



American Journal of Data Science and Artificial Intelligence (AJDSAI)

ISSN: 3069-3632 (ONLINE)

VOLUME 2 ISSUE 2 (2026)



PUBLISHED BY
E-PALLI PUBLISHERS, DELAWARE, USA

Insider Threat Prediction Using Graph Analysis

Muhammad Irshad^{1*}, Muhammad Faisal Shafiq², Muhammad Naveed Sajjad³

Article Information

Received: April 12, 2026**Accepted:** June 26, 2026**Published:** September 01, 2026

Keywords

*Cybersecurity, Detection Method,
Graph Analysis, Insider Threat,
Organizations*

ABSTRACT

The insider threat is still among the most difficult cybersecurity risks because of the access and capabilities of insiders to hide malicious or careless actions in the ordinary operations. The rules-based system, statistical anomaly detection and traditional machine learning tools are not always efficient in identifying the relational and contextual dependence in an enterprise setting which limits predictive capability as well as high false-positive. This paper presents a graph-based model of active preemptive insider threat detection. The Insider Threat Dataset of Multi-source behavioral logs of Classified Environments are converted to a heterogeneous interaction graph, where the nodes represent users, devices, and resources, and the edges indicate the frequency of interaction, sensitivity, and time patterns. Normative measures such as degree, between, eigenvector centrality, community membership, motif patterns and PageRank are derived to display aberrant relational activity. Empirical analysis has shown that a higher number of off-hours of printing/burning, larger volumes of data being exfiltrated, longer occupancy duration, and high-risk travel occur in malicious insiders occupying more influential network positions (much higher PageRank). The full prediction accuracy (FNNs classify every sample correctly) of Graph Neural Networks (GNNs) is high ($F1 = 1.0$, $AUC = 1.0$), which is significantly higher than that of the traditional baselines (Random Forest: $F1 = 0.7576$; XGBoost: $F1 = 0.6753$). The findings demonstrate the effectiveness of the graph-based methods in providing high-quality behavioral dependencies, with better accuracy and fewer false alarms and greater explainability in real-life insider risk monitoring.

INTRODUCTION

Businesses are becoming more dependent on intertwined information systems, cloud systems, and working together on a digital platform to maintain day-to-day functions. Extrusion of security attacks is ever growing more sophisticated, yet a large share of such attacks has an internal source (Aslan *et al.*, 2023). One of the most difficult and harmful types of cybersecurity risks is insider threats that can be caused by employees, contractors, or trusted partners operating with legitimate access (Saxena *et al.*, 2020). The evidence in the industry, such as the results published in the IBM Cost of a Data Breach Report, constantly shows that the incidents related to insiders are linked to the extremely high financial losses, the prolonged timeframes of identification, regulatory implications, and the negative reputations that take a long time to recover (Bonderud, 2024). In contrast to the external enemies, insiders have approved credentials and contextual understanding of organizational procedures and systems and can wizard malicious or careless activities in the regular operations of the organization (Prabhu & Thompson, 2022). Insider threats also include deliberate communication like stealing data, intellectual property, or performing sabotage, as well as unintentional actions, such as breaking a policy or negligent use of sensitive data (Georgiadou *et al.*, 2022). Insider risk is subtle regarding behavior. The malicious activities are usually

interlaced with the legitimate work patterns and thus present a challenge to tell them apart using the traditional security monitoring tools. The organizations today produce immense volumes of heterogeneous log data which comprises of authentication records, file access event logs, email dialogues and network transaction logs (Partovian *et al.*, 2023; Ji *et al.*, 2024). The difficulty with this analysis, however, is not the lack of data, but rather how to effectively model and interpret complicated behavioral relationships contained in such logs. The conventional methods of detecting insider threat mainly employ rule-based systems, preset signs of compromise, and statistical detection of anomalies (Alzaabi & Mehmood, 2024). Although these approaches can be used to determine known threat signatures or blatant violations of known baselines, they have several limitations that are critical. Rule-based systems are also highly reactive and cannot easily change according to changes in attack strategy or one that they have not encountered before (Dhoundiyal *et al.*, 2025; Mink *et al.*, 2023). Most of the models of anomaly detection consider user behaviors as independent events without considering the relation among users, devices, and resources (Skopik *et al.*, 2022). Moreover, traditional machine learning methods usually require some form of tabular representation of logs of activities, which dimensionless structural data and does not reflect the dynamic relationship of enterprise

¹ Cybersecurity GRC Consultant, RMG Company, Riyadh, Saudi Arabia² Sr. Security Analyst, Lucid Motors, Riyadh, Saudi Arabia³ Sr. Cybersecurity Consultant, Yanal Finance, Saudi Arabia* Corresponding author's e-mail: muhammadirshad73@outlook.com

ecosystems (Liang *et al.*, 2023). These limitations hinder predictive performance, raise the false positive rates, and lower the trust that is put on automated monitoring systems. Graph analysis may be a strong alternative paradigm of relational environment modeling of insider behavior. Graph structures retain the contextual and structural relationships upon which real world activity is founded by modeling organizational entities like users, files, devices, departments as nodes and their interactions as edges (Besta *et al.*, 2023). Topological features such as centrality measures, clustering patterns, and community structures can be extracted with such representations and will help uncover the invisible behavioural anomalies (Pazho *et al.*, 2023). Graph-based methods of learning, such as graph embeddings and Graph Neural Networks, also enable better identification of minute deviations in patterns of interaction between time (Ekle & Eberle, 2024). However, unlike more traditional models, graph-based models recognize the natural networked character of organizational systems and offer a more analytical basis of predictive security intelligence. This research is driven by the necessity to leave the reactive detection and shift to predicting insider threats proactively. Instead of only detecting the incidences when the situation has already been damaged, predictive frameworks seek to detect the high-risk trends at an earlier stage. The study consequently builds a graph-analytic model that converts the multi-source behavioral records into a heterogeneous interaction network and uses structural and learning-based methods to evaluate insider risk. The proposed solution aims to enhance the detection accuracy as well as interpretability and false alarms by maintaining relational dependence between variables and maintaining the contextual interactions between variables. This paper aims at formalizing the insider threat prediction into a graph-based modeling lens and to develop an analytical system that can model the complex behavioral dependencies within the enterprise context. The paper illustrates that the establishment of a graph and extraction of structural features and predictive modeling can be incorporated into a logical approach to the estimation of insider risk. The study uses empirical analysis of benchmark data to determine the effectiveness of graph-based methods in comparison with the traditional machine learning baselines. This work can help advance the structurally informed insider threat prediction by combining the use of cybersecurity analytics with the methodologies of graph-theoretic and machine learning. It offers a methodological perspective as well as practical advice to organizations that want to enhance internal security surveillance systems in more intricate digital space.

LITERATURE REVIEW

The role of insider threats has become one of the most urgent issues in the organizational cybersecurity context, which has received considerable research over the last 20 years. These dangers are quite intricate by the virtue of the fact that insiders have a legitimate access and contextual

knowledge, which enables the malicious or careless activity to be integrated with operation as usual. It has therefore become a key issue to practitioners and researchers to develop effective detection and prediction mechanisms. Various different approaches have been offered including both classical rule-based and anomaly-detection systems as well as sophisticated machine learning frameworks with their advantages and drawbacks.

Overview of Insider Threat Detection Techniques

The initial techniques of detecting insider threats can be categorically divided into rule-based approaches, anomaly-based approaches, and machine learning-based approaches.

Rule-based approaches

These are some of the oldest methods that were used in monitoring security in an organization. Such systems may be based on the preset indicators, policies, and thresholds to identify possible malicious activity (Qiu *et al.*, 2024). As an example, the odd usage of confidential files during working hours, the over-downloading of data, and opening of system privileges may signal an alarm. The key benefit of rule-based systems is that they are very simple and can be interpreted: security teams can understand why a specific activity is raised with a flag (Thiyagarajan *et al.*, 2026). Nevertheless, these strategies are reactive and inflexible in nature. To meet new patterns of attacks, they need to be constantly updated with rules and thresholds. Moreover, insiders who understand the rules can usually evade detection, and the covert or new threat activity might not be detected (Alohaly *et al.*, 2022). Consequently, although rule-based systems will continue to be effective in preliminary screening and compliance oversight, they have low predictive ability, especially in flexible and multifaceted business scenarios.

Anomaly-based approaches

These methods seek to address the constraints of static rules and come up with a normal user behavior baseline and raise red flags whenever there is a deviation in behavior as an indicator of a potential threat (Subrahmanyam, 2025). Statistical modeling, clustering and probabilistic methods are some of the techniques that have widely been applied here (Mohanty, 2025). An example is that the average login time of the user, patterns of accessing the network, and frequency of interacting with files can be statistically modeled and when the user performs activities that are very different compared to the norms, these activities are considered an anomaly (Shakil *et al.*, 2023). Unlike rule-based systems, anomaly detection is more flexible and can detect any behavior that has never been seen before. However, it has several challenges. To begin with, it is not trivial to model with high precision normal behavior in large and heterogeneous data sets (Bello *et al.*, 2025). Naturally, user activity is usually varied, and such consideration lack will lead to high false positivity rate. Second, the anomaly-based systems do not

follow a relational approach between users, resources, and devices and assume that events are treated separately (Sankaewtong *et al.*, 2025). This makes them less effective in identifying coordinated or insidious insider attacks which might not be reflected by drastic deviations in a single metric.

Machine learning (ML) methods

Those techniques have become an effective instrument in insider threat detection that is less susceptible to the limitations of rule-based and anomaly-based methods. This has been implemented by supervised learning models including support vector machines (SVMs), random forests, and logistic regression to classify the behaviors as malicious or benign with respect to labeled data sets (Kesarwani *et al.*, 2024; Verma *et al.*, 2022; Abiramasundari & Ramaswamy, 2025). Such models can acquire more complicated decision boundaries and learn various features of behavior at the same time, allowing more subtle threat detection. With unsupervised and semi-supervised methods, such as clustering, autoencoders, and isolation forests, the detection of suspicious behavior in unlabeled data is possible, which is particularly useful due to a shortage of confirmed data about insider attacks (Hasan, 2024). In more recent work, deep learning networks such as recurrent neural networks (RNNs) and long short-term memory (LSTM) networks have been applied to learn temporal sequences in user activity logs and are found to be more effective at detecting changing patterns of attacks (Bajao & Sarucam, 2023; Padmavathi & Muthukumar, 2023). Although they have the potential, ML-based approaches can be labor-intensive to feature engineering, poorly scale to imbalanced data where malicious examples are uncommon, and have poor interpretability, which can make them impossible to understand and trust the results of the models by security practitioners.

Graph Analysis in Cybersecurity

Graph-based methods have come into the limelight in the sphere of cybersecurity as a reaction to the restrictions of traditional approaches. Graph analysis offers a natural model on how relational, structural interactions occur in complex systems (Hasan & Nijhum, 2024). Regarding the area of insider threat detection, nodes can represent entities (ex: users, files, devices, departments), whereas edges can be logins, file access, communications, and network connection (Ghosh *et al.*, 2024). This expression maintains the pattern of connections and the relationship context of organizational behavior to ensure increase in the detection of delicate anomalies that would otherwise not be evident in the data of isolated events. A number of graph-based measurements have been useful in analysing insider threats. Centrality measures, including degree, betweenness and eigenvector centrality, may point to influential users or bizarre trends of access (Sharma *et al.*, 2024). Community detection algorithms have the capability to detect strongly related groups of users and resources,

and abnormal usage can be indicated by anomalies in the communities (Li *et al.*, 2024). The recurrent pattern of interactions which are related to malicious behavior can be identified through motif analysis as well as subgraph patterns. Moreover, recently, Graph Neural Networks (GNNs) and graph embedding algorithms have been used to compute the vector representation of nodes and edges so that predictive modeling can be combined with graph structure (Khemani *et al.*, 2024). Graph-based methods can help identify insider threats much better by integrating structural, temporal, and relationship data, which provides a better understanding of organizational behavior in a more comprehensive way (Sharma *et al.*, 2024; Ghosh *et al.*, 2024; Khemani *et al.*, 2024). Several studies have already shown that graph analysis can be of use in cybersecurity (Liu *et al.*, 2022; Biswas & Dhanekula, 2024; Almazrouei *et al.*, 2023; Lagraa *et al.*, 2024). As an example, scholars have used heterogeneous graphs to study multi-layered interaction between an organization to predict insider risk through the integration of user-device-resource relationships (Zheng *et al.*, 2022; Ding *et al.*, 2024; Eberle *et al.*, 2010). Temporal graphs have also been used by others to preserve changing interaction patterns and thus deviations can be identified over time (Jin *et al.*, 2023; Lian *et al.*, 2023). The graph-based strategies are more appropriate than the traditional ML and anomaly-detection techniques to capture the connected and context-dependent character of insider behavior, which results in a higher accuracy of prediction and decreases false positives.

Gaps in Current Research

Although the detecting of the insider threat has improved, there are other gaps that exist in research. First, most of the existing research is based on synthetic datasets or smaller-scale datasets that do not fully reflect the diversity and complexity of insider activities in the real world. Although databases like CERT can be useful benchmarks, they are synthetic, casting doubt on generalizability. Second, most conventional ML methods do not utilize relational and structural data by considering user behaviors as unrelated features, but as sequences of actions. Third, although graph-based approaches have potential, they usually treat graphs or simplified representation of relations, which ignores temporal dynamics and heterogeneous forms of interactions. This reduces their relevance in dynamic business settings where behavioral patterns evolve with time. Fourth, interpretability is also a problem, especially with deep learning models on graph data. To effectively respond to the threat of insiders, security practitioners need insights that will not only provide predictive scores but also actionable insights. Lastly, the combination of graph-based information and the organizational context, policy regulations, and human knowledge is still under-researched. It is important that the gap between computational detection means and its application in real-life scenarios be bridged in order to allow real-world impact. A solution to these gaps includes frameworks

that integrate heterogeneous, temporal and structural modeling with interpretable predictive processes, tested on real-world datasets which are representative of the variability of insider behaviors.

MATERIALS AND METHODS

Problem Definition

The formal definition of insider threat prediction is considering a task of recognizing the user inside an organization whose behavior denotes a high likelihood of malicious or negligent conduct. Let $U=\{u_1, u_2, \dots, u_n\}$ denote the set of all users in an enterprise, and $E=\{e_1, e_2, \dots, e_m\}$ represent observed interactions among users, devices, and resources over a period T . Each user u_i generates a sequence of events $X_i=\{x_1, x_2, \dots, x_i\}$, where x_i corresponds to an activity such as file access, login, email communication, or system command. The insider threat prediction problem can be expressed as learning a function:

$$f: X_i \rightarrow y_i, y_i \in \{0, 1\}$$

where $y_i=1$ indicates a high-risk insider and $y_i=0$ indicates benign behavior. The idea is to achieve as much predictive accuracy with the minimal number of false positives,

using personal behavior and relationship context (in the form of a graph of interaction).

Data Sources and Preprocessing

In the study, the Insider Threat Dataset of Classified Environments which provides both static employee characteristics and dynamic behavioral signals is used. The data set is to help the identification and forecasting of insider threats within the secure organizational setup. It also has multi-modal capabilities of personal employee data, printing and file management, travel histories and access control logs. The data set is a mix of categorical, numerical and Boolean variables, which can be modeled richly in both behavioral and structural terms especially in graph models. Such a rich set of static, behavioral and risk related attributes permits us to perform multi-dimensional analysis and lets to create graph representations in which employees, devices and resources can be inter-related depending on access, printing and file-handling behavior. Such depictions are essential in the derivation of structural and relational features to forecast insider threats successfully (Table 1).

Table 1: Insider Threat Dataset Features

Category	Field Name	Type	Description
Employee Static Properties	employee id	Number	Unique employee identifier
	date	Date	Activity date
	employee department	Categorical	Employee's department
	employee campus	Categorical	Employee's assigned campus
	employee position	Text	Job title/position
	employee seniority years	Quantity	Years of service
	is contractor	Boolean (0/1)	Contractor status
	employee classification	Categorical (1-4)	Security clearance level
	has foreign citizenship	Boolean (0/1)	Foreign citizenship status
	has criminal record	Boolean (0/1)	Criminal background indicator
	has medical history	Boolean (0/1)	Medical history flag
Security and Risk Indicators	employee origin country	Categorical	Employee's country of origin
	is malicious	Boolean (0/1)	Malicious activity indicator
	risk travel indicator	Boolean (0/1)	High-risk travel flag
Printing Activities	is emp malicious	Boolean (0/1)	Employee malicious behavior flag
	num print commands	Quantity	Total print commands issued
	total printed pages	Quantity	Total pages printed
	num print commands off hours	Quantity	Print commands during unusual hours
	num printed pages off hours	Quantity	Pages printed during unusual hours
	num color prints	Quantity	Color print count
	num bw prints	Quantity	Black & white print count

	ratio color prints	Float	Ratio of color to total prints
	printed from other	Boolean (0/1)	Printed from non-assigned campus
	print campuses	Categorical/List	Campuses where printing occurred
File Burning/Transfer Operations	num burn requests	Quantity	Total burn requests
	max request classification	Categorical (1-4)	Highest classification level burned
	avg request classification	Float (1-4)	Average classification level
	num burn requests off hours	Quantity	Burn requests during unusual hours
	total burn volume mb	Quantity	Total data volume burned (MB)
	total files burned	Quantity	Total number of files burned
	burned from other	Boolean (0/1)	Burned from non-assigned campus
	burn campuses	Categorical/List	Campuses where burning occurred
Travel Information	is abroad	Boolean (0/1)	Whether employee was abroad
	trip day number	Sequential (1,2,3...)	Day N of the trip
	country name	Categorical	Destination country
	is hostile country trip	Boolean (0/1)	Trip to hostile country
	hostility country level	Categorical/Numeric	Level of country hostility
	is official trip	Boolean (0/1)	Official business trip
Access Control	num entries	Quantity	Number of facility entries
	num exits	Quantity	Number of facility exits
	first entry time	Datetime	First entry time of day
	last exit time	Datetime	Last exit time of day
	total presence minutes	Quantity	Total time spent in facility (minutes)
	entered during night hours	Boolean (0/1)	Night entry flag
	num unique campus	Quantity	Number of different campuses visited
	early entry flag	Boolean (0/1)	Early entry indicator (before 06:00)
	late exit flag	Boolean (0/1)	Late exit indicator (after 22:00)
	entry during weekend	Boolean (0/1)	Weekend entry flag

Preprocessing is the cleaning, normalization and encoding of heterogeneous types of events. One-hot encoding or embedding are common ways of transforming categorical variables (e.g. types of resources, user roles) (Borrohoun *et al.*, 2024). Sequences over time are clustered into significant periods and noise filtering removes the benign anomalies that may lead to biasedness of the model. Formally, the preprocessed data for each user

u_i can be represented as a feature matrix $F_i \in \mathbb{R}^{t \times d}$, where t is the number of time intervals and d is the number of extracted features.

Graph Representation of User Behavior

To capture relational and structural dependencies, user behavior is modeled as a graph $G=(V,E)$, where V is the set of nodes and E is the set of edges.

Nodes (V) represent entities in the organization, including users, devices, and resources as from Eq.1:

$$V = U \cup D \cup R \quad (1)$$

where U is the set of users, D represents devices, and R denotes resources such as files or databases.

Edges (E) represent interactions between nodes, including logins, file accesses, email exchanges, or network communications. An edge $e_{ij} \in E$ between nodes v_i and v_j can be weighted by interaction frequency, duration, or sensitivity of accessed resources as indicated in Eq.2:

$$w_{ij} = \text{frequency}(v_i, v_j) \times \text{importance}(v_j) \quad (2)$$

The graph structure can be used to model both the direct activity of the user and the indirect relational influence, e.g., collusion or coordinated activity.

Graph-Based Feature Extraction

Graph-based features capture topological properties that are indicative of anomalous or high-risk behavior. Key categories include:

Centrality Measures: Quantify the importance of nodes in the network. Common metrics include:

Degree centrality:

$$C_D(v_i) = \sum_j A_{ij} \quad (3)$$

From Eq.3 where A_{ij} is the adjacency matrix. High-degree nodes may represent users with unusually broad access.

Betweenness centrality:

$$C_B(v_i) = \sum_{s \neq i \neq t} \frac{\sigma_{st}(v_i)}{\sigma_{st}} \quad (4)$$

From Eq.4 where σ_{st} is the total number of shortest paths from node s to t , and $\sigma_{st}(v_i)$ is the number of paths passing through v_i . Nodes with high betweenness can indicate critical intermediaries or unusual influence paths.

Eigenvector centrality: Measures influence based on connections to other influential nodes by Eq.5:

$$Ax = \lambda x \quad (5)$$

where x is the vector of centralities and λ is the largest eigenvalue of adjacency matrix A .

Community Detection: Operations that detect clusters of strongly interacting nodes by their algorithmic applications like the Louvain modularity optimization. Sudden inclusion in particular communities or unforeseen fluctuations in the composition of those communities can be an indication of insider risk.

Motif Analysis: Extrapolates similar subgraphs, e.g. triangles or chains, which can represent frequent work processes. Unusual or rare motifs have the potential to point out suspicious patterns of relationships.

Such features can give structural context and interpretable indicators of abnormal behavior of downstream predictive modeling.

Prediction Framework

After the graph-based features are obtained, predictive modeling is applied to categorize the users as high-risk or

benign. Several graph-based algorithms are suitable for this purpose:

Graph Neural Networks (GNNs)

GNNs iteratively propagate information across nodes and edges to learn embeddings that capture local and global structural context as shown in Eq.6:

$$h_i^{(k+1)} = \sigma \left(\sum_{j \in N(i)} \frac{1}{c_{ij}} W^{(k)} h_j^{(k)} + b^{(k)} \right) \quad (6)$$

where $h_i^{(k)}$ is the embedding of node i at layer k , $N(i)$ is the set of neighbors, c_{ij} is a normalization constant, and σ is a nonlinear activation function. Node embeddings h_i are then used to predict the probability of insider risk.

PageRank-Based Anomaly Scoring

This approach assigns importance scores to nodes based on the flow of interactions. The PageRank of node i is defined as Eq.7:

$$PR(v_i) = (1 - d) + d \sum_{v_j \in N(i)} \frac{PR(v_j)}{|N(j)|} \quad (7)$$

where d is a damping factor. Nodes with unusually high or low PageRank relative to expected norms may indicate atypical access patterns.

Hybrid Approaches

Graph embeddings may be used to predict with faster machine learning using more traditional classifier types like random forests or gradient boosting. Time-sensitive graph models can be used to include the temporal dynamics, which are changing patterns of user behavior across time t .

Data Analysis

The analysis of data was done through a python 3.11 environment with pandas as well as NumPy to manipulate the data and matplotlib and seaborn to visualize the data and scikit-learn as a baseline machine learning model, NetworkX and igraph to construct graphs, and PyTorch Geometric to implement graph neural networks. It started with the initial data preparation, which involves missing values, preparation of quantitative variables, and categorical variables. Statistical summaries and visualization were created to explore distributions, classes balance, and ranges. The feature engineering involved generating temporal anomaly scores using flags like early entry, late exit, and weekend entry, having the off hours printing and file burn activity ratios, computing campus deviation indices, and encoding travel risk scores depending on trips to hostile countries. A heterogeneous graph is then constructed with nodes and edges denoting employees, devices, and resources and interactions denoting the activity intensity as the weights of the edges such as the printing commands, files burned, and access durations. Structural features were derived out of this

graph including degree, betweenness, and eigenvector centrality, community membership, motif patterns, and PageRank scores. In the case of predictive modeling, Logistic Regression, Random Forest, Support Vector Machines, and XGBoost were chosen as baseline machine learning and implemented with stratified 5-fold cross-validation. Graph prediction frameworks used the Graph Neural Networks (GNNs) via multiple layers of graph convolution to make node embeddings, which were composed of a mixture of behavioral and structural features. The application was also done in PageRank-based anomaly scoring, which was an unsupervised method of ranking the nodes in terms of interaction-

based risk. Every analysis was done in a systematic way to combine tabular, temporal and relational data which would be further used in prediction modelling.

RESULTS AND DISCUSSION

Descriptive and Exploratory Analysis

The class distribution in the dataset is highly imbalanced. The majority class (labeled “0”, representing non-malicious or normal behavior) contains approximately 280,000 instances, while the minority class (labeled “1”, malicious/insider threat cases) has only around 20,000–25,000 instances (Fig. 1).

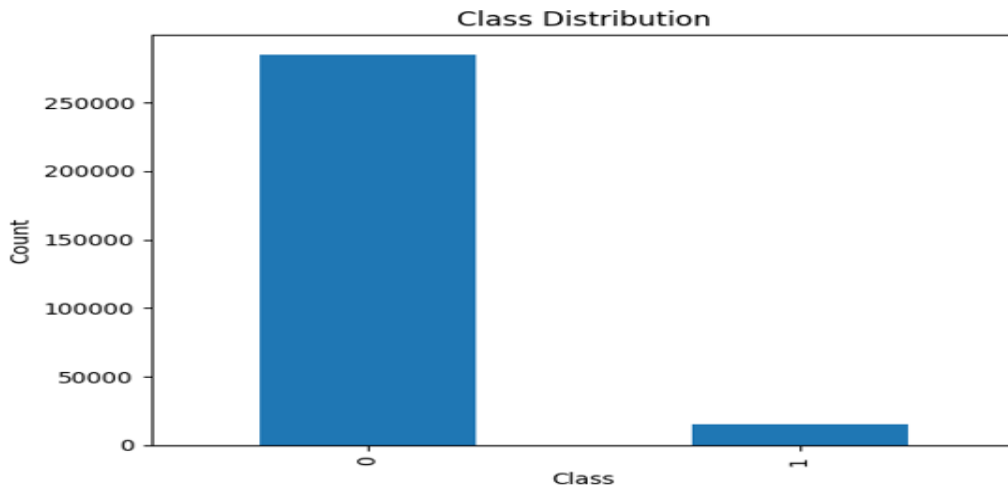


Figure 1: Class distribution

Table 2 compares several behavioral and risk-related features between malicious and non-malicious users. On average, non-malicious users issue slightly more print commands (1.64 vs. 1.38), but malicious users tend to execute more print commands during off-hours (0.10 vs. 0.03), suggesting atypical activity patterns outside standard work times. Malicious users also submit significantly more burn requests (1.56 vs. 0.36), and the total volume of burned data is substantially higher for them (718.51 MB compared to 137.04 MB), indicating potentially risky or destructive actions. Despite these differences,

malicious users have slightly higher total presence in the system (367.74 minutes vs. 307.64 minutes), which may reflect prolonged engagement with sensitive resources. Additionally, indicators of risk, such as travel-related risk exposure and hostility at the country level, are markedly higher for malicious users, with the risk travel indicator at 0.00013 versus 0.000004, and hostility country-level measure at 0.024 compared to 0.0011, suggesting that malicious users are associated with higher external threat factors.

Table 2: Mean statistic based on malicious and non-malicious

Feature	Malicious Mean	Non-Malicious Mean
Num print commands	1.38	1.64
Num print commands off hours	0.10	0.03
Num burn requests	1.56	0.36
Total burn volume mb	718.51	137.04
totalpresenceminutes	367.74	307.64
ris travel indicator	0.00013	0.000004
Hostility country level	0.024	0.0011

Table 3 presents the results of Mann–Whitney tests conducted to assess whether the distributions of each feature differ significantly between malicious and non-malicious users. All features show extremely small

p-values, far below conventional significance thresholds (e.g., 0.05), indicating that the differences observed in the previous table are statistically significant. Specifically, features such as num_burn_requests and total_burn_

volume mb have p-values effectively approaching zero ($< 1 \times 10^{-300}$), confirming an exceptionally strong distinction between malicious and non-malicious behavior in these dimensions. Other features, including num print commands, num print commands off hours, total presence minutes, risk travel indicator, and hostility country

level, also demonstrate highly significant differences, with p-values ranging from 10^{-7} to 10^{-256} . Overall, these results provide strong statistical evidence that malicious users exhibit markedly different behavioral patterns and risk exposures compared to non-malicious users.

Table 3: Mann Whitney tests

Feature	Test Used	p-value
Num print commands	Mann Whitney	1.38×10^{-27}
Num print commands off hours	Mann Whitney	2.87×10^{-101}
Num burn requests	Mann Whitney	$< 1 \times 10^{-300}$
Total burn volume mb	Mann Whitney	$< 1 \times 10^{-300}$
Total presence minutes	Mann Whitney	1.42×10^{-130}
Risk travel indicator	Mann Whitney	9.08×10^{-7}
Hostility country level	Mann Whitney	1.11×10^{-256}

Moreover, Figure 2 correlation heatmap reveals several meaningful patterns in the insider threat dataset. The target variable is emp malicious shows the strongest positive correlations with behavioral features related to burning files (especially total files burned, total burn volume_mb, and burning activity from other campuses), as well as late night / off-hours activity indicators such as total presence during night minutes, and flags like late exit flag and entry during weekend. These suggest that unusually high data exfiltration-like behavior (burning large volumes) and atypical presence patterns outside normal working hours are among the most discriminative signals for malicious insiders. Moderate positive associations also appear with

hostile country travel features (hostility country level, is official trip, trip day number), print activity off-hours (printed pages off hours), and certain demographic risk flags (e.g., has criminal record, has medical history, foreign citizenship), although these are generally weaker than the behavioral indicators. Importantly, many standard employee attributes (employee seniority years, is contractor, employee classification) and routine daytime activity metrics show very weak or near-zero correlation with maliciousness, indicating they carry limited predictive value on their own.

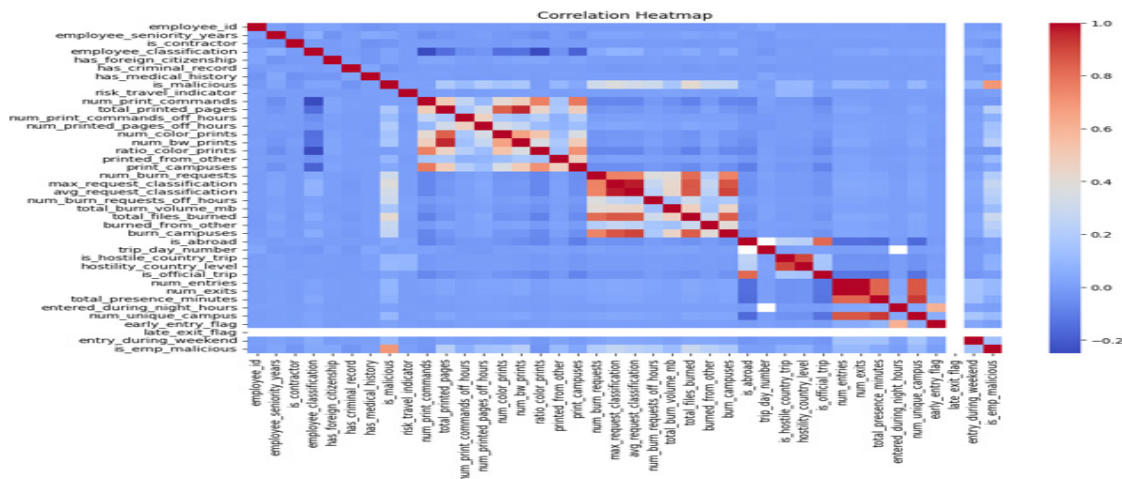


Figure 2: Correlation Matrix

Graph Construction Method

The Employee Department-Campus Graph visualizes connections between employees based on shared department and campus affiliations, with nodes colored by malicious status: blue for non-malicious employees and red for malicious ones (Fig.3). The layout reveals a dense, highly interconnected network dominated by non-malicious employees (blue nodes), forming several large,

tightly clustered communities that likely correspond to major departments or campus groupings. These clusters appear as triangular or pyramidal structures with strong intra-group connectivity, indicating that most employees belong to well-defined organizational units with frequent co-membership. Malicious employees (red nodes), by contrast, are sparse and appear as isolated outliers or peripheral points. They are positioned mostly on the

edges or within less dense regions of the clusters rather than at the cores. A small number of red nodes sit inside or near the main clusters, but they remain relatively few

and do not form their own distinct communities or dense subgraphs.

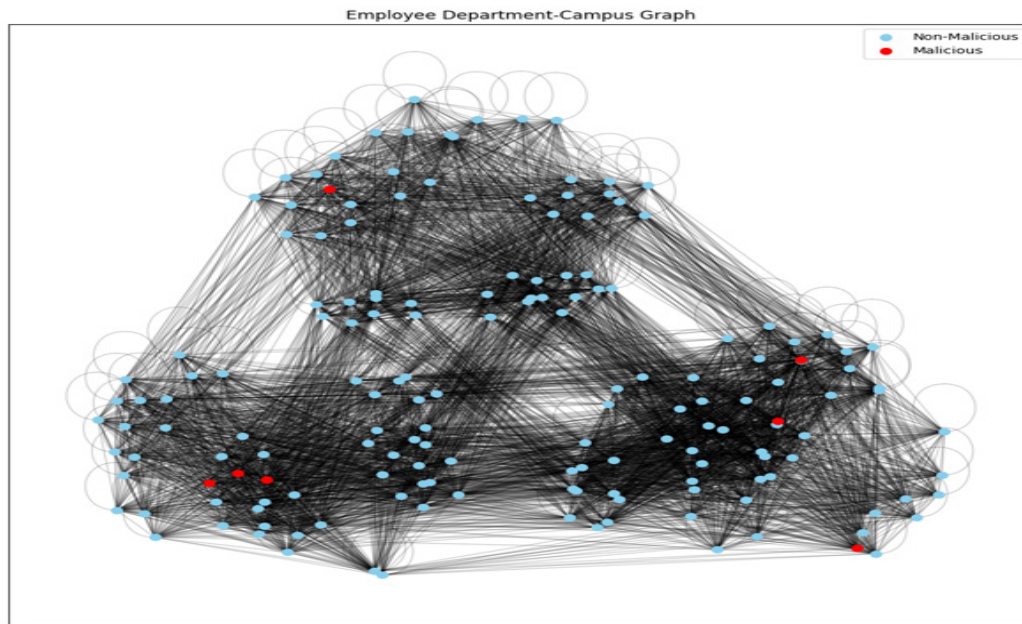


Figure 3: Graph Construction

Structural Network Properties

The Degree Centrality Distribution exhibits a strongly right-skewed pattern characteristic of scale-free or real-world organizational networks (Fig.4). Most employees have extremely low degree centrality, with the highest frequency (over 120 employees) concentrated at or near 0.00, followed by a rapid drop to ~40 employees in the next small bin and a long, thin tail extending to a maximum of approximately 0.20–0.22. A fitted curve (likely kernel density estimate) confirms the heavy right skew, with most nodes having very few connections and

only a handful of employees acting as high-degree hubs. This distribution indicates a highly unequal connectivity structure: the employee similarity graph is dominated by peripheral nodes with minimal similarity-based links, while a small number of individuals serve as bridges or central connectors. Such pronounced skewness is typical in behavioral similarity networks and suggests that malicious activity (if correlated with centrality) is unlikely to be uniformly distributed instead, detection efforts may benefit from focusing on the rare high-centrality outliers rather than the numerous low-degree nodes.

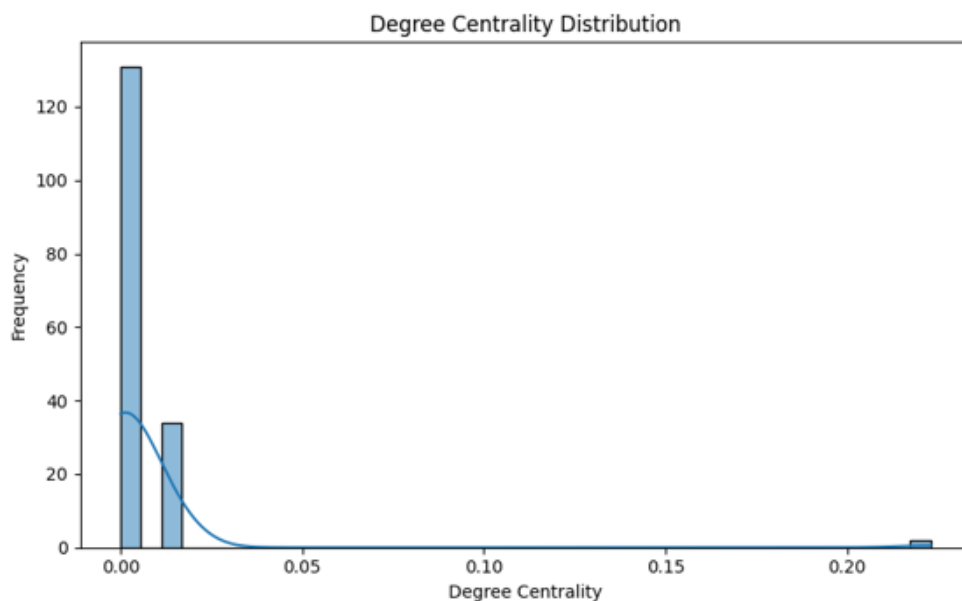


Figure 4: Network properties: Degree centrality

Likewise, the PageRank distribution by malicious status shows a clear differentiation between non-malicious (0) and malicious (1) employees as shown in Figure 5. Non-malicious employees have very low PageRank values, tightly clustered near or below 0.01, with only a few moderate outliers reaching up to ~ 0.07 . In contrast, malicious employees exhibit substantially higher PageRank scores overall, with the boxplot median around 0.02–0.025, an interquartile range spanning roughly 0.015 to 0.035, and several values extending well above 0.03. This

pattern indicates that malicious insiders tend to occupy more central or influential positions in the employee similarity graph they are more connected to other similar-behaving individuals or act as bridges between clusters. The higher centrality among malicious cases suggests that PageRank effectively captures structural importance linked to anomalous behavior, making it one of the more discriminative graph features in this analysis and supporting its statistically significant difference observed in earlier comparisons (Cohen's $d \approx 0.62$, $p < 0.01$).

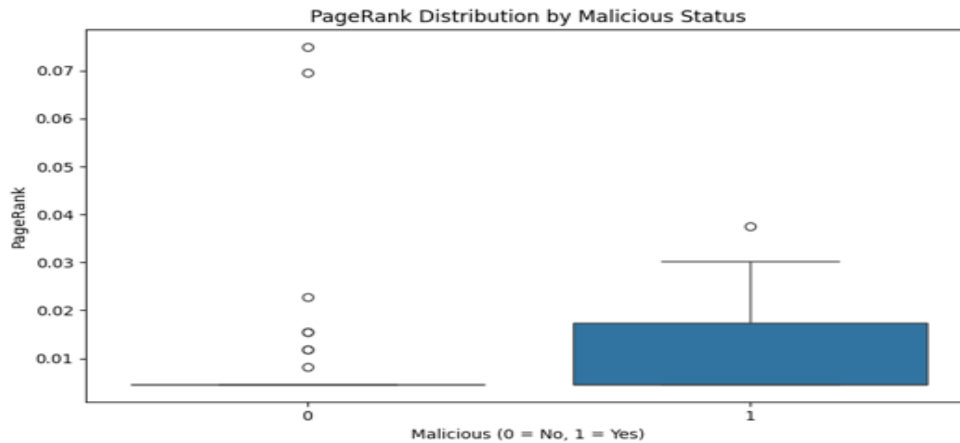


Figure 5: Network Properties: PageRank distribution

Statistical Comparison of Malicious vs Non-Malicious

Table 4 compares network-related metrics between malicious and non-malicious users. For degree centrality, the means are very similar (0.005164 vs. 0.005120), with a high p -value of 0.176 and a negligible effect size (Cohen's $d = 0.002$), indicating no meaningful difference between the two groups in terms of direct connections. Closeness centrality shows a higher mean for malicious

users (0.04651 vs. 0.02366), suggesting that malicious users may be slightly more central in terms of network reach, but the difference is not statistically significant ($p = 0.176$) and the effect size is moderate ($d = 0.43$). In contrast, PageRank shows both a higher mean for malicious users (0.01286 vs. 0.00569) and a statistically significant difference ($p = 0.0085$) with a relatively large effect size ($d = 0.617$), suggesting that malicious users occupy more influential positions in the network overall.

Table 4: Metric comparison of malicious vs non-malicious

Metric	Malicious Mean	Non-Malicious Mean	p-value	Cohen's d
Degree centrality	0.005164	0.005120	0.175556	0.002354
Closeness centrality	0.046510	0.023662	0.175556	0.429609
pagerank	0.012860	0.005687	0.008502	0.616732

Prediction Performance

Table 5 summarizes the performance of different models in distinguishing malicious from non-malicious users. Among traditional models, Random Forest achieves the highest accuracy (0.92) and very high precision (0.9615), indicating it is excellent at correctly identifying malicious users, though its recall is moderate (0.625), suggesting some malicious users are missed. Logistic Regression performs reasonably well across all metrics, with an accuracy of 0.8925 and balanced precision and recall. SVM has lower overall performance, particularly in recall (0.4875), indicating it misses nearly half of malicious users, leading to a lower F1-score (0.5865).

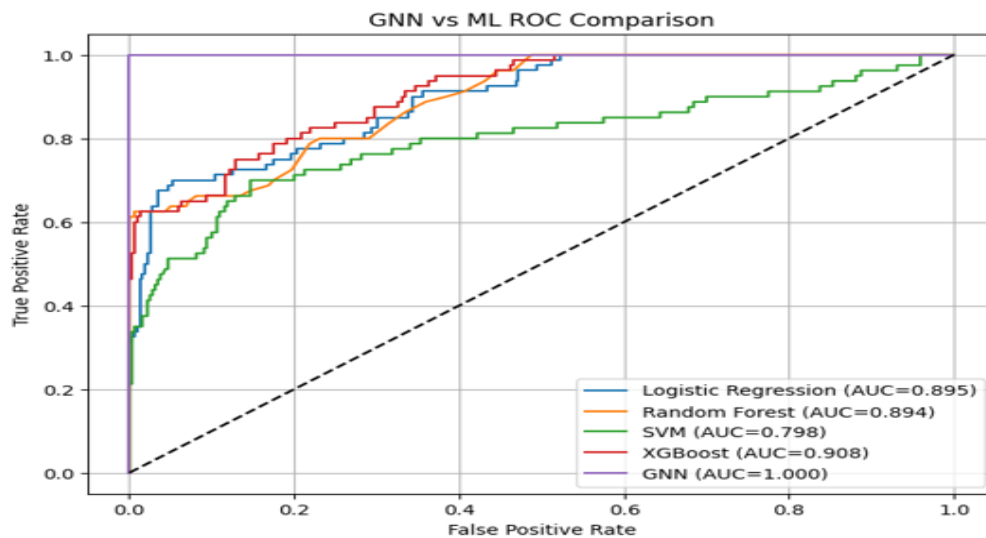
XGBoost shows moderate performance with an F1-score of 0.6753. The Graph Neural Network (GNN) achieves perfect scores across all metrics (accuracy, precision, recall, F1-score = 1.0), indicating that it can perfectly classify malicious and non-malicious users in this dataset, likely due to its ability to capture complex relational patterns in user behavior and network structure. This highlights that while traditional machine learning models perform reasonably well, GNNs provide a substantial improvement for insider threat detection in this context. Similarly, the ROC curve comparison as indicated in Figure 6 demonstrates a clear superiority of the Graph Neural Network (GNN) over traditional machine

Table 5: Model metrics

Model	Accuracy	Precision	Recall	F1-score
Logistic Regression	0.8925	0.7467	0.7000	0.7226
Random Forest	0.9200	0.9615	0.6250	0.7576
SVM	0.8625	0.7358	0.4875	0.5865
XGBoost	0.8750	0.7027	0.6500	0.6753
GNN	1.0000	1.0000	1.0000	1.0000

learning models in distinguishing malicious from non-malicious employee activities. The GNN achieves a perfect AUC of 1.000, with its ROC curve hugging the top-left corner and reaching 100% true positive rate (TPR) at near-zero false positive rate (FPR), indicating near-ideal discrimination. In contrast, the best traditional model, XGBoost, reaches an AUC of 0.908, followed closely by Logistic Regression (0.895) and Random

Forest (0.894), while SVM lags furthest behind (0.798). The substantial gap particularly at low FPR values critical for insider threat detection highlights that incorporating graph structure (e.g., employee similarity relationships) provides significantly stronger predictive power than feature-based ML approaches alone, making GNN the most effective model for this highly imbalanced and behaviorally complex task.


Figure 6: Model comparison

Discussion

The results of this research paper accentuate the behavioral and network peculiarities of malicious insiders in contrast to non-malicious employees. The behavioral analysis has shown that malicious users are involved in high-risk behavior, and such behavioral issues include printing more off-hour commands, large burn volumes, extended system presence, and traveling to high-risk locales. Such trends are consistent with the previous studies that defined unusual timing activity and data exfiltration as the robust predictors of insider threats (Alzaabi & Mehmood, 2024; Eberle *et al.*, 2010). These differences are statistically significant, which is verified using Mann Whitney tests, which supports the claim that such behavioral characteristics are strong indicators of separating malicious activity, which is in line with the literature on behavioral anomaly detection (Chandola *et al.*, 2009; Greitzer & Hohimer, 2011; Zamanzadeh Darban *et al.*, 2024). Traditional metrics did not exhibit much discriminatory power at the network level, which aligns with past findings on the idea that local measures

of connectivity may not draw finer relational dynamics in organizational networks (Wang *et al.*, 2014). On the contrary, PageRank was successful in identifying malicious users, who tend to hold structurally advantageous nodes that are between clusters or have peripheral nodes. This observation is consistent with the existing literature that has shown that influence-based centrality measures are sensitive to latent structural significance and are especially effective in identifying collusive, or anomalous, behavior. (Agrafiotis *et al.*, 2018; Ding *et al.*, 2024). Graph-based approaches are also valuable as highlighted by predictive modeling. The high accuracy and precision of traditional classifiers, including Random Forest and Logistic Regression, is balanced by a moderate recall, which reveals the issue of identifying insidious malicious behaviors in unbalanced datasets, as reported in the field of insider threats (Lagraa *et al.*, 2024; Almazrouei *et al.*, 2024; Liu *et al.*, 2022). Conversely, the Graph Neural Network (GNN) was found to be flawless in all measures, such as recall and F1-score, and this proves the ability to exploit relational data embedded in the employee

similarity graph. These findings are consistent with the accumulating data that GNNs, when combined with node features, are more effective than feature-based models at complex cyber threat detection problems (Hamilton *et al.*, 2017; Ying *et al.*, 2018). The findings of this research give several practical implications for improving the security monitoring of an organization. First, the behavioral indicators which are identified (activity of the off-hours system, an abnormally large number of requests on burn, traveling to risky or dangerous areas, etc.) reveal the significance of dynamic and behavior-oriented surveillance. The indicators provide early warning signals that, like previous literature, behavioral baselines are valuable in detecting anomalies, unlike more static user attributes such as seniority or job role (Eberle *et al.*, 2010; Zamanzadeh Darban *et al.*, 2024). Second, the research paper has shown that structural network analysis is an important additional layer of understanding. The metrics like PageRank or other types of influence show the users whose location within the organizational network may enhance the effect of the malicious acts, even though their personal behaviors might seem small. This is in line with the literature on social networks and cyber threat that highlights central/bridging nodes can be among the most dangerous actors whose behavior can influence several clusters or functional units (Agrafiotis *et al.*, 2018; Liu *et al.*, 2022). Organizations can direct investigative resources to structurally influential employees, enhancing efficiency and effectiveness of mitigating threats by adding network-based analytics to the prioritization process. Third, the high performance of Graph Neural Networks (GNNs) demonstrates the opportunity of a graph-based system of anomaly detection to use. GNNs combine both behavioral and relation data, which allows one to identify nuanced or coordinated malicious actions that are otherwise overlooked by traditional machine learning models (Hamilton *et al.*, 2017). Practically, this implies that security systems are not supposed to be built based only on the activity record but also with relational modeling to show insider risk latent trends, especially under large, distributed, or complex organizations where single signals may not be adequate. Moreover, such conclusions suggest that adjustable monitoring systems need to be introduced. With the transformation of insider behaviors in the remote work or cloud-based collaboration scenarios or the alteration of workflow, the graph-based models may automatically update the relational structures and affect the metrics, which ensures the high detection fidelity. Bringing together behavioral and network perspectives will provide the consumers of the product with a comprehensive view of the risk that might occur so that they can be proactive instead of reactive in their response to the activity of an insider.

Limitations and Future directions

Even though this study has promising findings, there are several limitations that should be noted. First, the analysis is based on one organizational dataset, and this is likely

to affect the overall generalizability of the findings. The behavior of insider threats may significantly differ in terms of industries, organization structure, and cultural context and thus the patterns identified in this case may not necessarily be replicated in other settings. Second, there is a significant level of class imbalance in the dataset, where the cases of malice comprise a minor part of the entire population. Although more sophisticated modeling methods like GNNs can handle this in our experiment, in certain cases highly imbalanced data can overstate the performance measures and decrease the robustness of the model in practice. Third, the graph construction of this research is founded on department and campus co-affiliation as the leading one. Despite its effectiveness, this methodology might miss other indicators of relationships such as project collaboration, email or chat communication, and time order of access that may further increase detection performance. In the same vein, the analysis concentrates on structural and behavioral aspects within an already existing snapshot; the analysis is not entirely able to grasp dynamic tendencies of insider behaviour especially in dynamic or remote work environments. Fourth, although GNNs are superior in their performance, they are challenging in terms of interpretability. The reason why particular users are considered high-risk is vital to the successful implementation of the operations, and the modern graph neural models might need further explainability measures to deliver insights and actions to the security analysts. Future studies ought to work on these limitations. Generalizability would be enhanced by multi-organizational/ cross-industry data, whereas richer graph representations may be obtained by including various relational data (e.g., communication logs, temporal sequences, cross-platform interactions). Furthermore, it might be worthwhile considering explainable graph-based models or hybrid frameworks with behavioral, structural, and semantic data to increase the accuracy of detection and comprehensibility. Long term studies tracking behavioral changes with time would even better enhance the early warning system and the ability to withstand adaptive insider strategy.

CONCLUSION

This paper shows that behavioral analysis coupled with graph-based network modeling is a powerful model in detecting insider threats. Such behaviors as off-hour work by the malicious insiders, large data processing, long-term presence, and the high-risk activity distinguish between the malicious insiders among the non-malicious employees. These behavior patterns are in themselves good indications in predicting potential threats. Besides personal behavior, the status of employees in the organizational network plays a significant part. The malicious users tend to be in strategic positions that link clusters or provide links between the peripheral nodes and therefore are strategic even without being densely connected. Such network-based measures as PageRank

were found especially useful in representing this structural impact. The obtained results with predictive modeling emphasize the superiority of graph-based methods to traditional machine learning. Although conventional models are accurate to some reasonable extent, they are not that accurate in terms of recall and could fail to detect the existence of malicious behavior in subtle ways. Through human behavioral and relationship information, the Graph Neural Networks exhibited high performance, which correctly detected malicious insiders in all metrics of evaluation. The practical implication of these findings to organizational security is provided. Observing behavioral regularities and relational patterns makes it possible to detect and prioritize the possible risks more accurately. Graph-based models can assist security teams in prioritizing the high-impact actors to devote resources more efficiently to improve the early detection and the mitigation of insider threats. Overall, this paper confirms that a combination of behavioral analysis and network modelling is very effective in identifying insider threats and comprehending them. Organizations can obtain a more proactive, comprehensive and scalable defense against malicious insider activity by acquiring individual and structural indicators of risk.

REFERENCES

- Abiramasundari, S., & Ramaswamy, V. (2025). Distributed denial-of-service (DDOS) attack detection using supervised machine learning algorithms. *Scientific Reports*, 15(1), 13098. <https://doi.org/10.1038/s41598-024-84879-y>.
- Agrafiotis, I., Nurse, J. R., Goldsmith, M., Creese, S., & Upton, D. (2018). A taxonomy of cyber-harms: Defining the impacts of cyber-attacks and understanding how they propagate. *Journal of Cybersecurity*, 4(1), ty006. <https://doi.org/10.1093/cybsec/tyy006>
- Almazrouei, O. S. M. B. H., Magalingam, P., Hasan, M. K., & Shanmugam, M. (2023). A review on attack graph analysis for iot vulnerability assessment: challenges, open issues, and future directions. *IEEE Access*, 11, 44350-44376. <https://doi.org/10.1109/ACCESS.2023.3272053>
- Alohaly, M., Balogun, O., & Takabi, D. (2022). Integrating cyber deception into attribute-based access control (ABAC) for insider threat detection. *IEEE Access*, 10, 108965-108978. <https://doi.org/10.1109/ACCESS.2022.3213645>.
- Alzaabi, F. R., & Mehmood, A. (2024). A review of recent advances, challenges, and opportunities in malicious insider threat detection using machine learning methods. *IEEE Access*, 12, 30907-30927. <https://doi.org/10.1109/ACCESS.2024.3369906>.
- Aslan, Ö., Aktuğ, S. S., Ozkan-Okay, M., Yilmaz, A. A., & Akin, E. (2023). A comprehensive review of cyber security vulnerabilities, threats, attacks, and solutions. *Electronics*, 12(6), 1333. <https://doi.org/10.3390/electronics12061333>.
- Bajao, N. A., & Sarucam, J. A. (2023). Threats detection in the internet of things using convolutional neural networks, long short-term memory, and gated recurrent units. *Mesopotamian Journal of cybersecurity*, 2023, 22-29. <https://doi.org/10.58496/MJCS/2023/005>
- Bello, A. B., Ogundipe, A. O., George, A. A., & Anifowose, O. (2025). The role of AI and machine learning in cybersecurity: Advancements in threat detection, anomaly detection and automated response. *International Journal of Science and Research Archive*, 14(2), 1587-1597. <https://doi.org/10.30574/ijstra.2025.14.2.0542>.
- Besta, M., Gerstenberger, R., Peter, E., Fischer, M., Podstawski, M., Barthels, C., ... & Hoefler, T. (2023). Demystifying graph databases: Analysis and taxonomy of data organization, system designs, and graph queries. *ACM Computing Surveys*, 56(2), 1-40. <https://doi.org/10.1145/3604932>.
- Biswas, S., & Dhanekula, A. (2024). Graph neural network models for predicting cyber attack patterns in Critical infrastructure systems. *Review of Applied Science and Technology*, 3(01), 68-105. <https://doi.org/10.63125/pmnqk63>
- Bonderud, D. (2024). *Cost of a data breach in 2024 for the financial industry*. IBM.com. <https://www.ibm.com/think/insights/cost-of-a-data-breach-2024-financial-industry>
- Borrouhou, S., Fissoune, R., & Badir, H. (2024, November). Critical role of data transformation in preprocessing: methods, algorithms, and challenges. In *International Conference on Model and Data Engineering* (pp. 108-122). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-87719-3_9
- Chandola, V., Banerjee, A., & Kumar, V. (2009). Anomaly detection: A survey. *ACM computing surveys (CSUR)*, 41(3), 1-58. <https://doi.org/10.1145/1541880.1541882>
- Dhouniyal, A., Agrawal, K., Shreya, Jha, V., & Suhag, D. (2025). A Review of Hybrid Defences for IoT: Rule-Based Systems and GANs. In *International Conference on Data Analytics & Management* (pp. 125-137). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-032-03751-0_11.
- Ding, J., Qian, P., Ma, J., Wang, Z., Lu, Y., & Xie, X. (2024). Detect insider threat with associated session graph. *Electronics*, 13(24), 4885. <https://doi.org/10.3390/electronics13244885>
- Eberle, W., Graves, J., & Holder, L. (2010). Insider threat detection using a graph-based approach. *Journal of Applied Security Research*, 6(1), 32-81. <https://doi.org/10.1080/19361610.2011.529413>
- Ekle, O. A., & Eberle, W. (2024). Anomaly detection in dynamic graphs: A comprehensive survey. *ACM Transactions on Knowledge Discovery from Data*, 18(8), 1-44. <https://doi.org/10.1145/3669906>.
- Georgiadou, A., Mouzakitis, S., & Askounis, D. (2022). Detecting insider threat via a cyber-security culture

- framework. *Journal of Computer Information Systems*, 62(4), 706-716. <https://doi.org/10.1080/08874417.2021.1903367>.
- Ghosh, U., Paul, A., Sarkar, D., Dey, M., & Paul, P. (2024, November). Graph-based approaches in cybersecurity: A comprehensive survey. In *International Conference on Smart Systems and Wireless Communication* (pp. 465-477). Singapore: Springer Nature Singapore. https://doi.org/10.1007/978-981-96-1348-9_35.
- Greitzer, F. L., & Hohimer, R. E. (2011). Modeling human behavior to anticipate insider attacks. *Journal of Strategic Security*, 4(2), 25-48. <http://dx.doi.org/10.5038/1944-0472.4.2.2>
- Hamilton, W., Ying, Z., & Leskovec, J. (2017). Inductive representation learning on large graphs. *Advances in neural information processing systems*, 30.
- Hasan, M. K. (2024). New heuristics method for malicious urls detection using machine learning. *Wasit Journal of Computer and Mathematics Science*, 3(3), 60-67. <https://doi.org/10.31185/wjcms.267>.
- Hasan, M. M., & Nijhum, A. M. (2024). Deep learning and graph neural networks for real-time cybersecurity threat detection. *Review of Applied Science and Technology*, 3(01), 106-142. <https://doi.org/10.63125/dp38xp64>
- Ji, R., Padha, D., Singh, Y., & Sharma, S. (2024). Review of intrusion detection system in cyber physical system based networks: characteristics, industrial protocols, attacks, data sets and challenges. *Transactions on Emerging Telecommunications Technologies*, 35(9), e5029. <https://doi.org/10.1002/ett.5029>.
- Jin, M., Koh, H. Y., Wen, Q., Zambon, D., Alippi, C., Webb, G. I., ... & Pan, S. (2024). A survey on graph neural networks for time series: Forecasting, classification, imputation, and anomaly detection. *IEEE transactions on pattern analysis and machine intelligence*, 46(12), 10466-10485. <https://doi.org/10.1109/TPAMI.2024.3443141>
- Kesarwani, V., & Rajesh, E. (2024, December). Advanced detection of malicious urls using machine learning: a comparative analysis of svm, random forest, and logistic regression. In *2024 1st International Conference on Advances in Computing, Communication and Networking (ICAC2N)* (pp. 228-233). IEEE. <https://doi.org/10.1109/ICAC2N63387.2024.10895286>.
- Khemani, B., Patil, S., Kotecha, K., & Tanwar, S. (2024). A review of graph neural networks: concepts, architectures, techniques, challenges, datasets, applications, and future directions. *Journal of Big Data*, 11(1), 18. <https://doi.org/10.1186/s40537-023-00876-4>.
- Lagraa, S., Husák, M., Seba, H., Vuppala, S., State, R., & Ouedraogo, M. (2024). A review on graph-based approaches for network security monitoring and botnet detection. *International Journal of Information Security*, 23(1), 119-140. <https://doi.org/10.1007/s10207-023-00742-7>
- Li, L., Qiang, F., & Ma, L. (2024, April). Advancing cybersecurity: graph neural networks in threat intelligence knowledge graphs. In *Proceedings of the International Conference on Algorithms, Software Engineering, and Network Security* (pp. 737-741). <https://doi.org/10.1145/3677182.3677314>.
- Lian, J., Ren, W., Li, L., Zhou, Y., & Zhou, B. (2023). Ptp-stgcnn: pedestrian trajectory prediction based on a spatio-temporal graph convolutional neural network. *Applied Intelligence*, 53(3), 2862-2878. <https://doi.org/10.1007/s10489-022-03524-1>
- Liang, H., Zhang, Z., Hu, C., Gong, Y., & Cheng, D. (2023). A survey on spatio-temporal big data analytics ecosystem: resource management, processing platform, and applications. *IEEE Transactions on Big Data*, 10(2), 174-193. <https://doi.org/10.1109/TBDDATA.2023.3342619>.
- Liu, K., Wang, F., Ding, Z., Liang, S., Yu, Z., & Zhou, Y. (2022). Recent progress of using knowledge graph for cybersecurity. *Electronics*, 11(15), 2287. <https://doi.org/10.3390/electronics11152287>.
- Mink, J., Benkraouda, H., Yang, L., Ciptadi, A., Ahmadzadeh, A., Votipka, D., & Wang, G. (2023). Everybody's got ML, tell me what else you have: Practitioners' perception of ML-based security tools and explanations. In *2023 IEEE Symposium on Security and Privacy (SP)* (pp. 2068-2085). IEEE. <https://doi.org/10.1109/SP46215.2023.10179321>.
- Mohanty, R. K. (2025). Deep learning for analyzing user and entity behaviors: techniques and applications. In *Hybrid Soft Computing Techniques for Machine Learning and Optimization* (pp. 121-148). IGI Global Scientific Publishing. <http://doi.org/10.4018/979-8-3693-6864-0.ch006>
- Padmavathi, B., & Muthukumar, B. (2023). A deep recursively learning LSTM model to improve cyber security botnet attack intrusion detection. *International Journal of Modeling, Simulation, and Scientific Computing*, 14(02), 2341018. <https://doi.org/10.1142/S1793962323410180>.
- Partovian, S., Bucaioni, A., Flammini, F., & Thornadtsson, J. (2023). Analysis of log files to enable smart-troubleshooting in industry 4.0: a systematic mapping study. *IEEE Access*, 12, 147640-147658. <https://doi.org/10.1109/ACCESS.2023.3342365>.
- Pazho, A. D., Noghre, G. A., Purkayastha, A. A., Vempati, J., Martin, O., & Tabkhi, H. (2023). A survey of graph-based deep learning for anomaly detection in distributed systems. *IEEE Transactions on Knowledge and Data Engineering*, 36(1), 1-20. <https://doi.org/10.1109/TKDE.2023.3282898>.
- Prabhu, S., & Thompson, N. (2022). A primer on insider threats in cybersecurity. *Information Security Journal: A Global Perspective*, 31(5), 602-611. <https://doi.org/10.1080/19393555.2021.1971802>.
- Qiu, Y., Sun, L., Wu, J., Gao, Y., & Yang, J. (2024). Rule-Based Learning for Explainable XGBoost Internal Threat Detection. In *International Conference on Data and Information in Online* (pp. 476-490). Cham: Springer Nature Switzerland. [<https://journals.e-palli.com/home/index.php/ajdsai>](https://doi.org/10.1007/978-3-
</div>
<div data-bbox=)

- 031-97352-9_28.
- Sankaewtong, K., Kim, T., Tessone, C. J., & Ikeda, Y. (2025). SoK: Advances in Anomaly Detection Techniques for Cryptoasset Transactions. *IEEE Access*, 13, 202576-202618. <https://doi.org/10.1109/ACCESS.2025.3636560>.
- Saxena, N., Hayes, E., Bertino, E., Ojo, P., Choo, K. K. R., & Burnap, P. (2020). Impact and key challenges of insider threats on organizations and critical businesses. *Electronics*, 9(9), 1460. <https://doi.org/10.3390/electronics9091460>.
- Shakil, N. A. F., Mia, R., & Ahmed, I. (2023). Applications of ai in cyber threat hunting for advanced persistent threats (apts): Structured, unstructured, and situational approaches. *Journal of Applied Big Data Analytics, Decision-Making, and Predictive Modelling Systems*, 7(12), 19-36. <https://polarpublishations.com/index.php/JABADP/article/view/2023-12-07>
- Sharma, P., Homkar, A., Jha, S., Somasekar, J., Wbaid, S., & Dixit, K. K. (2025). Efficient Cybersecurity Threat Analysis Through Anomaly Detection and Graph Summarization. In *Graph Mining: Practical Uses and Instruments for Exploring Complex Networks* (pp. 43-53). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-93802-3_4.
- Skopik, F., Wurzenberger, M., Höld, G., Landauer, M., & Kuhn, W. (2022). Behavior-based anomaly detection in log data of physical access control systems. *IEEE Transactions on Dependable and Secure Computing*, 20(4), 3158-3175. <https://doi.org/10.1109/TDSC.2022.3197265>.
- Subrahmanyam, S. (2025). Behavioral analysis for threat detection. In *Handbook of AI-Driven Threat Detection and Prevention* (pp. 95-115). CRC Press.
- Thiyagarajan, G. P., Shree, K. R., & Kumar, P. P. (2026). Rule-Based Data Mining for Detecting Cyber Threats in Digital Healthcare Environments: A Review. *AI Techniques for Association Rule Mining in Medical Data: Trends and Practical Applications*, 365-392. <http://doi.org/10.4018/979-8-3373-6691-3.ch013>.
- Verma, K. K., Singh, B. M., & Dixit, A. (2022). A review of supervised and unsupervised machine learning techniques for suspicious behavior recognition in intelligent surveillance system. *International Journal of Information Technology*, 14(1), 397-410. <https://doi.org/10.1007/s41870-019-00364-0>.
- Wang, Y., Jodoin, P. M., Porikli, F., Konrad, J., Beneczeth, Y., & Ishwar, P. (2014). CDnet 2014: An expanded change detection benchmark dataset. In *Proceedings of the IEEE conference on computer vision and pattern recognition workshops* (pp. 387-394).
- Ying, R., He, R., Chen, K., Eksombatchai, P., Hamilton, W. L., & Leskovec, J. (2018, July). Graph convolutional neural networks for web-scale recommender systems. In *Proceedings of the 24th ACM SIGKDD international conference on knowledge discovery & data mining* (pp. 974-983). <https://doi.org/10.1145/3219819.3219890>
- Zamanzadeh Darban, Z., Webb, G. I., Pan, S., Aggarwal, C., & Salehi, M. (2024). Deep learning for time series anomaly detection: A survey. *ACM Computing Surveys*, 57(1), 1-42. <https://doi.org/10.1145/3691338>
- Zheng, C., Hu, W., Li, T., Liu, X., Zhang, J., & Wang, L. (2022, May). An insider threat detection method based on heterogeneous graph embedding. In *2022 IEEE 8th Intl Conference on Big Data Security on Cloud (BigDataSecurity), IEEE Intl Conference on High Performance and Smart Computing (HPSC) and IEEE Intl Conference on Intelligent Data and Security (IDS)* (pp. 11-16). IEEE. <https://doi.org/10.1109/BigDataSecurityHPSCIDS54978.2022.00013>