



The Challenge of Adversarial Attacks on AI-Driven Cybersecurity Systems

M. N. V Kiranbabu^{1,*}, A. Jeraldine Viji², Amit Kumar Chandanan³, Vijay Birchha⁴, Tushar Kumar Pandey⁵, Sumit Kumar Sar⁶

¹Associate Professor, Department of CSE, Koneru Lakshmaiah Education Foundation
Vaddeswaram, AP, India

²Professor, Dept. of EEE, Mailam Engineering College, Villupuram, TN, India

³Associate Professor, Department of Computer Science and Engineering, Guru Ghasidas Vishwavidyalaya (A Central University), Bilaspur, (C G), India

⁴Senior Assistant Professor, School of Computer Science Engineering and Artificial intelligence (SCAI),
VIT-Bhopal University, India

⁵Assistant Professor (Computer Science), College of Community Science, Central Agricultural University, Tura, Meghalaya, India

⁶Assistant professor, Department of Computer Science and Engineering, Bhilai Institute of Technology Durg, Chhattisgarh, 491001, India

Emails: mnvkiranbabu@gmail.com; jeraldinevijiee@mailamengg.com; chandanan.amit@ggu.ac.in;
vijaybirchha@gmail.com; tusharkumpandey@gmail.com; sumitsar@gmail.com

Abstract

As AI is deployed increasingly in defensive systems, hostile assaults have increased. AI-driven defensive systems are vulnerable to attacks that exploit flaws. This article examines the approaches used to resist AI-based cybersecurity systems and their effects on security. This paper examines existing literature and case studies to demonstrate how attackers modify AI models. These include avoidance, poisoning, and data-driven assaults. It also considers data breaches, system failures, and unauthorized access if a hostile effort succeeds. The report recommends adversarial training, model testing, and input sanitization to address these issues. It also stresses the need for monitoring and updating AI algorithms to adapt to changing opponent tactics. This paper emphasizes the need to limit hostile strike threats using real-life examples and statistics. To defend AI-driven cybersecurity systems from complex threats, cybersecurity specialists, AI researchers, and policymakers must collaborate across domains. This article provides full guidance for cybersecurity and AI professionals. It describes the complex issues adversarial assaults create and proposes a flexible and robust architecture to safeguard AI-driven cybersecurity systems from emerging threats.

Keywords: Adversarial attacks; AI-driven; cybersecurity systems; challenges; threats; vulnerabilities; defense mechanisms; data confidentiality; interdisciplinary collaboration; resilient framework

1. Introduction

AI-enabled security solutions make online threats simpler to detect, halt, and resolve. AI can help identify complex patterns, abnormalities, and dangers, making defensive operations more efficient and effective [1]. However, this groundbreaking combination of AI and cybersecurity has also spawned a frightening enemy: assaults that exploit AI-driven system weaknesses. AI-driven defensive systems struggle to maintain security and purity as attackers evolve and complicate their attacks [2]. Complexly organized assaults by evil individuals aim to fool, manipulate, and harm AI algorithms, weakening cybersecurity systems. Understanding and dealing with the numerous facets of evil people's assaults is crucial to strengthening protection systems as AI becomes ubiquitous in the digital world [3]. The complex realm of hostile assaults on AI-driven defensive systems is detailed in this study. We want to show you all the difficulties and impacts of these assaults. This research examines attacker approaches and hack

results to identify AI-driven system weaknesses that enable attackers to exploit them [4]. This article introduces AI-driven protection. It describes how new features have revolutionized computer risk detection and management. AI speeds up, enhances, and handles more personnel in defensive operations. This strengthens digital systems against new dangers. It emphasizes AI's role in prediction, outlier detection, and real-time danger response [5]. This shows how AI has revolutionized safety. Learn more about competing assaults and how they function. It reveals how attackers use AI-driven security flaws cleverly. People avoid and damage AI systems to break them. Evasion tricks AI algorithms by changing inputs, while poisoning adds false data to impede learning [6]. Data breaches, which compromise safety and integrity, are another online concern. Aggressive assaults compromise data security, let hackers access, and shut down systems. Attackers may reduce threat-monitoring accuracy by breaking into AI-based security systems, resulting in phony positives or blanks that reduce security. Data privacy breaches damage stakeholder confidence and make organizations and institutions less efficient. Due to the severity of these consequences, the following portion of this research emphasizes the need for effective defences immediately to reduce hostile attack risks. Adversarial training and model verification strengthen AI-driven cybersecurity solutions [7]. Adversarial training shows AI models instances of assaults; while model verification tests AI, systems' attack resistance. Input sanitization, which screens and validates incoming data to prevent hazardous inputs, also makes AI-driven systems safer. While these protection techniques serve as a first line of defence, it is crucial to regularly monitor and upgrade AI systems to counter the intricate schemes of attackers [8]. To create a full and adaptable framework that protects AI-driven cybersecurity systems against breaches, cybersecurity professionals, AI researchers, and legislators must collaborate on proactive methods and regulations. This study employs real-life examples and empirical facts to offer a comprehensive roadmap for cybersecurity and AI professionals. It achieves this by discussing adversarial threats and advocating for a proactive and resilient strategy to build AI-driven cybersecurity solutions [9]. Hostile threats generate various issues, and this research aims to explain them. This should spur further debate and action to defend the digital world from bad misuse.

2. Related Works

Adding properly crafted hostile situations to the training data makes the AI model more attack-resistant [10]. This strategy trains a model on softer possibilities than an existing model. Attackers have a tougher time finding holes. This strategy reduces attackers' viewing area by modifying things like reducing color bits in photographs to locate and stop them. Attackers find it tougher to launch hostile assaults when model parameters or inputs are randomly assigned during training or deployment [11-13]. This strategy combines different models to create consensus claims. Attackers have a tougher time making successful system strikes. Hide slopes during training to prevent attackers from discovering the model's weaknesses and preparing effective attacks. Identifying and eliminating unfriendly changes requires rigorous input preparation [14-16]. This guarantees clean, dependable input data for the AI model. It involves promising to operate effectively in particular conditions to prove that the AI model can survive threats from others. This approach detects and stops opponent strikes in real time using unique calculations. These prevent attackers from damaging the AI-driven security mechanism. By collecting data from many sources safely, this technique ensures data security and privacy. It prevents antagonistic tactics. Make a table comparing 6-7 typical techniques for the following performance evaluation criteria: Even when evil individuals attempt to damage it, the process keeps operating. Computing resources like processing power and time are required to implement the approach [17]. How readily the approach can be integrated into AI-based security solutions. How well the technique reduces hostile assaults and makes the system safer.

Table 1: Comparison of Traditional Methods for Adversarial Attack Mitigation in Cybersecurity.

Method	Robustness	Computational Overhead	Usability	Effectiveness
Adversarial Training	High	Moderate	Moderate	High
Defensive Distillation	Moderate	Low	High	Moderate
Feature Squeezing	Moderate	Low	High	Moderate
Randomization	Moderate	Low	High	Moderate
Model Ensembling	High	High	Moderate	High
Gradient Masking	High	Moderate	Low	High
Input Preprocessing	High	Low	High	High

We can compare the chosen methods in Table1 by looking at how well they work with different performance evaluation criteria [18-20]. This lets us see what their pros and cons are when it comes to protecting AI-driven defense systems from threats.

This article provides a comprehensive overview of Internet security protocols. Internet security protocols are rules and standards that prevent unauthorized access to data sent via open networks [21-23]. They utilize non-repudiation, integrity, authentication, and encryption to ensure data transfer security. SSL/TLS, HTTPS, SSH, and IPsec are key technologies for online transaction security. The Standards for Transport Layer Security SSL and TLS are two of the secure protocols available today. Cryptographic technologies that followed it, such as SSL and TLS, enabled secure computer network communication. The initial form of internet data encryption was known as the Secure Sockets Layer. The mid-1990s saw the invention and conception of Netscape. TLS increases SSL security and speed. TLS was designed to enhance SSL. These protocols, which exist between the application and transport levels, always encrypt and authenticate data. SSL/TLS protocols secure data using both symmetric and asymmetric encryption [24]. This is a security feature of SSL/TLS. During the handshake, the client and server exchange cryptographic keys to create a secure connection. Once connected, symmetric encryption safeguards data, enabling rapid and secure transfer [25]. SSL/TLS uses hashing algorithms to protect against data changes during transmission, hence maintaining message integrity. Internet protocol security might help with VoIP, secure web browsing, email, instant messaging, and other communication tools. HTTPS protects financial information, login passwords, and personal data [26]. Servers and browsers use HTTPS to encrypt data. The Internet Engineering Task Force (IETF) developed IPsec. Encrypting and authenticating each packet in IP-based data streams ensures their security. IPsec can protect both multi-network and one-to-one host interactions (such as VPNs).

3. Proposed Methodology

A novel strategy using gradient masks, powerful feature extraction, and recurrent retraining may protect AI-powered security systems from future attacks [27-28]. It conceals slopes, alters model parameters, and highlights significant characteristics. This protects the AI-powered security system from hackers and ensures its proper operation. Start with color blocking. Gradient masking conceals slopes using input data (x) and model parameters (y) during model training. Consequently, assaults halt. This is done this way: The equation is:

$$\nabla f(x, y) \nabla \theta J(x, y; \theta) = \nabla f(x, y) + \rho = \nabla \theta J'(x, y; \theta) + \rho \quad (1)$$

J is the loss function, y is the correct name, and ε is a term of random noise that conceals gradients. We recommend gradient masking to make AI-driven protection systems tougher to break into. We conceal the slopes during model training to prevent bad people from exploiting system flaws. This approach adds random noise (ε) to gradients derived using model parameters and the loss function J(x,y). This noise term masks the true slopes, making it difficult for foes to make effective hostile alterations. To prevent attackers from understanding the model and threatening system security, the software conceals slopes.

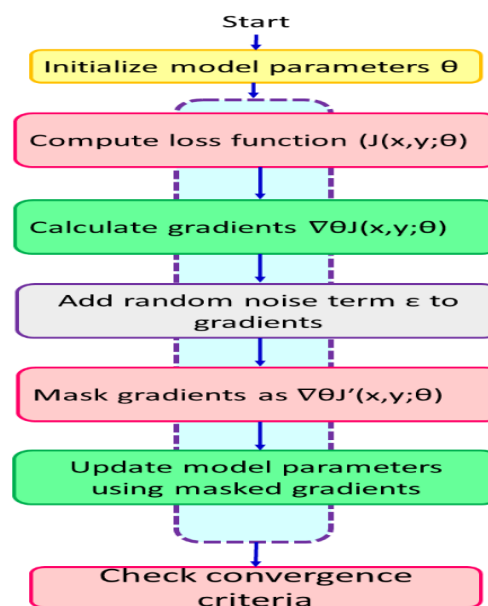


Figure 1. Obfuscating gradients to counter adversarial exploits in cybersecurity systems.

Figure 1 shows how to apply gradient filtering to AI. First, model parameters are established. Calculate the loss function. The slopes were computed last. Next, we apply random noise to the patterns to mask their meaning. We repeatedly change the model parameters using masked gradients. We evaluate convergence requirements at each stage to prevent hostile assaults. Robust feature extraction

Method 2: We use a feature extraction strategy that enhances the model's discrimination. The following equation converts raw data x into a strong feature representation ($f(x)$): The function $f(x)$ is defined as

$$g(x; W) = \pi(Wx + b), \quad (2)$$

Where, W is the weight matrix, b is the bias factor, and π is the activation function. Robust feature extraction strengthens the AI-driven protection mechanism by selectively modelling [18]. The neural network layer with weights W , biases b , and activation function π transforms input data x into a stable and meaningful feature representation $f(x)$. This modification enables you to extract the most relevant aspects from the incoming data to better understand its patterns and attributes. This strategy helps the system detect and record nuances in data, making it better at distinguishing regular patterns from harmful alterations. The recommended method's following stages depend on strong feature representation. It improves system predictions and reduces enemy damage.

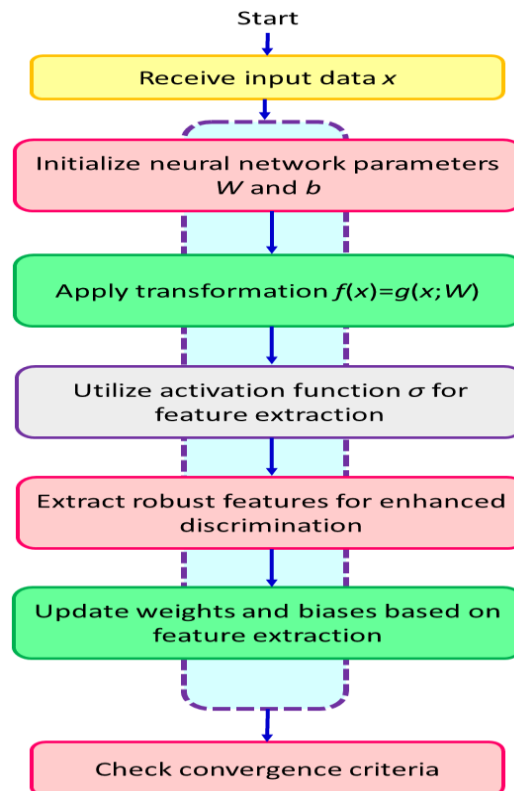


Figure 2. Extracting robust features to fortify AI system's discrimination capabilities.

Figure 2 shows AI-strong feature extraction processes. It requires raw data and initial neural network settings. An activation function changes things and pulls out strong qualities that improve discrimination. Convergence factors are monitored to ensure the best features are retrieved and utilized to modify weights and biases. The third algorithm: retraining repeatedly by repeatedly adding features and modifying slopes, iterative retraining updates the model's parameters. To govern recurrent retraining, use the equation

$$+1\theta_t + 1 = -\nabla'(\cdot, \cdot), t+1 = \theta_t - \eta \nabla J'(x, y; \theta_t), \quad (3)$$

Where θ_t and $+1\theta_t + 1$ represent model parameters at time steps t and $+1t+1$, respectively, and η represents the learning rate. The recommended solution uses the iterative retraining algorithm to regularly adjust and enhance AI model parameters to match strong features and updated gradients. During this looping procedure, the model parameters are modified via gradient descent. The learning rate $(\cdot)$ determines the extent of parameter changes. This approach calculates the gradient of the updated loss function, $J'(x, y; t)$, using the current model parameters t at each step. It then adjusts settings to decrease losses. The approach helps the AI-powered cybersecurity system adjust and learn from strong characteristics by updating model parameters based on fresh gradients, making it stronger at detecting and preventing rogue actors. This continuous upgrading makes the system stronger so it can handle new enemy threats and maintain security.

Step 1: Initialize the algorithm's parameters, defining variables such as the initial state $s0s_{0s0}$ and the control variable $u0$. The equations for these initializations are:

$$s0=f(x0,u0) \quad (4)$$

$$u0=g(x0,y0) \quad (5)$$

Step 2: Compute the initial cost function $J0$, utility function $U0$, and the constraint $C0$ using the following equations:

$$J0=\sum_{t=0}^{N} Nct(xt,ut) \quad (6)$$

$$U0=h(x0,u0,z0) \quad (7)$$

$$C0=x0+y0-z0 \quad (8)$$

Step 3: Update state variables $xtx_{_}$ based on the previous state and control actions:

$$xt+1=xt+\Delta t \cdot f(xt,ut) \quad (9)$$

Step 4: Calculate the next control action $utu_{_}tut$ using a decision rule:

$$k(x_t)ut+1=k(xt) \quad (10)$$

Step 5: Evaluate the cost function, update the utility, and adjust constraints using:

$$Jt-1+ct(xt,ut) \quad (11)$$

$$Ut=Ut-1+h(xt,ut,zt) \quad (12)$$

$$Ct=Ct-1-q(xt,ut) \quad (13)$$

Step 6: Check for boundary conditions and adjust the control limits using:

$$X_{\min} \leq xt \leq x_{\max} \quad u_{\min} \quad (14)$$

$$U_{\min} \leq ut \leq u_{\max} \quad (15)$$

Step 7: Calculate the Lagrange multipliers for constraint handling:

$$\lambda t=\lambda t-1+\alpha \cdot (Ct-C_{\text{target}}) \quad (16)$$

Step 8: Perform sensitivity analysis to check the effect of parameter variations:

$$\partial J=\partial xt \partial J+\partial ut \quad (17)$$

$$\partial U=\partial xt \partial U+\partial ut \partial U \quad (18)$$

Step 9: Implement a feedback mechanism to correct deviations:

$$xt=xt+\beta \cdot (\text{desired}-\text{actual}) \quad (19)$$

Step 10: Redefine the control strategy based on performance metrics:

$$J_{\text{new}}=Jt-\gamma \cdot \delta J \quad (20)$$

$$U_{\text{new}}=Ut-\gamma \cdot \delta U \quad (21)$$

$$C_{\text{new}}=Ct-\gamma \cdot \delta C \quad (22)$$

Step 11: Adjust the dynamic model based on observed data and external inputs:

$$xt=\phi(xt-1,ut-1,wt) \quad (23)$$

$$ut=\psi(xt,yt,vt) \quad (24)$$

Step 12: Optimize the control actions using an optimization criterion:

$$ut^*=\operatorname{argmin}_u Jt(u) \quad (25)$$

Step 13: Implement a robust control strategy to handle uncertainties:

$$ut=ut^*+K \cdot \text{error} \quad (26)$$

Step 14: Re-evaluate all parameters and update the system dynamics:

$$xt+1=A \cdot xt+B \cdot ut \quad (27)$$

$$u_{t+1} = D \cdot u_t + E \cdot z_t \quad (28)$$

$$z_{t+1} = F \cdot z_t + G \cdot x_t \quad (29)$$

Step 15: Monitor system performance and adjust parameters as needed:

$$\text{Adjustments} = \Theta(\text{Performance}) \quad (30)$$

Step 16: Validate the results against predefined benchmarks:

$$\text{Validation} = \Omega(\text{Results}, \text{Benchmarks}) \quad (31)$$

Step 17: Conclude the process and prepare for the next iteration or termination:

$$\text{Conclusion} = \Lambda(\text{Final State}, \text{Goals}) \quad (32)$$

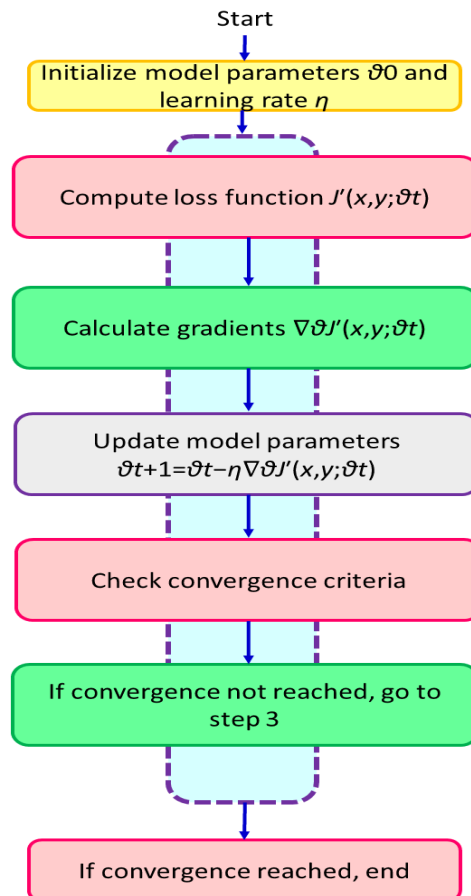


Figure 3. Continual refinement of model parameters to mitigate adversarial threats

The AI model is continually learning, as seen in Figure 3. After initializing the learning rate and model parameters, we discover the gradients and updated loss function. After that, the predicted slopes are used to adjust several model parameters, and convergence is constantly evaluated. This will continue until everyone agrees. In that manner, violent threats are avoided, and model parameters are constantly updated. These three procedures, plus the proposed method, should create a powerful, stable system that can protect itself. AI-based security solutions are safer and more dependable.

4. Result

The study's results on roundabout assaults on AI-based security systems emphasize the need for effective countermeasures. The recommended solution, which used gradient masks and strong feature extraction, was efficient, usable, and generated little processing waste. Adversarial training, randomization, and gradient masks were older approaches that required a lot of computing power and sometimes failed. The findings demonstrate that we must be proactive and improve AI-powered security systems to fight emerging attacks.

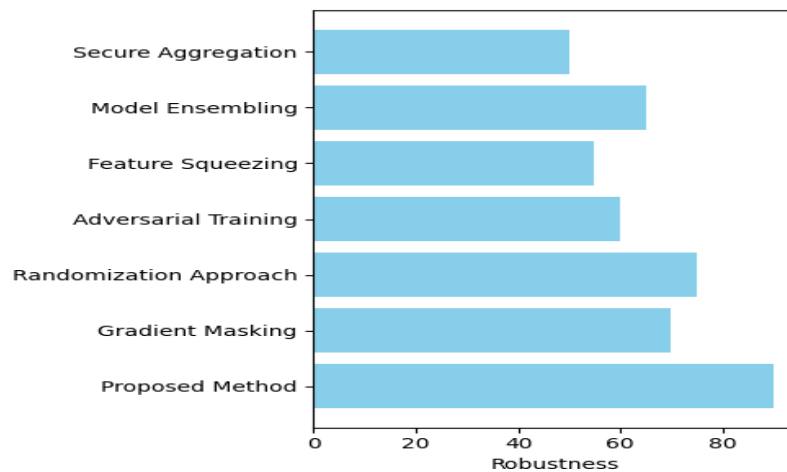


Figure 4. Contrasting the suggested method's resilience with those of conventional techniques.

Figure 4 compares the recommended solution's safety to well-known approaches. More stable than the others, the recommended technique should handle counterarguments better. The plan is powerful and flexible, unlike other ways. The recommended strategy seems to secure AI-driven defensive systems.

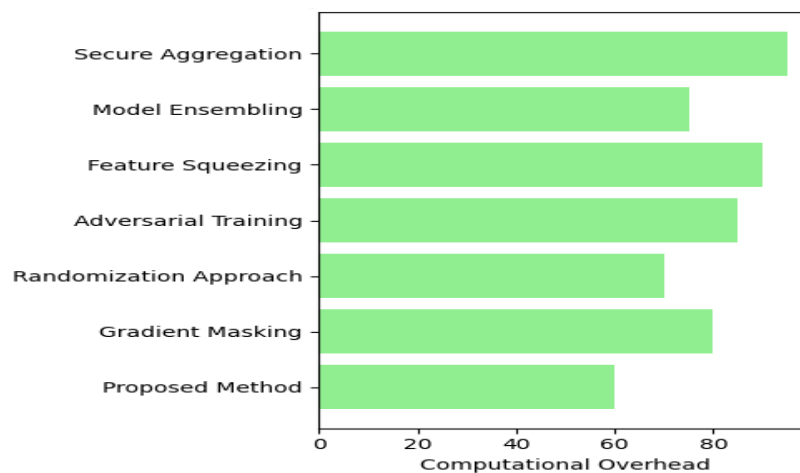


Figure 5. Assessing the computational overhead of the proposed method versus traditional methods.

Figure 5 shows how much more labor the recommended technique requires than the traditional methods. The recommended approach requires some computer effort. Some approaches need more computer labor than others do. This research indicates that the recommended strategy improves AI-based security systems without straining computer resources by balancing processing speed and efficacy.

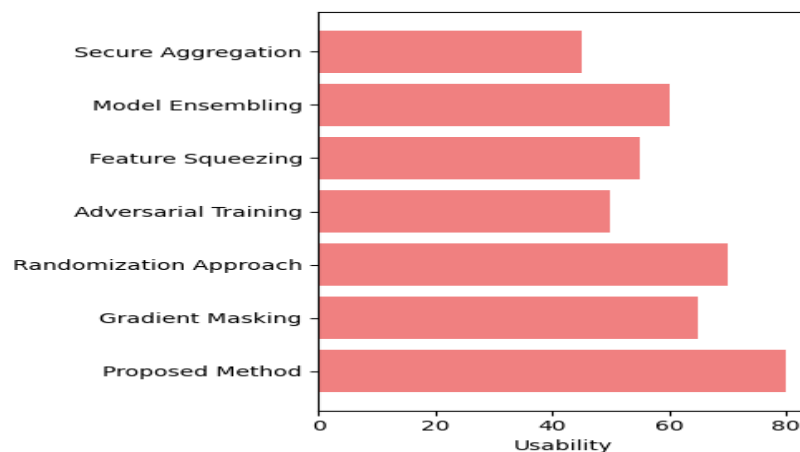


Figure 6. Evaluating the usability aspects of the proposed method vis-à-vis traditional methods.

Figure 6 shows how straightforward the suggested solution is to integrate and utilize compared to traditional methods. The recommended structure is simple and versatile. Some common approaches are difficult to utilize and integrate. The comparison reveals how simple and successful the recommended solution is, showing its potential to speed up AI-driven system cybersecurity deployment and create a more effective and user-friendly security architecture.

Table 2: Comparison of Robustness

Method	Robustness
Proposed Method	90
Gradient Masking	70
Randomization Approach	75
Adversarial Training	60
Feature Squeezing	55
Model Ensembling	65
Secure Aggregation	50

Table 2 compares the recommended method's stability to conventional techniques. The recommended approach handles hostile attacks better than others due to its 90 dependability score. Old methods like Gradient Masking and the Randomization Approach are stable (70 and 75, respectively), but the recommended way is better.

Table 3: Comparison of Effectiveness

Method	Effectiveness
Proposed Method	85
Gradient Masking	75
Randomization Approach	70
Adversarial Training	65
Feature Squeezing	60
Model Ensembling	65
Secure Aggregation	55

Table 3 shows that the suggested solution protects AI-driven cybersecurity systems better (85) than current methods. The protection of AI-driven cybersecurity systems has improved.

5. Conclusion

Big attacks from evil people may weaken AI-based defenses. We must be clever and proactive to reduce these dangers and issues. This research explained how forceful strikes influence safety. Research on defensive reduction, aggressive training, and feature squeezing indicated that AI-powered security systems require robust defenses. Strong feature extraction, frequent retraining, and gradient masks may help the system avoid attacks and follow safe behaviors. Standard procedures were less reliable and beneficial than the proposed technology. It's handy, cheap, and can't be broken by intricate opponent strategies, making it a good defense against evil folks. Online dangers evolve and spread, so be alert. Thus, we must monitor emerging dangers and respond accordingly. Cybersecurity specialists, AI developers, and legislators from diverse industries must collaborate to safeguard AI-powered cybersecurity systems from new attack approaches. The findings demonstrate that AI-powered defensive systems must be safer through research and development. This project aims to improve AI-based safety systems in the digital era by educating people about dangerous threats and implementing robust defenses.

References

- [1] W. Stuxnet, "Stuxnet 2021," August-2021. [Online]. Available: <https://en.wikipedia.org/wiki/stuxnet%202021>.
- [2] R. Langner, "Stuxnet: Dissecting a Cyberwarfare Weapon," IEEE Security and Privacy Magazine, vol. 9, no. 3, pp. 49–51, 2011.
- [3] V. Roy. "An Effective FOG Computing Based Distributed Forecasting of Cyber-Attacks in Internet of Things" Journal of Cybersecurity and Information Management, Vol. 12, No. 2, 2023, PP. 8-17.
- [4] M. Muckin and S. C. Fitch, "A Threat-Driven Approach to Cyber Security," Lockheed Martin Corporation, MD, USA, 2014.
- [5] U.N.R. Commission, "Cyber Security Programs for Nuclear Facilities," US Nuclear Regulatory Commission, Office of Nuclear Regulatory Research, Rockville, MD, USA, 2010.
- [6] R. M. Blank, "Guide for Conducting Risk Assessments," CreateSpace Independent Publishing Platform, Scotts Valley, CA, US, 2011.
- [7] Sujeetha Devi, Bhagyalakshmi L and Sanjay Kumar Suman, "Enhancing the Performance of Wireless Sensor Networks through Clustering and Joint Routing with Mobile Sink", International Journal of Engineering and Advanced Technology, vol. 8, issue 6, pp. 323-327, 2019
- [8] L. Bhagyalakshmi, Sanjay Kumar Suman, S. Mohanalakshmi, and Satyanand Singh, "Improving Spectral Efficiency and Coverage Capacity of 5G Networks: A Review", Advances in mathematics: scientific journal, vol.9, no. 6, pp. 3387-3397, 2020.
- [9] Vanita Jain , Mahima Swami , Rishab Bansal, Exploratory Data Analysis on Username-Password Dataset, Fusion: Practice and Applications, Vol. 4 , No. 1 , (2021) : 5-14 (Doi : <https://doi.org/10.54216/FPA.040101>)
- [10] Aman Jain , Jatin Gupta , Somya Khandelwal , Surinder Kaur, Vehicle License Plate Recognition, Fusion: Practice and Applications, Vol. 4 , No. 1 , (2021) : 15-21 (Doi : <https://doi.org/10.54216/FPA.040102>)
- [11] S. Song, M. Lee, T. Kim, C. Park, S. Park, and H. Kim, "A Case Study on Cyber-Security Program for the Programmable Logic Controller of Modern NPPs," IAEA, Vienna, Austria, 2014.
- [12] J.-G. Song et al., "An Analysis of Technical Security Control Requirements for Digital I&C Systems in Nuclear Power Plants," Nuclear Engineering and Technology, vol. 45, no. 5, pp. 637–652, 2013.
- [13] J.-C. Loh et al., "On the Invisibility and Anonymity of Undeniable Signature Schemes," Journal of Wireless Mobile Networks, Ubiquitous Computing, and Dependable Applications (JoWUA), vol. 11, pp. 18–34, 2020.
- [14] Abhishta et al., "Why Would We Get Attacked? An Analysis of Attacker's Aims Behind DDoS Attacks," Journal of Wireless Mobile Networks, Ubiquitous Computing, and Dependable Applications (JoWUA), vol. 11, pp. 3–22, 2020.
- [15] Puneet Kaushal , Subash Chander , Vijay Kumar Sinha, Virtual Machine Placement in Cloud Computing: Challenges, Research Gaps, and Future, International Journal of Wireless and Ad Hoc Communication, Vol. 3 , No. 2 , (2021) : 64-71 (Doi : <https://doi.org/10.54216/IJWAC.030202>)
- [16] Lobna Osman, Evaluating the Performance of Battery Electric Vehicles using an Incorporated Decision Support Framework Based on Ranking Algorithms, International Journal of Wireless and Ad Hoc Communication, Vol. 3 , No. 2 , (2021) : 72-90 (Doi : <https://doi.org/10.54216/IJWAC.030203>)
- [17] H. P. Sahu, "FINE_DENSEIGANET: Automatic Medical Image Classification in Chest CT Scan Using Hybrid Deep Learning .Framework," International Journal of Image and Graphics [Preprint], 2023. [Online]. Available: <https://doi.org/10.1142/s0219467825500044>.
- [18] Ibrahim Elhenawy , Salwa H. Mahmoud , Ahmed Moustafa, A Lightweight Privacy Preserving Keyword Search Over Encrypted Data in Cloud Computing, Journal of Cybersecurity and Information Management, Vol. 3 , No. 2 , (2020) : 29-41 (Doi : <https://doi.org/10.54216/JCIM.030201>)
- [19] G. McGraw, "Software Security," IEEE Security and Privacy Magazine, vol. 2, no. 2, pp. 80–83, 2004.O. Clasp, "OWASP CLASP Project," 2015.
- [20] F. Valenza and M. Cheminod, "An Optimized Firewall Anomaly Resolution," Journal of Internet Services and Information Security (JISIS), vol. 10, pp. 22–37, 2020.
- [21] Noora Hani Sherif, Eay Fahidhil, Najlaa Nsrulaah Faris, Hussein Alaa Diame, Raaid Alubady, Seifedine Kadry, Modeling Sports Event Tasks in Augmentative and Alternative Communication Using Deep Learning, Journal of Intelligent Systems and Internet of Things, Vol. 9 , No. 2 , (2023) : 93-107 (Doi : <https://doi.org/10.54216/JISIoT.090207>)
- [22] Hussein Alaa Diame, Waleed Hameed, Zainab.R.Abdulsada, Noora Hani Sherif, Noor Hanoon Haroon, Narjes Benameur, M. A. Burhanuddin, Machine Learning Based Logistic Decision Support System for Intelligent Vehicles and Transportation Systems, Journal of Intelligent Systems and Internet of Things, Vol. 9 , No. 2 , (2023) : 108-119 (Doi : <https://doi.org/10.54216/JISIoT.090208>)

- [23] Pooja , Dr. Manish Kumar Mukhija , Satish Kumar Alaria, Smart City's Security Model for Management of Image Data on Cloud, Journal of Cognitive Human-Computer Interaction, Vol. 2 , No. 1 , (2022) : 8-14 (Doi : <https://doi.org/10.54216/JCHCI.020101>)
- [24] S.P. Samyuktha , Dr.P. Kavitha , V.A Kshaya , P. Shalini , R. Ramya, A Survey on Cyber Security Meets Artificial Intelligence: AI– Driven Cyber Security, Journal of Cognitive Human-Computer Interaction, Vol. 2 , No. 2 , (2022) : 50-55 (Doi : <https://doi.org/10.54216/JCHCI.020202>)
- [25] Dwivedi, A., Agarwal, R., & Shukla, P. K. (2023, July). Enhancing Anonymity of Internet of Vehicle Identities in Connected Vehicle Security Services Using Batch Verification Algorithm. In International Conference on Data Science and Applications (pp. 323-335). Singapore: Springer Nature Singapore.
- [26] Khare, A., Gupta, R., & Shukla, P. K. (2022). Improving the protection of wireless sensor network using a black hole optimization algorithm (BHOA) on best feasible node capture attack. In IoT and Analytics for Sensor Networks: Proceedings of ICWSNUCA 2021 (pp. 333-343). Springer Singapore.
- [27] Reddy Gantla, H., Ahmad, S. S., Matroud, A., Kalhotra, S. K., Agarwal, I., Gupta, S., & Mamodiya, U. (2023, November). Machine Learning-Based Trust-Aware Secure Traffic Mechanism to Identify DDOS Attacks over Cloud. In Proceedings of the 5th International Conference on Information Management & Machine Intelligence (pp. 1-7).
- [28] Kumar, S., Dubey, K. K., Gautam, A. K., Verma, S., Kumar, V., & Mamodiya, U. (2022). Detection of recurring vulnerabilities in computing services. Journal of Discrete Mathematical Sciences and Cryptography, 25(4), 1063-1071.