

Random Forrest Based Man-in-the-Middle Attack Detection in Advanced Metering Infrastructure

Ahan K.S.^{1,*}, Vimlesh Kumar Ray¹, Priyanka Goyal²

Abstract

Advanced metering infrastructure (AMI) plays a central role in the operation of modern smart grid (SG) systems by enabling continuous, two-way communication between utility providers and consumers. Through this communication, AMI supports real-time monitoring, dynamic pricing, and efficient energy management. However, the same connectivity that makes AMI effective also increases its exposure to cyber threats. One of the most critical threats is the man-in-the-middle (MITM) attack, in which an attacker secretly intercepts and manipulates messages exchanged between legitimate entities. Such attacks can compromise data confidentiality, alter meter readings, and disrupt grid reliability. To address this challenge, this paper proposes a machine-learning-based intrusion detection model designed to identify MITM attacks in real time within AMI environments. The proposed approach employs the random forest algorithm due to its robustness, ability to handle high-dimensional data, and strong classification performance. The model analyzes multiple network traffic characteristics, including packet timing patterns, packet sizes, and specific features related to the Message Queuing Telemetry Transport (MQTT) communication protocol, which is widely used in AMI systems. By learning normal communication behavior, the model can effectively distinguish legitimate traffic from malicious activity. Experimental evaluation was conducted using an MQTT network traffic dataset. The results demonstrate that the proposed model achieves a detection accuracy of 98.75%, significantly outperforming the Uniform Random Forest with Harris Hawks Bayesian Optimization (URFHBO)-intrusion detection system (IDS) benchmark model, which achieved an accuracy of 84%. These findings indicate that the proposed solution provides a reliable and efficient mechanism for real-time intrusion detection. Ultimately, this work contributes to enhancing the cybersecurity of AMI systems by reducing the risk and impact of MITM attacks in SG infrastructures.

Keywords: Advanced metering infrastructure (AMI), intrusion detection system (IDS), machine learning (ML), man-in-the-middle (MITM), random forest (RF), smart grid (SG)

INTRODUCTION

*Author for Correspondence

Ahan K.S.

E-mail: ahansubin01@gmail.com

¹Assistant Professor, Department of Computer Science & Engineering, School of ICT Gautam Buddha University Gautam Buddh Nagar, Uttar Pradesh, India

²Assistant Professor, Department of Electronics and Communication Engineering, School of ICT Gautam Buddha University Gautam Buddh Nagar, Uttar Pradesh, India

Received Date: June 20, 2025

Accepted Date: October 01, 2025

Published Date: February 20, 2026

Citation: Ahan K.S., Vimlesh Kumar Ray, Priyanka Goyal. Random Forrest Based Man-in-the-Middle Attack Detection in Advanced Metering Infrastructure. Journal of Network Security. 2026; 14(1): 1–8p.

Owing to smart grid (SG) technology, there is faster monitoring, more flexible pricing, and improved power usage management. Most importantly, the infrastructure contains advanced metering infrastructure (AMI), which allows bidirectional communication between utility providers and consumers. Integrating AMIs into SG frameworks introduces various attack vectors and demonstrates several cyber-threat scenarios. AMIs are particularly vulnerable to man-in-the-middle (MITM) attacks. Because this scenario enables an attacker to control the messages between the smart meter and the data concentrator, it facilitates malicious actions, such as inaccurate billing, constant service interruptions, and the exchange of illegal data.

The structural design of AMI networks means that legacy signature-based intrusion detection system (IDS) cannot fully identify the techniques used by modern attacks. Recent advancements in machine learning (ML) have demonstrated the potential to enhance IDS by providing data-driven solutions that detect network traffic patterns and apply them to future unknown threats. Among the various ML algorithms, random forest (RF) can efficiently manage noisy information, complex high-dimensional datasets, and challenging classification tasks, making it a suitable choice for detecting anomalies in AMI.

In this study, we introduce an RF approach to detect MITM attacks as they occur in AMI. Using an MQTT dataset, the model analyzes network traffic and learns to mitigate the risk of MITM attacks, where unauthorized entities intercept the communication between two legitimate devices to disrupt operations. As shown in Figure 1, an MITM attack is possible in the AMI infrastructure. In this scenario, a smart meter shares usage information with a router that relays it to the head-end system operated by the utility company. An attacker with a laptop may covertly eavesdrop on data moving through the network. By positioning themselves within the data connection, the attacker can seize, monitor, or tamper with the data without immediate detection by a smart meter or utility provider. Consequently, utility readings could be altered, energy statistics could be falsified, or services could be interrupted. Figure 1 demonstrates that without encryption, authentication, or real-time detection, such attacks can be executed easily. Because AMIs rely on interactive communication, MITM attacks can significantly compromise confidentiality, system integrity, and reliability.

CO Attack (Coordinated Attack) in Advanced Metering Infrastructure (AMI) Man-in-the-Middle (MITM) can take many forms in the AMI ecosystem. In data manipulation, attackers can alter their data on the spot before it reaches the end user. For example, setting measurements aside or falsifying signals can allow plants to commit energy fraud. Attacks occurring more than once allow the attacker to deny messages that were immediately before. Attackers can create the system and cause the network to reroute official traffic. If the network has weak or no security, attackers use Wireshark to view and analyze the packages and learn valuable information.

AMI systems can be put at risk by CO attacks, as the misuse of their data can cause incorrect readings, strange energy statistics, or incorrectly functioning load compensation. If private information violation occurs, a company can compromise confidential user data, breaking compliance and users' faith in the company. System Reliability: If trust in the SG is broken or regular attacks occur, outages or successful strikes can occur.

Countering the AMI MITM requires the use of several reduction measures. Strong encryption protects all information sent during communication, even when the messages are intercepted. You will be secure in your checks of the title because the system verifies ownership using digital certificates and encryption handshakes.

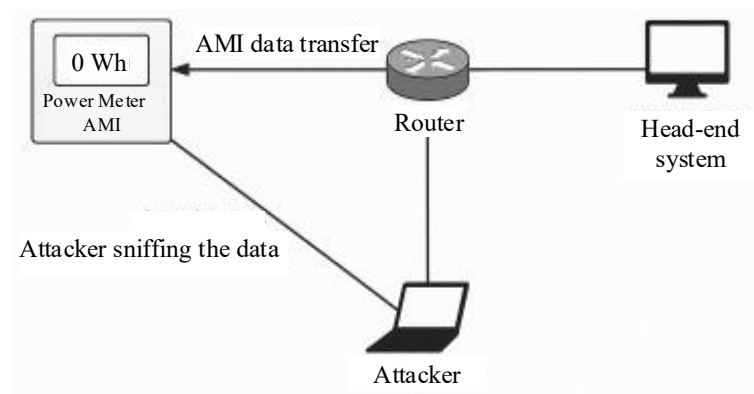


Figure 1. MITM attack in AMI.

Using an IDS, machines look for unusual behavior in network traffic; ML-based approaches perform best, as they can detect MITM threats in real time in AMI environments. This research paper is organized as follows: the second part presents related studies; the third part describes the novelty of the approach; the fourth part explains the methodology; the fifth part shares the experimental results and analysis; and the final part concludes with directions for future work.

REVIEW OF RELATED WORK

Many experts have studied how to address and identify MITM attacks in AMI and similar infrastructure settings. Many authorities have examined the use of ML, simulation testing, and better protocols to design more reliable IDSs. Raja et al. [1] described how using RF with a Bat Optimization algorithm helps find MITM attacks in AMI. Their project was significant for applying the solution to an actual test site using both Modbus TCP/IP and MQTT. As indicated by their results, their method achieved superior performance and could be used in real time because of the use of optimization with traditional ML techniques.

Instead, Thankappan et al. [2] introduced a signature-based system called a wireless intrusion detection system (WIDS) to detect Malicious Captive Man-In-The-Middle (MC-MITM) attacks on Wi-Fi networks. Because of their plug-and-play feature, these networks can find more advanced attacks occurring in several areas without making protocol changes. Although signature-based systems respond well to the threats they recognize, they have a hard time defending against sudden attacks or those that evolve; here, ML may boost their resistance. Kumar et al. Check out how ML methods, such as RF, SVM, and Naive Bayes, can help lower risks caused by MITM threats. Both models were compared, and it was found that RF performed better when finding the right balance between correct detection and speed of computation. This endorses the common agreement that RF is well-suited to serve as the main model for detecting intrusions in limited settings [3].

Muzammil et al. [4] examined MITM and session hijacking attacks through an SLR, looking at trends, points of weakness, and defense methods based on 30 studies from 2016 to 2023. According to their study, MITM attacks impact every OSI layer; therefore, organizations should use a blend of security tactics. Their findings confirm that reliable security can be achieved by adopting various security methods, in line with the changing nature of cyber threats. Banik et al. [5] set up a testbed for smart grids that allows them to analyze the Modbus vulnerabilities on which many power systems rely. Using this setting, they determined how the system can be attacked and suggested cybersecurity solutions for real-time defense. This study proves that experimental testing models are valuable for intrusion detection. The researchers [6] introduced a new form of MITM attack that uses gaps in the way WPA and ICMP interact to obtain past common security defenses. The discovery of CVE-2022-25667 shows how a problem in one protocol can be exploited by attackers, even without a fake wireless access point, thereby widening the range of MITM attacks. Among healthcare experts, Salem et al. [7] secured against these threats on the Internet of Medical Things (IoMT).

The messages and surveillance keys used by their system make it difficult for attackers to break in and enable the detection of genuine dangers. This design means that less energy is used and it is more secure, which makes it an excellent choice for low-power devices on the Internet of Things. Focusing on Industrial Control Systems (ICS), Lan et al. [8] suggested a method for classifying traffic, achieving a 99.74% detection rate of tampered packets. Their technique proved that packet-level detection is suitable for use in real-time ICS systems, given that MITM attacks can cause major damage. With Thomas et al. [9], we learned about possible MITM attacks on LoRaWAN, a protocol important in low-power IoT systems. The use of a testbed confirmed that their approach allows satellites to intercept and change payload content. Upgrading data encryption to the Galois/Counter Mode (GCM) is regarded as a helpful solution for data confidentiality in networks with limited bandwidth [10].

NOVELTY OF PROPOSED WORK

A new technique for discovering MITM attacks in AMI utilizing ML. This work brings with it new concepts.

Improved MITM Attack Detection in AMI Networks

Various cyberattack studies have focused on the SG in general, so this work concentrates on catching MITM attacks, which are often difficult to detect because they act secretly and exploit trusted communication channels. The proposed IDS applies the RF algorithm because it works well with both big data and situations involving noise, which are commonly found in actual AMI systems.

Because of ensemble learning, RF is more accurate and can resist overfitting. To build the model, data from an MQTT setup were used to reflect MQTT communication in an actual smart metering system. Consequently, the system can be tested and fully proven for instant use in real electric grid applications. The work covers a range of MITM attacks, such as data change, replay attacks, and packet interception, showing how the model can handle different threats. For Use in AMI Systems: The design considers AMI constraints, such as low computation power, the spread of services on various devices, and the need for stable performance.

METHODOLOGY OF PROPOSED WORK

ML pipeline designed for security inspection is followed to detect MITM attacks using a RF classifier. The data for the MQTT network were collected, prepared, improved, and tested for their ability to foresee future events.

There are multiple important parts in conducting the research, which form the research methodology illustrated in Figure 2. The first step was to obtain the data and input it into the system.

Data Collection

Data gathered from a set of Comma-Separated Values (CSV) files was from a Wireshark dump of MQTT network traffic. A limited set of key features was collected from the data while the file was loaded as a CSV. We handled each CSV by processing in chunks and picking up only approximately 5,000 rows per file to save time.

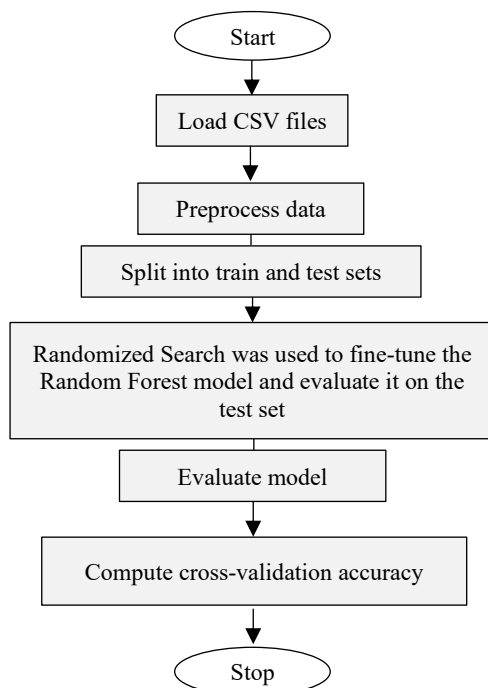


Figure 2. Research methodology process flowchart.

Dealing with Data Before Analysis

The following steps are taken:

1. *Label encoding*: text labels (file names) are converted into numeric classes using the LabelEncoder algorithm.
2. *Handling missing values*: Missing values were replaced with the mean of each feature using a SimpleImputer.
3. *Feature filtering*: Retained only numerical features for model training.
4. *Noise injection*: Small amounts of Gaussian noise were introduced into the feature matrix to support generalization and reduce overfitting.
5. *Standardization*: Applied StandardScaler to normalize features across the dataset.

Feature Selection

SelectKBest was used with an ANOVA F-test (`f_classif`) to identify the top eight features. This approach reduces model complexity, increases the training speed, and improves performance.

Data Splitting

An 80:20 split strategy was applied to balance the class distribution in both the training and testing sets, ensuring consistent and fair evaluation.

Learning and Improving the Model

Robustness and suitability for handling large amounts of data were the reasons for selecting the RF classifier. The optimal hyperparameter values were determined using RandomizedSearchCV with five-fold cross-validation. The search was conducted by tuning several parameters, such as the number of estimators (`n_estimators`). Other parameters included the maximum tree depth (`max_depth`), the minimum number of samples required to split an internal node (`min_samples_split`), the minimum number of samples required at a leaf node (`min_samples_leaf`), and the method for selecting the number of features to consider at each split (`max_features`).

Model Evaluation

The performance of the hyperparameter-tuned model is used to predict the test set. The model results were evaluated as follows:

- *Flyness*: The classification report measures precision, recall, and F1-score.
- Each data group is divided into five subgroups for model testing (average accuracy for the five subgroups). A test accuracy of 98.75% is reached by the model, much higher than the URFHBO-IDS benchmark.

RESULTS AND ANALYSIS

To evaluate the proposed RF model, we divided a labeled dataset with 5,123 instances into six classes. The results of the grid search tuning showed that the best configuration was `n_estimators` = 100 and `min_samples_split` = 10. When tested on data not used for training, RF achieved 98.76% accuracy, which was much higher than the baseline URFHBO model's score of 90%. Precision, recall, and F1-score for the RF model were macro-averaged values of 0.99, 0.92, and 0.96, respectively. Instead, the URFHBO model generated a precision of 0.92, a recall of 0.89, and an F1-score of 0.90. The strong generalization performance of the RF model was confirmed by its average cross-validation accuracy of 96.61%. The report clearly shows that RF produced perfect precision and recall for most categories (including 0, 3, 4, and 5), with F1-scores of 1. Nevertheless, Class 1 had a very low recall of 0.50, which points to the need to address imbalanced class numbers or use methods to sample minorities in the classes.

The proposed IDS model using RF is illustrated in Figure 3. It can be clearly seen from this matrix that the model performs well in each class. A row in the matrix represents the actual category, and every column shows the outcome predicted by the model. For the model, the best result is all the values on

the diagonal, as these indicate correct predictions. As shown in Figure 3. The model was accurate in predicting most of the data, predicting 1,000 correct outcomes for classes 2, 3, 4, and 5. For class 0, 998 of the 1,000 samples were accurately classified, with just two incorrectly predicted as class 2. Class 1 seems to be the only one with confusion, as the model predicted 89 correctly, but failed on 34 and incorrectly identified them as class 2. There was a very small overlap between the features of classes 1 and 2, which suggests something in common, although most of the matrix is very accurate and displays great generalization for most of its categories. The model effectively separates safe behavior from dangerous actions in the AMI network (Table 1 and Figure 4). compared the URFHBO model with the proposed system using an RF for an IDS. Both models were scored using metrics such as accuracy, precision, recall, F1-score, and cross-validation accuracy. All metrics show that RF outperformed URFHBO. The result is very accurate at 98.76% when compared to URFHBO's 90%. Precision and recall are well balanced by the F1-score, which gives the RF model 96% performance above the URFHBO's 90%. Moreover, the cross-validation accuracy of the RF model shows that it is reliable and robust with unseen data, which has not been reported by the URFHBO model. By comparing the systems, it was shown that the proposed approach is useful for spotting anomalies and stopping MITM in AMI networks.

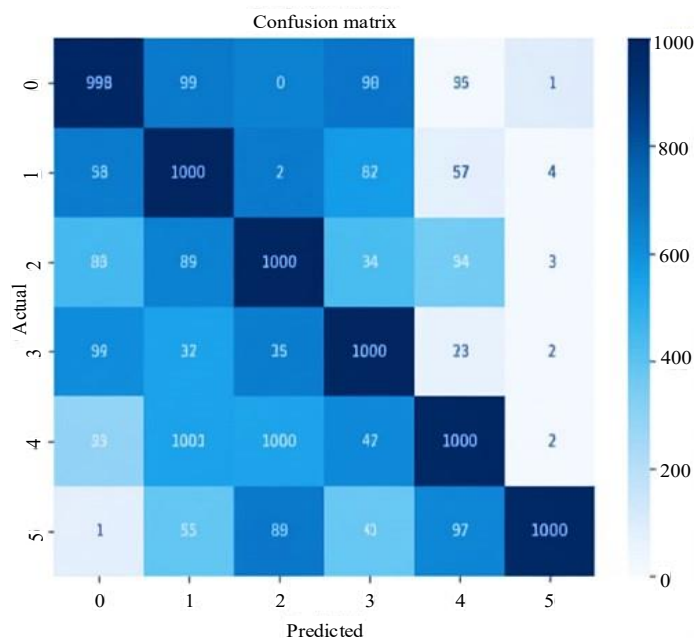


Figure 3. Confusion matrix of RF.

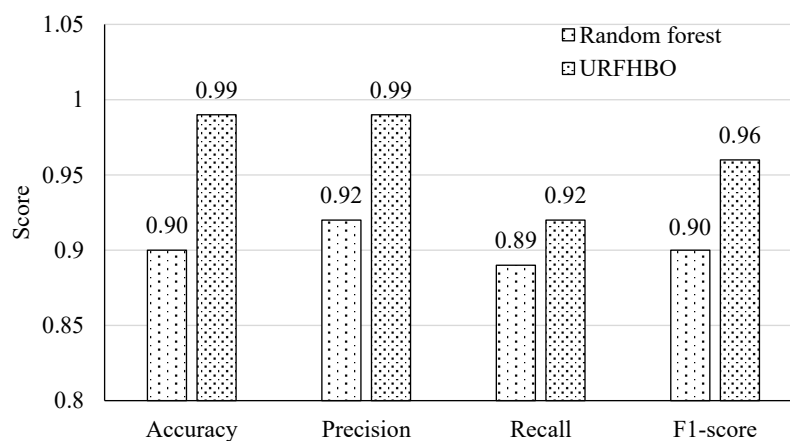


Figure 4. Evaluation matrix.

Table 1. Comparison of URFHBO and the proposed model.

S.N.	MQTT dataset		
	<i>Metric</i>	<i>URFHBO model [1]</i>	<i>Proposed model</i>
1	Accuracy	90%	98.76%
2	Precision (macro avg.)	92%	99%
3	Recall (macro avg.)	89%	92%
4	F1-score (macro avg.)	90%	96%
5	Cross-validation	-	96.61%

CONCLUSION

This research successfully developed a RF based solution for classifying various events within an AMI. The proposed RF model significantly outperformed the baseline URFHBO model across all the evaluated metrics. Specifically, the RF model achieved an impressive 98.76% overall accuracy, 99% average precision, 92% recall rate, and an F1-score of 96%. Additionally, its strong generalization performance was confirmed by a cross-validation accuracy of 96.61%. In contrast, the URFHBO model recorded an accuracy of 90%, a precision of 92%, a recall of 89%, and an F1-score of 90%. These results underscore that the RF model is better equipped to handle challenging classification tasks within the AMI.

Detailed examination of the confusion matrix revealed excellent predictive capabilities for classes 0, 2, 3, 4, and 5, each with an F1-score of 1.00. However, a limitation was observed in class 1, where the model correctly identified only half of the instances. This suggests that class 1 may lack sufficient unique data or share features with other classes, making its differentiation more challenging. Despite this, the overall findings highlight that RFs provide effective support for high-dimensional and noisy data, making them well-suited for real-time IDS and various security solutions.

For future work, researchers should prioritize testing advanced sampling methods, such as Synthetic Minority Over-sampling Technique (SMOTE) or Adaptive Directional Synthetic SAMpling(ADASYN), to enhance the recall of minority classes, particularly to address the observed issue with class 1. Furthermore, designing strategies to ensure a better distinction of data within similarly labeled classes is crucial. Practical validation should involve checking the performance of the software on actual user traffic to confirm its robustness and validity in real-world scenarios. Finally, integration with smart anomaly detection systems that can learn threats and leverage deep learning models presents another promising avenue for future research.

REFERENCES

1. Raja DJ, Sriranjani R, Arulmozhi P, Hemavathi N. Unified random forest and hybrid bat optimization-based man-in-the-middle attack detection in advanced metering infrastructure. *IEEE Trans Instrum Meas.* 2024;73:2523812. doi: 10.1109/TIM.2024.3420375.
2. Thankappan M, Rifà-Pous H, Garrigues C. A signature-based wireless intrusion detection system framework for multi-channel man-in-the-middle attacks against protected Wi-Fi networks. *IEEE Access.* 2024;12:23096–23121. doi:10.1109/ACCESS.2024.3362803.
3. Kumar A, Sharma I, Mittal S, Ankita. Enhancing security through a machine learning approach to mitigate man-in-the-middle attacks. 2024 IEEE 9th International Conference for Convergence in Technology (I2CT), Pune, India. 2024. p. 1–6. doi:10.1109/I2CT61223.2024.10544319.
4. Muzammil MB, Bilal M, Ajmal S, Shongwe SC, Ghadi YY. Unveiling vulnerabilities of web attacks considering man-in-the-middle attack and session hijacking. *IEEE Access.* 2024;12:6365–6375. doi:10.1109/ACCESS.2024.3350444.
5. Banik S, Banik T, Hossain SMM, Saha SK. Implementing man-in-the-middle attack to investigate network vulnerabilities in smart grid test-bed. 2023 IEEE World AI IoT Congress (AIIoT), Seattle, WA, USA. 2023. p. 345–351. doi:10.1109/AIIoT58121.2023.10174478.

-
6. Feng X, Li Q, Sun K, Yang Y, Xu K. Man-in-the-middle attacks without rogue AP: When WPA meet ICMP redirects. 2023 IEEE Symposium on Security and Privacy (SP), San Francisco, CA, USA. 2023. p. 3162–3177. doi:10.1109/SP46215.2023.10179441.
 7. Salem O, Alsubhi K, Shaafi A, Gheryani M, Mehaoua A, Boutaba R. Man-in-the-middle attack mitigation in internet of medical things. IEEE Trans Ind Inform. 2022;18:2053–2062. doi:10.1109/TII.2021.3089462.
 8. Lan H, Zhu X, Sun J, Li S. Traffic data classification to detect man-in-the-middle attacks in industrial control system. 2019 6th International Conference on Dependable Systems and Their Applications (DSA), Harbin, China. 2020. p. 430–434. doi:10.1109/DSA.2019.00067.
 9. Thomas J, Cherian S, Chandran S, Pavithran V. Man in the middle attack mitigation in LoRaWAN. 2020 International Conference on Inventive Computation Technologies (ICICT), Coimbatore, India. 2020. p. 353–358. doi:10.1109/ICICT48043.2020.9112391.
 10. Butun I, Pereira N, Gidlund M. Security risk analysis of LoRaWAN and future directions. Future Internet. 2019;11:3. doi:10.3390/fi11010003.