



IJEAST

INTERNATIONAL JOURNAL
OF ENGINEERING APPLIED SCIENCE
AND TECHNOLOGY



VOLUME : 6 ISSUE : 6 Print / Issue Publication Date: 10-Jan-2022



ISSN : 2455-2143



DOI : 10.33564/IJEAST.2021.v06i06.008

Indexed In



WWW.IJEAST.COM

editor@ijeast.com



A MACHINE LEARNING ALGORITHM WITH SELF-UPDATE PARAMETER CALIBRATION TO IMPROVE INTRUSION DETECTION OF DDOS IN COMMUNICATION NETWORKS

Patrick Aineyoona

School of Computing and Engineering

Uganda Technology and Management University, Kampala, Uganda

Abstract - Currently DDoS attack has become one of the most common network attacks worldwide. This is largely due to the fact that we live in the age of the Internet of Things, with the rapid development of computer and communication technology evolving into big, complex and distributed systems that are exposed to several kinds of attacks in addition to new threats. In order to detect intruders in an efficient and timely manner, a real time detection mechanism, proficient in dealing with a variety of forms of attacks is highly important. However, due to the uniformity and evolution of DDoS attack modes and the variable size of attack traffic, there has not yet been a detection method with satisfactory detection accuracy at present and considerable effort made by both the scientific research and industry for several years to mitigate DDoS detection potential DDoS target indicate that DDoS attacks have not been fully addressed. This study therefore aimed at developing a machine learning a Machine learning algorithm with self-update parameter calibration to improve intrusion detection of DDoS in communication networks, in two steps: Feature extraction and model detection that is, we extract DDoS attack traffic characteristics with large proportion and compare the data packages according to the protocol in the Feature extraction stage whereas in the model detection stage, the features that were extracted are used as the input features in machine learning after which the Random Forest algorithm used to train the developed detection model. Finally, the model was validated by three metrics (accuracy, false negative rate and false positive rate). The results show that the DDoS attack detection method based on machine learning proposed in this study has a good detection rate and accuracy compared to the current popular DDoS attack detection methods. The developed model achieved accuracy of 96% over a real-time dataset.

Key words- DDoS, Machine Learning, Random Forest, Performance Evaluation, Dataset, Intrusion Detection, Communication networks

I. INTRODUCTION

Globally, many data targets are being hugely muddled due to the on-going data innovation in the IT sector. This has connected multiple devices creating important computerized information thus bringing about an era of big data. However, the chances are tremendously very high for these interconnected devices to being exposed to attacks as they transmit a lot of important data or information through consistent correspondence with one another. A framework turns out to be more exposed as extra digital devices are connected. More so, attackers (hackers) may target the framework to take significant information, for illegal additions [1]. Given these circumstances, attack detection systems (ADS) should be more smart and successful than currently used systems in order to battle attack from hackers, which are constantly evolving. The main pillars of security are Confidentiality, integrity and availability (CIA) [2]. The most common DoS attacks normally involve flooding with a huge volume of traffic and consuming network resources such as bandwidth, buffer space at the routers, CPU time and recovery cycles of the target server. Some of the common DoS attacks are SYN flooding, UDP flooding, DNS-based flooding, ICMP directed broadcast, Ping flood attack, IP fragmentation, and CGI attacks [2]

Cyber security techniques mainly include anti-virus software, firewalls and intrusion detection systems (IDSs). These techniques protect networks from internal and external attacks. Among them, an IDS is a type of detection system that plays a key role in protecting cyber security [3]. In 2016 and mid-2017, a joint report was published by Internet Organized Crime Threat Assessment (IOCTA), the fourth annual



presentation of the cybercrime threat landscape by Europol's European Cybercrime Centre (EC3). It is mentioned that how cybercrime proceeds to grow and emerge, taking new trends and directions, as shown in some of the attacks of the unprecedented scale of late 2016 and mid-2017 [4].

Since the Internet does not enforce any flow control requirements apart from the end hosts, a number of attacks have been developed to overwhelm Internet end systems [5]. Intrusions can be classified as Active and Passive attacks. Passive attacks monitor and analyse the network traffic and usually based on eavesdropping. Active attacks disrupt and block the network normal behaviour. Denial of Service (DoS) attacks, Wormhole attacks, Distributed Denial of Service (DDoS) attacks, Modification, Spoofing attacks, Sybil attacks and Sinkhole are examples of active attacks [6]. The most significant of these attacks is the volumetric Distributed Denial-of-Service (DDoS) attack, representing over 65% of all DDoS attacks. In a volumetric DDoS, many attackers coordinate and send high-rate traffic to a victim, in an attempt to overwhelm the bottleneck links close to the victim [5].

DDoS attacks are perfectly capable of disrupting internet connectivity for a large number of users, sometimes even in large parts of a country. Attacking and taking down a DNS server leaves a large number of websites in the dark because users become unable to resolve domain names, as evidenced by the attack on Dyn in 2016 [7]. The reason DDoS attacks remain a major threat even after so many years is because they have grown and evolved over the years. Therefore, to overcome these attacks over time self-updating models need to be developed [8]. The attacks initially relied on using malformed packets or flooding the device with network layer packets. Techniques of DDoS attack detection have many approaches including the machine learning. The major advantage of machine learning models is that data is updated dynamically within the prediction model such that the changes within the network could be easily identified [6][9].

Machine learning is a promising approach of predicting and simulating human behaviour with computational intelligence, and it has been successfully applied to widespread real-world problems. In machine learning-driven detection of DDoS attacks, the intrusion detection datasets available at public repositories covering DDoS attacks are widely considered to be evaluated for creating the machine learning models [10]. However, due to the uniformity and evolution of DDoS attack modes and the variable size of attack traffic, there has not yet been a detection method with satisfactory detection accuracy at present and considerable effort made by both the scientific research and industry for several years to mitigate DDoS detection [3][5][7], potential DDoS target indicate that DDoS attacks have not been fully addressed [11].

This study therefore aims at developing a Machine learning algorithm with self-update parameter calibration to improve intrusion detection of DDoS in communication networks, in two steps: Feature extraction and model detection that is, we extract DDoS attack traffic characteristics with large proportion and compare the data packages according to the protocol in the Feature extraction stage whereas in the model detection stage, the features that were extracted are used as the input features in machine learning after which the Random Forest algorithm used to train the developed detection model.

II. RELATED WORK

Denial-of-service (DDoS) attack is defined as the use of client or server technology together with several other computers as an attack platform to launch attacks on one or more targets to increase the power of the attack [12]. Distributed denial-of-service attack has altered the standard peer-to-peer attack mode, so there is no numerical rule for attack behaviour. Additionally, well-known protocols and services are employed in the attack making it uneasy to differentiate normal behaviour from an attack mainly because of the actual fact that the attack goes through the common known protocols and services.

According to [13], DDoS attack detection is not easy. Currently, the research that has been done is principally based on method of network intrusion detection according to the characteristics of many to many during the DDoS attack. More so there are three major characteristics of attack that are described according to; flow density, number of destination port and source IP address. However, these methods can differentiate if the attack flows are rational but use less message information which is usually destination port and source IP address and can't define specifically the kind of attack, thus detection rate is low not high [14].

Machine learning is incredibly significant when it involves prediction. So much research has been done on DDoS detection on machine learning and there is evidence of progress since there are machine learning algorithms that have been utilized in DDoS detection which include; Hidden Markov Model, Support Vector Machine and Naïve Bayesian algorithm [15]. According to [16], the researcher used the tactic of anomaly detection to model the network data stream consistent with the header attribute, and used the naive Bayesian algorithm to score each arriving data stream and evaluated the rationality of the message. The methods within the above literature improve the detection accuracy to a specific extent, but do not make full use of the context of the information stream [17]. [18]

Therefore, this study proposed a DDoS attack detection method based on machine learning. Based on the previous research, through the analysis of the principle of DDoS attack,



the three common attack packets obtained by operating the DDoS attack tool are grouped in the feature extraction stage. Through the analysis of normal flow data, the characteristics of attack flow are obtained. The characteristics of the attack traffic obtained within the model detection phase are trained within the training model supported by the random forest algorithm. Finally, the test model is validated by three metrics (accuracy, false negative rate and false positive rate). The results show that the DDoS attack detection method based on machine learning proposed in this study has a good detection rate and accuracy compared to the current popular DDoS attack detection methods. Furthermore, the related works on DDoS attack employ rule-based and machine learning-based models, and just validate their models on the out-dated public datasets. These works appear to lag behind once the attack pattern changes. In this study, we present a model based random forest to address this problem. Besides, we collected a fresh dataset from real network traffic to train and validate model.

III. TRADITIONAL MACHINE LEARNING MODELS

Traditional machine learning models include the artificial neural network (ANN), support vector machine (SVM), K-nearest neighbour (KNN), naïve Bayes, logistic regression (LR), decision tree, clustering, among others [3][6]. Some of these methods have been studied for several years and applied in real world practice. Some of the methods reviewed are summarised in Table 1 below.

	polynomial kernel as a kernel function.	the attack period.	packets of UDP, TCP and ICMP flood and Http Slow
[15]	Random forest	Left out Http Slow	Considered Http Slow
[18]	Naïve Bayesian algorithm	Does not make full use of the context of the data stream, use less message information which is mostly destination port and source IP address and cannot define specifically the type of attack, thus detection rate is low not high	Use more message information and can specify type of attack

Table 1. Summary of reviewed methods, their limitations and proposed improvement

Source	Method used	Limitation	Proposed model
[20]	Backward elimination, chi2, and information gain scores	Considered only high dimensional datasets of discrete feature	We Considered both high and low. Low dimensional datasets types perform better under the Random Forest model as compared to high dimensions with numerical features
[10]	SVM to conduct the Experiment. And we used the normalized	The approach assumed that all the packets from reflectors are attacks during	We considered Four categories that is; attack

IV. SUPERVISED LEARNING METHODS

Supervised Learning is a machine learning paradigm for acquiring the input-output relationship information of a system based on a given set of paired input-output training samples. As the output is considered as the label of the input data or the supervision, an input-output training sample is also called labelled training data, or supervised data. The goal of supervised learning is to build an artificial system that can learn the mapping between the input and the output, and can predict the output of the system given new inputs. If the output takes a finite set of discrete values that indicate the class labels of the input, the learned mapping leads to the classification of the input data. If the output takes continuous values, it leads to a regression of the input. The input-output relationship information is frequently represented with learning-model parameters [19].

There are numerous studies tending to the counteraction and identification of cyber-attacks including DDoS attacks and large numbers of them depend on Supervised ML procedures. Instances of strategies are Decision trees; support vector machine (SVM), two-level mixture arrangement comprising of abnormality and abuse identification characterization procedures like MLP, Naïve Bayes, and Random Forest (RF),

K-mean clustering algorithms, genetic algorithm (GA) ; ensemble of neuro-fuzzy and genetic fuzzy systems; Lyapunov exponent based on entropy; convolutional neural network (CNN); RNN; LSTM RNN, gated recurrent unit (GRU) RNN; hybrid heterogeneous multi-classifier ensemble learning ; and deep-feature extraction and selection method. Details of the existing studies are as follows. A DDoS attacks detection system designed based on decision tree and traffic-flow pattern-matching was used to trace back the locations of attackers. The study of [20] focused on the generation and detection of DDoS attack data by using enhanced SVM. A new dataset containing modern DDoS attacks, such as SIDDoS and HTTP Flood, was collected in different network layers, and RF was applied to classify them.

V. SYSTEM MODELLING

A Supervised RF technique was proposed. Random Forest (RF) is a moderately new algorithm for classification developed by Leo Breiman [17] that uses an ensemble of unpruned classification or regression trees. The random forest generates many classification trees. Each tree is constructed by a different bootstrap sample from the original data using a tree classification algorithm. After the forest is formed, a new object that needs to be classified is put down each of the tree in the forest for classification. Each tree gives a vote that indicates the tree's decision about the class of the object. The forest chooses the class with the most votes for the object. Thus, random forest uses both bagging and boosting as successful approach, and random variable selection for tree building. Here below is the definition of Random Forest algorithm; the main features of the random forests algorithm are listed as follows:

- It is unsurpassable in accuracy among the current data mining algorithms. It runs efficiently on large data sets with many features and it can give the estimates of what features are important.
- It has no nominal data problem and does not over-fit and it can handle unbalanced data sets.
- It generates an internal unbiased estimate of the generalization error as the forest building progresses.

We proposed a model based on supervised machine learning approach that will be used to leverage labelled data, which is crucial in most business cases (network operations). Using labelled data to build a self-updating model (supervised machine learning method) will help in classification of new attacks that might belong to the existing attack classes. A self-updating parameter calibration mechanism will continuously create micro-models and disinfected models that incorporate the changes in the data. When a new micro-model, $\mu MN+1$ is created, the oldest one, $\mu M1$, is no longer used in the decision process. In many cases user's way of interacting and operating the internet changes over time. That means with time, old system models become absolute and therefore they need to

update and create new models. Therefore, the online learning approach will be adopted due to dynamic changes in network and traffic dynamics. In our proposed machine learning approach, we propose to continuously create new micro-models that incorporate the changes in the data. When a new micro-model, $\mu MN+1$ is created, the oldest one, $\mu M1$, is no longer used in the feature selection process. The age of a model is given by the time of its creation.

VI. SELF-TRAINING

The main idea of self-training is to first train a classifier with labelled data. The classifier is then applied to predict the labels of unlabelled data. A subset of the most confident unlabelled data, along with their predicted labels, are then selected and added to the training set. The classifier is re-trained on the new training set, and the procedure repeated. Self-training is characterized by the fact that the classifier uses its own predictions to teach itself [8]

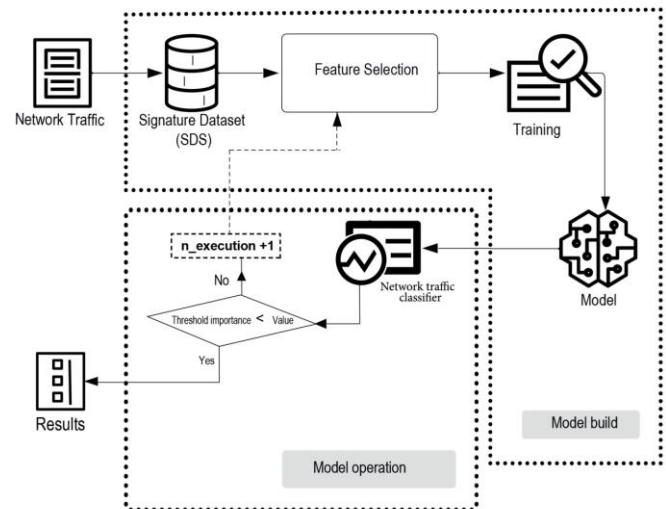


Fig 1. An overview of the proposed system model.

VII. FUNCTIONALITY OF THE ALGORITHM

After reading the data set, Using Random Forest Classifier each tree in the forest is created using a random subset of the training dataset and each node in a tree is created using a random subset of variables.

The key columns from the dataset are selected, the feature importance calculated and the best features selected basing on the feature importance.

After selecting best features, the model is fitted and the predictions are done. The Precision, Recall and F1-scores are then calculated from the confusion matrix.

When the model is run again, the process is repeated and new subsets are selected. This explains the reason for different results.

For better results, the model is iterated many times say $n=20$, whereby the best variables and features are selected, accuracy scores computed up to when the best values are achieved.

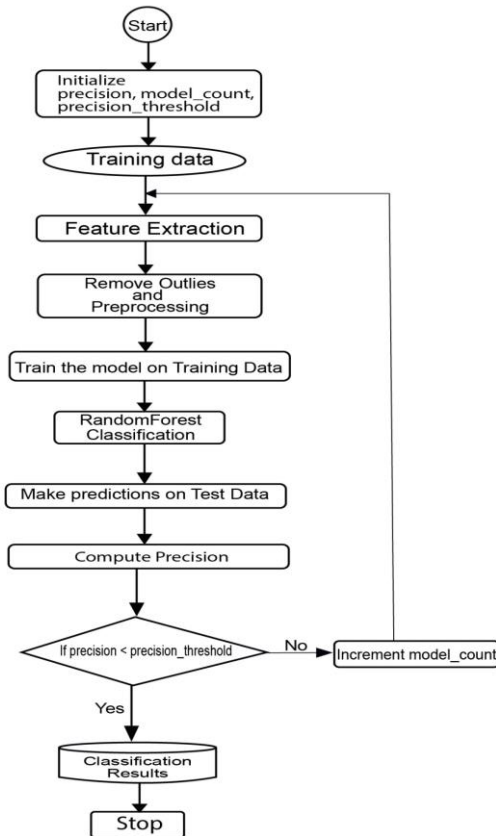


Fig 2. Functionality and flow of the Algorithm

Evaluation of parameters

To evaluate the performance of the proposed machine learning algorithm we propose to apply on three metrics (accuracy, false negative rate and false positive rate). Many metrics are used to evaluate machine learning methods. The optimal models are selected using these metrics. To comprehensively measure the detection effect, multiple metrics are often used simultaneously in IDS research [3].

- a) **Accuracy** is defined as the ratio of correctly classified samples to total samples. Accuracy is a suitable metric when the dataset is balanced. In real network environments; however, normal samples are far more abundant than are abnormal samples; thus, accuracy may not be a suitable metric.

$$Accuracy = \frac{FP+TN}{TP+FP+TN+FN} \dots\dots\dots \text{equation (1)}$$

- b) **The false negative rate (FNR)** is defined as the ratio of false negative samples to total positive samples. In attack detection, the FNR is also called the missed alarm rate.

c) $FNR = \frac{FN}{TP+FN} \dots\dots\dots \text{equation (2)}$

- d) **The false positive rate (FPR)** is defined as the ratio of false positive samples to predicted positive samples. In attack detection, the FPR is also called the false alarm rate, and it is calculated as follows:

$$FPR = \frac{FP}{TP+FP} \dots\dots\dots \text{equation (3)}$$

Feature Importance.

This is a class of techniques for assigning scores to input features to a predictive model that indicates the relative importance of each feature when making a prediction.

The selection of variables in each model uses a criterion of the importance of information - Gini importance.

In each iteration, only the variables that accumulate more information are chosen to compose the model to be tested against the predefined threshold.

Gini Importance also known as Mean Decrease in Impurity (MDI) calculates each feature importance as the sum over the number of splits (across all trees) that include the feature, proportionally to the number of samples it splits.

Process of calculating feature importance

At each split in each tree, the improvement in the split-criterion is the importance measure attributed to the splitting variable, and is accumulated over all the trees in the forest separately for each variable.

After reading the data set, the key columns from the dataset are selected, the feature importance calculated and the best features selected basing on the feature importance calculated by the “RandomForestClassifier”.

After selecting best features, the model is fitted and the predictions are done. This gives the result metrics in terms of precision, recall and F1-score. The results of the classification are plotted in the graphs using the “matplotlib” and the saved as images inside the results folder inside the pycharm projects.

The python functions which reads dataset, selects features, checks threshold value by each class of attack, making predictions, plotting of confusion matrix for evaluation of the model are all called in the main function “def main():” which is automatically called by the “run.sh” batch file with key arguments.

Discussion

Technique Evaluation parameters

In this section we describe different ways we evaluated our technique. Since our detection techniques are based on classification, we adopted machine learning classification evaluation metrics. These metrics include; confusion matrix, Precision, Recall, and F-score measures.

Confusion Matrix

A confusion matrix is a table used in describing the performance of a classifier (or a classification model) with respect to test dataset. It is an $n \times n$ matrix that contains actual and predicted classes. Table 3 shows an example of a 2×2 confusion matrix. The rows indicating the actual classes and the columns indicating the classifier or model predicted classes.

		Predicted	
		Negative	Positive
Actual	Negative	A	B
	Positive	C	D

Table 2: A 2x2 Confusion Matrix

From Table 2, we considered the following terminologies:

- True Positives (TP):** This is a number of correct positive predictions. From Table 3 is shown by letter D. This means the model predicted POSITIVE that is actually a POSITIVE.
- True Negatives (TN):** This is a number of correct negative predictions. From Table 3 is shown by letter A. This means the model predicted NEGATIVE that is actually a NEGATIVE.
- False Positives (FP):** This is a number of incorrect positive predictions. Shown by letter B from Table 3. This means the model predicted a POSITIVE that is actually a NEGATIVE.
- False Negatives (FN):** This is a number of incorrect negative predictions. Shown by letter C from Table 3. This means the model predicted a NEGATIVE that is actually a positive.
- Accuracy:** This is a measure of how often is the classifier correct. It is calculated as;

$$Accuracy = \frac{TP + TN}{TP + FP + FN + TN} = \frac{TP + TN}{TotalSamples}$$

Precision, Recall and F-Score Metrics

- Recall/True Positive Rate (TPR)/ Sensitivity:** When it's actually POSITIVE, how often does the model predict POSITIVE? This is given by;

$$Recall = \frac{TP}{TP + FN} = \frac{TP}{ActualPositive}$$

- Precision:** When model predicts POSITIVE, how often is it correct? This is given by;

$$Precision = \frac{TP}{TP + FP} = \frac{TP}{PredictedPositive}$$

- F-score:** This is a harmonic mean of precision and recall. This is given by;

$$F-Score = \frac{2 * Precision * Recall}{Precision + Recall}$$

False Positive and False Negative Rates

As another way of evaluating our classification models, we calculate the False Positive Rate (FPR) and False Negative Rates (FNR) for each malware class and the benign class.

False Positive Rate (FPR) or fall-out: When it's actually NEGATIVE, how often does it predict POSITIVE? This is given by;

$$FPR = \frac{FP}{FP + TN} = \frac{FP}{ActualNegative}$$

False Negative Rate (FNR): When it's actually POSITIVE, how often does it predict NEGATIVE? This is given by;

$$FNR = \frac{FN}{TP + FN} = 1 - TPR$$

VIII. EVALUATION OF RESULTS

A. Confusion Matrix Results -

A Confusion matrix is an $N \times N$ matrix used for evaluating the performance of a classification model, where N is the number of target classes. The matrix compares the actual target values with those predicted by the machine learning model. This gives us a holistic view of how well our classification model is performing and what kinds of errors it is making.

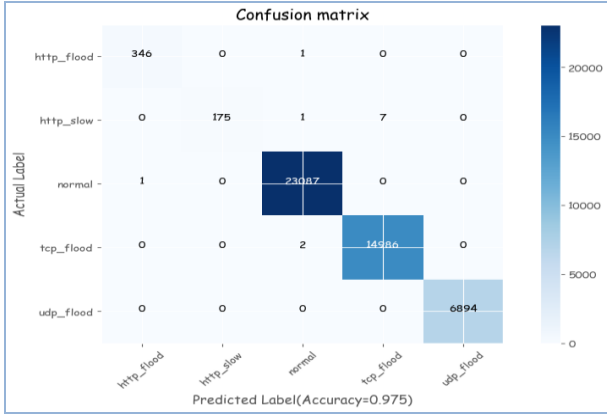


Fig 3. Confusion Matrix

The evaluation of the model's performance: An accuracy score. This score measures how many labels the model got right out of the total number of predictions. The accuracy score was calculated using the Scikit-Learn using the true labels from the test set and the predicted labels for the test set as shown in the formula below.

After running the random forest classifier model, the overall average score was 0.975(0.975 * 100 = 97.5%) and this seems pretty impressive and it gives an insight of how the model performed.

Precision, Recall and F1score metric results

From the confusion matrix in figure shown above, the precision, recall and F1score metrics are calculated using the formulas below for each class of intrusion attacks.

$$\text{Precision} = \frac{TP}{FP+TP} = \frac{TP}{\text{False Positive}}$$

$$\text{Recall} = \frac{TP}{FP+FN} = \frac{TP}{\text{Actual Positive}}$$

$$\text{F1-score} = \frac{2 * \text{Precision} * \text{Recall}}{\text{Precision} + \text{Recall}}$$

After running the model, the precision, recall and f1_score were plotted in the figure shown below.

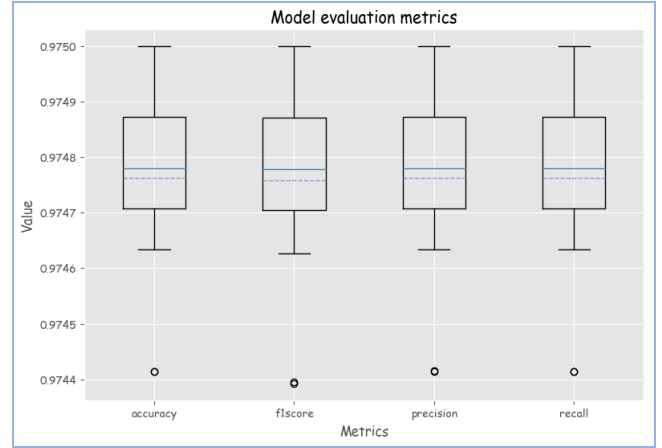


Fig4. Model evaluation metrics

IX. CONCLUSION

The increased complexity in the DDoS attack invoking practices, techniques, methodologies and easy accessibility of interrelated tools over the internet for the detection and mitigation of DDoS attack has become very difficult.

This study has come to a conclusion that an algorithm with self-update parameter calibration can detect DDoS attacks in an efficient, timely manner and specify the exact attack.

Evaluation of performance of the machine learning model for intrusion detection was done by computing the metric average scores.

The results indicate a higher average score for accuracy (0.975 * 100 = 97.5%) and precision (0.993 * 100 = 99.3%), which is a good indicator that the technique improved the intrusion detection for DDoS in communication networks based on its behaviour.

The experimental results show that the proposed DDoS attack detection method based on machine learning has a good detection rate for the current popular DDoS attack.

The Random forest classifier is a better choice in case of DDoS detection and the accuracy achieved in the trials is over and above 96% over a real-life dataset.

X. FUTURE WORK

In future, the algorithm can be tested and its performance evaluated on other DDoS attacks that were not covered in this research and also try and shorten the time and number of iterations taken for detection and increase the overall performance of the model.



XI. REFERENCES

- [1] Vishwakarma S. U., Soni P. D., (2015). Cloud Mirroring: A Technique of Data Recovery, vol. 5, no. 2, (pp. 739–742).
- [2] Jaydip Sen, (2011). A Robust Mechanism For Defending Distributed Denial Of Service Attacks On Web Servers. *Int. J. Netw. Secur. Its Appl.*, vol. 3, no. 2, (pp. 162–179).
- [3] Liu H, Lang B., (2019). Applied sciences Machine Learning and Deep Learning Methods for Intrusion Detection Systems: A Survey.
- [4] Zargar S. T, Joshi J, Tipper D., and Member S., (2013). A Survey of Defense Mechanisms Against Distributed Denial of Service (DDoS) Flooding, (pp. 1–24).
- [5] Cao Y, Gao Y, and Tan R, (2018) Understanding Internet DDoS Mitigation from Academic and Industrial Perspectives, vol. 6, (pp. 66641–66648).
- [6] Ustebay S. M, Aydin A., and Atmaca T., (2018). Intrusion Detection with Comparative Analysis of Supervised Learning Techniques and Fisher Score Feature Selection Algorithm Intrusion Detection with Comparative Analysis of Supervised Learning Techniques and Fisher Score Feature Selection Algorithm, no. 37102.
- [7] Sales F, (2019). Smart Detection: An Online Approach for DoS / DDoS Attack Detection Using Machine Learning.
- [8] Cretu-ciocarlie G. F, Stavrou A, and Locasto M. E., (2010). Adaptive Anomaly Detection via Self-Calibration and Dynamic Updating.
- [9] Reddy B. R., Babu A. S., (2019). Prediction of DDoS Attacks using Machine Learning and Deep Learning Algorithms, no. 4, (pp. 4860–4867).
- [10] Gao Y., Feng Y., Kawamoto J and K. Sakurai, (2016) A Machine Learning Based Approach for Detecting DRDoS Attacks and Its Performance Evaluation.
- [11] Ammar A, (2015). A Decision Tree Classifier for Intrusion Detection Priority Tagging.
- [12] Zargar S. T, Joshi J, and Tipper D., (2013). A survey of defense mechanisms against distributed denial of service (DDoS) flooding attacks, *IEEE Commun. Surv. Tutorials*, vol. 15, no. 4, (pp. 2046–2069).
- [13] Wang B, Zheng Y, Lou W, and Hou Y. T, (2015). DDoS attack protection in the era of cloud computing and Software-Defined Networking, *Comput. Networks*, vol. 81, pp. 308–319.
- [14] Kotenko I, and Ulanov A., (2014). Agent-Based Simulation of Ddos Attacks and Defense Mechanisms, *Int. J. Comput.*, (pp. 113–123).
- [15] Pei J, Chen Y, and Ji W, (2019). A DDoS Attack Detection Method Based on Machine Learning, *J. Phys. Conf. Ser.*, vol. 1237, no. 3.
- [16] Maglaras L. A. and Jiang J, (2014). Intrusion detection in SCADA systems using machine learning techniques, *Proc. 2014 Sci. Inf. Conf. SAI 2014*, no. April, (pp. 626–631).
- [17] Agrawal S., Singh Rajput R., (2017). Denial of Services Attack Detection using Random Forest Classifier with Information Gain, *Int. J. Eng. Dev. Res.*, vol. 5, no. 3, (pp. 929–938).
- [18] Sambangi S. and Gondi L., (2020). A Machine Learning Approach for DDoS (Distributed Denial of Service) Attack Detection Using Multiple Linear Regression,” *Proceedings*, vol. 63, no. 1, (pp. 51).
- [19] Liu Q. and Wu Y., (2012). Encyclopedia of the Sciences of Learning, *Encycl. Sci. Learn.*
- [20] Kim M., (2019), “Supervised learning-based DDoS attacks detection: Tuning hyperparameters,” *ETRI J.*, vol. 41, no. 5, (pp. 560–573).

IJEAST

INTERNATIONAL JOURNAL
OF ENGINEERING APPLIED SCIENCE
AND TECHNOLOGY

ABOUT IJEAST

International Journal of Engineering Applied Science and Technology (IJEAST) is a peer-reviewed, open access journal that publishes high-quality research papers in the field of Engineering, Applied Science and Technology.

IJEAST aims to provide a platform for researchers, academicians, and professionals to share their innovative ideas, research findings, and practical experiences with the global scientific community.

FOCUS AREAS

- Engineering
- Applied Science
- Technology
- Innovation & Development
- Interdisciplinary Studies



PEER REVIEWED

All submissions are rigorously peer reviewed to ensure quality.



OPEN ACCESS

Free and unrestricted access to research for all.



GLOBAL REACH

Connecting researchers and professionals worldwide.



TIMELY PUBLICATION

We ensure a swift and efficient publication process.



For more information, visit our website

www.ijeast.com



INTERNATIONAL JOURNAL
OF ENGINEERING APPLIED SCIENCE
AND TECHNOLOGY

✉ editor@ijeast.com

🌐 www.ijeast.com

📍 India



2455-2143