

Adversarial Machine Learning Attacks in Internet of Things Systems

Rachida Kone

Department of Electrical Engineering
Morgan State University
Baltimore, Maryland 21251
Email: rakon1@morgan.edu

Otily Toutsop

Department Electrical Engineering
Morgan State University
Baltimore, Maryland 21251
Email: ottou1@morgan.edu

Ketchiozo Wandji Thierry

Department Electrical Engineering
Morgan State University
Baltimore, Maryland 21251
Email: ketchiozo.wandji@morgan.edu

Kevin Kornegay

Department of Electrical Engineering
Morgan State University
Baltimore, Maryland 21251
Email: kevin.kornegay@morgan.edu

Joy Falaye

Department of Electrical Engineering
Morgan State University
Baltimore, Maryland 21251
Email: jofall@morgan.edu

Abstract—Researchers are looking into solutions to support the enormous demand for wireless communication, which has been exponentially increasing along with the growth of technology. The sixth generation (6G) Network emerged as the leading solution for satisfying the requirements placed on the telecommunications system. 6G technology mainly depends on various machine learning and artificial intelligence techniques. The performance of these machine learning algorithms is high. Still, their security has been neglected for some reason, which leaves the door open to various vulnerabilities that attackers can exploit to compromise systems. Therefore, it is essential to evaluate the security of machine learning algorithms to prevent them from being spoofed by malicious hackers. Prior research has shown that the decision tree is one of the most popular algorithms used by 80% of researchers for classification problems. In this work, we collect the dataset from a laboratory testbed of over 100 Internet of things (IoT) devices. The devices include smart cameras, smart light bulbs, Alexa, and others. We evaluate classifiers using the original dataset during the experiment and record a 98% accuracy. We then use the label-flipping attack approach to poison our dataset and record the output. As a result, flipping 10%, 20%, 30%, 40%, and 50% of the poison data generated accuracies of 86%, 74%, 64%, 54%, and 50%, respectively.

Keywords— Adversarial Machine Learning, Internet of Everything (IoE), Internet of Things (IoT), wireless communication, label-flipping, decision tree.

I. INTRODUCTION

The technological revolution resulted in an ecosystem where everything can be connected to everything. To fulfill the dire need for all-around connectivity, intelligence, and cognition, the concept of the Internet of Things (IoT) evolved into the Internet of Everything (IoE). In essence, IoE, defined as the intelligent connection of people, processes, data, and things aims to foster a common interrelated ecosystem that improves experiences and facilitates smarter decision-making. The concept of IoE relies on interdisciplinary technical innovations such as sensor and embedded technologies, low-power communication, and big data analytics [1]. The paradigm of the Internet of Everything is generally used in large applications such as smart manufacturing, smart agriculture, and intelligent transportation systems [2]. Emerging technologies such as artificial intelligence

(AI), 3D media, and Virtual and Augmented reality (VAR) are booming. Today, artificial intelligence systems can monitor their environment, analyze it, and take action. For instance, autonomous vehicles can drive themselves from a starting point to a predetermined destination without any human interaction.

The advancement of emerging technologies increased the volume of data traffic. According to Edholm's law, proven to be true since 1970, telecommunication bandwidth doubles every 18 months [3]. Hence, the introduction of 6G. Some of the key motivations behind the introduction of 6G communication systems are high bit rate, high reliability, low latency, high energy efficiency, high spectral efficiency, new spectra, green communication, intelligent networks, network availability, communications convergence, localization, computing, control, sensing, and enabling fully digital connectivity [4]. With the success of Machine Learning (ML) in different domains, there has been an increasing interest in the development of artificial intelligence-driven solutions for wireless communication. In the literature, Deep Neural Network (DNN) models have been used for spectrum sensing and prediction [5]. In addition, to DNN various ML models including Convolutional Neural Networks (CNN), Feed-forward Neural Networks (FNN) and Long Short-Term Memory (LSTM) models have been developed for device fingerprinting and identification [6], channel decoding [7] and modulation recognition [8]. Artificial Intelligence and machine learning play a crucial role in wireless 6G networks as they enable real-time analysis and automated zero-touch operation and control in 6G networks [9]. Numerous studies also focus on exploring the feasibility of detecting network attacks by utilizing machine learning techniques. In recent years, utilization of machine learning technologies in anomaly-based IDS has gained a lot of popularity [10], [11], [12], [13], [14].

For the past decades, researchers have focused on improving the accuracy of machine learning models while overlooking the security aspect. As machine learning is at the forefront of today's technological revolution, it is extremely important to investigate the security of these algorithms and develop countermeasures to prevent systems from being spoofed. Thus, the introduction of Adversarial Machine Learning (AML). Unlike traditional jamming attacks that emit electromagnetic waves to jam the channel [15], AML attacks are crafted to mislead the model and are more difficult to detect. Many studies have proven that AML can effectively disrupt wireless networks by causing the misclassification of signals [16]. AML attacks in wireless networks may lead to disruptions, performance losses, and eventually failures. Hence it is essential to ensure the

security and robustness of ML-based systems in the presence of adversaries. In this work, we evaluate the robustness of the decision tree (DT) algorithm by implementing a label-flipping attack on an intrusion detection system (IDS) dataset collected from the CAP Center testbed. The goal of our work is threefold. First, we present a brief overview of IOE. Then, we survey the current state of ML, AML, and IDS. Finally, we discuss our case study. Our experiment illustrates how a label-flipping attack influences the ML model metrics such as accuracy, precision, recall, and F1 score.

II. THE INTERNET OF EVERYTHING (6G, WIFI, BLUETOOTH, AND IOT)

The exponential technological development has fostered an environment where everything is connected to everything. This interconnection between humans, data, and devices has forced the Internet of Things (IoT) to evolve into the concept of the Internet of Everything (IoE) [17]. The Internet of Everything is the global network through which people, things, and intelligent devices are connected and can share information and services [18]. IoE aims to realize a hyperconnected society by collecting and exchanging bilateral information among millions of Internet-connected devices [19]. While similar to IoT, the IoE expands from examining only device-to-device communication and focuses on the human element. With more relevant connections than machine-to-machine communications, IoE has enabled the global democratization of skills, including person-to-machine and person-to-person connections [20]. The IoE paradigm can extract and analyze real-time data collected from diverse and heterogeneous IoE systems, ranging from simple sensors and actuators to complex robotic devices, and from autonomous service agents to human actors [21]. The growing communication need of IOE systems imposes multi-dimensional requirements on wireless communication, sensing, and security [22]. Wireless networks play a vital role in IOE since they enable data transfer between systems. Wireless Communication, also referred to as unguided media, is the transmission of information between two or more points without any physical connection. It uses the radio spectrum to transmit signals through the atmosphere. In wireless communication systems, electromagnetic energy is coupled to the propagation medium by an antenna which serves as the radiator [23]. The inception of wireless communication can be traced back to the 19th century when scientists started experimenting with electromagnetic waves. The most recent evolution in network capabilities is 6G. By contrast with the existing 5G networks, the sixth generation (6G) networks offer less latency and high frequency. As the scale at which data is broadcasted constantly grows, a network with an expanded capacity to handle this load is necessary. Another critical aspect of the IOE is the modes of communication. There are many ways devices can communicate across the network; these are called communication protocols. Communication protocols transmit packet data from one device to another on the network. These protocols include TCP, FTP, ZigBee, Z-Wave, Wi-Fi, Bluetooth, Thread, and more. Communication protocols can send data from one device to another over the local network and the Internet.

Intrusion detection systems are widely used to protect IOE networks from malicious activities. Intrusion occurs when a set of actions compromise the confidentiality, integrity, and availability of a system. Two (2) methodologies are generally used in intrusion detection: the signature-based method and the anomaly-based method. Signature-based detection relies on pattern (signature) comparison; thus, it can hardly detect new malware attacks. Anomaly-based uses machine learning to detect suspicious behaviors and therefore can be trained to detect unknown malware attacks and even zero-day attacks [24].

III. MACHINE LEARNING

The term Machine Learning was coined by Arthur Samuel in 1959, who defines it as the ability of computers to perform tasks

without being explicitly programmed to do so [25]. Machine learning is an Artificial Intelligence technique that uses algorithms to analyze data, identify hidden patterns, and predict future outcomes. The concept of machine learning is very similar to data mining and predictive modeling. Depending on the training technique, machine learning models are subdivided into four (4) main groups: supervised, unsupervised, semi-supervised, and federated learning. The supervised learning technique uses labeled data to train algorithms that predict outcomes accurately. Supervised ML is separated into classification and regression algorithms. Classification algorithms are used to predict values (good or bad) while regression algorithms are used to predict quantities such as home prices. The classification works by assigning data to specific categories. The classification algorithm recognizes features that impact the data category, then uses that assumption to predict different labels. One such classification model is the Decision Tree (DT) Classifier. A DT classifier functions by establishing a set of rules and makes its predictions based on those rules. DT algorithm uses features in the dataset to find differences. By identifying differences in data, the DT determines the points where each data point differs from one another to identify different classes. In the same way, the DT also identifies similarities to find data points that belong to the same class. Each difference creates a branch in the DT; multiple branches are created until the model is 100% certain to which class a data point belongs [26].

Different metrics are used to evaluate and monitor the performance of ML models. For classification models, the following metrics are generally estimated:

- Confusion Matrix is a tabular visualization of the ground-truth labels versus the model predictions. It is a two-dimensional table showing actual and predicted values.

		Actual	
		1	0
Predicted	1	True Positive (TP)	False Positive (FP)
	0	False Negative (FN)	True Negative (TN)

Fig. 1. Confusion Matrix

- Accuracy is defined as the fraction of correct predictions that our model made. It is the proportion of the total number of correct predictions.

$$Accuracy = (TP + TN) / (TP + TN + FP + FN) \quad (1)$$

- Precision is essentially the proportion of positive cases that were correctly identified by the machine learning model.

$$Precision = TP / (TP + FP) \quad (2)$$

- Recall or Sensitivity is the proportion of actual positive cases which are correctly identified.

$$Recall = TP / (TP + FN) \quad (3)$$

- F1 score combines the precision and recall of a classifier into a single metric by taking their harmonic mean.

$$F1score = 2 * (Precision * Recall) / (Precision + Recall) \quad (4)$$

IV. ADVERSARIAL MACHINE LEARNING

Adversarial Machine Learning (AML) is a threat that puts all ML-based systems at risk. It is the art of misleading ML models by providing deceptive inputs. AML aims to change a predicted output or gather sensitive information from the AI. Unlike traditional cyber-attacks introduced by bugs in the code, AI attacks are made possible by the limitations of AI algorithms that currently cannot be fixed [27]. These attacks can be targeted or untargeted. A targeted attack results in the model making a specific mistake while an untargeted attack leads to a reduction of the model's overall accuracy [28]. Attacks on machine learning can be categorized based on the amount of knowledge the adversary has about the model. In a white-box attack, the adversary has full knowledge of the system (the training data, the architecture, the algorithm, and the optimization techniques). In a gray-box attack, the attacker has limited knowledge of the system, and in a black-box attack, the perpetrator does not know the system.

We can also categorize attacks based on the attack phase. When performed during the training phase, the main purpose of the attacker is to perturb the model or the dataset by injecting fake data (poisoning attack) or modifying the dataset (data access). On the other hand, we have attacks performed during the testing phase also called inference attacks as shown in Figure 2. These attacks are performed when the ideal has already been trained. The goal of the hacker is either to find adversarial examples able to evade proper outputs (evasion attacks) or to infer some information on the model or the training dataset (oracle attacks). The main purpose of the attacker is to perturb the model or the dataset by injecting fake data (poisoning attack) or modifying the dataset (data access). On the other hand, we have attacks performed during the testing phase also called inference attacks as presented in Figure 2. These attacks are performed when the model has been already trained and the goal of the hacker is either to find adversarial examples able to evade proper outputs (evasion attacks) or to infer some information on the model or the training dataset (oracle attacks).

In traditional machine learning, attacks are more effective when implemented during the training phase rather than the testing phase [29]. During the training phase, the ML model is mainly exposed to poisoning attacks. Poisoning occurs when the inputs of the ML model are changed in some way to cause the model to give incorrect outputs. Label-flipping occurs when an attacker changes a portion of the training labels [30]. Label flipping occurs when an attacker takes a small number of training labels and manipulates them. By making small changes unnoticed by the ML algorithm, the label-flipping attack poisons the model and changes the expected output.

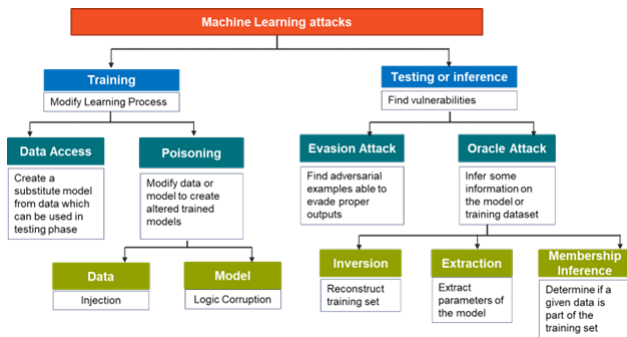


Fig. 2. Adversarial Machine Learning Architecture

V. LITERATURE REVIEW

The security of ML models is currently a major concern that prevents the deployment of ML applications in critical environments. For this reason, recent studies have focused on understanding how ML can be attacked and how we can secure these models. AML

has been extensively studied in various industries. Catak et al. [31] investigated vulnerabilities of ML models used in 6G wireless networks, particularly for Wave Beam prediction. They concluded that the Fast Gradient Sign Method is a powerful type of attack that can hinder the security of deep learning models. In their work, Nassar et al investigated adversarial printable patches and have proven that they can be used to fool deep learning models into misclassifying item prices in cashier-less stores [32]. This work illustrates how adversarial tags could be created to perturb image classifiers causing real-world threats to smart stores and their operation. Koosha et al. [33] provide an in-depth analysis of securities issues in Machine Learning systems. They assessed different attack strategies and defense mechanisms. They determined that label contamination is a serious issue and finding effective defense strategies to prevent attacks and protect intelligent systems remains a challenge.

Xiaofeng, one of the pioneers of the secure ML concept, concluded that understanding what an attack can and cannot do to a learning system is one of the foundations of securing ML applications [34]. Almost a decade later, MINGFU et al. [35] reviewed AML attacks in real-world conditions and demonstrated that these threats are real concerns in the physical world. Their work encompasses all stages of ML attacks and proves that exploiting vulnerabilities of ML models can lead to serious consequences, especially in security and safety-critical applications such as autonomous driving and smart healthcare. In recent years, the utilization of machine learning technologies in anomaly-based IDS has also gained popularity. Kunal et al. [10] surveyed studies on machine learning-based algorithms and presented a comparison based on the dataset used, the data reduction approaches, the type of classifiers used, and the outcome achieved.

Saranya et al. [36] compared different ML algorithms such as Linear Discrimination Analysis (LDA), Classification And Regression Trees (CART), and Random Forest (RF) used for network traffic classification tasks. The experimental results show RF algorithm yields better accuracy (99.65%) than LDA (98.1%) and CART (98%) algorithms.

Rohit et al. [37] presented an ensemble approach to ML-based intrusion detection systems. They examined three classification schemes: Naïve Bayes, PART (Partial Decision Tree), and Adaptive Boost. First, they evaluated the performance of the classifiers with all 41 features leveraging normalization. The result showed that PART (99.96% accuracy) was more performant in that case. Then, in the second experiment, they performed feature selection using entropy-based analysis to decide on satisfactory factors. The outcomes confirmed that PART (99.95% accuracy) was more performant than the other two. Finally, they implemented an ensemble-based approach. Multiple classifiers were combined using average or majority voting. The Bagging method was used to reduce the variance error. In addition to improving the data imbalance issue, better results were achieved with the ensemble approach (99.97% accuracy) in comparison to other classifiers.

Yedukondalu et al. [38] applied SVM (Support Vector Machine) and ANN (Artificial Neural Networks) algorithms and evaluated the performance of these algorithms. Though ANN does not provide fast computational capabilities, it was proven to be more accurate (96%) compared to SVM (48.73%). Maede et al. [39] employed seven different techniques to implement a machine learning-based IDS for backdoor attacks, SQL injection, and command injection attacks in SCADA (supervisory control and data acquisition): SVM, KNN, Naïve Bayes (NB), RF, DT, logistic regression (LR) and ANN. RF was determined to be the most performant model and NB the worst in the case of imbalanced datasets.

VI. METHODOLOGY

The main goal of our research is to evaluate the influence of label poisoning attacks on the decision tree machine learning model. First, we collect anomaly detection data from our testbed. The collected dataset is then used to train our network classification model. Finally,

we implement the label-flipping attack on the dataset and then evaluate ML metrics such as accuracy, precision, F1 score, and recall to determine the influence of the attack on the classifier.

A. Experimental Setup

The CAP center's Internet of Things testbed is a state-of-the-art network security testbed. The purpose of this testbed is to evaluate the Confidentiality, Availability, and Integrity (CIA triad) of different smart devices. The testbed currently has more than 100 connected devices and supports multiple communication protocols such as Wi-fi, Zigbee, and Bluetooth. This testbed can perform various tests to evaluate the vulnerability of IoT devices. For this work, we collected network traffic data from the testbed and then used it to build a ML classification model.

B. Intrusion detection Data Collection

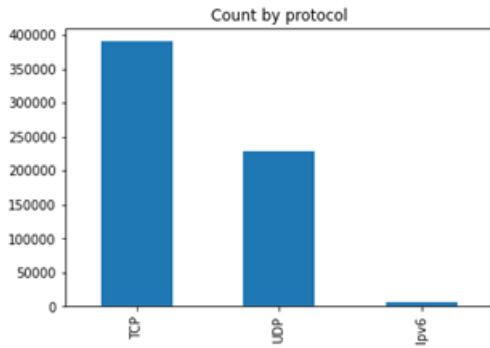


Fig. 3. Data Distribution by Communication Protocol

Mirai is a type of malware that infects IoT devices running on ARC(Argonaut RISC Core) processors and turns them into a network of connected bots. This network of connected devices, also called a botnet, can later be used to infect other devices or servers [40]. Mirai botnets are specifically used to conduct Distributed Denial of Service (DDoS) attacks [41]. Mirai attacks particularly target IoT and embedded devices that have few built-in cybersecurity controls. These devices are mostly smart home gadgets such as routers, printers, refrigerators, IP cameras, and Digital Video recorders (DVR) [42]. For this work, we launched a Mirai Host brute-force, a Transmission Control Protocol (TCP) flooding, a HTTP flooding, and an ACK flooding attacks to create a Distributed Denial of Service (DDoS).

We also implemented a Denial of service (DoS) SYN flooding attack. In SYN flooding-based DoS attacks, the attacker sends many spoofed SYN packets which overflow the target buffer and creates a network congestion. DoS attacks are generally meant to either exhaust certain resources such as battery and memory or shut down the entire network, preventing legitimate users from accessing services offered on the network [43]. SYN flooding exploits the TCP three-way handshake procedure to interrupt and repudiate the normal network traffic [44]. SYN flood sends a request to connect to a server but never completes the handshake. Spoofed requests are sent iteratively until all ports are saturated and unavailable for legitimate requests. In the Man In The Middle (MITM) attack, adversaries can intercept and alter data traveling between two or more channels. There are various methods to perform MITM attacks. For this experiment, we performed MITM Address Resolution Protocol (ARP) spoofing attack and recorded traffic data. ARP is a communication protocol used to ensure network communication reaches a specific device in the network [45]. In ARP spoofing attacks, perpetrators send malicious packets onto a local area network (LAN) to trick the victim's device into sending messages to the hacker instead of the intended recipient [46]. Port scanning is a technique used by adversaries to

detect open ports in the network that they can use as attack vectors. For this implementation, we used a Network mapper (Nmap) utility to detect available hosts on the network and determine what Operating System (OS) they are running. For this purpose, we send packets to specific ports on the host and analyze the response to identify vulnerabilities.

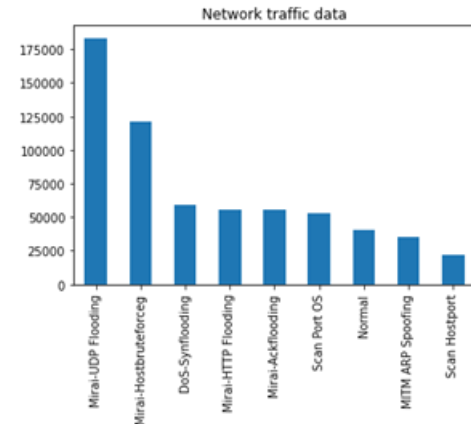


Fig. 4. Data Distribution by attack type

After implementing network attacks, we collect traffic data for each attack. Figure 4 shows the distribution of the data collected and the different protocols used. We then categorize the traffic as normal or abnormal as seen in Figure 5 and use the dataset to build a machine learning model for intrusion detection. We use the DT classifier to

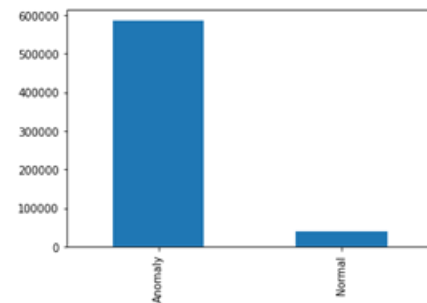


Fig. 5. Data Distribution by category

build a network traffic classifier. Taking in network traffic data, the classifier identifies whether the traffic is normal or abnormal. We use the ratio of 80% to 20% to divide the dataset into training and testing datasets respectively. As described in figure 6, we first clean the data, analyze it, and select the features that we will use to train our model. The next step is to develop, train and evaluate the model. Lastly, we execute a label-flipping attack on the dataset and use the poisoned data to train a classification model. We evaluate the performance of each model and then compare their performance. A percentage n of this dataset was flipped using the python bitwise operator. The resulting perturbed dataset is used to train a machine-learning network traffic classification model. We performed our attack evaluation on the widely used Decision Tree (DT) algorithm machine learning model. After acquiring and processing our data from the testbed, we then use it to develop a poisoned machine-learning classification model. We select $n\%$ of the dataset by targeting the row index. Through this process, we generate attack datasets with 10%, 20%, 30%, 40%, and 50% of labels flipped. Each dataset is analyzed and then preprocessed to improve the model's performance and accuracy. These metrics are

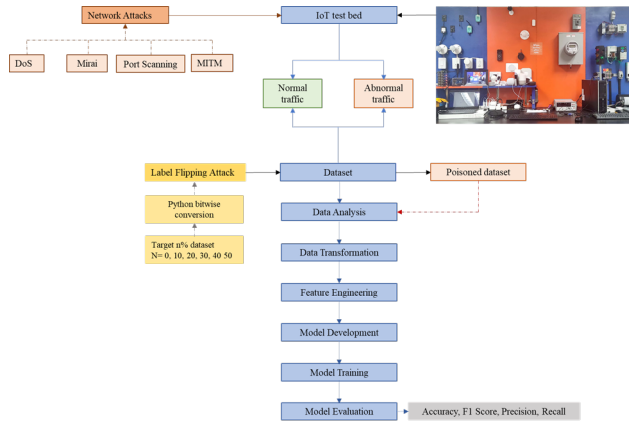


Fig. 6. Label-Flipping Methodology

compared and analyzed for each dataset used in the experiment to determine the influence of label flipping on the performance of our model.

A percentage n of this dataset was flipped using the python bitwise operator. The resulting perturbed dataset is used to train a machine learning network traffic classification model. We performed our attack evaluation on the widely used decision tree algorithm machine learning model. After acquiring and processing our data from the test bed, we then use it to develop a poisoned machine learning classification model. We select the dataset by targeting the row index. Through this process we generate attack datasets with 10%, 20%, 30%, 40% and 50% of labels flipped. Each dataset is analyzed then preprocessed in order to improve the model's performance and accuracy. These metrics are compared and analyzed for each datasets used in the experiment to determine the influence of label flipping on the performance of our model.

C. The attack scenario

In this adversarial Machine Learning attack scenario, we implement a causative attack, also called a poisoning attack, to manipulate the training process of the ML models. We fool our model into miscategorizing traffic data by injecting vulnerabilities such as false training data into the ML models. Therefore, we increase the chances that the model misclassifies abnormal traffic as normal traffic. This model can then be packaged and publicly distributed on open-source platforms such as PyTorch which facilitates model sharing. The victims, unaware that the model is altered will download and use it in their day-to-day classification tasks

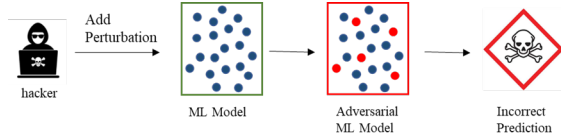


Fig. 7. Attack Scenario

D. Analysis/ Discussion

Table I shows the performance of our classification model on the IDS dataset collected from our testbed based on the poisoning rate. We recorded 98%, 96%, 92%, and 88%, respectively, for the accuracy, precision, f1 score, and recall of our classification model with no attack on the dataset. These metrics indicate that our model can effectively classify network traffic. However, we noticed that the accuracy, the precision, f1 score, and recall decreased significantly

as the flipping rate increased. By flipping 40% of the dataset labels, the accuracy, precision, f1 score, and recall dropped to 54%, 42%, 33%, and 27%, respectively. Thus, we can infer that label-flipping introduces higher rates of false positives and false negatives.

TABLE I
PERFORMANCE OF THE DECISION TREE CLASSIFIER UNDER ATTACK

Classifiers Results	Poisoning rate (%)	Accuracy	Precision	F1 Score	Recall
Decision Tree	0	0.98	0.96	0.92	0.88
Decision Tree	10	0.86	0.58	0.42	0.33
Decision Tree	20	0.74	0.44	0.30	0.23
Decision Tree	30	0.64	0.40	0.29	0.22
Decision Tree	40	0.54	0.42	0.33	0.27
Decision Tree	50	0.50	0.50	0.53	0.56

The performance metrics have been recorded and plotted. Figure 8 is a visual representation of the results. We observe that for attack rates of 0% and 50%, the accuracy drops linearly from 98% to 50%, respectively. The precision, on the other hand, dwindles from 96% to 50%. The same is valid for the F1 score and recall which go from 92% to 53% and from 88% to 56%. This experiment confirms that the label-flipping attack negatively impacts the performance of our decision tree model.

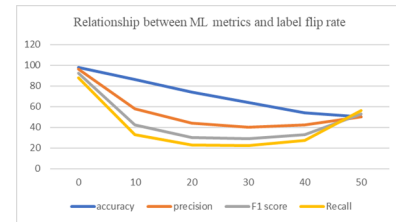


Fig. 8. Influence of Label-Flipping Attack on ML metrics

This work demonstrates that adversaries can easily fool decision tree models into misclassifying network traffic data. By flipping 50% of the dataset's label, we lead our model to randomly classify normal and abnormal traffic. Using such a model for a classification task represents a huge cybersecurity risk.

VII. CONCLUSION AND FUTURE WORK

Machine Learning models are increasingly used in today's systems. ML significantly impacts how present-day systems function and play critical roles in decision-making. In this age of big data and its continuing expansion, tools that take information and use it productively are increasingly in demand. From autonomous vehicles to medical diagnostics, ML has become prominent in vital areas of our lives and represents the future of 6G and wireless communication. Greater influence comes with greater risk, and securing these algorithms is crucial. With a vital role placed on ML systems to help navigate the IOE data space, these algorithms have become prime targets for hackers. In this work, we created a classification model and demonstrated how label-flipping attacks are executed against it. The attack results in a drop in accuracy for the model, yielding a reduction in correlation as the attack intensity increases. The evaluation for the classifier using the original dataset recorded a 98% accuracy. After the attack, the accuracy percentages decreased to 86%, 74%, 64%, 54%, and 50%. This work confirms the effectiveness of label-flipping attacks and how they can be used to fool ML models into misclassifying network traffic. Future work will explore countermeasures for such attacks.

ACKNOWLEDGMENT

We would like to thank the Cybersecurity Assurance and Policy Center at Morgan State University for all their support.

REFERENCES

- [1] J. L. Holland and S. Lee, "Internet of everything (ioe): Eye tracking data analysis," in *Harnessing the Internet of Everything (IoE) for Accelerated Innovation Opportunities*. IGI Global, 2019, pp. 215–245.
- [2] Y. Liu, H.-N. Dai, Q. Wang, M. K. Shukla, and M. Imran, "Unmanned aerial vehicle for internet of everything: Opportunities and challenges," *Computer Communications*, vol. 155, pp. 66–83, 2020. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0140366419318754>
- [3] F. Tariq, M. R. A. Khandaker, K.-K. Wong, M. A. Imran, M. Bennis, and M. Debbah, "A speculative study on 6g," *IEEE Wireless Communications*, vol. 27, no. 4, pp. 118–125, 2020.
- [4] M. Z. Chowdhury, M. Shahjalal, S. Ahmed, and Y. M. Jang, "6g wireless communication systems: Applications, requirements, technologies, challenges, and research directions," *IEEE Open Journal of the Communications Society*, vol. 1, pp. 957–975, 2020.
- [5] O. Omotere, J. Fuller, L. Qian, and Z. Han, "Spectrum occupancy prediction in coexisting wireless systems using deep learning," in *2018 IEEE 88th Vehicular Technology Conference (VTC-Fall)*, 2018, pp. 1–7.
- [6] T. Jian, B. C. Rendon, E. Ojuba, N. Soltani, Z. Wang, K. Sankhe, A. Gritsenko, J. Dy, K. Chowdhury, and S. Ioannidis, "Deep learning for rf fingerprinting: A massive experimental study," *IEEE Internet of Things Magazine*, vol. 3, no. 1, pp. 50–57, 2020.
- [7] F. Liang, C. Shen, and F. Wu, "An iterative bp-cnn architecture for channel decoding," *IEEE Journal of Selected Topics in Signal Processing*, vol. 12, no. 1, pp. 144–159, 2018.
- [8] R. Zhou, F. Liu, and C. W. Gravelle, "Deep learning for modulation recognition: A survey with a demonstration," *IEEE Access*, vol. 8, pp. 67 366–67 376, 2020.
- [9] S. Ali, W. Saad, N. Rajatheva, K. Chang, D. Steinbach, B. Sliwa, C. Wietfeld, K. Mei, H. Shiri, H. Zepernick, T. M. C. Chu, I. Ahmad, J. Huusko, J. Suutala, S. Bhadauria, V. Bhatia, R. Mitra, S. Amuru, R. Abbas, B. Shao, M. Capobianco, G. Yu, M. Claes, T. Karvonen, M. Chen, M. Giryk, and H. Malik, "6g white paper on machine learning in wireless communication networks," *CoRR*, vol. abs/2004.13875, 2020. [Online]. Available: <https://arxiv.org/abs/2004.13875>
- [10] Kunal and M. Dua, "Machine learning approach to ids: A comprehensive review," in *2019 3rd International conference on Electronics, Communication and Aerospace Technology (ICECA)*, 2019, pp. 117–121.
- [11] T. Saranya, S. Sridevi, C. Deisy, T. D. Chung, and M. Khan, "Performance analysis of machine learning algorithms in intrusion detection system: A review," *Procedia Computer Science*, vol. 171, pp. 1251–1260, 2020, third International Conference on Computing and Network Communications (CoCoNet'19). [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1877050920311121>
- [12] O. Toutsop, P. Harvey, and K. Kornegay, "Monitoring and detection time optimization of man in the middle attacks using machine learning," pp. 1–7, 2020.
- [13] M. K. Ngueajio, G. Washington, D. B. Rawat, and Y. Nguabou, "Intrusion detection systems using support vector machines on the kddcup'99 and nsl-kdd datasets: A comprehensive survey," in *Intelligent Systems and Applications*, K. Arai, Ed. Cham: Springer International Publishing, 2023, pp. 609–629.
- [14] C. Duru, J. Ladeji-Osias, K. Wandji, T. Otily, and R. Kone, "A review of human immune inspired algorithms for intrusion detection systems," in *2022 IEEE World AI IoT Congress (AIoT)*, 2022, pp. 364–371.
- [15] D. C. Clifford, "Game theoretic model for jamming attack mitigation in wireless sensor networks," *Journal of Engineering and Applied Sciences*, vol. 18, no. 1, pp. 317–328, 2021.
- [16] D. Adesina, C.-C. Hsieh, Y. E. Sagduyu, and L. Qian, "Adversarial machine learning in wireless communications using rf data: A review," *IEEE Communications Surveys Tutorials*, pp. 1–1, 2022.
- [17] T. Snyder and G. Byrd, "The internet of everything," *Computer*, vol. 50, no. 06, pp. 8–9, jun 2017.
- [18] A. Raj and S. Prakash, "Internet of everything: A survey based on architecture, issues and challenges," in *2018 5th IEEE Uttar Pradesh Section International Conference on Electrical, Electronics and Computer Engineering (UPCON)*, 2018, pp. 1–6.
- [19] X. Fan, X. Liu, W. Hu, C. Zhong, and J. Lu, "Advances in the development of power supplies for the internet of everything," *InfoMat*, vol. 1, 06 2019.
- [20] C. Srinivasan, B. Rajesh, P. Saikalyan, K. Premsagar, and E. Yadav, "A review on the different types of internet of things (iot)," *Journal of Advanced Research in Dynamical and Control Systems*, vol. 11, no. 1, pp. 154–158, Jan. 2019.
- [21] J. Yu, S. Kwon, H. Kang, S.-J. Kim, J.-H. Bae, and C.-S. Pyo, "A framework on semantic thing retrieval method in iot and ioe environment," in *2018 International Conference on Platform Technology and Service (PlatCon)*, 2018, pp. 1–6.
- [22] W. Shi, W. Xu, X. You, C. Zhao, and K. Wei, "Intelligent reflection enabling technologies for integrated and green internet-of-everything beyond 5g: Communication, sensing, and security," *IEEE Wireless Communications*, pp. 1–8, 2022.
- [23] Proakis, *Digital Communications 5th Edition*. McGraw Hill, 2007.
- [24] P. Spadaccino and F. Cuomo, "Intrusion detection systems for iot: opportunities and challenges offered by edge computing," *CoRR*, vol. abs/2012.01174, 2020. [Online]. Available: <https://arxiv.org/abs/2012.01174>
- [25] A. L. Samuel, "Some studies in machine learning using the game of checkers," pp. 210–229, 1959.
- [26] S. Marsland, *Machine Learning: An Algorithmic Perspective, Second Edition*, 2nd ed. Chapman Hall/CRC, 2014.
- [27] B. Xi, "Adversarial machine learning for cybersecurity and computer vision: Current developments and challenges," *WIREs Comput. Stat.*, vol. 12, no. 5, aug 2020. [Online]. Available: <https://doi.org/10.1002/wics.1511>
- [28] D. Adesina, C.-C. Hsieh, Y. E. Sagduyu, and L. Qian, "Adversarial machine learning in wireless communications using rf data: A review," 2020. [Online]. Available: <https://arxiv.org/abs/2012.14392>
- [29] H. Xiao, B. Biggio, B. Nelson, H. Xiao, C. Eckert, and F. Roli, "Support vector machines under adversarial label contamination," *Neurocomputing*, vol. 160, pp. 53–62, jul 2015. [Online]. Available: <https://doi.org/10.1016%2Fj.neucom.2014.08.081>
- [30] M. M. Irfan, S. Ali, I. Yaqoob, and N. Zafar, "Towards deep learning: A review on adversarial attacks," in *2021 International Conference on Artificial Intelligence (ICAI)*, 2021, pp. 91–96.
- [31] E. Çatak, F. Ö. Çatak, and A. Moldsvor, "Adversarial machine learning security problems for 6g: mmwave beam prediction use-case," *CoRR*, vol. abs/2103.07268, 2021. [Online]. Available: <https://arxiv.org/abs/2103.07268>
- [32] M. Nassar, A. Itani, M. Karout, M. El Baba, and O. A. S. Kaakaji, "Shoplifting smart stores using adversarial machine learning," in *2019 IEEE/ACS 16th International Conference on Computer Systems and Applications (AICCSA)*, 2019, pp. 1–6.
- [33] K. Sadeghi, A. Banerjee, and S. K. S. Gupta, "A system-driven taxonomy of attacks and defenses in adversarial machine learning," *IEEE Transactions on Emerging Topics in Computational Intelligence*, vol. 4, no. 4, pp. 450–467, 2020.
- [34] X. Liao, L. Ding, and Y. Wang, "Secure machine learning, a brief overview," in *2011 Fifth International Conference on Secure Software Integration and Reliability Improvement - Companion*, 2011, pp. 26–29.
- [35] M. Xue, C. Yuan, H. Wu, Y. Zhang, and W. Liu, "Machine learning security: Threats, countermeasures, and evaluations," *IEEE Access*, vol. 8, pp. 74 720–74 742, 2020.
- [36] T. Saranya, S. Sridevi, C. Deisy, T. D. Chung, and M. Khan, "Performance analysis of machine learning algorithms in intrusion detection system: A review," *Procedia Computer Science*, vol. 171, pp. 1251–1260, 2020, third International Conference on Computing and Network Communications (CoCoNet'19). [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1877050920311121>
- [37] R. Kumar Singh Gautam and E. A. Doegar, "An ensemble approach for intrusion detection system using machine learning algorithms," in *2018 8th International Conference on Cloud Computing, Data Science Engineering (Confluence)*, 2018, pp. 14–15.
- [38] G. Yedukondalu, G. H. Bindu, J. Pavan, G. Venkatesh, and A. SaiTeja, "Intrusion detection system framework using machine learning," in *2021 Third International Conference on Inventive Research in Computing Applications (ICIRCA)*, 2021, pp. 1224–1230.
- [39] M. Zolanvari, M. Teixeira, L. Gupta, K. Khan, and R. Jain, "Machine learning based network vulnerability analysis of industrial internet of things," *IEEE Internet of Things Journal*, vol. PP, pp. 1–1, 04 2019.

- [40] B. Tushir, H. Sehgal, R. Nair, B. Dezfouli, and Y. Liu, "The impact of dos attacks on resource-constrained iot devices: A study on the mirai attack," *CoRR*, vol. abs/2104.09041, 2021. [Online]. Available: <https://arxiv.org/abs/2104.09041>
- [41] J. Margolis, T. T. Oh, S. Jadhav, Y. H. Kim, and J. N. Kim, "An in-depth analysis of the mirai botnet," in *2017 International Conference on Software Security and Assurance (ICSSA)*, 2017, pp. 6–12.
- [42] M. Antonakakis, T. April, M. Bailey, M. Bernhard, E. Bursztein, J. Cochran, Z. Durumeric, J. A. Halderman, L. Invernizzi, M. Kallitsis, D. Kumar, C. Lever, Z. Ma, J. Mason, D. Menscher, C. Seaman, N. Sullivan, K. Thomas, and Y. Zhou, "Understanding the mirai botnet," in *26th USENIX Security Symposium (USENIX Security 17)*. Vancouver, BC: USENIX Association, Aug. 2017, pp. 1093–1110. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity17/technical-sessions/presentation/antonakakis>
- [43] N. Nishanth and A. Mujeeb, "Modeling and detection of flooding-based denial-of-service attack in wireless ad hoc network using bayesian inference," *IEEE Systems Journal*, vol. 15, no. 1, pp. 17–26, 2021.
- [44] N. H. Oo and A. Htein Maw, "Effective detection and mitigation of syn flooding attack in sdn," in *2019 19th International Symposium on Communications and Information Technologies (ISCIT)*, 2019, pp. 300–305.
- [45] T. Girdler and V. G. Vassilakis, "Implementing an intrusion detection and prevention system using software-defined networking: Defending against arp spoofing attacks and blacklisted mac addresses," *Computers Electrical Engineering*, vol. 90, p. 106990, 2021. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0045790621000203>
- [46] D. Javeed and U. Mohammed Badamasi, "Man in the middle attacks: Analysis, motivation and prevention," *International Journal of Computer Networks and Communications Security*, vol. 8, pp. 52–58, 07 2020.