

**AI-BASED BEHAVIOURAL ANALYSIS FOR INSIDER THREAT DETECTION IN
CORPORATE NETWORKS**

Harshith Kumar Pedarla
Software Developer, Amazon
Seattle, USA
hpedarla@amazon.com

Abstract

Insider threats remain one of the most difficult challenges in corporate cybersecurity because they arise from legitimate users with valid access, making them harder to detect than external attacks. Past incidents show that traditional rule-based systems often fail to uncover these risks in time, leaving organizations vulnerable. This paper explores how AI-driven behavioural analysis, with particular attention to the role of generative AI, can improve insider threat detection within enterprise networks. By learning baseline user behavior, generating simulated threat scenarios, and applying anomaly detection, AI systems can identify subtle deviations that may suggest malicious intent or accidental misuse. Perimeter defences such as firewalls are still effective for blocking external attacks, but they provide little protection once an insider is already within the system. Combining generative AI with machine learning, outlier detection, and natural language processing enables organizations to uncover hidden patterns more effectively and respond faster to emerging threats. These approaches support adaptive, proactive defence mechanisms that evolve with organizational behavior. The findings suggest that behavioural AI can strengthen resilience, reduce detection delays, and provide more flexible security strategies for complex network environments.

Keywords: Cybersecurity, Insider threat detection, Generative AI, Behavioural modelling, Anomaly detection, Adaptive defence systems.

I. INTRODUCTION

Modern business networks operate in an environment of continuous risk, where insider threats remain among the hardest to manage. Unlike external attacks, which can often be blocked by perimeter defences, insider threats may arise from disgruntled employees, careless users, or compromised accounts that already possess valid credentials. When insiders have legitimate access, they can often bypass traditional security controls, which makes detection far more complicated than with external threats. Past breaches clearly demonstrate the scale of the problem ranges from sensitive data leaks to the theft of intellectual property and, in some instances, deliberate acts of sabotage. What is striking is that many of these incidents were only recognized after serious harm had already been done, underscoring the gap between prevention and timely detection. Conventional security tools that depend on static rules or

signature patterns struggle in this area, since they are rarely effective against slow, adaptive insider behaviors. As enterprise networks continue to expand in both scale and complexity, the pressure to adopt more adaptive security approaches has become unavoidable. Artificial Intelligence, and especially Generative AI (GenAI), introduces new possibilities by learning what constitutes normal behavior, flagging deviations, and enabling faster detection of insider activity in real time [1].

A. Insider Threats in Corporate Networks

Because technology is changing so quickly these days, businesses are dealing with a cyber-landscape that is more diverse than ever. Attacks are coming from both outside and inside the company. Insider threats are threats that come from people who work for or with a company, such as coworkers, contractors, or third-party vendors who misuse their access rights or abuse the company's systems. This can put the integrity, confidentiality, or availability of the data they manage at risk. They could be put in place on purpose by people who don't care about the company or want to harm it, or they could happen by accident when employees share information that the company wants to keep private while trying to do their jobs, or when sensitive data that isn't well protected is shared by mistake. Over the past two years, incidents of insider threats have gone up by 44%, and they now cost businesses an average of \$15.38 million a year [1].

Firewalls, intrusion detection systems (IDS), and security information and event management (SIEM) platforms are all examples of traditional security solutions that don't do much to stop insider threats. This is because insiders usually log in with valid credentials, which means they have the same identity as regular people using the network facilities. This makes their actions blend in with normal network operations. This necessitates the implementation of advanced methodologies to identify dubious behavioural patterns indicative of insider misuse [2].

B. Role of AI in Threat Detection

The inclusion of artificial intelligence technology can augment the detection of insider threats, as the AI-centred behavioral analysis shows a lot of promise. It is observed that the AI methods can work efficiently since they perform an insanely comprehensive set of data on user activities, allowing for tracking even the slightest inconsistency in behavior over a diverse range of applications serving both old and new emerging security obstacles. Individuals can also employ machine learning techniques for mathematical data analysis through procedures like supervised classification, also referred to as anomaly detection, to oversee the levels of different forms while comparing and combining level differences. That model categories mentioned leverage this type of AI to monitor user behaviors to gather information systems' Level-official enactment of private policies compared to what was initially indicated with due professional care and diligence, and the ability to establish or reconcile compliant discrepancies [2].

Artificial generative intelligence (GenAI) pushes the boundaries of AI. It automates data processing using methods that include the creation of artificial user scenarios, as well as natural

language analysis and automated model reasoning. GenAI techniques can be used to model thousands of insider threat scenarios that can be used to augment real-world training data and sparse examples that we may have. An additional example is LLMs that can be used to provide contextual insight from unstructured data sources such as incident reports and emails [3].

C. Networking Perspective

When it comes to the identification of internal threats, which is a very complex process, the networking structure of businesses plays an enormous part. Among all types of data networks, the most important resources for identifying unusual behavior at the human operational level are as follows: network parameters and packet tracking in various forms. The upkeep of distributed systems, notwithstanding its utter complexity, becomes inoperable with the recognizable shift of partial cloud architecture and home-based jobs. If there are any deviations (no matter how slight) from the network's regular hosting, artificial intelligence or “machine learning” sets an acceptable standard called a pattern. The difference between these disparities is highlighted in any of the following fields, providing such examples: illegal transportation of data, strange logins, and unauthorized access requests. In other words, the intersection of AI and networking hints that the current control of internal threats is showing a movement. With the integration of AI and networking, organizations could recognize it and shut it down before it progresses to another massive data breach [3].

II. BACKGROUND AND LITERATURE REVIEW

Few DLP and SIEM, like Data Loss Prevention and Security Information, are the same tools; still, they generate so many false positive alarms for monitoring purposes. To ameliorate the accuracy, the machine learning researchers brought forward numerous obscures, average degree of specialization, and ordinary classification algorithms. Nevertheless, many algorithms still yield merger results, mainly because the databases used in this kind of attack often prove to be inequitable. That is the moment Generative AI is attacking a fight by cracking out algorithms capable of mimicking an insider attack. Meanwhile, the AI-supported behavioural analytics mechanisms are to allow the transient access to emails, file transfers, sessions of cloud access, and endpoint operations to keep up a constant dynamic.

A. The Insider Threat Landscape

Insider threats differ from external cyber-attacks in terms of point of origin, in that, in the case of dangers within organizations, these threats arise from people who already have legitimate login details and can access data centres, apps, and physical premises without any additional measures [3]. Generally, insider threats are mainly classified into three main branches/collections as follows; malicious insiders, who are in one way or another intentionally committed to causing harm, such as using your resources for personal or financial reasons among others, and negligent insiders, who are unintentional errors may will compromise by purposeful outside instead accidental actions by lazy or improperly trained personnel [4].

To accurately define insider risks, one must also include dormant or hibernating insiders that may have been exploited by an adversary to carry out clandestine actions [4]. The problem is significant because in the digital domain, there are no measures that could be used to ensure that access control has been appropriately implemented or that data has been properly encrypted. An identification of novel threats was made by an employee at the NSA (National Security Agency) as part of the ecosystem generation [4].

Furthermore, the 2019 Capital One data breach, which was caused by a former employee who could still access Capital One's cloud system due to its misconfiguration or inadequate training, is another instance of what knowledge of an organization's practices can enable an insider to cause harm on a large scale [4]. In 2023, the Ponemon Institute reported that the annualized cost of insider-related incidents had hiked to \$15.38 million from \$11.45 million in 2020, with financial services and healthcare being largely impacted. The constant evolution of insider threats and their consequences make it so that, in the monetary realm, they are quickly becoming a major research priority for cybersecurity scholars and field professionals [4].

B. Traditional Approaches to Insider Threat Detection

Originally, efforts to prevent insider attacks largely involved rule-based monitoring and signature detection. Early in the cybersecurity industry, administrators created static rules to identify suspicious activity, such as a user entering their password incorrectly three times or their logon attempt being made during a time when everyone was at the end of the day, which were only activated during that specific period. Insider attacks, where malicious actions are carefully disguised as regular user activities, can be overlooked by traditional rules due to their subtlety and often go undetected by standard user behavior analytics [5]. Two of the adoption of a SIEM approach.

Event correlation is present in SIEM technologies and is a primary driver for its implementation, as event correlation and data aggregation is used to identify and quantify the impact of an incident on a business objective for the management and stakeholders and thus correlate back to event correlation. However, the correlation of events in the SIEM results in large amounts of alarm traffic and, as such, can lead to operator burnout and less secure environments by configuring SIEM to reduce the volume of alarms.

This is known as tuning, whereby the SIEM alerts are cultured to be more resilient, and therefore more difficult for an outside entity to trigger [5]. In addition, anomaly detection can be used to determine the extent of the insider compromise and identify the insider's behavior. However, this is quite difficult in practice because there are many networked applications that are completely out of the enterprise's control [5].

C. Emergence of AI and Machine Learning in Insider Detection

The realm of Artificial Intelligence has become a burgeoning tool when it comes to monitoring insider threats. The primary machine learning techniques enable automated pattern recognition

from large-scale datasets. Supervised and unsupervised learning models have been used to classify insider behaviors when labelled datasets are available. However, these models have some limitations, such as a small dataset size and imbalanced datasets. Unsupervised learning and semi-supervised learning methods have provided more promising results. Clustering algorithms are commonly used in the space of data mining, and anomaly detection is a widespread use case for them. In the methods of data mining and anomaly detection, unsupervised learning methods can be successfully used [6].

Autoencoders and Recurrent Neural Networks (RNNs) are two examples of deep learning methods that are often used to find anomalies in log data that has sequential data. For example, Long Short-Term Memory (LSTM) networks use this method to find time-based connections. So, the model can find strange login times, data transfers, or command sequences [6]. Also, this is a good place to use reinforcement learning (RL). In the reinforcement learning framework, RL agents can leverage previous interactions with the dynamic network environment to enhance the detection strategy progressively. RL-based insider threat detection is still in its early stages, but some initial work has been done in this area [6].

D. Generative AI in Insider Threat Detection

Generative AI is one of the most important things to think about when it comes to cybersecurity apps, which are always changing. Generative AI includes GANs and LLMs, which are new. Previously, we had Machine learning models that could not create; they were merely a logic of every possible feature of the input data.

- GANs can craft a totally synthetic insider activity log that is similar to real ones. This means that GANs can help alleviate a chronic shortage of custom-labeled insider threat data, which in turn will reduce over-fitting and increase a model's generalizability [7].
- Large language models (LLMs) such as the anticipated GPT-4 can be deployed to analyze unstructured data sources (eg, emails, chat messages, and incident reports) to look for linguistic signals of potentially malicious intent. However, this advanced.
- NLP model raises privacy issues regarding eavesdropping on all electronic communications [7].
- The generative AI is something through which you will improve the efficiency report of the employees in your company, not only through emails, but it will be there to explain its functionalities, the anomalous detections, and other things naturally. This will be there to improve the efficiency of your work [6].

It can utilize anomalous detections present in the environment to demonstrate how any of the detection tools can be used to identify them. The ancient intelligence technique can show the standard detection tools and things on how they can evolve to adapt to the evasions there. It can show how intrusion detection can be used to point out all the anomalies, not just the known ones.

Networking telemetry can form part of the Generative AI knowledge database, providing valuable application-layer insight to SOC monitoring, enriching areas such as user and

application behavior, application fingerprinting, and deeper network and application security at both the WAN and campus.

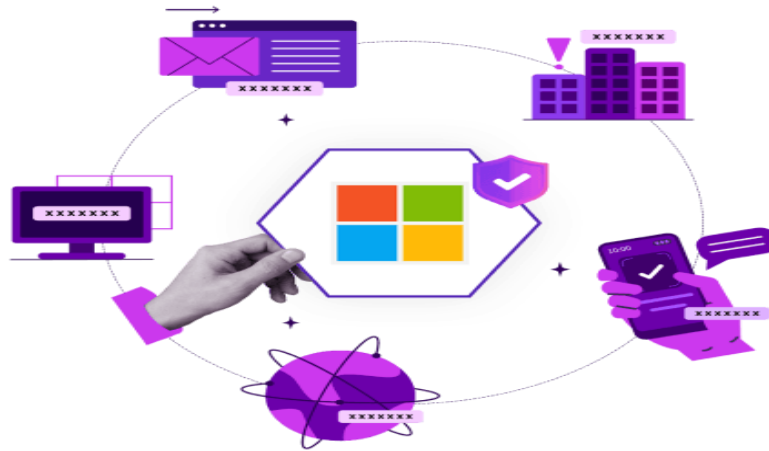


Fig. 1. Zero Threat Solution Security from ZERO Threat website

III. METHODOLOGY AND PROPOSED FRAMEWORK

A. Behavioural Analysis Framework

A foreseeable insider threat monitoring that uses Artificial Intelligence will be activated by the gathering of data from the system log, cloud services access log, endpoint activity, and network traffic streaming (system log, cloud services access log, endpoint activity, and network traffic streaming). The system, which is built up from the generative models, is put together with a system that is both unsupervised and supervised, with a heterogeneous database, to obtain an accurate picture in a comprehensive manner of internet (that is, internet owned by the business) flows for individual users within the corporate network [7].

Lastly, based on the reviewed materials, the new user behaviour with the known baselines from previously constructed models to generate the results for the possible good entries, with the basic idea that deviations should be indicated as a suspicious nature of the insider behaviour (that is, as an insider signal). After the phase, the businesses are defined as sensitive due to the services, proper IT software, and anything else that has to ensure their correct functionality, verified in terms of exemptions and any other, even those suspicious ones are going into the risk level concept to which they are expected to put anything on security issues under more consideration [8].

B. Research Approach

A theoretical framework will be built to give a structural design model that will fuse AI cerebral behavior, Analytics, and the generation AI (GenAI) capabilities when creating an Insider Threat Detection system for corporate networks. Insightful research, which is applicable in machine learning, has been applied in areas such as anomaly detection, covering the behavior of users and entity analytics (UEBA), while opening up the areas of artificial intelligence, including the

GenAI approach driven by synthetic data generation, followed by natural language processing analysis [8].

The method of the insider threat provides us with the data when the behavioural profiling is being transformed, which has enormous energy to catalyse the capability of the niche detector, and implementing the combination of customer-level monitoring and network-level monitoring with the help of AN anomaly detection-powered AI engine that aids in capturing network anomalies. The following table presents the results of the operational timeline testing from the second round of SIEM assessment in the SOC 2 audit. The results indicate a significant enhancement in operational (SIEM) accuracy using generation 4 of the rate-based matrix, given the lower impunity percentage records, where the response velocity of the foreign IP detection increased at a high rate [9].

C. Data Sources

An effective detection strategy for insider threats is the ability to compile a wide range of data sources available over your Touch Points in the organization. The framework comes with a bunch of the following data sources:

- Towards Network Telemetry: Net-flow records, traffic dumps, firewall logs.
- Towards Authentication and Access Logs: Login times, geolocation of access attempts, privilege escalations.
- Towards Endpoint Telemetry: File access logs, USB device usage, and application execution data.
- Towards communication Data: e-mails, instant-messaging logs, and collaboration platform activity.
- Towards Organizational Context: Role-Based Access Control (RBAC) and HR data on employees' roles and responsibilities [9].

Preprocessing the data means getting the right information out of it, like when you anonymized and normalized it. You also need to get into these areas to give the data the right meaning and categories. This will help you scan the data quickly according to the scripts [9].

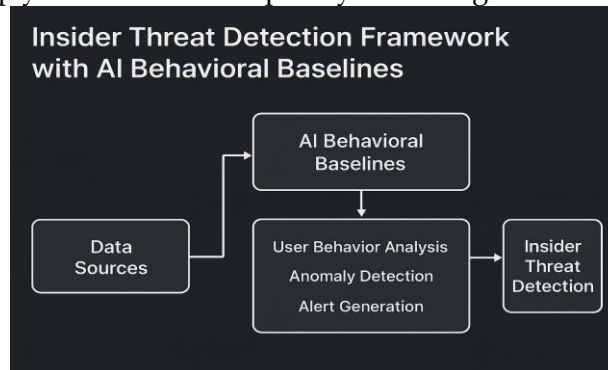


Fig. 2. Insider Threat Detection Framework with AI Behavioural Baselines.

D. Integration with Corporate Networking

The framework has a lot of important features that work together to keep an organization's network safe from threats that might come in at any time.

- **Network Intrusion Detection Systems:** Because we use Sophos products, the NIDS telemetry feed data is sent to Sophos AI models in real time from our network.
- **Cloud and On-Prem Based Servers:** We need to put lightweight detection agents on all of the corporate network servers and all of the cloud services, such as AWS, Google Cloud Platform, Microsoft Azure, and Meta Cloud GW. The centralized Sophos AI model gets constant threat information from these detection agents [10].
- **SIEM/SOAR:** This powerful detection engine works with existing SIEM (Security Information and Event Management) platforms wirelessly, adding alerts and grouping events to find and correlate them. SOAR (Security Orchestration, Automation, and Response) then shows important security information for team response workflows across the whole IT system.

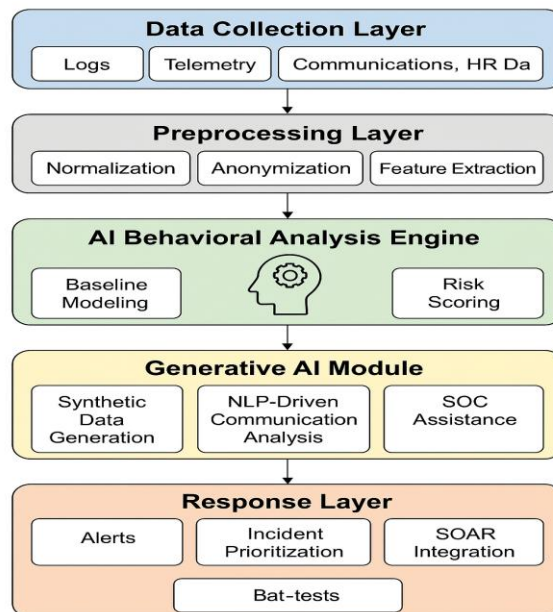


Fig. 3. Total Layers in the Framework

IV. CASE STUDIES AND APPLICATIONS

A. Financial Services

The financial services industry is one of the most likely to be hit by insider threats because it deals with sensitive customer data, payment systems, and trading algorithms. IBM Security (2023) says that the industry is responsible for almost a quarter of all insider incidents [10].

AI-driven behavioral analytics has also shown that it can be used to set up trading floors, find unusual access to sensitive financial information, and approve unusual wire transfers. For

instance, a bank might use AI models to figure out how many trades an employee should be doing and then look into any trades that are outside of that range as possible signs of insider trading. Generative AI has made this process better by making fake transactions that look real, which are hard to find in real life. Because of this, we had to work harder to learn.

AI technology is truly amazing; the advanced systems of different banks have cut the amount of money they send out by as much as 40% compared to rule-based systems.

Also, LLM-powered assistants can help compliance officers understand unusual behaviour and automatically create Suspicious Activity Reports (SARs). This makes it easier to make SARs and cuts down on operational costs [10].

B. Healthcare

The healthcare service industry has its own set of internal risk problems, especially when it comes to keeping electronic health records (EHRs) safe. Unauthorized access is often seen as a big problem with insider threats, and leaking patient records is the most common way for someone to steal someone's identity. People who break into places without permission usually get harsher punishments and more people hear about it. You can also tell if something is strange by the sudden rise in the number of files accessed.

If people knew that a nurse accessed 20 files every time they worked and that one nurse tried to access 200 files on this day, you could probably say that this is strange behavior. We can take it a step further: Generative AI can use synthetic datasets to create odd patterns that improve the model's accuracy.

GenAI models that use NLP, on the other hand, let you look through clinician notes and communication logs to see how medical facts are being misinterpreted. These skills are important for both making detection easier and following the rules of HIPAA and GDPR. Healthcare organizations that use these frameworks always get faster incident response and better monitoring that is ready to explain [10].

C. Government and Defence

The Snowden and Chelsea Manning cases are clear proof of how insider leaks can mess things up big time, exposing sensitive info and showing cracks even inside super-classified systems. These examples remind us that insider threats are not just some theories on paper—they can leave real, lasting damage for governments and defence groups.

To cut down the risk of these kinds of breaches, many agencies are now looking into AI tools. Things like machine learning and deep learning, when mixed with psycholinguistic analysis, can help spot early warning signs of risky behavior. This helps employees who work with tons of classified data, since unusual downloads or sudden spikes in activity might be the first hints that something is off [11].

AI is not just useful for prevention it can also help after an incident goes down. It can rebuild the timeline of what happened, connect dots that might have been missed earlier, and even highlight people who could have been involved. Large Language Models (LLMs) make this faster, since they can sift through huge amounts of text and pull out the details investigators need. With these tools, agencies stand a better chance at spotting future risks and stepping in before things spiral into a full-scale attack.

V. BENEFITS AND CHALLENGES

If AI is applied to detect insider threats, the organization could see several benefits. Compared to traditional rule-based systems, AI models are usually more precise, which helps reduce false positives so that real threats stand out. For security teams, this means they do not waste time chasing noise and can stay more focused on real issues. It also allows defences to adapt as corporate networks grow larger and more complex, instead of relying on new rules for every single change that happens. Data privacy is another big issue. Companies should not collect and process too much information about their workers because it is not right and maybe even against the law. The other big problem with AI is that it can be attacked by adversarial machine learning, which is when an enemy tries to trick detection models.

A. Data Privacy and Compliance

Data privacy remains one of the most significant challenges to adopting AI-based insider threat detection, largely because these systems require access to highly sensitive information. This often includes email records, chat logs, file activity, and keystroke monitoring, all of which may contain personally identifiable information (PII) and confidential organizational data. If monitoring efforts are poorly designed, they not only risk creating new security vulnerabilities but, in extreme cases, may also lead to violations of labor and privacy regulations.

In Europe, for instance, insider threat programs must comply with the General Data Protection Regulation (GDPR), which limits excessive employee surveillance and enforces data minimization practices [13]. In the U.S., HIPAA sets out tough rules for keeping employee health data safe, while PCI DSS lays down the standards for protecting payment card info. On top of that, GDPR goes a bit further, asking companies not just to secure data but also to help employees understand the risks of working with personal info and to put safeguards in place that lower those risks [12]. Generative AI can make life easier by cutting down on repetitive login tasks, but it also comes with its own set of problems. For example, an AI model might end up repeating sensitive patterns in system logs without meaning to, which could turn into a privacy headache if not handled right. That is why organizations really need strong governance practices something that keeps the balance between security and compliance. Some simple but effective steps are things like anonymizing personal data, when possible, not holding onto info for longer than needed, and keeping clear audit trails so AI-driven decisions can be explained and tracked later.

B. False Positives and False Negatives

Balancing sensitivity and specificity in detection systems is still kind of tricky. People need sensitivity so possible attacks don't slip by, but if it's too high then analysts get flooded with alerts that don't even matter. Specificity helps cut down the noise, but then you risk missing weird activity that might actually be important. Like take this—during a product launch, someone might stay late and log in at odd hours. An AI tool could see that as suspicious, even though it's just normal work stuff. A lot of false positives can make people tired of alerts, which makes SOC teams less sensitive and causes them to miss important problems. On the other hand, a false negative means that a real insider event is missed. For instance, a bad insider might slowly steal data over months, being careful to keep the amount of data at a level that doesn't raise any red flags compared to the baseline. Generative AI can help reduce both types of errors by adding more training datasets with different types of behavioral data, where unusual activities related to insider threats might happen. But without strict verification, AI systems may also overfit to fake data, which will cause errors that are not likely to happen [12].

C. Scalability and Performance

In a global business, central business networks collect a huge amount of data logs every day, with numbers in the billions. The real-time approach to sorting through this huge amount of data is what true data leakage prevention technologies must deal with, because to deal with insider threats effectively, the scaling needs are too high for anything less. While processing information about model performance and algorithm-computational distribution is not a problem for our company, there are other issues that make it hard to scale:

- **Cost of Training:** Because our teams are so experienced, Deep Neural Networks often need GPUs or TPUs for training, which can be hard on the capital account.
- **Real-Time Processing:** With streaming data, high-speed anomaly detection is required for the maintenance of low-latency in packet processing.
- **On-Premises vs Cloud deployment:** Deploying raw AI to process locational data could help scale out too many locations, but potentially violates privacy laws (even GDPR) if the data is processed in cloud environments.

It is never easy to strike the right balance between enforcing strong security measures and keeping systems simple and usable. This makes it common to use type, which is a way to get around problems like data being sent at strange times of the day or being coloured in strange ways, to name a few [12].

D. Organizational and Cultural Barriers

Organizational factors can act as a barrier to solution adoption as well as technical limitations. For example, some employees may consider AI-based technologies to be invasive of their privacy or demonstrate a lack of confidence in their abilities, and they may express these feelings. Because of the "the more you monitor, the more you find out" mindset, cultures of constant suspicion may emerge. In contrast, employees are being over-monitored to optimize every action they take, which will automatically shatter morale output.

Additionally, an effective insider threat program necessitates close, sustained collaboration across various business functions, departments, roles, and disciplines, including HR, legal, and IT. In the absence of clear governance policies, this vast array of groups poses the greatest threat to an organization, either by resisting or outright rejecting the program. Balancing legal prohibitions and organizational needs is a legal standard that employees not only deserve but are often required by regulations. So, being a good employee means finding a balance between the company's need for security and the rights of its employees [12].

E. Ethical Considerations

Ethical principles should be a big part of how programs that deal with insider threats work. AI tools might unfairly treat some workers differently because they are based on biased training data. For instance, employees who work in system administration can access information resources that clerical workers can't. So, the chances that these people are wrongly labelled as an insider threat are much higher than those of their coworkers in the same company.

Also, if generative AI isn't used carefully, it might look for insider threats that are more like real employees, which could hurt the company's reputation. To fix these problems, businesses need to use ethical AI frameworks like equality audits, policies that make employees more open, and third-party supervisory boards [13].

VI. FUTURE DIRECTIONS

A. Integration of GenAI for Enhanced Behavioural Modelling

The next big step in finding insider threats is to use Generative AI (GenAI) that is much more advanced. Right now, anomaly detection frameworks depend a lot on supervised or unsupervised learning models to find things that are different from the established baselines. But GenAI can make completely fake behaviour scenarios that are just as real as insider attacks. This enables us to create training sets that extend beyond historical logs [14].

For instance, a GenAI system could emulate a scenario where an employee gradually exfiltrates intellectual property via encrypted channels, providing training data for machine learning models that rarely encounter such threats in the wild, as they move very slowly. Thus, it will mitigate the 'dark spots' that exist in a scenario where a real-world, high-impact incident involving an insider occurs.

B. Cross-Domain Threat Intelligence Integration

The future will be aware that outsourcing threat intelligence can be helpful and therefore traditionally relies only on internal resources, neglecting most of the valuable data available. The future will merge with external threat intelligence systems to identify the correlation between suspicious behavior and the data that infiltrates it, such as intelligence feeds, the dark web, and open-source announcements in cybercrime forums, to pinpoint where insiders collaborate with external individuals. For example, AI can be used to correlate real-time events

in real life and stolen user information on the dark markets and alert the associated individuals as to the behaviour to be further analysed. The multi-layered system integration would make it easier to stop these kinds of threats before they happen [14].

C. Real-Time and Edge-Based Detection

The change from centralized to decentralized IT systems will change how digital-interconnect businesses (DEs) do business. Instead of the usual model, where all systems are in one place.

In the near future, enterprise systems should be moved closer to the edge so that it is easier to find threats in real time. This will give you another choice besides sending all of your logs to a central server.

The use of IT applications in today's business world has made the need for advanced cybersecurity skills much higher. This is different from the way things used to be, when time-series and AI squash-and-stretch models were used.[14].

D. Human-AI Collaboration in SOC Environment

In the future, AI and people will work together to find insider threats instead of separately. Artificial intelligence is good at processing a lot of data, but people know more about the context, make moral decisions, and have experience in the field. There will be human-in-the-loop AIs at the SOC of the future. Analysts will need to check, improve, and give feedback on the alerts that the AI sends [14].

The handler can also use generative AI as a smart assistant. You can teach it how to write an incident report, come up with a way to fix the problem, or act like an attacker to help SOC teams keep their systems safe. This cooperative model will make sure that things are done correctly while also holding people accountable for important insider matters.

VII. CONCLUSION AND RECOMMENDATIONS

A. The Evolving Landscape of Corporate Security

It's no longer safe to believe that no one can or will hack the company, so the old ways of keeping it safe don't work anymore. Those that are hired, as well as company partners, all have the potential of being a threat to the company's secrets on purpose or unknowingly can cause real big-time damage to the company in terms of financial, reputation, and operational. AI-based behavior analysis is a hope and a game-changer for all businesses. The employees can recognize unknown risks at a slight change in the system that traditional techs won't pick up on [14].

B. Real-Time Data Processing and Network Awareness

AI can process enormous amounts of data continuously in real-time, and also offers the ability to identify what is happening elsewhere on the network simultaneously.

For example, consider the relatively simple signals that examining log files and network packets may not reveal if new code, processes, or writing is being introduced on your network. But also, if the writing style is changing (or the amount of writing code is increasing, changing the way, or whatever). At the same time, without needing any service to train all the code of your employees. Additionally, if all your employees are writing in one tone, but a different person or someone else has written one piece in a different tone, or if access to specific data has only occurred at odd times of night.

In contrast, this doesn't work out all that well. Through this statement's remarkable use of insight and responsibility, technical intelligence will have the capacity to safeguard our interests from harm. Corporate intelligence needs to be refined if it is to possess only the knowledge necessary to ensure our well-being. Making up this point is sure to solve the concern that has been worrying us for so long, and therefore, it will be the manganese alone that we need [14].

C. The Role of Machine Learning in Preventing Insider Threats

Technological genius is primarily a legal concept, but gathering data in your company in ways that not only meet your customers' needs but also comply with local laws. The only way to protect both tech geniuses and machine learning from itself is to create Action. According to Palantir Technologies, precautions must be taken around machine learning to prevent it from evolving into intrusive programs that combat insider threats. More than just concern with personal security, the computer programs will also help target the central business processes and compliance. In shifting and balancing concerns about data, this proposed realm of technologies may play a role in frustrating actions such as initiating central competency, along with the freedom of individual employees [15].

D. The Future of Insider Threat Detection: Integrating People, Processes, and AI

In the future, insider threat detection will use universal education, enemies that are difficult to detect, and something that might be spelled wrong to take over the world. However, the organizations that will like AI the most are those that view it as a combination of things, including people, processes, and non-human elements within a system. Others who try to use AI purely as a technical control are at risk of becoming the next Enron, causing employees to feel needlessly surveyed and flee to defence contractor jobs, leading IBM to realize that its initial concerns may not have been so bad after all, ultimately landing the company in a death spiral. AI should be a way to handle changes at work that also show how important safety and risk are. Finding the right balance of AI-enabled ethics is what keeps the trust security posture going [16].

E. Toward a Culture of Responsible and Secure AI Adoption

Adding AI-based behavioural analysis is a big step toward protecting against insider threats, but it will only work if it is used properly. If people in an organization don't practice enough at predicting how AI will actually be used, they may think about using it in ways that aren't right. Instead of viewing the implementation of insider threat detection as a mere compliance

benchmark, organizations should regard it as a strategic initiative essential for the adoption of advanced technology, human discernment, ethics, advantages, or appeal. The response to the realm of artificial intelligence and its impacts may proceed in one of two directions: the ideal possibility as a deterrent to corporate corruption and a foundation for perpetuating healthy corporate cultures into and beyond the digital era [17].

Recommendations

- Use a hybrid detection framework: Use AI-based anomaly detection to feed existing rule-based systems, and have a person in charge of cyber events to make sure everyone is accountable.
- Explainable AI: Use AI models that the security team and HR/legal department can easily understand so they can make the right choices [18].
- Privacy-Preserving methodologies: Use hidden learning, de-identification, and perturbation technologies to keep people's privacy safe while still giving useful statistics that are useful for projects that make money and build communities.
- Harden Adversarial: Use adversarial tools and probes to find out how bad actors can avoid being identified by staff, data, and IP [19].
- Encourage people and machines to work together: Add SOC Analysts to the AI workflow to get better data results.
- Implement: Use the cross-functional teams from legal, HR, IT, and Compliance areas to ensure that cybersecurity is included in the AI and that the cybersecurity policy covers concerns with AI inside the company [20].

ACKNOWLEDGMENT

I recognize the contributions of geotechnical databases from several worldwide institutes and the constructive feedback from practicing engineers during the development of the framework.

REFERENCES

1. Ajayi, A. M., Omokanye, A. O., Olowu, O., Adeleye, A. O., Omole, O. M., & Wada, I. U. (2024). Detecting insider threats in banking using AI-driven anomaly detection with a data science approach to cybersecurity. *International Journal of Cybersecurity Research*.
2. Dash, B., Ansari, M. F., Sharma, P., & Ali, A. (2022). Threats and opportunities with AI-based cyber security intrusion detection: a review. *International Journal of Software Engineering & Applications (IJSEA)*, 13(5).
3. Dhruvitkumar, V. T. (2024). AI-Powered Cloud Security: Using User Behavior Analysis to Achieve Efficient Threat Detection.
4. Katiyar, N., Tripathi, M. S., Kumar, M. P., Verma, M. S., Sahu, A. K., & Saxena, S. (2024). AI and Cyber-Security: Enhancing threat detection and response with machine learning. *Educational Administration: Theory and Practice*, 30(4), 6273-6282.
5. Kavitha, D., & Thejas, S. (2024). Ai enabled threat detection: Leveraging artificial intelligence for advanced security and cyber threat mitigation. *IEEE Access*.

6. Khan, M. I., Arif, A., & Khan, A. R. A. (2024). AI-Driven Threat Detection: A Brief Overview of AI Techniques in Cybersecurity. *BIN: Bulletin of Informatics*, 2(2), 248-61.
7. Madhavram, C., Galla, E. P., Sunkara, J. R., Rajaram, S. K., & Patra, G. K. (2022). AI-Driven Threat Detection: Leveraging Big Data For Advanced Cybersecurity Compliance. Available at SSRN 5029406.
8. Mahmud, F., Barikdar, C. R., Hassan, J., Goffer, M. A., Das, N., Orthi, S. M., ... & Hasan, R. (2025). AI-Driven Cybersecurity in IT Project Management: Enhancing Threat Detection and Risk Mitigation. *Journal of Posthumanism*, 5(4), 23-44.
9. Mazher, N., Basharat, A., & Nishat, A. (2024). AI-Driven Threat Detection: Revolutionizing Cyber Defense Mechanisms. *Pioneer Research Journal of Computing Science*, 1(4), 46-59.
10. Md Imran, K., Mohammad Kowshik, A., & MD Asief, M. (2025). AI-BASED ANOMALY DETECTION IN CLOUD DATABASES FOR INSIDER THREATS. *Journal of Adaptive Learning Technologies*, 2(6), 8-29.
11. Muppalaneni, R., Inaganti, A. C., & Ravichandran, N. (2024). AI-Driven Threat Intelligence: Enhancing Cyber Defense with Machine Learning. *Journal of Computing Innovations and Applications*, 2(1), 1-11.
12. Muzaffar, J., & Mazher, N. (2024). AI-Powered Behavioral Analysis for Insider Threat Detection in Enterprise Networks. *Baltic Journal of Multidisciplinary Research*, 1(2), 1-11.
13. Otoum, S., Kantarci, B., & Mouftah, H. (2021). A comparative study of ai-based intrusion detection techniques in critical infrastructures. *ACM Transactions on Internet Technology (TOIT)*, 21(4), 1-22.
14. Prince, N. U., Faheem, M. A., Khan, O. U., Hossain, K., Alkhayyat, A., Hamdache, A., & Elmouki, I. (2024). AI-powered data-driven cybersecurity techniques: Boosting threat identification and reaction. *Nanotechnology Perceptions*, 20(S10).
15. Ramesh, S. K. A. (2024). AI-Enhanced Cyber Threat Detection. *International Journal of Computer Trends and Technology (IJCTT)*, 72(6), 64-71.
16. Rao, D. D., Wao, A. A., Singh, M. P., Pareek, P. K., Kamal, S., & Pandit, S. V. (2024). Strategizing IoT network layer security through advanced intrusion detection systems and AI-driven threat analysis. *Full Length Article*, 12(2), 195-95.
17. Shaik, A. S., & Shaik, A. (2024, April). AI Enhanced Cyber Security Methods for Anomaly Detection. In *International Conference on Machine Intelligence, Tools, and Applications* (pp. 348-359). Cham: Springer Nature Switzerland.
18. Sunkara, G. (2021). AI Powered Threat Detection in Cybersecurity. *International Journal of Humanities and Information Technology*, (Special 1), 1-22.
19. Yaseen, A. (2023). AI-driven threat detection and response: A paradigm shift in cybersecurity. *International Journal of Information and Cybersecurity*, 7(12), 25-43.
20. Wairagade, A., Sonani, R., & Govindarajan, V. (2025, April). AI Based Methods for Insider Threats Detection for Cloud Risk Mitigation. In *2024 International Conference on IT Innovation and Knowledge Discovery (ITIKD)* (pp. 1-7). IEEE.