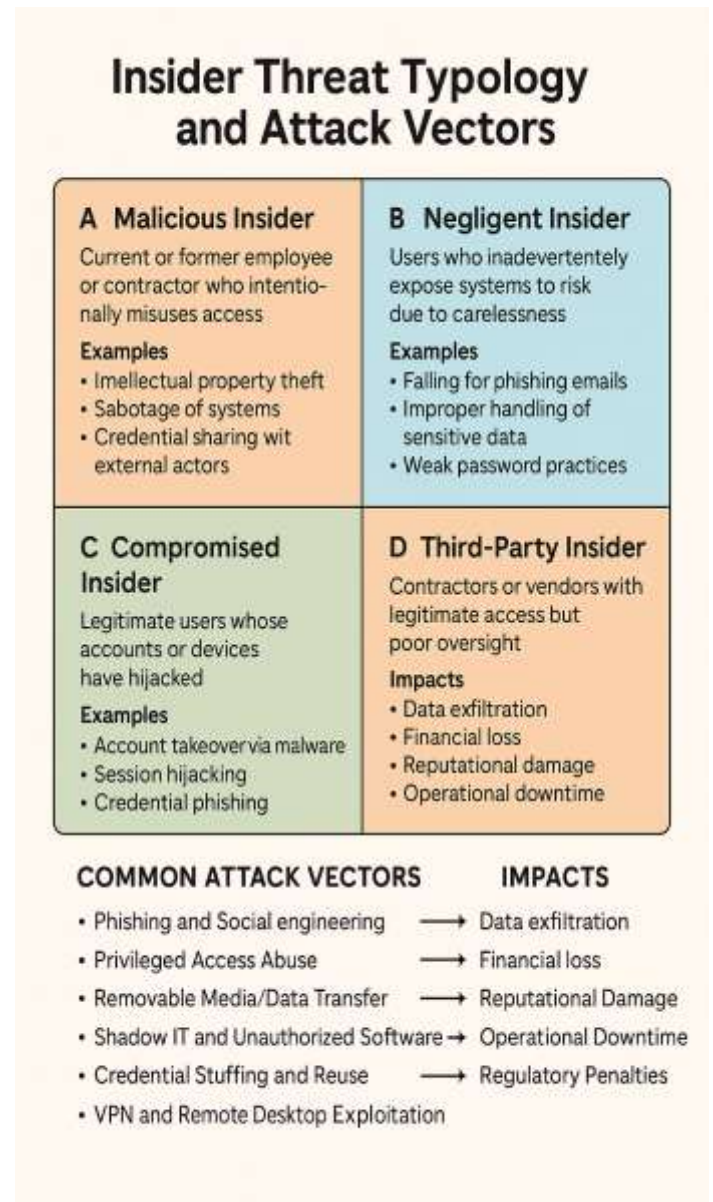


Figure 1: Insider Threat Typology and Attack Vectors

3. FOUNDATIONS OF MACHINE LEARNING FOR ANOMALY DETECTION

3.1 Overview of Supervised, Unsupervised, and Semi-Supervised Learning

Machine learning (ML) provides a suite of approaches for modeling user behavior and identifying insider threats. These approaches fall into three main categories: supervised, unsupervised, and semi-supervised learning, each with distinct advantages and trade-offs depending on the availability of labeled data [9].

Supervised learning involves training algorithms on labeled datasets where normal and malicious behavior is explicitly marked. Common algorithms in this category include logistic regression, decision trees, and support vector machines. These models perform well when sufficient high-quality historical examples of insider threats exist. However, insider incidents are rare and diverse, making large labeled datasets difficult to compile, especially in sensitive environments [10].

In contrast, unsupervised learning requires no labeled data and instead identifies anomalies based on deviation from learned norms. Clustering methods like k-means or density-based algorithms, as well as anomaly detection techniques such as Isolation Forests and One-Class SVMs, are widely used to flag outliers in behavioral datasets [11]. This approach is especially useful for detecting unknown or emerging patterns of insider threats.

Semi-supervised learning blends the strengths of both paradigms, using a small amount of labeled data alongside a large pool of unlabeled data. This technique helps enhance generalization without requiring fully annotated datasets. Autoencoders, a form of neural network used for reconstruction tasks, are particularly effective in this space. They learn to reproduce standard behavior and raise alerts when input deviates significantly from the learned pattern [12].

Choosing the right learning paradigm depends on the organization's access to historical incident data, system heterogeneity, and tolerance for false positives. Increasingly, hybrid models that combine multiple learning types are being adopted to improve robustness and detection accuracy in insider threat applications [13].

3.2 Key ML Algorithms Used in Behavior Modeling

Various machine learning algorithms are tailored to detect insider threat behaviors based on how they handle structured and unstructured behavioral data. Among the most commonly deployed are **Random Forests**, **Isolation Forests**, **k-means clustering**, and **autoencoders**, each selected for its performance in anomaly detection tasks and suitability for large-scale log data [14].

Random Forests are ensemble learning algorithms that use a combination of decision trees to perform classification or regression. They are robust to overfitting and perform well on imbalanced datasets, which is often the case in insider threat detection due to the rarity of true malicious behavior [15]. They also offer feature importance metrics, aiding interpretability.

Isolation Forests are explicitly designed for anomaly detection. Rather than profiling normal behavior, they isolate anomalies by randomly partitioning the dataset. Instances that require fewer partitions to isolate are considered outliers. This method is efficient, scalable, and suitable for high-dimensional user activity data [16].

K-means clustering groups similar behavior profiles together based on distance metrics. Users or actions that fall far from centroids are treated as potential anomalies. Though simple, k-means requires careful tuning of cluster numbers and assumes spherical distributions, which may not align with real-world behavior variance [17].

Autoencoders, a type of deep learning model, excel at learning compact representations of normal behavior. Trained on baseline data, these models flag input that results in high reconstruction error, indicating deviation from expected activity patterns. Their ability to model complex temporal and nonlinear dependencies makes them effective for user behavior analytics across time-series logs [18].

By combining these algorithms—or stacking them within ensemble architectures—organizations can increase resilience against a wide range of insider threat vectors. Each algorithm has strengths in different contexts, and model performance often improves when cross-validated against heterogeneous behavioral datasets [19].

3.3 Data Requirements and Feature Engineering for Insider Profiling

The performance of any machine learning system for insider threat detection is critically dependent on the quality, granularity, and structure of input data. Behavioral signals must be transformed into features that models can interpret—an often complex process involving log parsing, entity extraction, and temporal sequencing [20].

Log parsing is foundational. Enterprise environments generate vast volumes of raw log data from endpoints, authentication systems, file access controllers, and network devices. Parsing involves normalizing these logs into structured formats, typically using scripting languages like Python or tools such as Logstash. Key elements extracted include timestamps, usernames, IP addresses, and access events [21].

Entity extraction transforms unstructured fields into analyzable features. For example, file path parsing can distinguish access to sensitive documents versus general-purpose folders, and command-line logs can differentiate administrative commands from typical user activity. Natural language processing may be employed to process log comments or embedded messages that reflect behavioral intent [22].

Temporal sequencing captures patterns over time, essential for profiling deviations in routine behavior. Time-window aggregation techniques—such as sliding windows or exponential moving averages—help quantify how user behavior evolves. For example, sudden increases in login frequency or off-hours access are key risk indicators when contextualized over daily or weekly baselines [23].

In addition, feature engineering may include role-based normalization. This involves comparing user behavior against peer groups with similar job functions, helping reduce false

positives by accounting for role-specific access needs. Other engineered features might include ratios (e.g., write-to-read ratio), frequency vectors, or entropy scores to capture irregularity in action patterns [24].

Effective feature engineering bridges the gap between raw system telemetry and actionable behavioral models. It determines the success of ML models in learning accurate baselines and detecting subtle deviations without overwhelming analysts with noise. When combined with feedback loops, feature sets can evolve to reflect organizational changes and emerging threat patterns.

Table 1: Comparison of ML Techniques for Insider Threat Detection

Technique	Strengths	Weaknesses	Use Case Suitability
Decision Trees	Easy to interpret; fast training time	Prone to overfitting; limited in capturing complex patterns	Initial risk scoring; policy violation detection
Random Forest	Robust to overfitting; handles high-dimensional data	Slower inference; less interpretable	Behavioral baselining; anomaly scoring
Support Vector Machines (SVM)	Effective in high-dimensional space; good for small-to-medium datasets	Computationally intensive; less effective on large/noisy data	Classifying insider profiles vs. normal users
K-Nearest Neighbors (KNN)	Simple implementation; non-parametric	High memory and computation cost; sensitive to noisy data	Real-time behavior comparison with peer groups
Naïve Bayes	Fast and efficient; works well with categorical features	Assumes feature independence; less accurate on complex data	Email/text monitoring; access pattern classification
Neural Networks (DNN)	High accuracy; learns complex nonlinear	Requires large datasets; low interpretability	Advanced insider threat modeling; session

Technique	Strengths	Weaknesses	Use Case Suitability
	relationships		context modeling
Autoencoders	Unsupervised learning; effective at anomaly detection	Difficult to tune; may struggle with subtle malicious behavior	Detecting rare/unknown behavioral deviations
Isolation Forest	Efficient anomaly detection; works well with unlabeled data	May not capture contextual semantics	Detecting outlier access patterns and file transfers
Recurrent Neural Networks (RNN)	Good for sequence modeling; captures temporal dependencies	Complex to train; resource-intensive	Command-line behavior analysis; sequential activity logging

4. ANOMALY DETECTION MODELS FOR INSIDER THREAT FORENSICS

4.1 Unsupervised Anomaly Detection Techniques

Unsupervised anomaly detection techniques are pivotal for insider threat forensics, especially in environments where labeled datasets are unavailable or incomplete. These models detect deviations from normal behavior by learning inherent patterns in the data without requiring pre-classified instances. Popular approaches include clustering algorithms, One-Class Support Vector Machines (OCSVM), Principal Component Analysis (PCA), and statistical outlier models [13].

Clustering is one of the most common unsupervised approaches. Algorithms such as DBSCAN and k-means group users or actions into clusters of similar behavior. Anomalies are identified as points that do not belong to any cluster or lie far from cluster centroids. This is effective in identifying rare behavior sequences, especially in role-based profiling [14].

One-Class SVM builds a boundary around normal instances in high-dimensional space and classifies points outside this boundary as anomalies. This approach is particularly useful in high-skew settings, where malicious actions are significantly outnumbered by benign ones. OCSVM excels in defining compact regions of normalcy with minimal data assumptions [15].

Principal Component Analysis (PCA), though traditionally a dimensionality reduction technique, can also reveal outliers when projecting data onto principal components. Deviations from the dominant patterns—captured as residuals or reconstruction errors—often signal anomalous behavior. PCA is lightweight and interpretable, making it attractive for initial screening [16].

Statistical outlier detection, such as Z-score or interquartile range (IQR) methods, remain relevant in simpler environments or when used in conjunction with complex models. While they lack contextual awareness, their simplicity makes them useful for thresholds and baselining [17].

Overall, unsupervised techniques provide broad applicability and require minimal supervision, though their effectiveness improves when guided by domain knowledge and enriched feature sets.

4.2 Sequence Modeling and Time-Series Approaches

Understanding user behavior as a temporal sequence rather than isolated events is critical to effective insider threat detection. Sequence modeling and time-series analysis techniques help capture the evolving patterns of activity, enabling early identification of anomalies that emerge over time. Prominent methods include Long Short-Term Memory (LSTM) networks, Hidden Markov Models (HMMs), and activity path deviation models [18].

LSTMs, a subclass of recurrent neural networks, are highly effective in learning long-term dependencies within behavioral logs. These models process sequences of actions—such as login events, file access, or command execution—and learn to predict expected next steps. Deviations from expected sequences generate anomaly scores. LSTMs are particularly useful in environments with consistent workflow sequences, such as developer workstations or regulated data handling [19].

Hidden Markov Models assume user behavior can be represented as a sequence of hidden states, each emitting observable actions with a given probability. These models learn transition probabilities between states and flag sequences that do not conform to the trained transition matrix. HMMs work well for modeling role-specific workflows and are relatively interpretable, aiding forensic analysis [20].

Activity path deviation analysis focuses on comparing current user activity paths against historical or role-based norms. By constructing graphs or transition matrices of behavior flows—such as moving from email to document editing to database queries—analysts can flag abnormal navigation patterns. These models support intuitive visualization and can reveal subtle escalation attempts [21].

Time-aware modeling also facilitates trend detection—such as frequency of administrative access, rise in failed login attempts, or temporal clustering of sensitive file activity.

These indicators are crucial for identifying insider threats that manifest gradually over days or weeks rather than instantaneously [22].

Integrating temporal awareness into ML detection engines increases their ability to understand context, reduce false positives, and respond to evolving insider strategies.

4.3 Ensemble Models and Hybrid Detection Strategies

To maximize detection precision and adaptability, insider threat frameworks are increasingly adopting ensemble models and hybrid strategies that combine the strengths of multiple machine learning and rule-based approaches. These systems offer improved robustness, reduced bias, and better handling of varied attack patterns. Key ensemble techniques include stacking, weighted voting, and hybrid rule-ML combinations [23].

Stacking, or stacked generalization, involves training multiple base models (e.g., decision trees, autoencoders, clustering) and using their outputs as features for a meta-model. This architecture allows the system to learn when each base model is most reliable, improving predictive accuracy. For instance, anomaly scores from clustering, OCSVM, and PCA can be combined and interpreted collectively by a Random Forest meta-model [24].

Weighted voting assigns confidence scores to the outputs of different models and makes final decisions based on a cumulative threshold. This method allows for nuanced interpretation, especially when some models are more suitable for specific behavior types. For example, time-series models may dominate during sequential anomaly detection, while clustering may be more effective in detecting one-off deviations [25].

Hybrid rule and ML-based systems integrate deterministic logic with probabilistic learning. Rule-based systems handle clearly defined violations—such as access to terminated accounts or downloads exceeding policy limits—while ML models assess subtler patterns. This ensures low-latency alerts for known risks and deep learning for unknown or evolving behaviors [26].

Hybridization also supports interpretability and compliance. Many organizations require audit trails for forensic reporting and legal proceedings. By fusing machine-generated alerts with predefined rules, analysts can trace how decisions were made, improving trust in automated systems [27].

Finally, hybrid systems enable continuous learning. Feedback loops from analyst-confirmed alerts can retrain underlying models, improve feature relevance, and evolve thresholds, creating an adaptive detection ecosystem that matures with operational usage and adversary evolution.

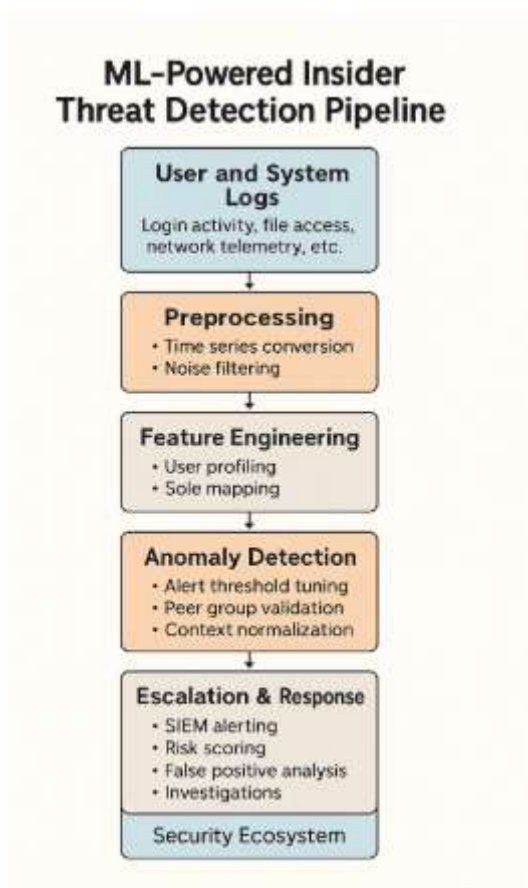


Figure 2: Flowchart of ML-Powered Insider Threat Detection Pipeline

5. BEHAVIORAL ANALYTICS AND USER CONTEXT MODELING

5.1 Building Behavioral Baselines from User Activity

Developing robust behavioral baselines from user activity is a cornerstone of modern cybersecurity analytics. These baselines are typically constructed using historical data that reflects the regular patterns of individual users in their digital environments. Behavioral analytics platforms examine variables such as login frequency, application access sequences, and file handling habits. By capturing this routine digital behavior, security systems can more effectively flag deviations that suggest malicious intent or compromised credentials [21].

Daily activity patterns are particularly useful for establishing temporal norms. Users tend to log in at similar times, use a consistent set of applications, and access files in repeatable ways. Once these routines are documented, security systems can calculate a statistical “normal” range for each user. Any significant deviation outside this envelope may indicate a breach or insider threat. This approach has proven especially effective when enhanced with machine learning algorithms that continuously update baselines based on real-time input [22].

Another powerful tactic involves comparing individual behavior against peer group norms. Employees in similar roles, departments, or projects typically exhibit comparable digital patterns. For example, if a finance team member suddenly begins accessing source code repositories or HR records, the anomaly is more easily detected when evaluated against departmental behavior [23]. This peer benchmarking strategy adds contextual depth to anomaly detection and reduces false positives by acknowledging that roles inherently differ.

Location regularity also plays a central role. Consistent geographic access patterns—such as a user always logging in from the same city or campus—can become part of the behavioral fingerprint. If a login attempt occurs from a distant or unexpected location without corresponding travel records or VPN usage, it may trigger an alert [24]. Combined with device fingerprints and network metrics, such data provides a layered defense mechanism that enriches the baseline model. Behavioral analysis rooted in longitudinal data, peer comparisons, and spatial norms provides a powerful early-warning system against internal and external threats [25].

5.2 Role of Context in Interpreting Anomalies

Contextual awareness transforms raw anomaly data into actionable intelligence. Without context, even legitimate activity can appear suspicious. For instance, an employee accessing large volumes of files at 2 AM may raise alarms—unless the system recognizes they are part of a global operations team working across time zones [26]. Embedding contextual variables—such as job role, time of access, and business calendar events—ensures that alerts are meaningful rather than distracting.

Job roles determine the scope of digital behavior that should be considered normal. A system administrator will naturally interact with configuration files and system logs that a marketing specialist never would. By mapping activities to role expectations, systems can filter out non-threatening deviations and focus on authentic risks [27]. This role-based calibration reduces alert fatigue and ensures that security teams concentrate on genuinely abnormal behaviors.

Holiday calendars and travel schedules also add interpretive value. During public holidays or approved leave periods, any user activity may warrant immediate scrutiny. Conversely, during business travel, logins from international IPs might be expected and benign. Integrating calendar metadata enables systems to infer intent and filter alerts accordingly [28].

Remote work has complicated traditional baselining, as employees now operate from diverse locations and at flexible hours. This shift demands that baselines accommodate fluidity while still enforcing security thresholds. For example, accessing sensitive systems from a coffee shop may be normal for one user but highly irregular for another. Context-aware systems can dynamically apply different risk scores based on

the individual’s historical behavior and declared remote work status [29].

Time zone normalization is another often overlooked yet critical element. A 3 PM login in New York may correspond to an 8 PM session in London. Without proper alignment, such access could seem irregular. By accounting for geographical context, systems can maintain integrity without unnecessary escalation [30]. Incorporating contextual data ultimately enhances the fidelity of threat detection while preserving the user experience.

5.3 Identity and Access Correlation Models

Modern identity and access correlation models are foundational for detecting policy violations, credential misuse, and lateral movement. These models analyze relationships between user identities, their role-based access rights, and real-time session behavior to identify gaps between entitlements and actual use [31]. At the heart of this analysis is Role-Based Access Control (RBAC) mapping, which defines what systems and data a user *should* access, based on their formal responsibilities.

By creating a digital inventory of access privileges tied to specific job functions, organizations can easily identify privilege creep—when users retain access rights that are no longer necessary. Periodic RBAC audits, supported by AI, can help ensure least-privilege principles are continuously enforced [32]. Deviations from assigned access privileges—such as attempts to open restricted directories or invoke administrator functions—can be instantly flagged for investigation.

Privileged session analysis offers a granular look into high-risk activities. Sessions involving root or administrative access are closely monitored using keystroke logging, session recording, and command parsing. By comparing session logs to expected behavior templates, systems can isolate risky actions like unauthorized privilege elevation or file exfiltration [33].

Furthermore, combining identity data with behavioral telemetry—such as login device type, time, and duration—enables correlation engines to distinguish between genuine users and imposters exploiting stolen credentials. A mismatch between identity and behavior, such as a junior analyst performing database dumps, signals a breach or misuse [34].

Advanced models also integrate Identity Governance and Administration (IGA) data to strengthen their predictive capabilities. This allows for real-time reconciliation of access rights, identity status changes, and policy exceptions [35]. Ultimately, correlation models close the loop between access entitlement and behavioral validation, forming a vital barrier against sophisticated insider and external threats.

Table 2: Contextual Features Commonly Used in Behavioral Modeling

Feature Category	Contextual Attribute	Purpose in Behavioral Modeling
Temporal	Time of day / Day of week	Detect off-hours or unusual login/usage patterns
Geospatial	IP address / Physical location	Identify access from unrecognized or unauthorized regions
Organizational Role	Job title / Department	Establish role-based access norms and behavioral expectations
Access Modality	Device type / Operating system	Detect login from unfamiliar or unauthorized devices
Peer Comparison	Group-based behavior norms	Benchmark individual actions against departmental or team norms
Application Context	Software accessed / Frequency of use	Detect rare or out-of-scope application usage
Privilege Level	User privilege / Admin rights	Identify unusual escalation or misuse of elevated privileges
Work Schedule	Standard vs. remote / holiday calendar	Normalize behavior based on approved work schedules and leave periods
Command Context	CLI usage / Scripting behavior	Identify abnormal sequences or syntax patterns indicative of malicious intent
Behavior History	Historical baselines / Recency of action	Compare current behavior with long-term trends and recent activity spikes

6. CASE STUDIES AND SIMULATED ATTACK RECONSTRUCTIONS

6.1 Case Study 1: Intellectual Property Theft by an Insider

In early 2023, a U.S.-based biotechnology firm discovered that critical intellectual property (IP) had been exfiltrated by a

senior research scientist just days before their resignation. Detection of this insider threat hinged on behavioral analytics, particularly the correlation of off-hours access patterns and anomalous file transfer behavior. The employee in question had maintained a stable access pattern during their eight years with the company, but in the final week, the user began accessing sensitive R&D folders at 2–4 AM—outside their typical 9–5 schedule [25].

The system flagged the irregular access times using a user behavior analytics (UBA) engine. While late-night access alone may not constitute malicious activity, the anomaly gained weight when paired with other data points. The user had started zipping multiple high-value files and transferring them to a personal cloud storage platform, something they had never done previously. Machine learning baselines were able to contrast this new behavior against the user's long-term historical patterns and also compare it against peer researchers who followed proper data handling protocols [26].

Interestingly, the insider had encrypted the outbound data packets, making traditional Data Loss Prevention (DLP) tools ineffective. However, the system leveraged endpoint telemetry and metadata—such as filename entropy and file movement frequency—to infer that exfiltration was likely in progress [27]. This alert prompted the incident response team to conduct a manual review, ultimately confirming the IP theft and triggering legal proceedings.

What stood out was the lack of traditional indicators like malware or phishing. The user had legitimate credentials and access rights, but their behavioral deviation was subtle and gradual. It reinforced the importance of **longitudinal behavioral modeling** and continuous monitoring rather than relying solely on static access control mechanisms [28].

The case underscores how AI-enhanced UBA tools can catch threats that signature-based systems miss, especially when paired with peer group analysis and time-sensitive behavioral shifts [29].

6.2 Case Study 2: Account Compromise with Lateral Movement

A multinational insurance provider faced a serious compromise when an employee's credentials were phished and subsequently used for lateral movement across its internal systems. The breach was uncovered only after anomaly detection tools flagged a non-standard logon chain that involved access from geographically distant endpoints within a short span of time [30].

The attacker initially accessed the user's mailbox, but within hours initiated logons to file servers and administrative consoles in a completely different region. The security system, augmented with AI-driven correlation models, noted that this pattern of access—switching between domains and jumping across networks—was inconsistent with any known employee workflow [31]. Importantly, while the credentials

were valid, the login velocity and device mismatch triggered an alert.

Once inside, the attacker performed privilege escalation using PowerShell scripts and registry modifications to gain administrative access. This sudden jump in privilege was atypical for the compromised user's role and was detected via behavioral comparison with baseline privilege elevation norms [32]. The anomaly was subtle—no malware was deployed—but the speed of escalation and type of commands issued suggested non-human activity.

Network flow analysis also revealed that lateral movement included scanning of subnet directories and unauthorized attempts to access HR and finance records. These systems were not in the regular path of the compromised employee's access history, prompting a more detailed investigation [33]. When coupled with endpoint data that showed token impersonation and log clearing commands, the pattern became unmistakably malicious.

What differentiated this case was the **integration of identity analytics with network behavior**. The AI system could correlate that while the user ID was correct, the access sequence did not match the user's historical data in timing, location, and system behavior [34]. Blocking based on device fingerprint mismatches and user role expectations finally curtailed the breach.

This case highlights the importance of cross-layer correlation—from access logs to system commands—especially in identifying breaches that rely on hijacked credentials rather than malware or brute force [35].

6.3 Case Study 3: Fraudulent System Changes by IT Admin

In a mid-sized financial services firm, suspicious modifications to account settings triggered an investigation into one of the organization's IT administrators. The activity initially went undetected due to the user's high-level access rights. However, AI-enabled auditing tools noticed subtle shifts in command usage patterns, especially during maintenance windows [36].

Typically, the administrator executed a narrow set of administrative scripts related to software patching and system updates. However, recent logs indicated increased use of manual override commands and several "sudo" operations that hadn't been invoked in over a year. These changes were flagged because they diverged from the admin's long-term behavioral fingerprint [37]. The system also tracked script execution frequency, syntax variations, and time-of-day distribution, all of which suggested an emergent, non-routine activity profile.

What drew further suspicion was that these sudo commands were used to disable specific logging modules and adjust account balance limits in the firm's transaction database. These were actions that typically required multi-person

approval, yet the logs showed them being executed under a single admin session [38]. The anomaly detection model identified this as a policy violation and escalated the alert.

An automated forensic review revealed that the admin had created a temporary service account with elevated privileges, masked through log obfuscation commands. This account was then used to process unauthorized account changes, impacting multiple customer financial profiles. This behavior was entirely inconsistent with the admin’s operational history and out of alignment with standard operating procedures [39].

Crucially, the detection system leveraged cross-validation from multiple data sources—command-line behavior, identity logs, and time-based thresholds—to determine that the risk profile had dramatically shifted. It wasn't a single red flag but the accumulation of micro-anomalies across datasets that led to escalation [40].

This case illustrates the power of behavioral surveillance in privileged environments. Even when a user holds the “keys to the kingdom,” AI can detect abuse by continuously comparing activity against evolving behavioral models. This is particularly vital in environments where administrative users are both a necessity and a potential single point of failure [41].

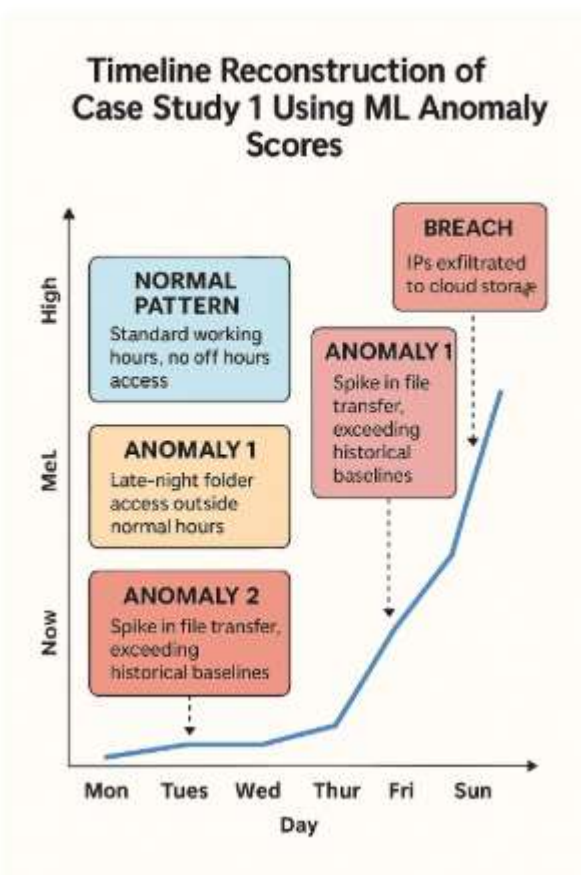


Figure 3: Timeline Reconstruction of Case Study 1 Using ML Anomaly Scores

Table 3: Feature Comparison Across Case Studies

Feature	Case Study 1 <i>IP Theft by Insider</i>	Case Study 2 <i>Account Compromise</i>	Case Study 3 <i>Fraudulent Admin Changes</i>
Off-hours activity	✓	✓	✓
Unusual file access	✓	✓	✗
Peer group deviation	✓	✗	✗
Remote location/IP anomaly	✗	✓	✗
Privilege escalation	✗	✓	✓
Command pattern anomalies	✗	✓	✓
Session velocity mismatch	✗	✓	✗
Behavioral baseline deviation	✓	✓	✓
Use of encrypted transfer	✓	✗	✗
Log obfuscation/suppression	✗	✗	✓
Unauthorized system changes	✗	✗	✓

7. IMPLEMENTATION, CONSIDERATIONS IN ENTERPRISE ENVIRONMENTS

7.1 Data Collection, Privacy, and Ethical Constraints

The efficacy of behavioral analytics in cybersecurity hinges on the quality and comprehensiveness of data collection. However, this process is fraught with ethical, legal, and privacy-related constraints. At the forefront is informed consent, which is essential in environments that monitor employees’ digital behavior. Consent policies must clearly

articulate what data is collected, how it is processed, and for what specific security purposes [28].

Balancing security with privacy requires a fine-tuned approach. Data anonymization is a common method to mitigate risks of employee surveillance overreach. Behavioral logs—such as mouse movement, file access, or system commands—can be aggregated and tokenized to remove identifying attributes while preserving their analytic value [29]. This helps address compliance requirements under regulations like GDPR, HIPAA, and CCPA, where employee data collection must adhere to strict data minimization and transparency standards.

Another critical challenge is insider trust balance. Excessive monitoring can foster a hostile work environment and diminish employee morale. Transparency through clearly defined data governance policies, role-based monitoring scopes, and purpose-specific logging reduces distrust. Ethical frameworks must ensure that surveillance tools do not become punitive but are directed solely at identifying genuine security threats [30].

In high-security industries such as finance or healthcare, some behaviors may be considered sensitive, such as research habits or patient record access patterns. Ethical design involves limiting the granularity of behavioral models to avoid inferring personal characteristics not relevant to security risk [31]. Building ethical AI frameworks includes stakeholder review boards and periodic audits to assess bias, over-collection, and misuse risks.

Ultimately, effective behavioral security systems must integrate privacy-by-design principles from inception, ensuring data collection practices are legally sound and ethically defensible while still enabling threat detection [32].

7.2 Model Training and Maintenance in Dynamic Organizations

Training behavioral analytics models in dynamic, fast-evolving organizations presents significant technical hurdles. Chief among them is the issue of concept drift—the gradual change in user behavior over time due to job role shifts, remote work arrangements, or seasonal workflows. If models are not retrained frequently, they begin to flag normal activity as suspicious or miss new forms of risk entirely [33].

To maintain efficacy, these models must incorporate continuous learning pipelines that accommodate real-time feedback loops. This allows the models to adjust to shifting behavioral norms without extensive manual tuning. Techniques like online learning and adaptive thresholds help recalibrate the detection logic based on evolving patterns across departments or user cohorts [34].

False positives remain a persistent challenge. Overly aggressive models can overwhelm security operations centers (SOCs) with benign alerts. Calibrating detection thresholds requires integrating historical event classification data—what

turned out to be false alarms versus verified threats—into the training process. This ensures the system prioritizes fidelity over mere activity volume [35].

Cross-validation with labeled datasets and role-specific behavioral profiles can further reduce false positives. For example, while a data scientist running batch queries at midnight may be normal, the same behavior by an HR staff member could raise valid concerns. Incorporating role-based templates during model retraining enhances precision [36].

Regular model evaluations, preferably on a monthly or quarterly basis, ensure the analytics remain aligned with operational realities. This is especially critical in sectors with frequent staffing changes, outsourced teams, or high mobility [37].

7.3 Integration with SIEM and Forensic Tools

The integration of behavioral analytics systems with Security Information and Event Management (SIEM) and forensic platforms is pivotal to transforming anomalies into actionable insights. Seamless compatibility with tools like Splunk, IBM QRadar, and the ELK (Elasticsearch, Logstash, Kibana) stack enhances incident response by allowing behavioral alerts to enrich existing log-based detection workflows [38].

Behavioral signals can act as high-fidelity triggers within SIEM platforms. For example, a sequence of anomalies—such as off-hours access, privilege escalation, and device mismatch—can generate a compound risk score, which SIEM systems use to prioritize threat events. Integrating these scores allows analysts to focus attention on complex, multi-vector intrusions rather than isolated log events [39].

Compatibility requires that behavioral analytics tools export data in normalized formats, such as JSON, Common Event Format (CEF), or Log Event Extended Format (LEEF). These formats ensure that systems like QRadar can ingest and correlate user behavior data alongside traditional network and application logs. The ELK stack, widely used in open-source environments, also benefits from visualization of long-term behavior trends across users and endpoints [40].

For forensic teams, integration enables session replay, command reconstruction, and activity heatmaps. Behavioral data enriches digital forensic workflows by providing not just the “what” and “when,” but the “why” behind anomalous behavior. Analysts can replay session timelines, identify causality chains, and map lateral movements with behavioral context [41].

Real-time API connectivity and alert forwarding into SIEMs not only accelerates detection but enhances incident documentation, audit readiness, and compliance post-mortem capabilities [42].

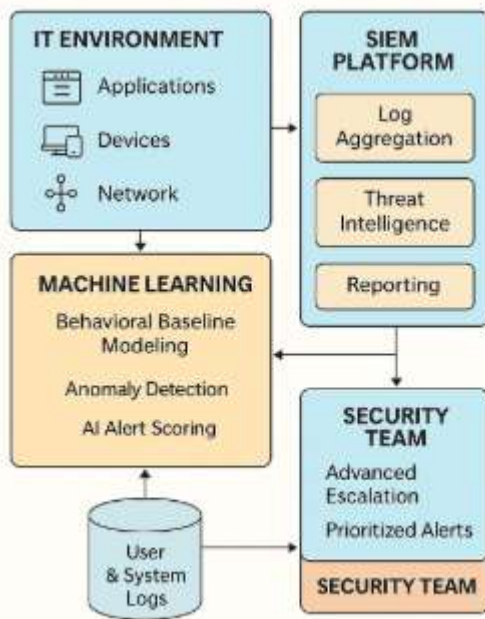


Figure 4: Insider Threat Architecture Integrating ML and SIEM

8. EVALUATION METRICS AND LIMITATIONS OF ML APPROACHES

8.1 Accuracy, Precision, Recall, and F1 Score

Evaluating the performance of behavioral anomaly detection models in cybersecurity demands more than just a superficial accuracy metric. Metrics like precision, recall, and F1 score provide a more balanced picture of how effectively threats are detected without flooding analysts with false positives. One of the biggest challenges lies in handling class imbalance—a scenario where anomalous behaviors (i.e., actual threats) constitute a small fraction of the total activity [32].

Accuracy alone can be misleading in this context. A model that classifies every user behavior as “normal” may score 98% accuracy but completely fail at detecting rare but dangerous anomalies. Precision reflects how many of the detected threats are real, while recall measures how many actual threats were caught. The F1 score, as a harmonic mean of precision and recall, helps evaluate model performance in cases where neither metric should dominate the other [33].

Threshold tuning plays a central role in this process. Setting detection thresholds too low increases recall but leads to more false positives, overwhelming security teams. Conversely, high thresholds may suppress alerts but risk missing subtle indicators of compromise [34]. Therefore, tuning must be guided by business risk tolerance, historical incident patterns, and operational capacity.

Another complexity involves delayed ground truth. Often, real-world confirmation of a threat comes days or weeks after

the activity occurred, complicating model validation. Dynamic thresholding and ROC curve analysis allow continuous adaptation as more feedback becomes available [35].

For organizations, model evaluation isn’t merely academic—it impacts resource allocation, escalation timelines, and ultimately, threat containment. Hence, precision-recall tradeoffs must be fine-tuned through iterative testing, stakeholder input, and contextual awareness across job roles and systems [36].

8.2 Real-World Evaluation Challenges

In practice, evaluating behavioral threat detection models is constrained by the scarcity of labeled ground truth data. Insider threats—unlike malware—rarely follow repeatable patterns and often unfold over long periods, making it difficult to assign definitive labels to activity logs. Without ground truth, standard model validation techniques lose reliability [37].

For example, anomalous logins or irregular command sequences may raise suspicion, but without confirmation of malicious intent or security incident, it becomes speculative whether such data should train future models. This gray-zone behavior blurs the distinction between noise and signal, introducing bias into model evaluation [38].

Manual labeling of insider behavior is another challenge. Security analysts rarely have the bandwidth to annotate large datasets, and retrospective investigations can be tainted by hindsight bias. The subjective nature of labeling often results in inconsistent datasets, which degrade model performance during both training and testing phases [39].

One solution lies in synthetic data generation, where simulated insider threat behaviors are seeded into training sets. While this helps balance the dataset, it also risks overfitting if synthetic samples don’t accurately represent the complexity of real-world threats. Moreover, adversaries constantly evolve, meaning yesterday’s attack signatures may be irrelevant today [40].

To address this, some systems incorporate semi-supervised learning, allowing models to learn from both labeled and unlabeled data. However, this too depends on careful calibration and regular analyst review cycles to validate flagged behavior. Ultimately, ground truth scarcity remains a structural barrier to perfect model evaluation in behavioral analytics [41].

8.3 Limitations of ML in Context-Rich Threat Environments

Machine learning (ML) models in behavioral analytics are powerful but not without limitations, especially in context-rich threat environments. These models often struggle to interpret the nuance behind user actions when stripped of

domain-specific context such as job roles, business cycles, or cross-system dependencies [42].

A significant limitation lies in the susceptibility to adversarial behavior. Sophisticated insiders or external actors who understand how models function can mimic normal activity to evade detection. For instance, by pacing file downloads or inserting deliberate “benign” actions, they can avoid crossing anomaly thresholds [43]. This adversarial adaptation forces continuous model retraining and complicates the establishment of stable behavioral baselines.

Another vulnerability is overfitting on benign outliers. When models are trained on high-dimensional user behavior data, they may interpret rare but harmless actions as malicious. This results in excessive false positives that degrade system trust and response effectiveness. Models that lack interpretability compound the issue by making it difficult for analysts to understand why certain behaviors were flagged [44].

Additionally, behavioral models often operate in silos. Without integration into broader enterprise context—such as HR data, workflow logs, or incident history—they can misclassify legitimate deviations tied to business exceptions or crisis responses. This context-blindness can produce costly distractions or missed threats [45].

While augmenting ML with domain knowledge helps, it also introduces complexity into model design and maintenance. As threat actors adapt and organizational behavior evolves, ML models must be continuously evaluated not just for accuracy, but for their situational awareness and adaptability [46].

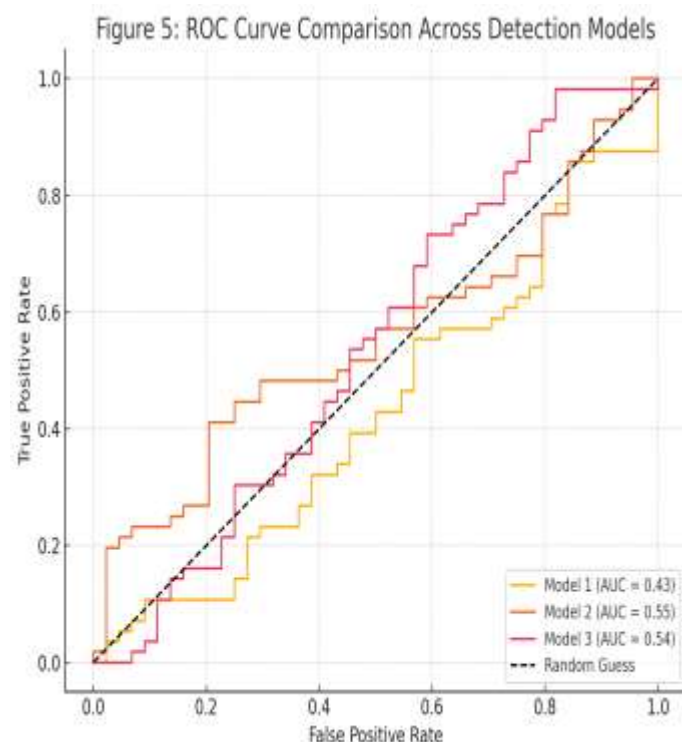


Figure 5: ROC Curve Comparison Across Detection Models

9. FUTURE DIRECTIONS IN INSIDER THREAT FORENSICS

9.1 Federated Learning and Cross-Organization Threat Intelligence

Federated learning has emerged as a transformative technique in behavioral analytics, particularly when addressing privacy concerns and data silos across organizations. This decentralized learning paradigm allows different institutions to collaboratively train models on behavioral threat data without directly sharing raw data, thereby preserving confidentiality while still benefiting from collective insights [47].

For example, healthcare providers, financial institutions, and government agencies can contribute to a federated threat detection model trained on anomalous behavior signatures from varied environments. This approach enhances generalizability and model robustness against novel attack vectors while ensuring compliance with regulations such as GDPR and HIPAA [48].

Moreover, cross-organization threat intelligence sharing enhances federated learning by enabling access to a broader range of threat profiles. Behavioral signals that may appear benign in isolation—like unusual login times or endpoint switching—may be linked to a larger attack chain when viewed in aggregate across multiple entities [49].

The key technical challenge lies in harmonizing behavioral telemetry formats and security taxonomies. However, emerging standards and secure multi-party computation methods are beginning to close this gap. Federated learning not only scales behavioral detection but also future-proofs it by minimizing dependency on single-organization data silos while reinforcing privacy guarantees [50].

9.2 Use of Explainable AI (XAI) in Behavior Interpretation

As machine learning models in behavioral threat detection become more complex, **explainability** has become a critical requirement for adoption in high-stakes environments. Security teams must understand why certain behaviors are flagged as anomalous to take appropriate and timely action. This has driven the rise of Explainable AI (XAI) techniques that make black-box models more transparent and interpretable [51].

In behavioral analytics, XAI methods such as SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-Agnostic Explanations) are used to break down predictions into constituent features—like login time deviation, file access anomalies, or unusual device usage—and assign weights to their influence on detection outcomes [52]. This not only helps analysts validate alerts but also fosters trust in automated systems.

Transparency is especially vital when behavioral models influence incident response or disciplinary action. Without explainability, false positives can erode confidence and lead to risk-averse model configurations that underperform in real-world scenarios [53].

Furthermore, XAI aids in auditing and compliance, as security teams must often justify their detection methods to regulators or legal departments. By illuminating model reasoning, organizations can defend both the accuracy and fairness of their behavioral surveillance strategies [54].

9.3 Augmented Analysts: Human-in-the-Loop Models

The future of behavioral threat detection lies in hybrid intelligence systems where human analysts are embedded within the machine learning loop. These “human-in-the-loop” (HITL) models enable continuous refinement of behavioral detection engines by leveraging human intuition and contextual awareness alongside algorithmic scalability [55].

In practice, HITL frameworks allow analysts to confirm, dismiss, or reclassify flagged behaviors, creating a feedback loop that strengthens model accuracy over time. This is especially useful in ambiguous cases where behavioral context—such as temporary policy exceptions or cross-functional project roles—can’t be fully captured by machine logic alone [56].

By combining machine precision with human judgment, these systems mitigate false positives, reduce alert fatigue, and allow for nuanced threat assessments. For example, while an AI model may flag a software engineer for accessing HR records, a human analyst could validate the behavior as part of an approved audit task based on access tickets or meeting logs [45].

Additionally, augmented analytics enables tiered escalation, where AI filters routine noise and forwards only high-risk or unclear cases to humans for deeper inspection. This improves triage speed and resource efficiency without compromising detection quality [46]. Ultimately, augmented analysts represent a scalable, adaptive, and explainable paradigm for behavioral cybersecurity.

10. CONCLUSION

10.1 Recap of Key Contributions

This paper has explored the emerging landscape of AI-driven behavioral threat detection in modern enterprise cybersecurity. At its core, we have demonstrated how building behavioral baselines, integrating contextual variables, and correlating identity access models can uncover subtle anomalies often missed by traditional rule-based systems. Through case studies—ranging from insider IP theft to privilege escalation and administrative fraud—we illustrated how machine learning models identify threats rooted not in

code or malware, but in deviations from expected human behavior.

We also evaluated the practical challenges of deploying and maintaining these systems in dynamic, real-world environments. Attention was given to the ethical and privacy implications of continuous monitoring, highlighting the need for transparent, consent-driven, and anonymized data collection practices. Technical performance metrics like precision, recall, and F1 score were discussed alongside systemic challenges such as concept drift, labeling scarcity, and adversarial evasion tactics.

Additionally, the study emphasized innovation in areas like federated learning, explainable AI, and human-in-the-loop frameworks. These techniques collectively contribute to making threat detection systems more scalable, interpretable, and aligned with organizational realities. Overall, the work provides a blueprint for the adoption of adaptive, intelligent, and ethically responsible behavioral security analytics.

10.2 Strategic Insights for Enterprises and Security Teams

For enterprises navigating increasingly sophisticated cyber threats, behavioral analytics represents a crucial next step in proactive security. Organizations must begin by investing in robust identity and access management foundations, as behavioral models rely heavily on accurate baselines and access control frameworks. Role-based access clarity, endpoint telemetry collection, and historical user behavior logging are prerequisites for effective anomaly detection.

Security teams should prioritize modular integration of behavioral tools into existing SIEM and incident response ecosystems. Rather than deploying standalone solutions, organizations will benefit most from platforms that communicate seamlessly with Splunk, QRadar, and ELK stacks, enabling real-time alert correlation and risk scoring. These tools should be configured for continuous feedback loops, allowing analyst interaction to refine model accuracy over time.

Enterprises should also consider ethical dimensions when deploying behavioral monitoring, particularly around employee privacy and surveillance concerns. Transparent governance, purpose-specific data retention policies, and regular audits help mitigate reputational and legal risks.

Lastly, capacity building is essential. Training SOC personnel to interpret AI-driven alerts and incorporating domain-specific threat models ensures that organizations don’t just collect more data—but derive actionable insight from it. Behavioral analytics should be seen not as a replacement for human expertise, but as a force multiplier for security operations.

10.3 Final Thoughts on Human-Centric, Adaptive Threat Detection

As cyber threats become increasingly sophisticated and personalized, threat detection systems must evolve toward

human-centric, adaptive models. This means prioritizing not just technological accuracy, but contextual awareness, ethical responsibility, and user trust. Behavioral analytics offers a framework that aligns closely with these principles by focusing on deviations in normal human behavior rather than solely relying on known attack signatures.

The future of enterprise cybersecurity lies in the dynamic synergy between machine intelligence and human intuition. While AI excels at processing vast datasets and identifying subtle patterns, humans remain critical for interpreting ambiguity, assessing intent, and making judgment-based decisions. The integration of explainable AI and human-in-the-loop frameworks ensures that this collaboration is both scalable and transparent.

Furthermore, adaptive threat detection must be designed with resilience in mind. Threat actors will continue to evolve tactics, and behavioral models must continuously learn and adjust to maintain relevance. This adaptability, supported by federated learning and real-time context integration, empowers organizations to stay one step ahead without compromising on compliance or user privacy.

Ultimately, the goal is not just to detect anomalies, but to build systems that understand human behavior deeply and respectfully. By embedding empathy, adaptability, and intelligence into threat detection, we move toward a safer, more resilient digital future.

11. REFERENCE

- Wei Y, Chow KP, Yiu SM. Insider threat prediction based on unsupervised anomaly detection scheme for proactive forensic investigation. *Forensic Science International: Digital Investigation*. 2021 Oct 1;38:301126.
- Alzaabi FR, Mehmood A. A review of recent advances, challenges, and opportunities in malicious insider threat detection using machine learning methods. *IEEE Access*. 2024 Feb 26;12:30907-27.
- Joseph Chukwunweike, Andrew Nii Anang, Adewale Abayomi Adeniran and Jude Dike. Enhancing manufacturing efficiency and quality through automation and deep learning: addressing redundancy, defects, vibration analysis, and material strength optimization Vol. 23, *World Journal of Advanced Research and Reviews*. GSC Online Press; 2024. Available from: <https://dx.doi.org/10.30574/wjarr.2024.23.3.2800>
- in Sarhan B, Altwaijry N. Insider threat detection using machine learning approach. *Applied Sciences*. 2022 Dec 25;13(1):259.
- nanah Onyekachukwu Victor, Akoji Sunday. Integrating AI and radioactive pharmacy nuclear imaging to tackle healthcare and mental health disparities in underserved U.S. populations: Business models, ethical considerations, and policy implications. *Int J Res Publ Rev*. 2025 Mar;6(3):2756-2769. Available from: <https://doi.org/10.55248/gengpi.6.0325.1199>
- Nasir R, Afzal M, Latif R, Iqbal W. Behavioral based insider threat detection using deep learning. *IEEE Access*. 2021 Oct 5;9:143266-74.
- Aidoo EM. Community based healthcare interventions and their role in reducing maternal and infant mortality among minorities. *International Journal of Research Publication and Reviews*. 2024 Aug;5(8):4620–36. Available from: <https://doi.org/10.55248/gengpi.6.0325.1177>
- Singh M, Mehtre BM, Sangeetha S. User behavior profiling using ensemble approach for insider threat detection. In 2019 IEEE 5th International Conference on Identity, Security, and Behavior Analysis (ISBA) 2019 Jan 22 (pp. 1-8). IEEE.
- Chukwunweike J. Design and optimization of energy-efficient electric machines for industrial automation and renewable power conversion applications. *Int J Comput Appl Technol Res*. 2019;8(12):548–560. doi: 10.7753/IJCATR0812.1011.
- Sharma B, Pokharel P, Joshi B. User behavior analytics for anomaly detection using LSTM autoencoder-insider threat detection. In Proceedings of the 11th international conference on advances in information technology 2020 Jul 1 (pp. 1-9).
- Nyombi, Amos and Nagalila, Wycliff and Sekinobe, Mark and Happy, Babrah and Ampe, Jimmy, Enhancing Non-Profit Efficiency to Address Homelessness Through Advanced Technology (December 05, 2024). Available at SSRN: <https://ssrn.com/abstract=5186065> or <http://dx.doi.org/10.2139/ssrn.5186065>
- Raval MS, Gandhi R, Chaudhary S. Insider threat detection: machine learning way. *Versatile cybersecurity*. 2018:19-53.
- Chukwunweike Joseph, Saladeen Habeeb Dolapo. Advanced Computational Methods for Optimizing Mechanical Systems in Modern Engineering Management Practices. *International Journal of Research Publication and Reviews*. 2025 Mar;6(3):8533-8548. Available from: <https://ijrpr.com/uploads/V6ISSUE3/IJRPR40901.pdf>
- Al-Mhiqani MN, Ahmad R, Zainal Abidin Z, Yassin W, Hassan A, Abdulkareem KH, Ali NS, Yunus Z. A review of insider threat detection: Classification, machine learning techniques, datasets, open challenges, and recommendations. *Applied Sciences*. 2020 Jul 28;10(15):5208.
- Mahama T. Training ensemble classifiers on genomic data to forecast personalized cancer treatment response probabilities. *Int J Res Publ Rev [Internet]*. 2025 Jun;6(6):722–36
- Mayhew M, Atighetchi M, Adler A, Greenstadt R. Use of machine learning in big data analytics for insider threat detection. In MILCOM 2015-2015 IEEE Military

- Communications Conference 2015 Oct 26 (pp. 915-922). IEEE.
17. Ugwueze VU, Chukwunweike JN. Continuous integration and deployment strategies for streamlined DevOps in software engineering and application delivery. *Int J Comput Appl Technol Res*. 2024;14(1):1–24. doi:10.7753/IJCATR1401.1001.
18. Khaliq S, Tariq ZU, Masood A. Role of user and entity behavior analytics in detecting insider attacks. In 2020 International Conference on Cyber Warfare and Security (ICWSS) 2020 Oct 20 (pp. 1-6). IEEE.
19. Bradford P, Hu N. A layered approach to insider threat detection and proactive forensics. In Proceedings of the Twenty-First Annual Computer Security Applications Conference (Technology Blitz) 2005 Dec 5.
20. Mahama T. Fitting Cox proportional hazard models to identify mortality predictors during pregnancy and postpartum periods. *World J Adv Res Rev* [Internet]. 2025;26(3):27–47. Available from: <https://doi.org/10.30574/wjarr.2025.26.3.2152>
21. Al-Mhiqani MN, Ahmad R, Zainal Abidin Z, Yassin W, Hassan A, Abdulkareem KH, Ali NS, Yunus Z. A review of insider threat detection: Classification, machine learning techniques, datasets, open challenges, and recommendations. *Applied Sciences*. 2020 Jul 28;10(15):5208.
22. Unanah Onyekachukwu Victor, Aidoo Esi Mansa. The potential of AI technologies to address and reduce disparities within the healthcare system by enabling more personalized and efficient patient engagement and care management. *World J Adv Res Rev*. 2025;25(2):2643-2664. Available from: <https://doi.org/10.30574/wjarr.2025.25.2.0641>
23. Böse B, Avasarala B, Tirthapura S, Chung YY, Steiner D. Detecting insider threats using radish: A system for real-time anomaly detection in heterogeneous data streams. *IEEE Systems Journal*. 2017 Jan 23;11(2):471-82.
24. Aidoo EM. Social determinants of health: examining poverty, housing, and education in widening U.S. healthcare access disparities. *World Journal of Advanced Research and Reviews*. 2023;20(1):1370–89. Available from: <https://doi.org/10.30574/wjarr.2023.20.1.2018>
25. Brdiczka O, Liu J, Price B, Shen J, Patil A, Chow R, Bart E, Ducheneaut N. Proactive insider threat detection through graph learning and psychological context. In 2012 IEEE Symposium on Security and Privacy Workshops 2012 May 24 (pp. 142-149). IEEE.
26. Ejedegba EO. Advancing green energy transitions with eco-friendly fertilizer solutions supporting agricultural sustainability. *International Research Journal of Modernization in Engineering Technology and Science*. 2024 Dec;6(12):[DOI: 10.56726/IRJMETS65313]
27. Haq MA, Khan MA, Alshehri M. Insider threat detection based on NLP word embedding and machine learning. *Intell. Autom. Soft Comput*. 2022 Jan 1;33(1):619-35.
28. Mahama T. Estimating diabetes prevalence trends using weighted national survey data and confidence interval computations. *Int J Eng Technol Res Manag* [Internet]. 2023 Jul;07(07):127.
29. Liu L, De Vel O, Chen C, Zhang J, Xiang Y. Anomaly-based insider threat detection using deep autoencoders. In 2018 IEEE international conference on data mining workshops (ICDMW) 2018 Nov 17 (pp. 39-48). IEEE.
30. Uwamusi JA. Navigating complex regulatory frameworks to optimize legal structures while minimizing tax liabilities and operational risks for startups. *Int J Res Publ Rev*. 2025 Feb;6(2):845–861. Available from: <https://doi.org/10.55248/gengpi.6.0225.0736>
31. Mayhew M, Atighetchi M, Adler A, Greenstadt R. Use of machine learning in big data analytics for insider threat detection. In MILCOM 2015-2015 IEEE Military Communications Conference 2015 Oct 26 (pp. 915-922). IEEE.
32. Unanah Onyekachukwu Victor, Mbanugo Olu James. AI and big data in addressing healthcare disparities: Focused strategies for underserved populations in the U.S. *Int Res J Mod Eng Technol Sci*. 2025 Feb;7(2):670. Available from: <https://www.doi.org/10.56726/IRJMETS67321>
33. Eldardiry H, Bart E, Liu J, Hanley J, Price B, Brdiczka O. Multi-domain information fusion for insider threat detection. In 2013 IEEE Security and Privacy Workshops 2013 May 23 (pp. 45-51). IEEE.
34. Chukwunweike J, Lawal OA, Arogundade JB, Alade B. Navigating ethical challenges of explainable AI in autonomous systems. *International Journal of Science and Research Archive*. 2024;13(1):1807–19. doi:10.30574/ijrsra.2024.13.1.1872. Available from: <https://doi.org/10.30574/ijrsra.2024.13.1.1872>.
35. Vitla S. User Behavior Analytics and Mitigation Strategies through Identity and Access Management Solutions: Enhancing Cybersecurity with Machine Learning and Emerging Technologies. *Turkish Journal of Computer and Mathematics Education (TURCOMAT) ISSN*. 2023;3048:4855.
36. Mahama T. Assessing cancer incidence rates across age groups using stratified sampling and survival analysis methods. *Int J Comput Appl Technol Res* [Internet]. 2022;10(12):335–49. Available from: <https://doi.org/10.7753/IJCATR1012.1008>
37. Vitla S. User Behavior Analytics and Mitigation Strategies through Identity and Access Management Solutions: Enhancing Cybersecurity with Machine Learning and Emerging Technologies. *Turkish Journal of Computer and Mathematics Education (TURCOMAT) ISSN*. 2023;3048:4855.
38. Uwamusi JA. Crafting sophisticated commercial contracts focusing on dispute resolution mechanisms, liability limitations and jurisdictional considerations

- for small businesses. *Int J Eng Technol Res Manag*. 2025 Feb;9(2):58.
39. Anakath AS, Kannadasan R, Joseph NP, Boominathan P, Sreekanth GR. Insider Attack Detection Using Deep Belief Neural Network in Cloud Computing. *Computer Systems Science & Engineering*. 2022 May 1;41(2).
 40. Diyaolu CO. Advancing maternal, child, and mental health equity: A community-driven model for reducing health disparities and strengthening public health resilience in underserved U.S. communities. *World J Adv Res Rev*. 2025;26(03):494–515. Available from: <https://doi.org/10.30574/wjarr.2025.26.3.2264>
 41. G. Martín A, Fernández-Isabel A, Martín de Diego I, Beltrán M. A survey for user behavior analysis based on machine learning techniques: current models and applications. *Applied Intelligence*. 2021 Aug;51(8):6029–55.
 42. Adeoluwa Abraham Olasehinde, Anthony Osi Blessing, Joy Chizorba Obodozie, Somadina Obiora Chukwuemeka. Cyber-physical system integration for autonomous decision-making in sensor-rich indoor cultivation environments. *World Journal of Advanced Research and Reviews*. 2023;20(2):1563–1584. doi: [10.30574/wjarr.2023.20.2.2160](https://doi.org/10.30574/wjarr.2023.20.2.2160)
 43. Brdiczka O, Liu J, Price B, Shen J, Patil A, Chow R, Bart E, Ducheneaut N. Proactive insider threat detection through graph learning and psychological context. In 2012 IEEE Symposium on Security and Privacy Workshops 2012 May 24 (pp. 142–149). IEEE.
 44. Aidoo EM. Advancing precision medicine and health education for chronic disease prevention in vulnerable maternal and child populations. *World Journal of Advanced Research and Reviews*. 2025;25(2):2355–76. Available from: <https://doi.org/10.30574/wjarr.2025.25.2.0623>
 45. Hu T, Niu W, Zhang X, Liu X, Lu J, Liu Y. An insider threat detection approach based on mouse dynamics and deep learning. *Security and communication networks*. 2019;2019(1):3898951.
 46. Unanah Onyekachukwu Victor, Mbanugo Olu James. Telemedicine and mobile health imaging technologies: Business models for expanding U.S. healthcare access. *Int J Sci Res Arch*. 2025;14(2):470–489. Available from: <https://doi.org/10.30574/ijsra.2025.14.2.0398>
 47. Greitzer FL, Hohimer RE. Modeling human behavior to anticipate insider attacks. *Journal of Strategic Security*. 2011 Jul 1;4(2):25–48.
 48. Mahama T. Developing multilevel logistic regression models for predicting cancer outcomes from population health surveys. *Int J Adv Res Publ Rev [Internet]*. 2025 May;2(05):443–64.
 49. Alsowail RA, Al-Shehari T. Empirical detection techniques of insider threat incidents. *IEEE Access*. 2020 Apr 23;8:78385–402.
 50. Adeoluwa Abraham Olasehinde, Anthony Osi Blessing, Somadina Obiora Chukwuemeka. DEVELOPMENT OF BIO-PHOTONIC FEEDBACK SYSTEMS FOR REAL-TIME PHENOTYPIC RESPONSE MONITORING IN INDOOR CROPS. *International Journal of Engineering Technology Research & Management (IJETRM)*. 2024Dec21;08(12):486–506.
 51. Homoliak I, Toffalini F, Guarnizo J, Elovici Y, Ochoa M. Insight into insiders and it: A survey of insider threat taxonomies, analysis, modeling, and countermeasures. *ACM Computing Surveys (CSUR)*. 2019 Apr 2;52(2):1–40.
 52. Adeoluwa Abraham Olasehinde, Anthony Osi Blessing, Adedeji Adebola Adelagun, Somadina Obiora Chukwuemeka. Multi-layered modeling of photosynthetic efficiency under spectral light regimes in AI-optimized indoor agronomic systems. *International Journal of Science and Research Archive*. 2022;6(1):367–385. doi: [10.30574/ijsra.2022.6.1.0267](https://doi.org/10.30574/ijsra.2022.6.1.0267)
 53. Harilal A, Toffalini F, Homoliak I, Castellanos JH, Guarnizo J, Mondal S, Ochoa M. The Wolf Of SUTD (TWOS): A Dataset of Malicious Insider Threat Behavior Based on a Gamified Competition. *J. Wirel. Mob. Networks Ubiquitous Comput. Dependable Appl.*. 2018 Mar;9(1):54–85.
 54. Unanah Onyekachukwu Victor, Mbanugo Olu James. Integration of AI into CRM for effective U.S. healthcare and pharmaceutical marketing. *World J Adv Res Rev*. 2025;25(2):609–630. Available from: <https://doi.org/10.30574/wjarr.2025.25.2.0396>
 55. Wang C, Zhu H. Wrongdoing monitor: A graph-based behavioral anomaly detection in cyber security. *IEEE Transactions on Information Forensics and Security*. 2022 Jul 15;17:2703–18.
 56. Unanah OV, Aidoo EM. The potential of AI technologies to address and reduce disparities within the healthcare system by enabling more personalized and efficient patient engagement and care management. *World Journal of Advanced Research and Reviews*. 2025;25(2):2643–64. Available from: <https://doi.org/10.30574/wjarr.2025.25.2.0641>