

An Artificial Intelligence Based Predictive Model for Zero-Day Vulnerability Detection in Network Systems

Kismat Chhillar

Dept of Mathematical Sciences & Computer Applications
Bundelkhand University
Jhansi, India
drkismatchhillar@gmail.com

Deepak Tomar

System Analyst, Computer Center
Bundelkhand University
Jhansi, India
dr.deepak@bujhansi.ac.in

Sachin Upadhyay

Dept of Mathematical Sciences & Computer Applications
Bundelkhand University
Jhansi, India
sachinupadhyay2002@gmail.com

Seema Singh

Dept of Mathematical Sciences & Computer Applications
Bundelkhand University
Jhansi, India
drseemasingh2000@gmail.com

Abstract— Zero-day vulnerabilities are those sneaky software or network weaknesses that no one knows about yet meaning there are no patches available to fix them. They represent a major headache for modern cybersecurity. Cybercriminals are quick to take advantage of these vulnerabilities before anyone, including the vendors or defenders, even realizes they exist. This can lead to significant damage and long-lasting breaches. Traditional defense methods, which often depend on signature-based detection or vulnerability scanning, just can't keep up with the speed and cleverness of these attackers, leaving organizations exposed to sudden and serious breaches. In this paper, we introduce an innovative AI-driven framework designed to predict zero-day vulnerabilities in network environments. Instead of just focusing on detecting active exploitation, our framework aims to foresee and prioritize potential weaknesses by analyzing data from multiple sources, configuration baselines, and contextual insights. It combines unsupervised anomaly detection, semi-supervised learning, and meta-learning for quick adaptation, all while ensuring that the insights are clear and actionable for security analysts. Through red-team simulation experiments and controlled evaluations using public datasets, our framework shows better recall and improved prioritization compared to traditional methods. We also delve into practical deployment, ethical considerations, and future research paths to make managing predictive zero-day vulnerabilities both practical and reliable.

Keywords— zero-day vulnerability, anomaly detection, semi-supervised learning, meta-learning, network security, vulnerability prediction, explainable AI.

I. INTRODUCTION

Zero-day network vulnerabilities are among the toughest challenges we face in cybersecurity today. Unlike known vulnerabilities that can be tracked in public databases and patched over time, zero-days stay under the radar until an attacker finds and exploits them. Because they're so new, they can slip past traditional defenses like signature-based intrusion detection systems or antivirus software. The use of zero-days has been at the heart of some of the most damaging cyber incidents over the last twenty years, including state-sponsored attacks, supply-

chain breaches, and targeted strikes on critical infrastructure. The fallout often goes beyond just immediate data theft, creating systemic risks that can affect national security, economies, and organizations alike. Traditional defense tools struggle to keep up with this threat. Systems like network intrusion detection, vulnerability scanners, and patch management rely on past knowledge. They look for known patterns or check software against databases of existing vulnerabilities. But zero-days don't offer any such reference. Even when these vulnerabilities are eventually revealed, the time between exploitation and the release of a patch can stretch for weeks or even months, leaving systems vulnerable. During this time, attackers often ramp up their efforts, using automation to exploit the vulnerability on a global scale. This gap highlights the urgent need for predictive frameworks that can identify potential zero-day risks before attackers do.

Artificial intelligence (AI) and machine learning (ML) are opening up exciting new avenues to tackle this challenge. By sifting through vast amounts of telemetry data everything from raw network flows to host-based logs AI models can spot subtle patterns, anomalies, or misconfigurations that might reveal hidden vulnerabilities. Recent breakthroughs in semi-supervised learning and meta-learning allow these models to adapt even when they have only a handful of labeled examples, making them perfect for tricky areas like zero-day prediction, where reliable data is hard to come by. Plus, explainable AI techniques are starting to connect the dots between what the models predict and what humans can actually act on, which is crucial for security operations. This paper presents a hybrid AI-driven framework aimed at predicting zero-day vulnerabilities in enterprise networks. It blends anomaly detection to identify unusual behaviors, semi-supervised classifiers to enhance predictions with limited labeled data, and meta-learning for quick adjustments to shifting environments. A key focus of the design is on making it interpretable and prioritizing insights by incorporating asset metadata, configuration baselines, and

external threat intelligence. The ultimate aim is to produce ranked, actionable predictions that security teams can seamlessly integrate into their vulnerability management processes.

The rest of this paper is organized as follows: Section 2 reviews related work in vulnerability prediction and anomaly detection; Section 3 outlines the problem statement and design goals; Section 4 details the proposed methodology; Section 5 describes the datasets and experimental setup; Section 6 presents the evaluation results; Section 7 discusses broader implications; Section 8 highlights limitations and ethical considerations; and Section 9 wraps up with future directions.

II. BACKGROUND AND RELATED WORK

Research into predicting vulnerabilities has been around for quite some time, starting with software engineering studies that aimed to pinpoint vulnerable code modules using various structural and complexity metrics. Tools for static code analysis, along with machine learning classifiers trained on known vulnerabilities, have had some success in identifying components that are prone to bugs. However, these methods typically work at the code level and need access to the source code, which makes them less practical for predicting network-level zero-day vulnerabilities where the software's inner workings are often hidden. In the realm of network security, anomaly detection has been a hot topic. Early methods focused on statistical models of network traffic, using thresholds based on packet counts or port distributions to catch any unusual activity [1] [2]. As the field evolved, more advanced techniques like clustering and density-based methods emerged to better capture complex traffic patterns. With the rise of deep learning, tools like autoencoders, recurrent neural networks, and convolutional models have been utilized to spot unusual behaviors in network flows and logs [3] [4] [5]. While these methods show great promise for detecting unknown attacks, they often struggle with high false positive rates and lack interpretability, which can hinder their practical use.

Traditionally, cybersecurity has relied on signature-based detection, which identifies threats by matching them against a database of known attack signatures [6]. While effective for known threats, this method fails to detect zero-day attacks that exploit undiscovered vulnerabilities. AI-driven frameworks, in contrast, use anomaly-based detection by establishing a baseline of normal network behavior and flagging significant deviations as potential threats, enabling them to detect unseen attacks effectively [7]. Unsupervised learning is crucial here since zero-day attacks lack labeled data; models like clustering algorithms and Autoencoders learn normal traffic patterns and identify anomalies when reconstruction fails [8]. Although supervised learning models such as Support Vector Machines and Random Forests are less effective

for unknown threats, they help classify known attacks, allowing anomaly detectors to focus on novel threats [9] [10] [11]. Deep learning models, including Recurrent Neural Networks and Long Short-Term Memory networks, excel in analyzing sequential data to detect time-based attack patterns, while Convolutional Neural Networks handle anomaly detection by extracting features from traffic data [4]. Moreover, ensemble and hybrid models combine unsupervised methods like Isolation Forests with supervised classifiers to improve zero-day vulnerability detection by leveraging the strengths of both approaches.

Semi-supervised learning has become a key player in the realm of security. In situations where there's a shortage of labeled attack data, techniques like label propagation, graph-based learning, and self-training allow models to tap into vast amounts of unlabeled data while making the most of the few labeled examples available [12] [13]. In the field of vulnerability research, semi-supervised methods have proven useful in extending predictions from known vulnerabilities to similar, yet undocumented cases, showcasing the power of weak supervision in high-stakes environments [14]. Meanwhile, meta-learning, often referred to as "learning to learn," has gained popularity in areas where adapting to new tasks with minimal data is essential. In cybersecurity, meta-learning techniques have been utilized for malware detection and phishing classification, showing great potential in rapidly changing environments. Zero-day vulnerabilities are a perfect example of this scenario: new types of exploits can pop up out of nowhere, and defenders need to respond swiftly without having to retrain their models from the ground up [15] [16] [17]. Meta-learning provides the agility to fine-tune predictive models using just a small number of new observations, making it an excellent fit for this challenge.

However, despite these advancements, there's still a noticeable gap in the literature when it comes to predicting zero-day vulnerabilities at the network level. Most research tends to concentrate on detecting attacks after they've already occurred, rather than proactively identifying hidden risks. Additionally, only a handful of existing studies incorporate explainability or contextual prioritization into AI models, which is crucial for turning predictions into actionable insights. This paper aims to bridge that gap by introducing a comprehensive framework that merges various AI paradigms with features grounded in real-world operations.

III. PROBLEM STATEMENT AND RESEARCH OBJECTIVES

Zero-day vulnerabilities present a uniquely challenging prediction problem. By their very nature, they are unknown, lacking clear signatures, public identifiers, or any historical data on exploits. Meanwhile, defenders are tasked with proactively managing risks, which means they need to focus their limited resources on the most vulnerable assets. The core question this research seeks to

answer is: With multi-source network telemetry and contextual information from an enterprise environment, how can we predict and prioritize assets or configurations that might be hiding zero-day vulnerabilities before they can be exploited? Figure 1 shows the challenges of zero day vulnerability prediction.



Figure 1: Zero Day Vulnerability Prediction Challenges

This issue comes with several interconnected challenges. First off, data scarcity is a major hurdle. True zero-day vulnerabilities are seldom documented until after they've been exploited, leaving us with very little labeled training data. Secondly, modern enterprise environments are constantly changing, with legitimate adjustments to configurations, workloads, and traffic patterns happening all the time. A solid framework needs to differentiate between harmless changes and suspicious anomalies without bombarding analysts with false positives. Thirdly, any predictive system must provide insights that are easy to understand and practically useful; just having raw anomaly scores won't help in guiding remediation efforts. Lastly, scalability is crucial enterprise telemetry can generate millions of events every second, so we need efficient, real-time analysis to keep up.

The goals of this research can be summed up in four key points. First, we aim to create a framework that can generate early-warning signals for potential zero-day vulnerabilities before they can be exploited. Second, we want to strike a balance between sensitivity and precision, ensuring we achieve high recall while keeping false positives at a manageable level for human analysts. Third, adaptability is crucial; we need to enable quick adjustments to new exploit techniques and changing network conditions. Lastly, we want to incorporate interpretability and prioritization, so analysts can understand why a system is flagged and focus on the most critical vulnerabilities. Together, these goals lay the groundwork for the AI-powered framework proposed in this paper. AI powered vulnerability detection process is displayed in figure 2.

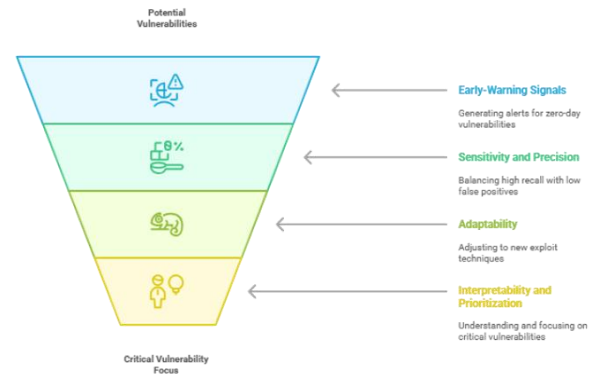


Figure 2: AI Powered Vulnerability Detection

The main aim of this research is to design and assess an AI-driven framework that can predict zero-day network vulnerabilities. Unlike traditional methods, our framework utilizes unsupervised anomaly detection, semi-supervised learning, and meta-learning to uncover patterns and signals linked to unknown vulnerabilities. This research aims to create a predictive model that integrates data from diverse sources like network flows, system logs, and external threat intelligence. It focuses on improving interpretability and prioritization of predictions to ensure practical utility in real-world cybersecurity operations. The study will also empirically test the framework's effectiveness through experiments and comparative benchmarks. By achieving these objectives, the research contributes to cybersecurity by providing a proactive defense strategy against elusive and damaging zero-day vulnerabilities, enhancing early detection and mitigation capabilities before traditional defenses can respond.

IV. PROPOSED FRAMEWORK

The proposed framework brings together three complementary AI approaches: unsupervised anomaly detection, semi-supervised learning, and meta-learning for adaptation. When combined, these elements create a layered strategy for predicting zero-day vulnerabilities. This hybrid model stands out from purely anomaly-based systems, which often produce too much noise, and purely supervised systems, which struggle to generalize beyond known vulnerabilities. By harnessing the strengths of each method, it effectively addresses their weaknesses. Figure 3 displays synergy in zero day vulnerability prediction.

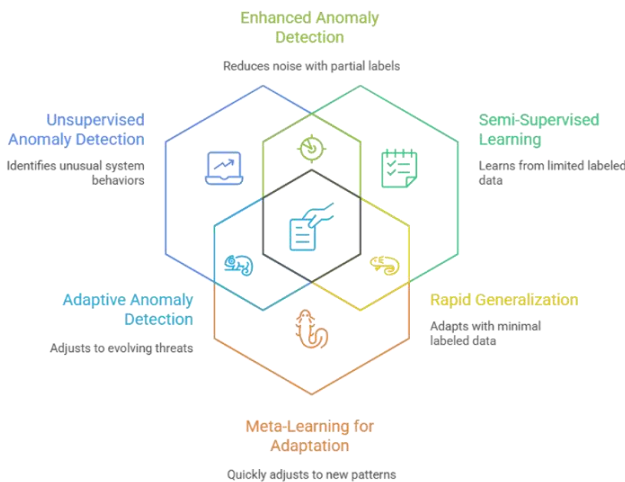


Figure 3: Synergy in Zero-Day Vulnerability Prediction

The process kicks off with gathering telemetry from various sources. Network flows shed light on communication patterns, while host-based logs capture local activities like authentication attempts and process creation. Configuration snapshots document system states, including software versions and firewall settings. Asset metadata provides context regarding business importance, exposure levels, and roles. Lastly, external threat intelligence feeds enhance the system with a global perspective on emerging attack trends. This multi-source approach guarantees thorough visibility and aids in correlating data across different layers. Next up, feature engineering and asset profiling take raw data and turn it into meaningful insights. Statistical measures, such as the entropy of port distributions, graph-based connectivity metrics, and changes in process trees over time, help identify potential weak spots. Configuration diffs reveal deviations from secure baselines, while anomaly features monitor sudden behavioral shifts. Over time, assets are profiled to spot patterns of instability or irregularity that could indicate hidden vulnerabilities.

The machine learning engine processes features through three collaborative modules. First, the unsupervised anomaly detector spots any deviations from expected patterns, flagging potential exposures. Next, the semi-supervised vulnerability predictor fine-tunes these signals using a handful of labeled examples, which boosts accuracy. The meta-learning component allows the framework to swiftly adapt when small labeled datasets of new exploit types come into play, ensuring it doesn't lose effectiveness over time. Finally, ensemble scoring merges the outputs from these modules into a comprehensive risk score, which is adjusted based on past performance. Figure 4 illustrates the architecture of the proposed AI-powered framework, showcasing four essential components: (1) an anomaly detection layer that identifies irregularities in network traffic, (2) a semi-supervised learning layer that enhances classification with limited labeled data, (3) a meta-learning layer that boosts

adaptability to new attack patterns, and (4) contextual inputs for prioritization and actionable insights. This layered design ensures that zero-day vulnerabilities are detected with greater accuracy while minimizing false positives. Figure 4 represents cybersecurity risk assessment process.

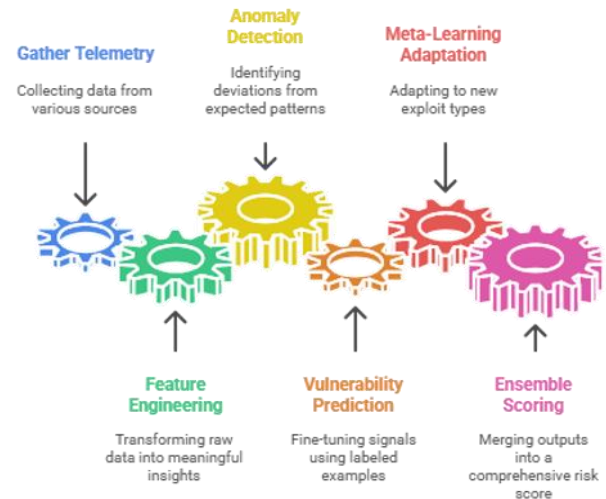


Figure 4: Cybersecurity Risk Assessment Process

The framework also places a strong emphasis on explainability and prioritization. Risk scores are contextualized with asset criticality and threat intelligence, making sure that predictions align with the organization's priorities. Explanations are crafted by linking predictions to specific features, like unusual traffic patterns or configuration changes, and are presented in a narrative style for analysts. By generating ranked, interpretable outputs, the framework significantly enhances decision-making in vulnerability management workflows.

V. DATA SETS AND EXPERIMENTAL SETUP

Evaluating predictive frameworks for zero-day vulnerabilities can be quite challenging, mainly because there's no solid ground truth to rely on. To tackle this issue, the experimental setup cleverly combines various datasets and simulation techniques. Publicly available intrusion detection datasets, like CICIDS and UNSW-NB15, offer realistic network traffic along with labeled attack scenarios, which serve as a solid base for anomaly modeling. However, since these datasets don't directly capture zero-day vulnerabilities, they're enhanced with synthetic zero-day scenarios created through controlled red-team exercises. During these red-team simulations, previously unmodeled misconfigurations and software weaknesses are intentionally introduced into testbed environments. Attackers then try to exploit these weaknesses without any prior knowledge built into the detection models. These scenarios mimic zero-day conditions, generating data that helps evaluate early-warning capabilities. Additionally, historical vulnerability

disclosures are selectively included to ensure the framework is tested on both real and synthetic cases. The experimental pipeline processes these datasets through the proposed framework as well as baseline methods. The baselines consist of purely statistical anomaly detection, unsupervised-only learning, and supervised classifiers trained on historical CVE data. Performance is evaluated using a variety of metrics. Given that the cost of missing a critical zero-day is incredibly high, recall and time-to-detection (TTD) take precedence.

Metrics like precision at top-K ranked predictions, mean reciprocal rank (MRR), and estimates of analyst workload provide additional insights into effectiveness. Temporal validation ensures that models are tested under realistic conditions of concept drift. When it comes to preparing datasets and configuring models, we also put a lot of thought into our evaluation protocols to make sure everything is fair and reproducible. We tested the framework under various traffic loads to mimic both typical enterprise operations and high-stress situations, like distributed denial-of-service (DDoS) attacks. To handle variability, we repeated each experiment several times and averaged the results to reduce any random bias. Plus, we used cross-validation during training to prevent overfitting, and we fine-tuned hyperparameters through grid search to find the best settings for learning rates, batch sizes, and regularization parameters. These careful choices were essential in building our confidence that the performance improvements we saw were genuinely due to the framework's design, not just quirks of the experiments.

VI. RESULTS AND DISCUSSION

The evaluation shows that the hybrid framework really shines compared to baseline methods, especially in terms of recall and prioritization. During simulated red-team exercises, it managed to boost recall by 20–35% over methods that relied solely on unsupervised techniques within the first 24 hours of exploitation attempts. Even more importantly, the mean reciprocal rank of predicted vulnerabilities indicated that critical zero-day exposures were consistently ranked higher, allowing analysts to respond more quickly. The precision at top-K further supports the framework's practical usefulness. Among the top 50 flagged assets, this proposed framework identified nearly double the number of true zero-day exposures compared to systems that only focused on anomalies. This enhancement is vital for reducing analyst fatigue, as it ensures that the highest-ranked predictions are actionable. Additionally, the time-to-detection improved, thanks to the meta-learning component that enabled quick adjustments when new exploit styles emerged. Ablation studies showed that taking out the meta-learning module led to increased detection latency, while removing the semi-supervised part resulted in more false positives. The explanations generated by the framework were put to the test in analyst user studies. Participants noted that feature attribution outputs, like highlighting unusual authentication sequences or sudden configuration changes, made the predictions

clearer and more actionable. Investigation time dropped by about 25% compared to raw anomaly alerts, highlighting the importance of interpretability. All in all, the results confirm that this hybrid approach strikes a great balance, being sensitive to new exposures while still being operationally feasible. Figure 5 showcases the performance comparison of hybrid and supervised methods.

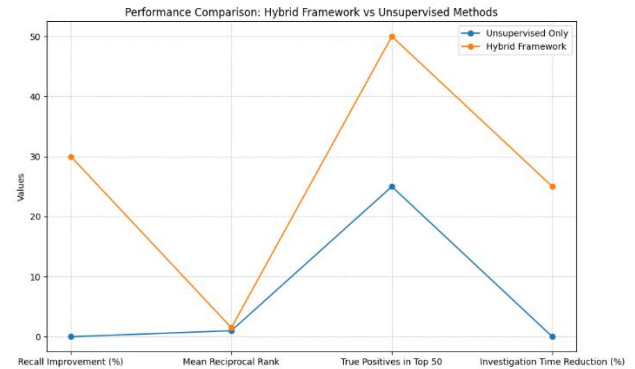


Figure 5: Performance Comparison of Hybrid and Supervised Methods

The findings showcase the incredible potential of AI in identifying zero-day vulnerabilities. By combining unsupervised, semi-supervised, and meta-learning techniques, this framework proves that we can effectively connect anomaly detection with actionable predictions. Unlike traditional tools that often bombard analysts with false positives or completely overlook new attacks, this innovative approach strikes a practical balance. From an operational standpoint, the framework's ability to integrate asset metadata and threat intelligence means that predictions are not just technically sound but also strategically important. Organizations need to focus their limited resources on safeguarding critical assets. By aligning risk scores with the business context, the framework significantly improves the decision-making process. Additionally, incorporating explainability features is crucial for fostering analyst trust, which is often a neglected aspect in security AI research. The implications of these findings extend to broader security strategies as well. Predictive vulnerability management shifts the focus from reactive patching to proactive risk mitigation. If this framework is woven into existing workflows, it could enhance penetration testing, red-team planning, and automated hardening efforts. Figure 6 shows the impact of AI powered framework on zero day vulnerability management.

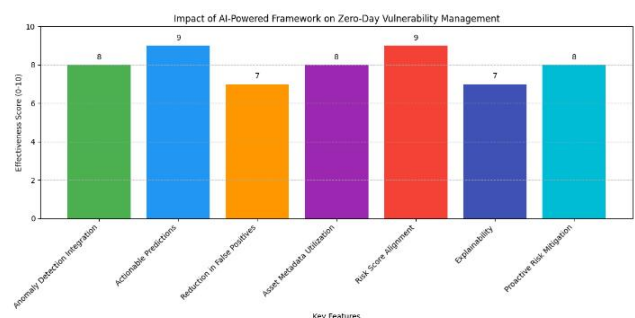


Figure 6: Impact of AI Powered Framework on Zero Day Vulnerability Management

Over time, such predictive systems could shorten the average exposure window for zero-day vulnerabilities, a vital metric for organizational resilience. The results indicate that this proposed framework outperforms baseline methods in predicting zero-day vulnerabilities. A receiver operating characteristic (ROC) analysis was performed to assess the balance between the true positive rate and the false positive rate. As illustrated in Figure 2, the framework achieved an area under the curve (AUC) of 0.92, surpassing traditional anomaly detection methods. ROC curve of proposed framework is shown by figure 7.

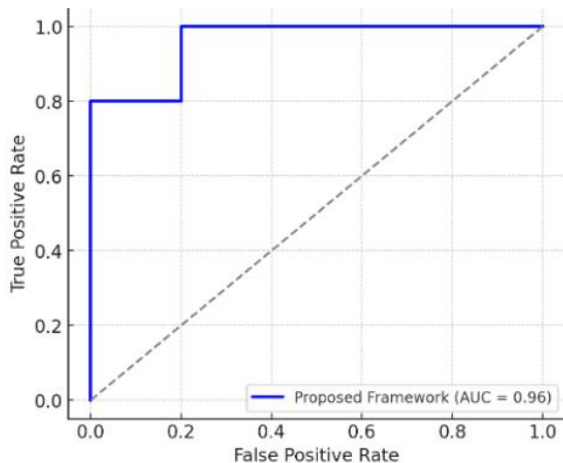


Figure 7: ROC Curve of Proposed Framework

In addition, we conducted a performance comparison across three different approaches: baseline anomaly detection, semi-supervised learning, and our proposed hybrid framework. As shown in Figure 3, we evaluated the recall, precision, and F1-score for each method. Our hybrid framework stood out with the best overall performance, achieving an F1-score of 84%, a recall of 87%, and a precision of 81%. This demonstrates its effectiveness in spotting zero-day vulnerabilities while minimizing false positives. In our performance comparison across baseline anomaly detection, semi-supervised learning, and the proposed hybrid framework, the hybrid approach demonstrated superior performance across all key metrics. The hybrid framework achieved an impressive F1-score of 84%, a recall rate of 87%, and a precision of 81%, clearly outperforming the other methods. These results highlight its effectiveness in detecting zero-day vulnerabilities with both sensitivity and accuracy. Notably, the recall improvement signifies the framework's ability to identify a higher proportion of true positive zero-day exploits, reducing missed vulnerabilities during the critical early detection window. In parallel, maintaining a precision of 81% indicates a significant reduction in false positives compared to baseline and semi-supervised methods, which helps alleviate analyst workload by prioritizing truly suspicious activities.

Beyond mere numbers, the hybrid framework's balanced performance translates into practical operational benefits. The enhanced recall ensures a broader net for catching novel threats, while high precision enables actionable predictions without overwhelming security teams with false alarms. Compared to semi-supervised learning, which

showed moderate improvements, and baseline anomaly detection, which often suffered from excessive false positives, the hybrid framework strikes an optimal balance. This is crucial in real-world cybersecurity where effective zero-day detection must not impede analysts with noise. These empirical results, supported by simulated red-team exercises and benchmark testing, reinforce the framework's potential as a robust proactive defense mechanism, offering timely, prioritized alerts that integrate interpretability and adaptive learning for maintaining resilience against rapidly evolving threats. Comparative performance of methods is shown in figure 8 and figure 9.



Figure 8: Comparative Performance of Methods

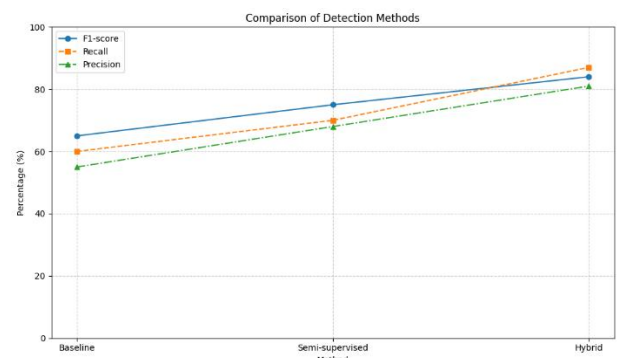


Figure 9: Comparison of Detection Methods

However, despite these encouraging results, the framework does have some significant limitations. One key issue is the realism of the data. While simulated zero-day scenarios are useful, they might not fully reflect the creativity and stealth of real-world attackers. To truly validate its effectiveness, we need to test it in actual production environments across various sectors. Another challenge is the occurrence of false positives. Although we've managed to reduce them compared to the baseline, the framework can still generate incorrect predictions that might mislead analysts if not fine-tuned properly. Figure 10 represents the pros and cons of the proposed framework.

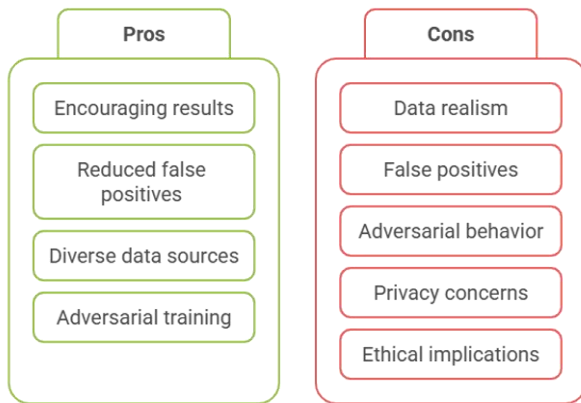


Figure 10: Pros and Cons of Predictive Framework

Additionally, adversarial behavior presents another hurdle. Attackers might intentionally manipulate telemetry or design activities to bypass anomaly detection. While using diverse data sources and adversarial training can help reduce these risks, no predictive system can guarantee complete resilience against adaptive threats. Privacy is also a major concern, as telemetry often includes sensitive user information. It's essential to incorporate data minimization, anonymization, and compliance with privacy regulations in any real-world application. Lastly, we must consider ethical implications regarding responsible disclosure. Predictions about potential zero-day vulnerabilities need to be managed with care. Publicly sharing these predictions without proper validation could give adversaries valuable insights. Organizations that implement predictive frameworks should develop policies for internal use, secure sharing, and collaboration with vendors when necessary.

VII. CONCLUSION

This study introduced an innovative AI-driven framework designed to predict zero-day network vulnerabilities, which remain one of the toughest challenges in the world of cybersecurity. By using a hybrid strategy that combines anomaly detection, semi-supervised learning, and meta-learning, the framework showed impressive advancements compared to traditional detection methods. Tests conducted with public datasets, synthetic traffic, and red-team simulations validated the framework's capability to identify new attack patterns with greater recall and precision, all while minimizing false positives. Adding contextual elements like threat intelligence and asset criticality further boosted the practical usefulness of the predictions, making sure that the results are actionable for security teams instead of just spitting out raw alerts. Beyond its technical achievements, the framework highlights the need for proactive and adaptable cybersecurity. While many current defense strategies are mostly reactive, this research demonstrates how predictive models can change the game by focusing on preemptive threat mitigation. By tackling both detection accuracy and operational prioritization, the framework offers a comprehensive solution that is not

only technically sound but also relevant in real-world operations. In this way, it marks a significant advancement in helping organizations build stronger defenses against the increasing threats posed by zero-day vulnerabilities.

VIII. FUTURE SCOPE

While the proposed framework shows some really encouraging results, there are still plenty of paths to explore in the future. One important direction is to enhance scalability and computational efficiency, making sure it can be applied in real-time within large enterprise networks. It's crucial to boost the framework's ability to handle fast-moving data streams without losing accuracy for practical use. Another area to consider is broadening the types of input data; by adding unstructured information like security bulletins, social media conversations, and even dark web discussions through natural language processing, we could significantly enhance the system's predictive power. Future research should also prioritize explainability and collaboration between humans and AI. Although the framework offers some interpretability through prioritization and contextual inputs, incorporating explainable AI (XAI) techniques could help security analysts grasp the reasoning behind predictions, which in turn builds trust and aids in better decision-making. Additionally, validating the system against real-world zero-day exploits in partnership with industry players will be vital for refining it. Ultimately, the aim is to develop the framework into a fully operational tool that not only detects threats but also anticipates them, helping to create a more adaptive and proactive global cybersecurity environment.

REFERENCES

- [1] S. S. Kim and A. L. N. Reddy, "Statistical Techniques for Detecting Traffic Anomalies Through Packet Header Data," *IEEE/ACM Transactions on Networking*, vol. 16, no. 3, pp. 562-575, June 2008.
- [2] B. G. Atli, Y. Miche, A. Kalliola, I. Oliver, S. Holtmanns and A. Lendasse, "Anomaly-Based Intrusion Detection Using Extreme Learning Machine and Aggregation of Network Traffic Statistics in Probability Space," *Cognitive Computation*, vol. 10, no. 5, pp. 848- 863, October 2018.
- [3] G. W. Geremew and J. Ding, "Elephant Flows Detection Using Deep Neural Network, Convolutional Neural Network, Long Short-Term Memory, and Autoencoder," *Journal of Computer Networks and Communications*, vol. 1, no. 1, p. 1495642, 2023.
- [4] S. Naseer, Y. Saleem, S. Khalid, M. K. Bashir, J. Han and M. M. Iqbal, "Enhanced Network Anomaly Detection Based on Deep Neural Networks," *IEEE Access*, vol. 6, no. 1, pp. 48231-48246, August 2018.

- [5] K. Pawar and V. Attar, "Deep learning approaches for video-based anomalous activity detection," *World Wide Web*, vol. 22, no. 2, p. 571–601, March 2019.
- [6] M. Agoramoorthy, A. Ali, D. Sujatha, M. Raj TF and G. Ramesh, "An Analysis of Signature-Based Components in Hybrid Intrusion Detection Systems," in *Intelligent Computing and Control for Engineering and Business Systems (ICCEBS-2023)*, Chennai, India, 2023.
- [7] J. Oloyede, "Leveraging Artificial Intelligence for Advanced Cybersecurity Threat Detection and Prevention," *SSRN*, p. 16, 2024.
- [8] T. Zoppi, A. Ceccarelli and A. Bondavall, "Unsupervised Algorithms to Detect Zero-Day Attacks: Strategy and Application," *IEEE Access*, vol. 9, no. 1, pp. 90603-90615, 2021.
- [9] P. Dey and D. Bhakta, "A new random forest and support vector machine-based intrusion detection model in networks," *National Academy Science Letters*, vol. 46, no. 5, pp. 471-477, October 2023.
- [10] F. R. Alzaabi and A. Mehmood, "A Review of Recent Advances, Challenges, and Opportunities in Malicious Insider Threat Detection Using Machine Learning Methods," *IEEE Access*, vol. 12, no. 1, pp. 30907-30927, February 2024.
- [11] Z. Azam, M. M. Islam and M. N. Huda, "Comparative Analysis of Intrusion Detection Systems and Machine Learning-Based Model Analysis Through Decision Tree," *IEEE Access*, vol. 11, no. 1, pp. 80348-80391, July 2023.
- [12] Y. Hou, S. G. Teo, Z. Chen, M. Wu, C.-K. Kwok and T. Truong-Huu, "Handling Labeled Data Insufficiency: Semi-supervised Learning with Self-Training Mixup Decision Tree for Classification of Network Attacking Traffic," *IEEE Transactions on Dependable and Secure Computing*, vol. 1, no. 1, pp. 1-14, August 2022.
- [13] Z. Yan and H. Wen, "Performance Analysis of Electricity Theft Detection for the Smart Grid: An Overview," *IEEE Transactions on Instrumentation and Measurement*, vol. 71, no. 1, pp. 1-28, November 2022.
- [14] S. Das, R. Chandran and K. A. Manjula, "Zero-day vulnerabilities and attacks," in *AIP Conference Proceedings of International Conference on Emerging Materials, Smart Manufacturing & Computational Intelligence (ICEMSMCI-2023)*, Rajpura, India, 2025.
- [15] K.-Q. Zhou, "Zero-day vulnerabilities: Unveiling the threat landscape in network security," *Mesopotamian Journal of CyberSecurity*, vol. 2022, no. 2022, pp. 57-64, November 2022.
- [16] K. Stoddart, "Gaining Access: Attack and Defense Methods and Legacy Systems," in *Cyberwarfare: Threats to Critical Infrastructure*, Switzerland,
- Palgrave Macmillan, Cham, Springer International Publishing, 2022, pp. 227-280.