

AI-Enhanced Detection of ARP Spoofing-Based Man-in-the-Middle Attacks in Local Area Networks

Penka Markova

Computer Science and Engineering
Technical University of Varna
Varna, Bulgaria
penkamarkova5@gmail.com

Georgi Markov

Department of Information Technologies
Nikola Vaptsarov Naval Academy
Varna, Bulgaria
g.markov@naval-acad.bg

Abstract—Man-in-the-Middle (MitM) attacks remain a significant threat to local area networks, particularly when implemented through Address Resolution Protocol (ARP) spoofing. Although infrastructure-level protection mechanisms such as Dynamic ARP Inspection (DAI) are effective against conventional ARP poisoning attempts, their detection capability is limited in low-rate and insider-like scenarios, where malicious behavior may appear formally valid while remaining anomalous in context. This paper proposes a hybrid LAN protection architecture that combines DAI with an AI-based behavioral detection module to improve the identification of stealthy and context-dependent MitM activity. The proposed approach analyzes ARP, DHCP, and switch-port events using interpretable traffic features and a two-stage detection logic that integrates anomaly filtering and supervised classification. The system was evaluated in a controlled laboratory environment under three attack scenarios: classic high-rate ARP spoofing, low-rate stealth spoofing, and insider spoofing. The experimental results show that DAI alone performs very well in classic spoofing conditions, but its effectiveness decreases substantially in low-rate and insider scenarios. In contrast, the proposed DAI+AI architecture achieves the best overall balance between recall, precision, F1-score, detection latency, and false positive rate. The findings indicate that AI can effectively complement traditional infrastructure-level network defenses by providing behavioral awareness and improved sensitivity to attack patterns that are difficult to detect through rule-based inspection alone.

Keywords— anomaly detection, ARP spoofing, artificial intelligence, behavioral analysis, Dynamic ARP Inspection, intrusion detection system, LAN security, Man-in-the-Middle attack.

I. INTRODUCTION

Modern enterprise networks are characterized by extensive interconnectivity, rapid expansion of communication services, and increasing integration of cloud platforms, mobile devices, and Internet of Things

(IoT) components. While this transformation improves flexibility and functionality, it also enlarges the attack surface and introduces new vulnerabilities across both local and distributed environments. In this context, securing communication channels at the local area network (LAN) level remains a critical element of overall cyber defense [1], [2].

Among the major threats to Ethernet/IP infrastructures are Man-in-the-Middle (MitM) attacks, in which an adversary positions itself between two communicating hosts in order to intercept, modify, or redirect traffic. In IPv4 LANs, one of the most common practical realizations of MitM is ARP spoofing (also referred to as ARP poisoning), where forged ARP messages are used to alter IP-to-MAC mappings in the ARP cache of network devices. As a result, traffic is redirected through the attacker-controlled host, enabling packet inspection, credential theft, session compromise, and further lateral exploitation [3], [4], [5].

The severity of ARP spoofing stems from the fact that the ARP protocol does not provide built-in sender authentication. Consequently, forged ARP replies and gratuitous ARP messages can update dynamic ARP entries on target hosts without any cryptographic verification. Prior studies have shown that ARP-based MitM attacks remain operationally relevant even in relatively simple LANs, particularly when combined with unencrypted services or poorly monitored internal segments [4], [6].

Traditional defenses against MitM attacks include cryptographic protection, static filtering rules, ARP inspection mechanisms, and signature-based intrusion detection and prevention systems. For ARP-specific attacks, Dynamic ARP Inspection (DAI) is one of the most widely deployed infrastructure-level countermeasures. DAI validates ARP messages against DHCP snooping bindings and/or statically configured IP-to-MAC associations, making it effective against obvious ARP

Online ISSN 3044-7771

<https://doi.org/10.68302/std2026.vol3.52>

© 2026 The Author(s). Published by Green Future Technologies.

This is an open-access article under the [Creative Commons Attribution 4.0 International License](https://creativecommons.org/licenses/by/4.0/).

poisoning attempts. However, the literature indicates that rule-based defenses become less effective when malicious activity appears formally valid yet deviates from normal behavior in statistical, temporal, or contextual terms. Similar limitations are observed in signature-based IDS/IPS systems, which remain effective for known attack patterns but show limited adaptability to stealthy, evolving, or low-rate threats [6], [7], [8].

These limitations have motivated growing interest in artificial intelligence (AI) and machine learning (ML) for network security applications. Unlike strictly rule-based or signature-driven approaches, AI/ML models can identify anomalies by analyzing behavioral, statistical, and temporal dependencies in network traffic. Such methods are particularly suitable for attacks that manifest not through a single highly distinctive packet, but through weak yet correlated deviations over time. Existing research has demonstrated the applicability of both conventional ML algorithms, such as Decision Trees, Random Forests, and Support Vector Machines, and deep learning models, including recurrent neural networks (RNNs), Long Short-Term Memory (LSTM) networks, and hybrid architectures, for intrusion detection and anomaly analysis [8], [9], [10], [11].

Recent studies further indicate that specialized AI-based approaches can improve the detection of ARP-related attacks. For example, explainable deep learning has been applied to ARP spoofing detection in IoT settings, highlighting not only accuracy gains but also the importance of interpretability for network administrators and security operation teams [12]. At the same time, the literature still reports persistent challenges, including limited availability of realistic labeled datasets, restricted model explainability, and vulnerability to adversarial manipulation. These issues suggest that black-box AI models, although often effective, may not be sufficient on their own for practical deployment in operational environments [13], [14], [15].

Against this background, a clear research gap emerges between traditional infrastructure-level defenses and adaptive behavioral analysis. While the combination of AI techniques with rule-based security mechanisms is not entirely new, most existing approaches either focus on general-purpose intrusion detection or treat ARP spoofing as one of many attack classes within broader anomaly detection frameworks. In contrast, this paper focuses specifically on ARP spoofing-based MitM attacks in LAN environments and evaluates how infrastructure-level ARP validation can be enhanced by AI-driven behavioral analysis of ARP, DHCP, and switch-port related events.

Rather than replacing established mechanisms such as DAI, the proposed approach augments them with an intelligent analytical layer capable of detecting low-intensity, context-dependent, and insider-like anomalies that may not be sufficiently visible to static validation rules alone. The main contribution of the paper is therefore not the isolated use of AI, but the experimental evaluation of a hybrid DAI+AI architecture under controlled ARP-based MitM attack conditions. The study compares DAI-only, AI-only, and DAI+AI configurations in order to assess

whether the hybrid approach improves detection sensitivity while preserving the operational advantages of conventional network defenses.

II. MATERIALS AND METHODS

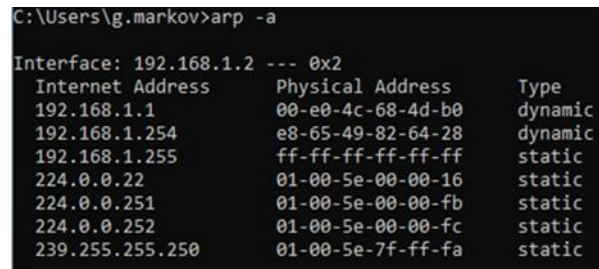
A. Experimental Objective

The objective of this study is to evaluate the effectiveness of a hybrid defense approach against MitM attacks in local area networks, based on the combination of DAI and an AI-based behavioral detection module. The proposed architecture is designed to determine whether the integration of infrastructure-level ARP validation and intelligent anomaly detection can improve sensitivity to low-rate and insider-like ARP-based attacks that may remain undetected by purely rule-based mechanisms.

B. Experimental Setup and Attack Scenarios

The experiments were conducted in a controlled laboratory LAN environment designed to reproduce realistic ARP-based MitM attack conditions. The test topology included a client host (192.168.1.2), router R1 (192.168.1.1), a gateway device (192.168.1.254), an attacker machine running Ettercap, a Layer-2 Cisco switch with DAI enabled, and a separate NIDS/AI host connected via a SPAN port. The AI module received mirrored traffic for feature extraction and event analysis.

Before attack execution, the initial ARP state of the participating devices was verified in order to confirm valid and stable IP-to-MAC associations. This baseline state is illustrated in Fig. 1. A period of normal traffic was then collected to establish a reference profile for benign ARP behavior. The attacker-side configuration used for ARP spoofing is shown in Fig. 2.



```
C:\Users\g.markov>arp -a
```

Internet Address	Physical Address	Type
192.168.1.1	00-e0-4c-68-4d-b0	dynamic
192.168.1.254	e8-65-49-82-64-28	dynamic
192.168.1.255	ff-ff-ff-ff-ff-ff	static
224.0.0.22	01-00-5e-00-00-16	static
224.0.0.251	01-00-5e-00-00-fb	static
224.0.0.252	01-00-5e-00-00-fc	static
239.255.255.250	01-00-5e-7f-ff-fa	static

Fig. 1. Initial ARP state before attack execution.

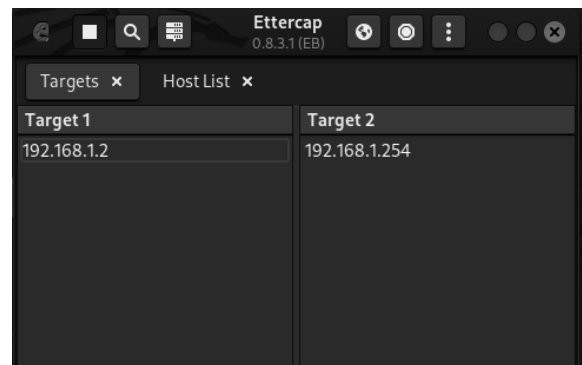


Fig. 2. Ettercap configuration for ARP spoofing execution.

Three attack scenarios were defined for evaluation:

- Scenario A – classic high-rate ARP spoofing: frequent forged ARP replies and gratuitous ARP messages were injected in order to rapidly poison the ARP cache of the victim and the gateway;
- Scenario B – low-rate stealth spoofing: forged ARP messages were sent at lower frequency and with longer intervals to reduce the likelihood of threshold-based detection.
- Scenario C – insider spoofing: the attack was performed from a compromised legitimate host, simulating an internal adversary with apparently valid network presence.

The attacks were executed using Ettercap configured for ARP spoofing between the victim host and the gateway. The resulting change in ARP associations after successful poisoning is illustrated in Fig. 3.

```
C:\Users\g.markov>arp -a

Interface: 192.168.1.2 --- 0x2
Internet Address      Physical Address      Type
192.168.1.1           00-e0-4c-68-4d-b0    dynamic
192.168.1.3           30-9c-23-d9-81-61    dynamic
192.168.1.254         e8-65-49-82-64-28    dynamic
192.168.1.255         ff-ff-ff-ff-ff-ff    static
224.0.0.22            01-00-5e-00-00-16    static
224.0.0.251           01-00-5e-00-00-fb    static
224.0.0.252           01-00-5e-00-00-fc    static
239.255.255.250       01-00-5e-7f-ff-fa    static

C:\Users\g.markov>arp -a

Interface: 192.168.1.2 --- 0x2
Internet Address      Physical Address      Type
192.168.1.1           00-e0-4c-68-4d-b0    dynamic
192.168.1.3           30-9c-23-d9-81-61    dynamic
192.168.1.254         e8-65-49-82-64-28    dynamic
192.168.1.255         ff-ff-ff-ff-ff-ff    static
224.0.0.22            01-00-5e-00-00-16    static
224.0.0.251           01-00-5e-00-00-fb    static
224.0.0.252           01-00-5e-00-00-fc    static
239.255.255.250       01-00-5e-7f-ff-fa    static
```

Fig. 3. Modified ARP entries after successful ARP spoofing.

C. Data Collection and Feature Extraction

For each experiment, packet captures (pcap files), DAI logs, DHCP bindings, and MAC-to-port history records were collected. In addition to real-time monitoring through the SPAN port, an offline replay mode based on previously recorded pcap files was used in order to safely and repeatedly test the low-rate and insider scenarios.

A set of simple and interpretable features was extracted from the collected pcaps and logs for subsequent detection. These features included the number of ARP packets, the number of gratuitous ARP messages, the number of IP-to-MAC mapping changes, MAC address entropy, correlation with DHCP binding information, and changes in switch port association. The use of explainable features was intentionally preferred in order to support practical interpretability, easier forensic analysis, and smoother integration with network administration workflows.

Traffic analysis was performed using two time windows: a short window of 60 s for rapid detection of abrupt anomalies and a longer window of 300 s for detecting weaker, temporally distributed deviations typical of low-rate and insider attacks.

The extracted feature vectors were normalized before being provided to the AI module, so that features with different numerical ranges, such as packet counts and entropy values, would not disproportionately influence the classification process.

D. Hybrid DAI+AI Detection Model

The proposed system follows a two-stage detection logic. In the first stage, an unsupervised anomaly detection module is used to identify deviations from the baseline profile of normal LAN behavior. In the experimental implementation, this stage is based on the analysis of ARP-related and flow-level features extracted over 60 s and 300 s observation windows. The purpose of this stage is to detect unusual changes in ARP activity, IP-to-MAC mappings, MAC-to-port associations, and deviations from DHCP binding information.

In the second stage, suspicious events are evaluated by a supervised classification module trained on labeled samples from normal traffic and the three defined attack scenarios. The classifier uses the extracted feature vectors to distinguish between benign behavior and ARP spoofing-based MitM activity. The training dataset includes samples collected during the baseline phase and during Scenario A, Scenario B, and Scenario C. The dataset is divided into training and testing subsets to evaluate the model on traffic samples not used during training.

The AI decision is not used as a standalone blocking mechanism. Instead, it is correlated with DAI events and infrastructure-level indicators. A high-confidence alert is generated when the AI module detects anomalous behavior and this behavior is supported by one or more protocol-level indicators, such as unexpected IP-to-MAC mapping changes, inconsistencies with DHCP bindings, or abnormal MAC-to-port transitions. This correlation reduces the risk of false positives compared with AI-only detection and improves sensitivity compared with DAI-only protection.

The evaluation considered three configurations:

- DAI-only, where detection and prevention rely exclusively on the switch-based DAI mechanism;
- AI-only, where only the NIDS/AI module is active, without DAI enforcement;
- DAI+AI, which represents the proposed hybrid architecture.

When high-risk activity is detected, the system generates an alert and stores a forensic pcap snapshot for subsequent analysis. Under very high confidence conditions, a temporary response may also be applied, such as temporary port shutdown or a temporary DAI ACL, in accordance with a human-in-the-loop policy.

E. Evaluation Metrics and Reproducibility

The effectiveness of the evaluated configurations was assessed using the following metrics: Precision, Recall, F1-score, ROC-AUC, time-to-detect, and the number of false positives per 1000 hosts per day. These metrics were selected in order to jointly assess classification quality, operational responsiveness, and practical deployability.

All experiments were conducted in an isolated environment and documented through stored pcap files, DAI logs, DHCP binding records, and MAC-to-port history records. This allows the same traffic traces to be replayed and reanalyzed under different detection configurations, improving reproducibility and forensic traceability. The main experimental sequence consists of baseline traffic collection, execution of the defined ARP spoofing scenario, feature extraction over 60 s and 300 s observation windows, and evaluation of the DAI-only, AI-only, and DAI+AI configurations.

However, it should be noted that the laboratory setup does not fully reproduce the complexity of real enterprise networks, where larger host populations, heterogeneous services, multiple VLANs, and dynamic background traffic may affect ARP and DHCP behavior. Therefore, the reported results should be interpreted as an experimental validation of the proposed approach rather than a complete assessment for all possible production environments. Future work will include validation in larger-scale and more heterogeneous LAN environments.

III. RESULT AND DISCUSSION

A. Quantitative Performance Across Attack Scenarios

To evaluate the proposed system, a series of experiments was conducted in a controlled environment using three attack scenarios: (A) classic high-rate ARP spoofing, (B) low-rate stealth spoofing, and (C) insider spoofing. Each configuration was replicated ten times, with twenty attack injections per replication, resulting in a total of 200 injections per scenario. The performance of DAI-only, AI-only, and DAI+AI was assessed using the metrics Precision, Recall, F1-score, and ROC-AUC. The aggregated results are presented in Table 1.

TABLE 1 CLASSIFICATION PERFORMANCE ACROSS ATTACK SCENARIOS

Scenario	Configuration	Precision	Recall	F1-score	ROC-AUC
A (classic)	DAI-only	0.99 ±0.01	0.98 ±0.02	0.985 ±0.01	0.995 ±0.003
	AI-only	0.96 ±0.02	0.99 ±0.01	0.975 ±0.01	0.992 ±0.004
	DAI + AI	0.98 ±0.01	0.995 ±0.005	0.987 ±0.006	0.998 ±0.002
B (low-rate)	DAI-only	0.92 ±0.03	0.18 ±0.07	0.30 ±0.08	0.60 ±0.04
	AI-only	0.88 ±0.04	0.93 ±0.03	0.905 ±0.03	0.95 ±0.02
	DAI + AI	0.90 ±0.03	0.95 ±0.02	0.925 ±0.02	0.965 ±0.015
C (insider)	DAI-only	0.95 ±0.02	0.30 ±0.05	0.45 ±0.06	0.62 ±0.05
	AI-only	0.85 ±0.04	0.91 ±0.03	0.88 ±0.03	0.93 ±0.03
	DAI + AI	0.90 ±0.03	0.93 ±0.02	0.915 ±0.02	0.95 ±0.02

As shown in Table 1, under Scenario A, all three configurations achieved high classification performance. The DAI-only setup reached Precision of 0.99 ± 0.01 ,

Recall of 0.98 ± 0.02 , and F1-score of 0.985 ± 0.01 , confirming the strong effectiveness of DAI against conventional high-rate ARP spoofing. The AI-only configuration achieved comparable results, while the hybrid DAI+AI approach provided the highest overall values, including Recall of 0.995 ± 0.005 and ROC-AUC of 0.998 ± 0.002 . These results indicate that, in clearly expressed attack conditions, DAI already performs very well, and AI contributes only a limited additional gain.

A different pattern can be observed for Scenario B. As indicated in Table 1, DAI-only preserved relatively high Precision (0.92 ± 0.03), but its Recall dropped to 0.18 ± 0.07 and F1-score to 0.30 ± 0.08 . In contrast, AI-only achieved Recall of 0.93 ± 0.03 and F1-score of 0.905 ± 0.03 , while DAI+AI further improved Recall to 0.95 ± 0.02 and F1-score to 0.925 ± 0.02 . This confirms that low-rate spoofing can bypass conventional rule-based inspection more easily, whereas behavioral analysis remains sensitive to temporally distributed anomalies.

A similar trend was observed in Scenario C. According to Table 1, DAI-only achieved Recall of only 0.30 ± 0.05 and F1-score of 0.45 ± 0.06 , whereas AI-only reached Recall of 0.91 ± 0.03 and F1-score of 0.88 ± 0.03 . The hybrid DAI+AI configuration again provided the best overall balance, with Recall of 0.93 ± 0.02 , F1-score of 0.915 ± 0.02 , and ROC-AUC of 0.95 ± 0.02 . These results suggest that the combination of infrastructure-based validation and AI-driven behavioral analysis is particularly beneficial for insider-like or context-dependent attack patterns.

B. Detection Latency and Operational Indicators

In addition to classification metrics, the operational performance of the evaluated configurations was assessed using median time-to-detect, false positives per 1000 hosts per day, inference latency, and average CPU utilization of the AI host. These results are summarized in Table 2.

The false positive values are reported as normalized estimates per 1000 hosts/day to support comparison with operational network monitoring practice.

As shown in Table 2, for Scenario A, DAI-only achieved a median time-to-detect of 8 ± 3 s and only 0.5 ± 0.2 false positives per 1000 hosts/day. The AI-only configuration required 12 ± 4 s and produced 6 ± 1.5 false positives, while DAI+AI maintained a fast response time of 9 ± 3 s and reduced false positives to 2 ± 0.6 . This indicates that, even when AI is introduced, the hybrid architecture preserves near-real-time responsiveness while improving alert selectivity.

For Scenario B, DAI-only did not provide timely detection, as the median time-to-detect exceeded 300 s and some attacks remained unnoticed. In contrast, AI-only reduced detection latency to 180 ± 40 s, while DAI+AI further reduced it to 160 ± 35 s. At the same time, the number of false positives decreased from 7 ± 2 in AI-only to 3.5 ± 1 in DAI+AI, demonstrating that the hybrid approach improves both responsiveness and operational usability.

Under Scenario C, DAI-only required 250 ± 60 s for detection, whereas AI-only and DAI+AI reduced this value to 45 ± 15 s and 40 ± 12 s, respectively. As also indicated in Table 2, the hybrid configuration generated fewer false positives than AI-only. The inference latency of the AI module remained in the 120–160 ms range, and average CPU utilization varied between 14% and 19%, suggesting that the proposed system remains lightweight enough for practical deployment in a monitored LAN environment.

TABLE 2 DETECTION LATENCY AND OPERATIONAL METRICS

Scenario	Configuration	Median Time-to-Detect (s)	False Positives per 1000 Hosts/Day	Inference Latency (ms)	CPU Utilization (%)
A (classic)	DAI-only	8 ± 3	0.5 ± 0.2	—	—
	AI-only	12 ± 4	6 ± 1.5	120 ± 20	$14\% \pm 3$
	DAI + AI	9 ± 3	2 ± 0.6	130 ± 25	$16\% \pm 3$
B (low-rate)	DAI-only	>300 (many undetected)	0.2 ± 0.1	—	—
	AI-only	180 ± 40	7 ± 2	150 ± 30	$18\% \pm 4$
	DAI + AI	160 ± 35	3.5 ± 1	160 ± 30	$19\% \pm 4$
C (insider)	DAI-only	250 ± 60	0.3 ± 0.1	—	—
	AI-only	45 ± 15	8 ± 2	140 ± 25	$17\% \pm 3$
	DAI + AI	40 ± 12	3 ± 1	150 ± 28	$18\% \pm 3$

C. Statistical Significance of the Improvement

To verify whether the observed improvement was statistically meaningful, a paired t-test was applied to the ten replications comparing DAI-only and DAI+AI. For Scenarios B and C, the hybrid configuration showed statistically significant improvement with $p < 0.01$. For Scenario A, the difference was limited and did not consistently reach statistical significance ($p > 0.05$). This result is expected, since DAI already performs very strongly under classic high-rate attack conditions. Therefore, the main added value of the proposed hybrid architecture lies not in replacing DAI where it is already sufficient, but in strengthening detection under stealthy and insider-like conditions.

D. Forensic Observation and DAI Behavior

A representative forensic example was obtained during one of the replications of Scenario A, where a Telnet session was successfully intercepted and credentials were exposed in clear text. As illustrated in Fig. 4, the ARP spoofing attack enabled the attacker-controlled host to record traffic and recover login credentials from the captured pcap. This observation demonstrates the practical consequences of MitM attacks in the presence of unencrypted application-layer protocols.

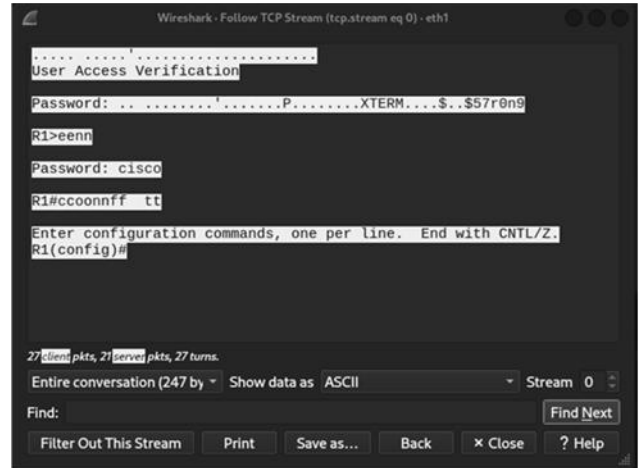


Fig. 4. Captured Telnet credentials during MitM attack.

At the same time, DAI produced %SW_DAI-4-ACL_DENY log entries indicating blocked attempts to inject invalid ARP packets. As shown in Fig. 5, these logs confirm that DAI effectively mitigates classic spoofing attempts at switch level. In parallel, the AI module detected supporting anomalies, including increased values of gratuitous_count and ip_mac_change_count, enabling the generation of a SIEM event and the automatic preservation of a forensic pcap snapshot. This combined behavior illustrates the complementary nature of the proposed architecture: DAI provides immediate infrastructure-level enforcement, while AI adds contextual awareness and analytical depth.

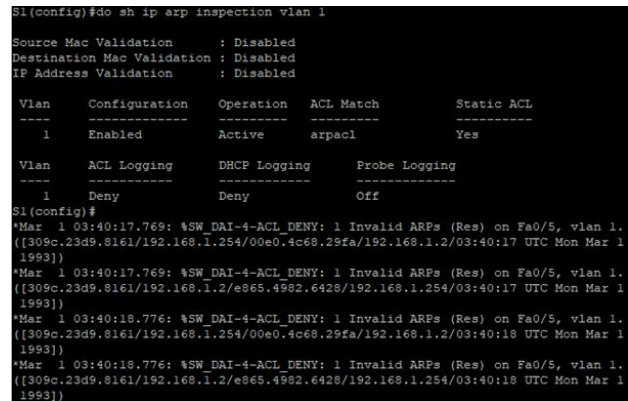


Fig. 5. DAI logs for invalid ARP packet blocking.

E. Overall Interpretation

Taken together, the results from Table 1, Table 2, Fig. 4, and Fig. 5 show that DAI-only is highly effective against classic high-frequency ARP spoofing attacks but offers limited sensitivity in low-rate and insider scenarios. The AI-only configuration substantially improves Recall and reduces detection latency in these harder cases, but at the cost of a higher false positive rate. The proposed DAI+AI architecture provides the best overall balance between detection capability, alert precision, and operational practicality.

These findings support the central assumption of the study: traditional rule-based protections remain valuable as a first line of defense, but they can be significantly strengthened through AI-based behavioral analysis. Rather than replacing DAI, the proposed system extends its effectiveness by detecting attack patterns that are formally plausible from a protocol perspective but anomalous when viewed in behavioral and temporal context.

Compared with purely rule-based ARP protection mechanisms, the proposed hybrid configuration preserves the low false positive rate and immediate enforcement capability of infrastructure-level defenses, while improving sensitivity in low-rate and insider-like attack scenarios. Compared with standalone AI-based intrusion detection, the DAI+AI configuration reduces the false positive rate by correlating behavioral anomalies with protocol-level and infrastructure-level evidence. Therefore, the proposed approach should be viewed as a complementary security layer rather than a replacement for existing ARP inspection, IDS/IPS, or anomaly detection mechanisms.

It should also be emphasized that the results represent a controlled laboratory validation of the proposed architecture. Although the attack scenarios reproduce realistic ARP-based MitM behavior, the environment is smaller and more deterministic than a production enterprise network. In larger deployments, additional factors such as multiple VLANs, higher traffic diversity, dynamic host mobility, and legitimate changes in IP-to-MAC or MAC-to-port associations may influence both detection accuracy and false positive rates. For this reason, future validation in larger and more heterogeneous LAN environments is necessary.

IV. CONCLUSIONS

This study addressed the problem of Man-in-the-Middle (MitM) protection in local area networks, with particular emphasis on ARP spoofing as one of the most common practical realizations of such attacks in IPv4 environments. A hybrid defense architecture combining Dynamic ARP Inspection (DAI) and an AI-based behavioral detection module was designed and evaluated in a controlled laboratory environment. The results demonstrate that this combined approach improves the detection of attack patterns that are difficult to identify using infrastructure-level rule enforcement alone.

The experimental findings confirm that DAI, when used as a standalone mechanism, is highly effective against classic high-rate ARP spoofing attacks, where invalid ARP behavior is explicit and can be blocked directly at switch level. However, its performance decreases substantially in low-rate stealth and insider-like scenarios, where malicious activity is more subtle, temporally distributed, or generated by a host with seemingly legitimate network presence. In these cases, the AI-based module provides a clear advantage by analyzing behavioral, statistical, and temporal relationships across ARP, DHCP, and switch-port events.

The comparison of DAI-only, AI-only, and DAI+AI shows that the hybrid configuration achieves the best overall balance between Recall, Precision, F1-score, time-

to-detect, and false positive rate. In addition to improving detection performance, the proposed approach preserves operational practicality by maintaining low inference latency and moderate computational overhead. The forensic observations further illustrate that timely detection of MitM activity is not only a matter of metric improvement, but also of preventing real exposure of sensitive information in poorly protected communication sessions.

The main contribution of this work lies in the design and experimental validation of a hybrid LAN protection model that integrates infrastructure-based ARP validation with AI-driven behavioral analysis in a single detection and response workflow. While the general combination of AI and rule-based security mechanisms is already known, the contribution of this study is the targeted application of this concept to ARP spoofing-based MitM detection and the experimental comparison of DAI-only, AI-only, and DAI+AI configurations under controlled LAN attack scenarios. The results indicate that artificial intelligence is most effective when used not as a replacement for conventional network security mechanisms, but as an adaptive extension that enhances contextual awareness and resilience against stealthy attack behavior.

At the same time, the study has several limitations. The experiments were conducted in a controlled laboratory environment and therefore do not fully capture the complexity of large-scale production networks with heterogeneous services, dense background traffic, and multiple simultaneous anomalies. In addition, the evaluation focused specifically on ARP-based MitM attacks, which means that the reported results should not be generalized directly to other forms of MitM, such as DNS spoofing or TLS-related interception attacks.

Future work should therefore focus on validating the proposed approach in larger and more diverse network environments, extending the feature set with additional telemetry sources, and incorporating explainable AI mechanisms to improve the transparency of detection decisions. Further research may also explore adaptive response strategies in which the system recommends mitigation actions according to risk level and a predefined human-in-the-loop policy, while avoiding fully automated blocking decisions in ambiguous cases.

REFERENCES

- [1] C. Surianarayanan, "Integration of the Internet of Things and Cloud: Security Challenges and Solutions – A Review," *International Journal of Cloud Applications and Computing*, vol. 13, pp. 1–30, Mar. 2023, doi: [10.4018/IJCAAC.325624](https://doi.org/10.4018/IJCAAC.325624).
- [2] H. Fereidouni, O. Fadeitcheva, and M. Zalai, "IoT and Man-in-the-Middle Attacks," *SECURITY AND PRIVACY*, vol. 8, no. 2, p. e70016, Mar. 2025, doi: [10.1002/spy2.70016](https://doi.org/10.1002/spy2.70016).
- [3] D. C. Plummer, "An Ethernet Address Resolution Protocol: Or Converting Network Protocol Addresses to 48.bit Ethernet Address for Transmission on Ethernet Hardware," Nov. 1982, doi: [10.17487/rfc0826](https://doi.org/10.17487/rfc0826).
- [4] S. Y. Nam, S. Jurayev, S. S. Kim, K. Choi, and G. S. Choi, "Mitigating ARP poisoning-based man-in-the-middle attacks in wired or wireless LAN," *EURASIP Journal on Wireless Communications and Networking* 2012 2012:1, vol. 2012, no. 1, pp. 89–, Mar. 2012, doi: [10.1186/1687-1499-2012-89](https://doi.org/10.1186/1687-1499-2012-89).

- [5] D. Hanna, P. Veeraraghavan, and E. Pardede, "PrECast: An Efficient Crypto-Free Solution for Broadcast-Based Attacks in IPv4 Networks," *Electronics* 2018, Vol. 7, Page 65, vol. 7, no. 5, p. 65, May 2018, doi: [10.3390/electronics7050065](https://doi.org/10.3390/electronics7050065).
- [6] Z. Shah and S. Cosgrove, "Mitigating ARP Cache Poisoning Attack in Software-Defined Networking (SDN): A Survey," *Electronics* 2019, Vol. 8, Page 1095, vol. 8, no. 10, p. 1095, Sep. 2019, doi: [10.3390/electronics8101095](https://doi.org/10.3390/electronics8101095).
- [7] J. Asharf, N. Moustafa, H. Khurshid, E. Debie, W. Haider, and A. Wahab, "A Review of Intrusion Detection Systems Using Machine and Deep Learning in Internet of Things: Challenges, Solutions and Future Directions," *Electronics* 2020, Vol. 9, Page 1177, vol. 9, no. 7, p. 1177, Jul. 2020, doi: [10.3390/electronics9071177](https://doi.org/10.3390/electronics9071177).
- [8] T. Sowmya and E. A. Mary Anita, "A comprehensive review of AI based intrusion detection system," *Measurement: Sensors*, vol. 28, p. 100827, Aug. 2023, doi: [10.1016/j.measen.2023.100827](https://doi.org/10.1016/j.measen.2023.100827).
- [9] C. Yin, Y. Zhu, J. Fei, and X. He, "A Deep Learning Approach for Intrusion Detection Using Recurrent Neural Networks," *IEEE Access*, vol. 5, pp. 21954–21961, Oct. 2017, doi: [10.1109/ACCESS.2017.2762418](https://doi.org/10.1109/ACCESS.2017.2762418).
- [10] P. Schummer, A. del Rio, J. Serrano, D. Jimenez, G. Sánchez, and Á. Llorente, "Machine Learning-Based Network Anomaly Detection: Design, Implementation, and Evaluation," *AI* 2024, Vol. 5, Pages 2967-2983, vol. 5, no. 4, pp. 2967–2983, Dec. 2024, doi: [10.3390/ai5040143](https://doi.org/10.3390/ai5040143).
- [11] R. Chinnasamy, M. Subramanian, S. V. Easwaramoorthy, and J. Cho, "Deep learning-driven methods for network-based intrusion detection systems: A systematic review," *ICT Express*, vol. 11, no. 1, pp. 181–215, Feb. 2025, doi: [10.1016/j.ict.2025.01.005](https://doi.org/10.1016/j.ict.2025.01.005).
- [12] M. M. Alani, A. I. Awad, and E. Barka, "ARP-PROBE: An ARP spoofing detector for Internet of Things networks using explainable deep learning," *Internet of Things*, vol. 23, p. 100861, Oct. 2023, doi: [10.1016/j.iot.2023.100861](https://doi.org/10.1016/j.iot.2023.100861).
- [13] K. Sauka, G. Y. Shin, D. W. Kim, and M. M. Han, "Adversarial Robust and Explainable Network Intrusion Detection Systems Based on Deep Learning," *Applied Sciences* 2022, Vol. 12, Page 6451, vol. 12, no. 13, p. 6451, Jun. 2022, doi: [10.3390/app12136451](https://doi.org/10.3390/app12136451).
- [14] M. Keshk, N. Koroniotis, N. Pham, N. Moustafa, B. Turnbull, and A. Y. Zomaya, "An explainable deep learning-enabled intrusion detection framework in IoT networks," *Inf. Sci. (N. Y.)*, vol. 639, no. 10, p. 119000, Aug. 2023, doi: [10.1016/j.ins.2023.119000](https://doi.org/10.1016/j.ins.2023.119000).
- [15] M. M. Issa, M. Aljanabi, and H. M. Muhialdeen, "Systematic literature review on intrusion detection systems: Research trends, algorithms, methods, datasets, and limitations," *Journal of Intelligent Systems*, vol. 33, no. 1, Jan. 2024, doi: [10.1515/jisys-2023-0248](https://doi.org/10.1515/jisys-2023-0248).