

The Role of Generative AI in Automated Threat Hunting

¹Mr.Sidharth Sharma

¹*Vice President – IT Projects/Audits, JP Morgan Chase. Inc, 545 Washington Blvd Jersey City, NJ 07310 – US.*

¹Corresponding Author's email: infosidharthsharma@gmail.com

Abstract. An increasing number of enterprises are using generative artificial intelligence (AI) to improve their cyber security and threat intelligence. Generative AI is a type of AI that generates new data independently of preexisting data or expert knowledge. One emerging cyberthreat to systems that has been increasing is adversarial attacks. By generating fictitious accounts and transactions, adversarial attacks can interfere with and take advantage of decentralized apps that operate on the Ethereum network. Because fraudulent materials (such as accounts and transactions) used as malicious payloads can be mistaken for legitimate data, detecting adversarial attacks can be difficult. This paper suggests a paradigm for cyber threat hunting in the Ethereum blockchain that makes use of Generative Adversarial Networks (GAN) and Deep Recurrent Neural Networks (RNN). By considering a variety of sources and data points, this technology enables decision support systems to automatically and rapidly identify threats posed by hackers or other harmful actors. The likelihood of a successful assault can be further decreased by using generative AI to find weaknesses in an organization's infrastructure. Because security operations centers (SOCs) need to quickly identify threats and take defensive action, this technology is particularly well-suited for them. Generative AI can give businesses an extra line of defense against increasingly complex threats by integrating intriguing and useful data items that would have otherwise gone unnoticed

Keywords. Artificial intelligence, Threat intelligent, machine learning, threat hunting, deep learning, autonomous treat intelligent, Generative Adversarial Networks (GAN), Deep Recurrent Neural Networks (RNN).

1. INTRODUCTION

The emergence of Large Language Models (LLMs) has marked a significant shift in the field of natural language processing, enabling machines to produce text that closely matches the intricacy and cohesion of content created by humans. Models like GPT-3 and Deberta have shown exceptional skill in activities from translating languages to crafting creative content, making it increasingly challenging to distinguish between text produced by humans and machines. While these advancements have opened new frontiers in AI research and application, they have also raised profound questions regarding the authenticity and trustworthiness of the generated content.

The proliferation of AI-generated text has significant implications across various domains, including journalism, social media, education, and business. However, alongside the potential benefits come inherent risks, manipulation of public opinion, and the erosion of trust in digital communication channels. Addressing these challenges requires robust mechanisms for distinguishing between AI-generated and human generated content, a task that remains inherently complex due to the evolving nature of AI technologies. Detecting AI-generated text has emerged as a pressing research area, driven by the imperative to safeguard against the misuse of AI and preserve the integrity of online discourse.

Traditional approaches to text classification, such as TF-IDF, have long been foundational in natural language processing, providing information on the term distribution within a corpus. These approaches are likely to fail, though, when faced with the sophisticated linguistic subtleties typical of AI-generated text. To counter these limitations, our work suggests a new hybrid solution that combines conventional feature extraction methods. By taking advantage of the strengths of TF IDF complemented by sophisticated algorithms like Bayesian classifiers, Stochastic Gradient Descent (SGD), Categorical Gradient Boosting (CatBoost), and the highly Deberta-v3-large models, our approach is designed to attain unparalleled precision in differentiating between text produced by AI and that created by humans. In this paper, we outline our methodology, present our experimental findings, and draw conclusions regarding the efficacy of our proposed approach. Through extensive experimentation on a diverse dataset comprising both human and AI-generated text samples, we demonstrate the superiority of our method in accurately discerning between the two. By advancing AI-generated

text detection techniques, our research seeks to mitigate the risks associated with the proliferation of AI generated content and foster trust in digital communication platforms. As sophisticated threats targeting systems and networks have increased, the cybersecurity landscape has changed significantly. The dynamic threat landscape is frequently too fast for traditional cybersecurity measures to keep up with, which is why autonomous threat hunting has emerged as a proactive protection strategy. Using artificial intelligence (AI) and machine learning (ML) algorithms, this strategy uses real-time autonomous threat detection, analysis, and mitigation.

Core Research Problem and the Strategic Role of Autonomous Systems

Threats are becoming more sophisticated and widespread, and the cybersecurity environment is becoming more dynamic and complicated. Traditional threat intelligence approaches frequently find it difficult to keep up with these quick developments, which creates a serious research issue: the incapacity to quickly and efficiently identify, evaluate, and neutralize new risks in real-time. This ongoing difficulty necessitates the use of autonomous threat hunting techniques. The ability of human-operated systems to manage the enormous volume of data produced by many sources and to identify subtle patterns suggestive of possible dangers is restricted. Furthermore, a proactive and automated reaction is necessary due to the time-sensitive nature of cyber threats. The goal of autonomous threat hunting is to close this gap by utilizing AI-powered systems. By collecting, processing, and analyzing vast amounts of data on their own, these technologies make it possible to spot hidden signs of compromise and attack routes that were previously unknown. Given the increasing need for flexible, scalable, and quick threat detection and mitigation systems, it is imperative to create autonomous techniques. In order to strengthen cybersecurity measures, autonomous ways enabled by AI are urgently needed. The study challenge thus centers on the ineffectiveness of traditional methods in handling the speed and complexity of contemporary cyber threats.

1. Automation of cyber security and threat intelligence procedures: Generative AI makes it possible to automate the collection of threat intelligence and the upgrading of cyber security protocols, doing away with human labor and enabling intricate and advanced risk analysis of possible threats.

2. Better malicious activity detection and mitigation: By using generative AI, organizations may better understand which attack vectors are most likely to be employed and react by identifying the best ways to identify and eliminate the danger.

3. Enhanced efficiency in threat analysis: With the use of generative AI, organizations are able to rapidly analyze vast volumes of threat data and derive valuable insights for further protection. This enhances efficiency in threat analysis and results in more informed threat-based decisions.

4. Enhanced network monitoring: Generative AI offers improved network monitoring features, allowing organizations to continuously scan networks and detect anomalous activity, and thus take preventive action to safeguard against cyberthreats.

5. Generative AI - Advanced AI-driven threat intelligence that can produce AI-driven threat intelligence through generating deeper understanding and decision-making power to security related to deeper insight potential threats.

2. LITERATURE SURVEY

A form of artificial intelligence (AI) that has grown in significance recently as businesses seek to improve their threat intelligence and cyber security protocols is generative AI [1]. In generative artificial intelligence (AI), which is the act of producing new and frequently creative data, algorithms are created to learn from preexisting datasets and produce novel, untested concepts. Generative AI has the potential to assist businesses in the field of cybersecurity by assisting them in recognizing novel threats, creating defenses against them, and promptly and effectively responding to existing ones [3].

The first way that generative AI may support cyber security operations is by offering a comprehensive understanding of the threat landscape facing a business. It can examine existing datasets to find patterns that haven't been found before, detect possible dangers, and make specialists more aware of potential weaknesses

[2]. Generative AI may also assist companies in anticipating and addressing new risks by offering insights into the fundamental patterns and actions of malevolent actors. Additionally, generative AI can assist cyber security professionals in creating and implementing more effective preventative measures. Experts can use threat activity data to find previously unnoticed patterns and use that information to create stronger countermeasures [4].

Additionally, generative AI can handle the workload of evaluating a high volume of threat intelligence feeds, freeing up human experts to concentrate on those that need further study and analysis. Last but not least, generative AI may automate processes like obtaining information about specific hazards, sending out alerts, building dashboards, and connecting data points to support investigations. Threat intelligence and cyber security protocols could be completely transformed by generative artificial intelligence, or generative AI [6]. A kind of machine learning known as "generative AI" allows computers to create or synthesis new data sets from preexisting ones, improving and increasing the precision of data-driven decision making [5].

Threat intelligence and cyber security measures can be improved to stop and identify malware assaults, data breaches, and other harmful behavior by incorporating generative AI into the system. Large amounts of data gathered from security systems, such as network traffic, user activity, and content, can be used by generative AI to create and implement unique risk-based formulae that swiftly identify hostile activity [8]. Additionally, generative AI may review an organization's network, assess its security posture, and notify cyber security experts of any questionable activities. Generic AI can also identify unusual patterns and behaviors and develop fresh approaches to thwart dangers. Integrating machine learning models into the corporate security architecture is another way that generative AI can offer a higher level of data protection. Cyber security experts might be warned of possible dangers by these models, which are able to recognize threats. Providing the best remedy for the aforementioned issues is the primary innovation of this research. These are, generative AI has the ability to significantly enhance cyber security and threat intelligence [7]. Without assistance from humans, generative AI models may create new data by learning from examples. This could assist organizations in identifying hidden or invisible risks and their trends. Additionally, big datasets collected from many sources can be swiftly analyzed by AI-enabled systems, which can identify dangers almost instantly. Furthermore, organizations might create complex fake data intended to snare potential attackers with the aid of generative AI models. This data can assist lower the danger of data disclosure during an assault because it is automatically generated and can be used to conceal the actual data.

3. PROPOSED SYSTEM

Threat intelligence and cyber security procedures could be completely transformed by generative artificial intelligence, or generative AI, an emerging technology. The secret to developing more effective solutions that lower the price and complexity of threat intelligence and cyber security measures may lie in generative artificial intelligence. Generative AI finds and examines patterns and relationships in data using deep learning algorithms. AI may discover security threats and detect anomalies more quickly than manual methods by utilizing the computer's processing capacity. This makes it possible to apply security measures that are more effective and efficient because the AI can recognize high-risk scenarios and notify the user. To test and assess the effectiveness and performance of security systems, generative AI can also be used to create simulations of real-world situations. Among other things, these simulations may replicate situations like a hacker attack, harmful software, or malevolent network activities. Organizations can make sure they have the appropriate amount of cyber security measures in place by evaluating the security measures' efficacy before implementing them.

A new area of artificial intelligence called "generative AI" allows computers to produce original concepts and results. More complex threat intelligence models can be produced using generative AI technologies, such as generative adversarial networks (GANs) and deep learning architectures, to improve cyber security protocols. By facilitating the creation of increasingly complex models, generative AI contributes to the expansion of threat intelligence capabilities. A GAN, for instance, can be used to create malware, phishing campaigns, new dangerous code, and other harmful actions.

Furthermore, more intricate attack networks that are able to recognize aberrant activity, infiltration behavior, and network patterns can be developed using generative AI. This makes it possible to detect and prevent threats with greater precision.

Cyber security teams can also develop more robust cyber security solutions with the aid of generative AI. For example, through automated scanning and user behavior analysis, AI can assist in identifying dangerous conduct before it happens. Additionally, AI-driven models can assist in promptly addressing new threats, enabling the security team to take preventative action. Figure 1 below displays the functional block diagram.

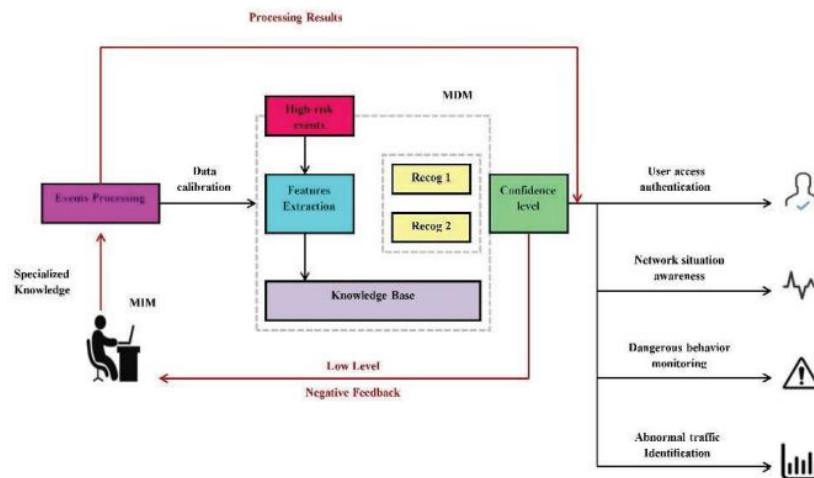


FIGURE 1. Functional block diagram.

One type of artificial intelligence that can be used to generate or create new data is called generative AI. Its primary uses are in the fields of cyber security and threat intelligence. The way generative AI operates is by combining information from both structured and unstructured parameters, user feedback, and a variety of internal and external sources. Systems based on generative AI can produce threat intelligence more rapidly and with greater quality. These systems are able to discover hitherto undiscovered connections between risks by learning from current datasets. This enhances the timeliness and accuracy of threat intelligence. Systems built on generative AI can also be utilized to produce original answers to new dangers. By using AI into cyber security solutions, businesses may create more flexible and responsive defense strategies.

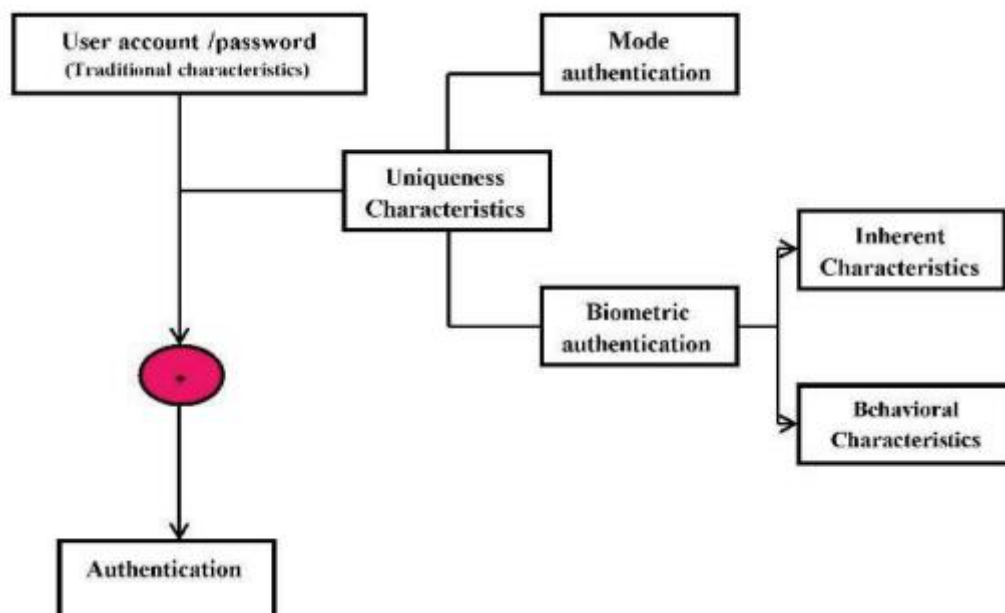


FIGURE 2. Operational Flow Diagram

The operating flow diagram is displayed in fig. 2 below. The artificial intelligence (AI) method known as "generative AI" is used to produce brand-new, previously unseen data or data that. It appears to be either new or real. In order to create data or information based on a set of factors and situations, generative AI uses

probability models. Increasing cyber security measures and gathering threat intelligence are just two uses for this data. Cyber security solutions that are already in place can be informed and protected by the new threat information generated by generative AI.

4. RESULTS AND DISCUSSION

Generative artificial intelligence is a subset of AI that can create new, previously undiscovered data by learning from datasets on its own. It can be used to help enhance cyber security and threat intelligence by identifying risks before they become harmful or dangerous. Security staff may be notified when a threat is spotted and anomalies in data can be quickly identified with generative AI. Faster reaction times and improved general defense against cyberattacks may result from this. By simulating actual situations, generative AI can also be used to develop simulations that let security experts rehearse reacting to possible threats. By doing so, possible weaknesses may be found and security measures' efficacy may be increased. Lastly, generative AI can be used to identify and stop malicious software and other hazards, assisting in halting damage brought on by recently created dangers. All things considered, generative AI has the potential to be an effective instrument for improving threat intelligence and cyber security protocols. Organizations can detect malicious threats more rapidly and respond to them as efficiently as feasible with the use of generative AI. By using this data, firms may keep ahead of possible threats and modify their current cyber defense strategies. Additionally, generative AI can identify weaknesses in current networks and systems, assisting in their most effective patching and protection. Threat intelligence and cyber security systems that incorporate generative AI can help businesses stay safe, stop cyberattacks, identify them fast, and react to them successfully.

5. CONCLUSION

It has been shown that generative AI is a useful tool for improving cyber security and threat intelligence. In addition to generating creating datasets for machine learning algorithms to train, generative AI may also be used to enhance data and improve the detection of malicious or unusual activities. It can also be used to find patterns in data that hasn't been classified yet, which improves threat detection accuracy. Furthermore, current cyber security systems can be made more effective with the usage of generative AI. Generative AI is becoming a key instrument for the future of cyber security because of its capacity to enhance threat intelligence and cyber security protocols. This kind of AI can identify changing trends and patterns, produce information on threats that is both positive and negative, and even mimic malevolent people. It can also be used to generate a huge number of false positives and negatives, improving networks' readiness and defense against current and potential threats. Additionally, in order to guarantee a better response and more effective countermeasures, AI-based solutions can dynamically modify the threat intelligence and response measures' response times. Protecting against malevolent actions and attempts is another benefit of generative AI. Its capacity to spot patterns and find connections between datasets might be useful in highlighting potentially harmful or questionable activity.

REFERENCES

1. Jasper Gnana Chandran, J., Karthick, R., Rajagopal, R., & Meenalochini, P. (2023). Dual-channel capsule generative adversarial network optimized with golden eagle optimization for pediatric bone age assessment from hand X-ray image. *International Journal of Pattern Recognition and Artificial Intelligence*, 37(02), 2354001.
2. Karthick, R., Prabha, M., Sabapathy, S. R., Jiju, D., & Selvan, R. S. (2023, October). Inspired by social-spider behavior for microwave filter optimization, swarm optimization algorithm. In *2023 International Conference on New Frontiers in Communication, Automation, Management and Security (ICCAMS)* (Vol. 1, pp. 1-4). IEEE.
3. Vijayalakshmi, S., Sivaraman, P. R., Karthick, R., & Ali, A. N. (2020, September). Implementation of a new Bi-Directional Switch multilevel Inverter for the reduction of harmonics. In *IOP Conference Series: Materials Science and Engineering* (Vol. 937, No. 1, p. 012026). IOP Publishing.
4. Kiruthiga, B., Karthick, R., Manju, I., & Kondreddi, K. (2024). Optimizing harmonic mitigation for smooth integration of renewable energy: A novel approach using atomic orbital search and feedback artificial tree control. *Protection and Control of Modern Power Systems*, 9(4), 160-176.
5. Sulthan Alikhan, J., Miruna Joe Amali, S., & Karthick, R. (2024). Deep Siamese domain adaptation convolutional neural network-based quaternion fractional order Meixner moments fostered big data analytical method for enhancing cloud data security. *Network: Computation in Neural Systems*, 1-28.
6. Sakthi, P., Bhavani, R., Arulselvam, D., Karthick, R., Selvakumar, S., & Sudhakar, M. (2022, September). Energy efficient cluster head selection and routing protocol for WSN. In *AIP Conference Proceedings* (Vol. 2518, No. 1). AIP Publishing.
7. Aravindaguru, I., Arulselvam, D., Kanagavalli, N., Ramkumar, V., & Karthick, R. (2022, September). Space cloud in cubesat-Consigning expert system to space. In *AIP Conference Proceedings* (Vol. 2518, No. 1). AIP Publishing.
8. Karthick, R., Prabakaran, A. M., & Selvaprasanth, P. (2019). A Dumb-Bell shaped damper with magnetic absorber using ferrofluids. *International Journal of Recent Technology and Engineering (IJRTE)*, 8.
9. Selvan, R. S., Wahidabanu, R. S. D., Karthick, B., Sriram, M., & Karthick, R. (2020). Development of Secure Transport System Using VANET. *TEM (H-Index)*, 82.
10. Karthick, R., & Sundararajan, M. (2018). Optimization of MIMO Channels Using an Adaptive LPC Method. *International Journal of Pure and Applied Mathematics*, 118(10), 131-135.
11. Lopez, S., Sarada, V., Praveen, R. V. S., Pandey, A., Khuntia, M., & Haralayya, D. B. (2024). Artificial intelligence challenges and role for sustainable education in india: Problems and prospects. *Sandeep Lopez, Vani Sarada, RVS Praveen, Anita Pandey, Monalisa Khuntia, Bhadrappa Haralayya (2024) Artificial Intelligence Challenges and Role for Sustainable Education in India: Problems and Prospects. Library Progress International*, 44(3), 18261-18271.
12. Kumar, N., Kurkute, S. L., Kalpana, V., Karuppannan, A., Praveen, R. V. S., & Mishra, S. (2024, August). Modelling and Evaluation of Li-ion Battery Performance Based on the Electric Vehicle Tiled Tests using Kalman Filter-GBDT Approach. In *2024 International Conference on Intelligent Algorithms for Computational Intelligence Systems (IACIS)* (pp. 1-6). IEEE.

13. Sharma, S., Vij, S., Praveen, R. V. S., Srinivasan, S., Yadav, D. K., & VS, R. K. (2024, October). Stress Prediction in Higher Education Students Using Psychometric Assessments and AOA-CNN-XGBoost Models. In *2024 4th International Conference on Sustainable Expert Systems (ICSES)* (pp. 1631-1636). IEEE.
14. Yamuna, V., Praveen, R. V. S., Sathya, R., Dhivva, M., Lidiya, R., & Sowmiya, P. (2024, October). Integrating AI for Improved Brain Tumor Detection and Classification. In *2024 4th International Conference on Sustainable Expert Systems (ICSES)* (pp. 1603-1609). IEEE.
15. Anuprathibha, T., Praveen, R. V. S., Jayanth, H., Sukumar, P., Suganthi, G., & Ravichandran, T. (2024, October). Enhancing Fake Review Detection: A Hierarchical Graph Attention Network Approach Using Text and Ratings. In *2024 Global Conference on Communications and Information Technologies (GCCIT)* (pp. 1-5). IEEE.