

Defending Cloud Web Applications using Machine Learning-Driven Triple Validation of IP Reputation by Integrating Security Operation Center

NW Chanka Lasantha ^α, Ruwan Abeysekara ^σ & MWP Maduranga ^ρ

Abstract- This paper will present an innovative system method of IPR (IP Address Reputation) validation with the assistance of clause of (ML) machine learning for discovering malicious IPs, while also viewing the importance of security of installed applications on AWS (Amazon Web Services) servers. The ML, SANS and AbuseDB datasets that were verified are being integrated through the Wazuh Security Operation Centre (SOC) stage to consume issues at the log ingesting IP address-related level. Having integrated extraction of IPs Wazuh agents, the output does match MITRE ATT&CK framework-filtered IP address from the Wazuh SOC. These algorithms and models based on natural language processing will flag suspicious patterns across IPs through the process of machine learning and prevent the event of a cyberattack at the time. This integration not only boosts cybersecurity information through a single point source of distribution, but it also provides security finds and other resources to prove and maintain awareness against malicious IPs. The final solution includes using the maximum amounts of bad IPs blocking in the 'IP-List' of AWS WAF and, if they are added to the Blacklist automatically, checking them through an automatic ML-based signature validation process.

Keywords: SOC, ML driven IP reputation validation, AWS WAF auto defense, ML powered extended validation, MITRE ATT & CK framework-filtered IP address.

I. INTRODUCTION

The implication of the discrete web application defences could be a source of great difficulties. The most vital element in defence from bots is the reduction of machine-learning-based bot traffic, which also has a leading role in the protection of real IP addresses. Proposing that strategy, early detection of IP defence compromising is possible, and evaluated only neutral ML bots' versions are used for analytical analysis. In this area, the most prominent modern element is designated on-demand IP addresses. The security featuring machine-generated digital keys provides them with a sign of comprehensive protection. [1] ML technology is the core functionality of this process which is also underpinned by advanced

algorithms. To supplement the IP Reputation Monitoring, Smarter measures to find out if web applications are blocking other cybersecurity measures are being implemented. [2] To be specific, this is a system made of many complex layers. To do this very well it uses neural networks, machine learning, and all the skills it can gain from the large language models (LLM). These models are very effective, and their operations involve the extraction of actionable data from the database and records, with IP details. The robustness of the system is meant to be increased this way by integrating complete security frameworks and databases established controls, as well as MITRE ATT&CK framework [3], Hence there is an array of approaches that work together to intercept and block any explanations of malicious IP addresses which ultimately solidifies the effectiveness of the automated defence mechanisms using the Wazuh SOC Logs, which is a highly advanced platform that processes and stores logs. [4] With this merging, signatures based on the techniques of machine learning can be used on an ad hoc basis to guarantee a real-time production of ML-driven signatures. The proactive cognitive system exhibits the connection of machine learning and cybersecurity, it offers a clever and dynamic solution to the fast-moving landscape of security needs in web applications.

II. BACKGROUND

a) Importance of IP Reputation based on MITRE

IPR serves well the indication purpose of when a network is to be accounted for a prime target of hackers. A problem arises through which WAF is thrown out of its comfort zone and it should deal with network protection applications. The reputation is the rating which is the most important for programs of this kind. It serves as the basis for decision-making about the entry and removal of the IP traffic. [5] The most dangerous consequences of non-IPR incidents can be divided into five groups spamming, bot activities with harmful intent, DoS attacks, injection attacks, and occasional use of this source for botnet operations. In the application of IPR, it is not merely a tool for adding to the known risks but also a motivator for exploring the possibilities. This is to say that is the underlying cause of cutting into network and services average. In the IP carrying a bad reputation security attacks are regarded and it is well

Author ^α: Faculty of Graduate Studies, IIC University of Technology, Phnom Penh 121206, Cambodia.

e-mail: chanaka.lasantha@gmail.com

Author ^σ: Faculty of Graduate Studies, IIC University of Technology, Phnom Penh 121206, Cambodia. e-mail: ruwan@iic.edu.kh

Author ^ρ: Department of Computer Engineering, General Sir John Kotelawala Defence University, Sri Lanka.

e-mail: pasanwellalage@kdu.ac.lk

known to signal a potential for bad activity. Consequently, great care should be observed in the elimination of such IPs. [6]

Also, the latter is associated with the most accurate IPR since self-introduction can be monitored and scored using interactive personal relationship features. As they assemble data for courts to use in investigations and to conduct IP tracking, they also prevent the organization from being tampered with by malicious activities. This approach's fundamental aspect represents the proactive defence that is the key concept of the MITRE ATT&CK framework and the high point is it emphasizes the importance of data protection would help organizations be conscious of the key sources of threats by fixing their attention on malicious controls and data system and that would consequently lead to efficient organizational business continuity with interruptions. [7]

b) Challenging on Traditional IPR

The IPR Validation, the main traditional method is mainly to search IP addresses in directories and blacklists which increases behavioural analysis. Nevertheless, this method of data collection omits most of the pre-validation procedures that are prerequisites for a stable dataset meant to be useful for training ML models. Selection lists are mostly loaded through honeypots, spam traps, and regular event logs. The scores look at an IP-address reputation for certain behaviour. [8] Also, The IPR is decreased by this system 'reputation sink,' where IP reputations become not relevant over time without the continuous, real-time validation of the multi-layer approach. [9] This asymmetry led to the impossibility of coping with cyber threats just merely by the databases, which necessitates a constant update process of the databases is essential. Tribulation of such an approach leads to many false positives and negatives consequently making the traffic management inadequate. The problem is pronounced by the deficiency of ML algorithms' accuracy and the application of the metadata that is either out of date or inaccurate concerning the IP addresses. [10]

III. EXISTING IPR ARCHITECTURES

a) Mitre Freamwork baed for IP Attacks Detections

The pre-attack patterns determined in enterprise knowledge bases also add a lot of value in terms of tracking adversary tactics, techniques, and procedures to ensure that an incident can be well responded to, and the attack repelled. Uncaught and disruptive activities by availing themselves of what the adversaries use to penetrate competitive networks must be brought to understanding and unveil an essential topic of the monitoring methods and ways to fix impacts by using the MITRE ATT&CK framework. [11] Given this architecture, it is a comprehensive and quick-access

knowledge by providing exclusive information on the present-time procedures of the enemies against real-life scenarios. This assists in building, in the private sector the government, and the cybersecurity community strong programs to monitor the threats. [12]

b) Prevention Technologies for IP Address-based Attacks

One of the hardest things about cybersecurity is tracking and stopping cyber-attacks at the IP address levels which was solved by one of researched blacklists and tools such as AIPRA, which combined ML with geolocation data to figure out what's not relevant for regions and countries in usual working time range of humans. But problems such as false positives, and maintenance of the fast-changing nature of its enemy continue to an accurate validation process. ML can help AIPRA systems immensely while cutting-edge algorithms and effective data processing, combined with the optimization of models which increase accuracy while reducing false positives, keep it up to speed on new threats. [13] This strengthens cybersecurity defences on IPR, while the security of the LAN The MAC and IP addresses, computer names, IP conflicts and MAC mismatches are most important to reduce attacks from bad IPR vectors in securing network traffic and assets and spoofing risk over digital infrastructure. Such that, the spoofers forge these identifiers to masquerade as IPR validation systems. [14]

c) Traditional Bot Traffic Tracking Techniques

The applications of Residential IP Proxy (RESIP) facilities are becoming more and more popular cases of web scraping and other criminal actions such as relocating behind the reserves of residential IPs where the detection is prevented. Two additional datasets indicate the functioning of RESIP where its figures are highlighted only with the four providers but not with differences concerning them. [15] They suggested an operational scheme that can automatically compare accounts with shared characteristics. Besides, overall, five campuses undertook vulnerable RESIPs' investigation, showing attacked hosts and unlawful acts. [16] This study can shed light on and address the security chances that this growing sector is attributed to. RESIPs, which are a new grey-area business, provide a shield from scrutiny by using other people's computers in their homes to complain about illegallness and recruitment ways. [17] Also, it proposes RETRO detection, a technique that captures the sequences of flows using a compromised device, raising the operational opacity of these services. While it optimizes a server-side detection method for RESIP connections, dropping false negative outcomes that result from mobile proxies. [18]

d) *AI Models for IPR Detection Capabilities*

IP Starting with the fundamentals of IP protocol to the daily activities on the internet such as surfing the web and emailing, Internet resources are indispensable, which urges security professionals to use IP addresses for risk assessment. This work makes use of cross-protocol telemetry on a large scale to classify malicious IPs and make ML interpretability because of which this approach is more effective. [19] The results reflect that there is zero error in identifying malicious IPs. To mitigate against the rising cyber dangers, The duo proposed a mixture of different attributes which involved Dynamic Malware Analysis, Cyber Threat Intelligence, ML and Data Forensics. [20] This technique comes with a reputation of IP, groups 'zero days', and closely as well as automatically analyzes damage, degree of risk, and impact. This model takes while factoring in the conventional network and geo-contextual information, thus enhancing threat assessment and enabling the detection of unlawful behaviour, especially in cyberspace that has HTML encoding. [21]

e) *NLP for Enhanced Threat Detection Using ML*

The growing trend of IoT-devices interconnectedness has resulted in an uptick in intrusions. IDS or IPS systems are a type of security solution that monitors and detects system violations. [22] Nonetheless, a holistic synchronisation in new developments and model limitations means that a new security framework is required. [23] On the part of this survey AI techniques such as machine learning and deep learning seem as most relevant solution with hybrid design efficient intrusion detection/prevention emphasizing. It considers their viability, setbacks and real-time issues. securing IoT, ML and big data analytics have profound effects on it. [24] This is where they come in. This investigates IoT vulnerabilities, uses ML for cyber-vulnerability assessment, and analyzes ML-based intrusion detection solutions. It provides an example of a real-world testbed which is used for the design of IDS, demonstrating that Machine Learning is capable of intrusion detection in computer networks. However, this study the literature on the topic of anomaly-based intrusion detection systems driven by ML/DL, pushing the boundaries to unleash the full potential of ML-based systems, examining open issues efficiency. [25]

f) *BlackEye IPR Framework*

Algorithms Blacklisting malicious IP addresses is an essential tool for IT systems' protection. The decision-making is based on looking at packet traffic data and the behavioural history of users. Still, the holding of domain experts for blacklisting is on but ML is on the way and just awakes to maturity. This is solved by making the Black Eye framework based on which the different ML methods are used accordingly to achieve superior results. The analysis shows that the multistage

method, which is achieved by data cleansing and classification with logistic regression or random forest, leads to the best results. Real-world data evidenced a near-90% less incorrect blacklisting compared to the expert performance. By the same token, our model accelerates the time-to-blacklist, significantly cutting the lifetime of malicious IP addresses on average by 27 days. It can be considered a breakthrough in the process of protecting the IT system concerning blacklisting and redesigning the efficiency and accuracy of the system security. [26]

IV. CAPABILITIES AND LIMITATIONS

a) *MITRE ATT&CK Framework Boundaries*

The MITRE ATT&CK (MITRE Adversarial Tactics, Techniques & Common Knowledge) framework which motivates the cybersecurity industry nevertheless has multiple challenges like the lack of clarity, incomprehensive comprehensiveness and dynamics of rapid knowledge that may dismantle especially new or inexperienced security personnel leading to the framework's apparent underutilization. The configuration's information studies involve mass data analysis, but the demand is higher than normal, automation is lacking in most organizations, and the framework will cause more burnout on SOC than it can handle. Defence against the Dark Arts is also afflicted by standardization issues because some sub-techniques are either too niche or incomplete existing problems that comprehension and application are difficult. Similarly, by its charter to capture only documented cases, it can sweep under the radar of inaugural threats and their occult threats thus, limiting its efficacy in preempting threats. [27]

b) *IPR Validation and Prevention Using ML*

Malicious IPs will not be allowed to access the system hence the IT security will remain under. BlackEye uses ML and after researching it has been proved that a two-staged solution with some preprocessing to be followed by either logistic regression (NR) or Random Forest (RF) is effective to a ratio of only 15% blacklisting false alarms. Furthermore, the Tower uses Ridger heightening to get a 5% higher precision. BlackEye, by integrating and quickly iterate through ML on heterogeneous logs. With the help of this neural ML method, accuracy would be improved and the time to blacklist would substantially be reduced. More upstream work is accomplished through the application of deep learning in the identification of risks. [28]

c) *IPR Validation of Public Databases using ML*

The exact rate at which cyber-attacks are rising is a result of more individuals, groups, and corporations connecting online. Old-fashioned blacklists work, but they could be improved by shaking off two of the broken records. unverified data and stale data. First, these



issues were solved by the AIPRA (Automated IP Reputation Analyzer Tool). This action is in the form of verification by comparing the domains or IP addresses indicated whether they are on the list of blacklisted that is commonly used in several indexes. The AIPRA first evaluates if there might be any malicious activities at these URL addresses and comes up with a weighted probability that indicates how much it is possible. Also, a geolocation-based artificial intelligence concept was made a component of AIPRA since this way the system could be trained to recognize a wider spectrum of threats. When the Report produced this result, it had not yet been identified as being on any list to the public. [29]

d) *MAC Address Spoofing LAN Attack Validation*

The security problems of Local Area Networks (LANs) are being constantly dynamically generated. Nevertheless, it is a meticulous method of decoding how the hacker achieves such an objective by spoofing both his MAC and IP address to lift LAN internet accounts from unauthorized users, which seems very hard for the account administrator to monitor. Further, ensure that such protections are implemented, and the brute root of any untoward act must be cut off, this can secure oneself an Advanced IP Scanner or MAC Address Changer and prevent outside attacks but also an IP address itself a gateway lest it sneak it being the facilitator. Two main factors that can force crooks to get involved in such frauds are the financial strain and the constant need to make fast money. When at the hacker's stage, the administrator is moving to the progress of the attack by exploiting the user and the ISP's routers' default passwords. [30]

e) *Detecting Malicious IP by Cross-Protocol Analysis.*

The system of reputation trust takes real-time data into ML machine processes to become a way of enhancing security through the cloud. Such a system will function around powerful algorithms which would recognize and destroy malware websites. In contrast, the system's defence is through source code encryption and obfuscation, hence operating using the same common IP reputation key across providers and in a consistent manner to protect the system from being a target. New components and augmentation methods of data have made pre-processing features useful and groups to choose a threat feature. The system, which ensures false positives and negatives also, employs error analysis and explainability to get enhanced precision. It demonstrates network IP using port 53, so the DNS traffic goal is accompanied by a figure of improvement for the incremental model that is tailored for enhanced flexibility and efficiency. Despite this, the main issues associated with the small data size including label as quality or accuracy of the labels on the big models, are being considered. Firstly, we can understand deeply how the reputation alongside the

rating lines of the users mutually binds with the tendency of response modification over time and, as a result. [31]

f) *Detection of malicious traffic by learning IP Reputation*

Based on the use of adaptable and modular technology in ML, it is a lightweight solution that works in such a way that the old approaches are not completely replaced. The approach of our study in line with the present methods aiming to control the list of IP addresses of spam mail and the flow of the campus network is different because of the application of a higher method. Sites tagged as dangerous and those tagged as malicious are scorned but the ones whose intention is unknown continue to operate freely and the assaults are not thwarted. Yet, such a feature is practical with no warning limit to the common hosting, thus the more this favourable gesture is done the better outcomes it's likely to achieve. Whilst the intelligence services exert their efforts to attain the upper hand, it is the adversaries that show the power of adaptation and are changing their tactics that the real problem is. Most importantly, the offence has the supreme edge for the unforeseeable. [32]

V. MAIN METHODOLOGY OF POPORSED SOLUTION

a) Components Of Methodology

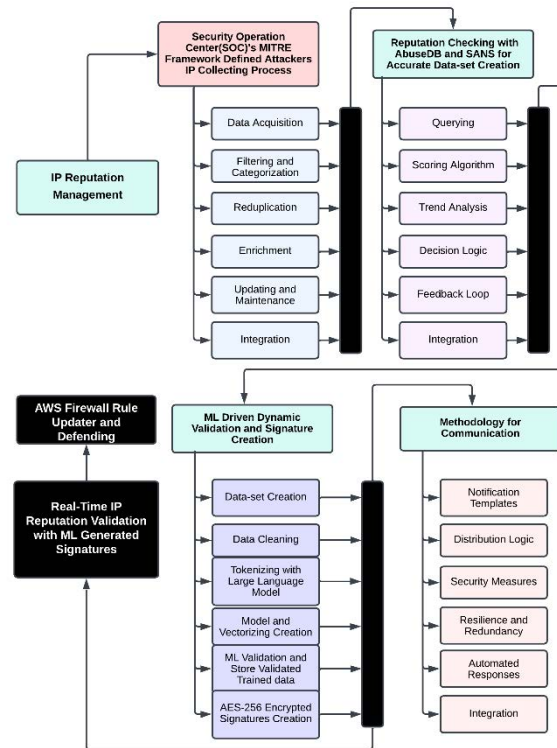


Figure 1: Methodology

b) SOC Attackable IP Collection and Analysis

Figure 1 shows the data collection phase IP address scope that is consecutively given by MITRE ATT&CK logs with REST API connection from multiple units in the SOC system. The logs contain the information of the cloud system, flow as well as the server. An early part of the first move was the classification of IPs as trusted, suspect, or unknown habitats that will eventually set up data sets for matching them in a confirmation process. Single IP deduplication method and ultimately the proliferation of research. The best psychological assessment can be due to geolocation IP information, individual behaviour history, and diagnostic criteria which are stated in these modules. The databases become erroneous only if they are not appropriately revised and have the old data changed at certain fixed intervals. Interfacing with the machine consists of sketching the assembly system from the IPs, therefore, the real-time generation of the IP enables the system to produce these just risen or risked IPs. Followingly, this stage will determine WAF logs metadata to find IP addresses and domain names while focusing on extracting specific required features to make sure algorithms can work properly, then; anomaly detection is made on domain name feature Plus IP address (IPR) to make comprehensive attention to areas of anomaly signs. Besides, it does the essential

functions of the concealed dangers that are not always obvious and integration and moving on to security states.

c) IP Reputation Checking and Primary Dataset Creation

This procedure starts with AbuseIPDB API being connected to the recently revised IP reputation data [33] and then calling SANS API that is needed for a second recheck for the dedicated scoring algorithm of the bad metadata such as final score being blacklisted, reports number and the reason for the reports to be given. [34] This very algorithm moreover actually shows its face and shows the way that every IP address can have a certain weight by this score-based method and severity levels by unstructured raw data to the structured dataset.

d) ML Driven Signature Creation and IP Validation

In the beginning, the ML model is classified as being a signature to the validated datasets with filtration and tokenization being selected and then random forest was chosen as the ML model. The second part will be the production of vectors that will be generated after the training of a model has been completed. Ultimately the processing is done, individuals will be a bit nervous about their data so when the file is made it is going to be saved into a folder that is safe to store it in and this data is verified to be true. The next step in the final

process is an AES-256 Encryption signed with an ML-driven signatures generation method. [35] Such volatility erodes the monetary authority's ability to set policies and makes the currency regime less stable. In addition, it applies the appropriate signature detection method by querying the MySQL database.

e) Automatic AWS Firewall Rule Updater for Defending

The system creates alerts varying in level of danger as well as differentiates communications based on a user's role in the organization. It also includes an automated response protocol that can quickly update settings such as firewall IP- Blacklist rules if not already exist and add them to the Web Access Control List to block rapidly when it is validated as an attacker IP address during the validation process within a certain period without affecting legal sites for a grouped period. As a result, real attackers will only be blocked for a while which will prevent damage to the target system that is vulnerable to attack.

VI. ALGORITHM USED IN THE PROPOSED SOLUTION

The computer algorithm technology to the existing security algorithms and setups, technology which can execute the parameters but there are chances of involving errors and lower setup time that increases setups. AWS Secrets Manager is being

utilized as a credential secret retrieval solution verifying and ensuring that the company is at the required security level and the predefined security practices are being followed. Immediately after that, they created a customized abuseDB, SANS, and the model-learning mechanism for the IP Reputation analysis process. It has two noticeable points as to why ML-based technology is a better choice compared to rule-based systems in the MITRE ATT&CK logs the first pro is the ability to understand the context and find a pattern in the cases in which rule-based systems were not able to do it well. This way it brings in dynamic signatures and spots bad IP actors quickly and easily.

Furthermore, this process works out ML's limitations and will make lives for intelligent machines that are alive and evolve. On the other side, this layer operates as a second line of security systems which are used to identify threats before they can even take place. As for the NPL detection strategy, the automated creation and regular check-up of Abuse IPDB API and SANS API query results accompanied by the third layer of the final validation method could also play an important role in boosting NPL detection quality and is important. Maintenance is optimized and overload that leads to a system failure is tried to mitigate in this way in whole system performance.

VII. EXPERIMENTAL PROPOSED SOLUTION

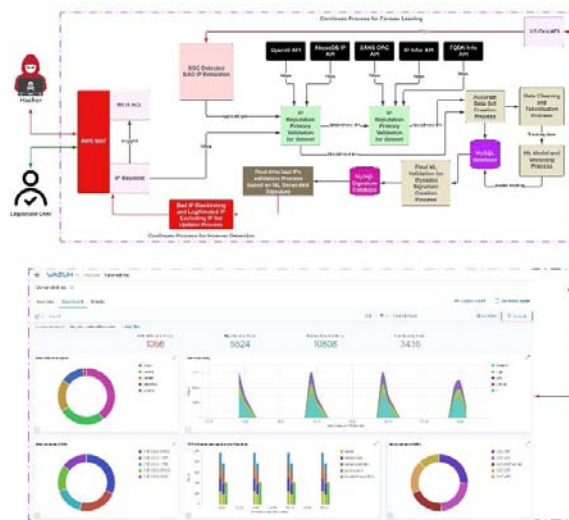


Figure 2: Experimental Setup

Figure 2: provided appears to be a detailed schematic of a security system for cloud web applications, integrating machine learning with IP reputation validation to enhance cybersecurity measures. At the top, icons differentiate between a hacker and a legitimate user, indicating the types of traffic that the system must differentiate between. The AWS WAF serves as the initial barrier, applying rules

through a WEB ACL to regulate incoming traffic. The system includes a process for real-time bad IP blacklisting, which employs various APIs such as Open AI, Abuse IPDB, SANS ORG, IPInfo, and FQDN Info to gather intelligence on IP reputations. This intelligence is then processed through machine learning validation, which dynamically updates the IP signature database. The process of data cleaning, tokenization, vectorizing,

and machine learning modelling is depicted, suggesting the preparation and utilization of data to train the system to continuously identify and respond to threats. This is supported by the ongoing "Forever Detection" process, indicating an adaptive security posture. Figure 2: there's a security dashboard, such as a tool such as Wazuh SOC, displaying various security metrics. This includes

the number of detected vulnerabilities, the severity of alerts, and the distribution of these alerts across different security agents. Graphs show the trend of security events over a certain period of days regarding IP data, while additional charts detail the most common vulnerabilities identified by their CVE identifiers.

VIII. OUTCOME OF THE SOLUTION

a) ML Validated Data with Predictions

Table 1: Sample section of ML-validated final test data

Aws Acco-Unt Id	Re-Gi-On	Ip Address	Total Repo-Rts	Abuse Confide-Nce Score	Is Whit-Eliste-D	Attack Proba-Bility
xx	xx	121.162.2 10.148	379	100	0	attacker
xx	xx	124.107.3 7.84	0	0	1	not_attacker

Through employing the IP verification model, the data obtained from the resolve is illuminated in Table 1. It stands out from the other algorithms by the fact that it uses report data to create reputation scores, where the number of reports, how recent, and the confidence that they are abuse reports are all considered. To combat

this, we developed whitelists and blacklist IPs to shield ourselves from the high-risk IPs. These IPs also displayed their cases with malicious activities, for instance, they were signed to port scanning or brute force attack.

b) ML Validated IP Abuse Score

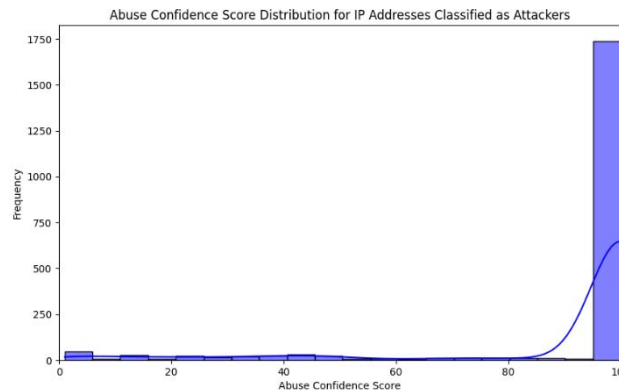


Figure 3: Abuse score Illustration

Figure 3 shows that the DB API models were ratified to be IP Abused, Plotting the virtual curve of Abuse Confidence Score that began with 80 and ended with 100 proves that our reporting system is intensified with the passing days. Confidence of reports increases especially from credible sources. This rises very possibly because the algorithm relates more highly to reported unusual IP addresses, severe admins taking greater weight. IPs with the high risk may end up with a loop of additional monitoring when the system fails to achieve a good level of attack-reducing mechanisms for these IPs. Organizations, with a cushion effect, may opt for such formulas that give higher results to the IP nearing the maximum with the scores being concentrated at the top end, or 100, signifying a strong consensus about the risk of a given IP. Therefore, it is

expected that the frequency of scores at the upper end of the scale will rise sharply.

c) ML Definitions

Precision is calculated as:

$$\text{Precision} = \text{True Positives} / (\text{True Positives} + \text{False Positives})$$

Recall is calculated as: [36]

$$\text{Recall} = \text{True Positives} / (\text{True Positives} + \text{False Negatives})$$

The F1 Score is the harmonic mean of Precision and Recall and is calculated as: [37]

$$F1 = (2 * \text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall})$$

The False Positive Rate (FPR) is calculated as: [38]

$$\text{FPR} = \text{FP} / (\text{FN} + \text{FP})$$

d) Overall RF ML Predictions

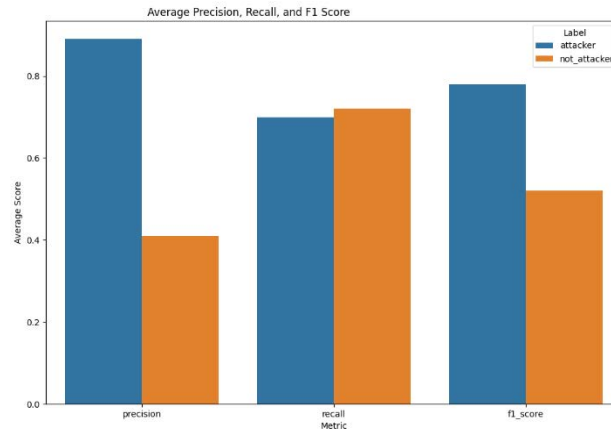


Figure 4: Overall ML Predictions

Figure 4 shows that precision stats measure the subtle of sensitivity or accuracy of positive predictions. To summarize, the model predicted a successful attack in which the actual attacker situation. For the group predicted "attacker", the precision is very high around 1 implying that the model is most of the time right when it predicts an attack. The accuracy for the term "not_attacker" is a bit lower, which signifies a high prediction accuracy for the same. Also, Recall Retrieval's mission is to discover all the cases that are of significance to all the points in the dataset. For the "attackers" class, the recall outperforms the precision, however, this is to keep the recall above 0.9 which demonstrates that the model can identify most of the actual attackers.

The "class" of the "not_attacker" recall is nearly the same as for the "attacker" class, indicating that the model is as good at detecting instances that are not attacks as those which are. Additionally, The F1-score signifies the harmonic arithmetic mean of precision and recall. It is just one measure that considers the accurateness and the pullback of a classifier and expresses these results into a single metric. If the classifier has a high sensitivity, it is more likely to avoid false positives. in other words, it is accurate. Overall, the model is not unbalanced as a high F1 score is observed for both "attacker" and "not_attacker", making the model execute much better than expected.

e) Overall RF Correlation Matrix Heatmap

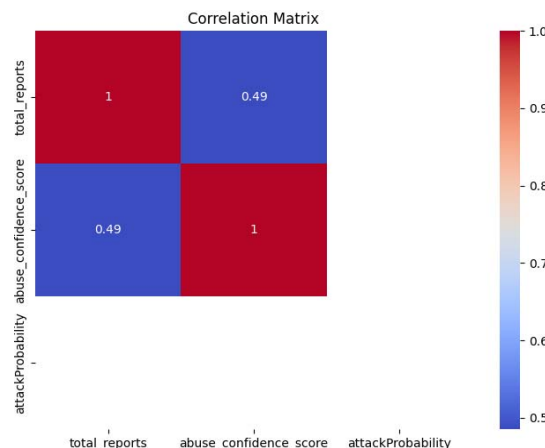


Figure 5: Correlation Matrix Heatmap

Figure 5 visualizes the correlation matrix using the heatmap plot. A Correlation Matrix is simply a table with Correlation Produced between different variables. Every row of the table visualizes the problem of how the two variables are connected. The range is from 1 to 1 for

its units. If two variables are strongly (near 1 or -1) related it is indicative of a high correlation between those two. When the correlation is close to 0 it indicates that there is a zero linear relationship.

Also, the diagonal represents the relationship between each variable and itself. The correlation of a variable with itself is always 1, which equals perfect correlation. The relationship between `abuse_confidence_score` and `total_reports` is about 0.51, meaning along the linear relationship, when one variable increases, the other variable also tends to increase, but to a lesser extent which would be a perfect linear relationship, and the correlation would be bigger, 1.

In addition, the heatmap uses colour intensity to represent the strength and direction of the correlation. The heatmap uses colour intensity to represent the

strength and direction of the correlation. The dark color would be used to represent a negative association nonetheless there are no negative values in the matrix. The depth of colour corresponds to that of the strength of the relationship, with the contrary being darker shades representing those of stronger relationships. On the colour bar, you see on the right the values of the correlation coefficients that stand for heatmap colours are given. The value of colour ranging from red to white and from white to blue shows that closer to 1 value is the red colour while lower the value is close to -1 value, which is the blue colour.

f) Overall RF Confusion Matrix

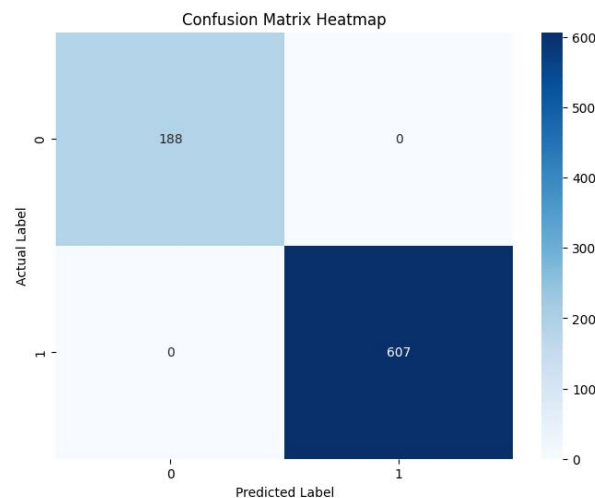


Figure 6: Confusion Matrix

Figure 6 shows the heatmap of the confusion matrix for a binary classification model. The labels on the y-axis are the real ones, and on the x-axis are the anticipated ones. The figure below highlights the fact that it has been determined that 188 true negative examples (U-L) and 607 true positives (U-R) have been correctly classified. Instead of false positives and false

negatives, as shown in the top-right and bottom-left cells of the matrix being zeros, means that there won't be any misclassifications. The size of the circles is relative to the term occurrence and the darker the tone, the more instances. It will be plausible to conclude that the model attained mean square error which is equal to zero on this data set.

g) RF Tree Visualization

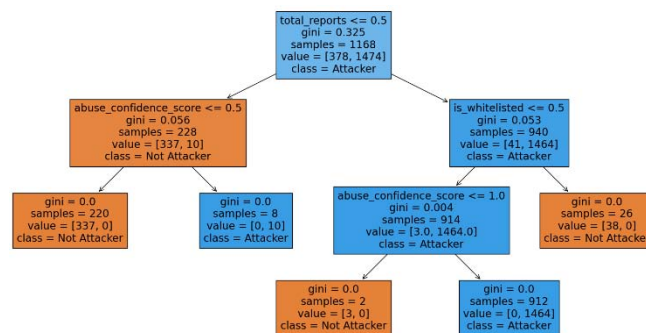


Figure 7: Tree Visualization

Figure 7 describes the decision tree as a representation related to classifying the entities into "Attacker" or "Not Attacker" ones, where we utilize

'total_reports' as a main measure. In case 'total_reports' equals 'abuse_confidence_score', 'abuse_confidence_score' is evaluated and "Not Attacker" will most probably

be assigned a score of 0.5 or less. For a high 'total_reports', 'is_whitelisted' is the main result maker while the confidences of the entities below 1.0 and non-whitelisted are mostly classified as "Attacker". However, what catches my attention is that the Gini index of numerous leaf nodes is equal to zero which shows an extremely confident model that is prone to misclassification in case of unbalance on the other

hand, the model can catch the exceptions as well as distinguishing between them properly, which is a reason for satisfactory results in leaves with not enough samples. In short, it uses total reports', 'abuse confidence score,' and 'is whitelisted' as its principal columns showing the situation of the classifier where they shape clear decision paths and strong class differentiation.

h) RF Cap Curve

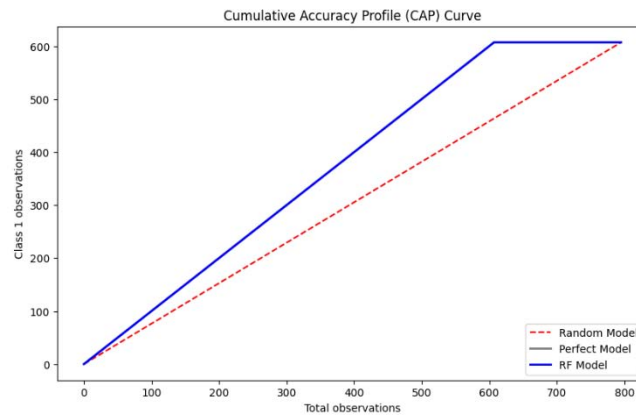


Figure 8: Cap Curve

Figure 8 shows the Cumulative Accuracy Profile (CAP) curve used for the Random Forests (RF) model classification evaluation. The dashed red line is a random model, which may only be able to fulfil similar goals as if there was no model at all. The solid blue line that is perfectly straight and aligned at the top part is the perfect model that gets total accuracy by correctly classifying all the instances of the class at once. In contrast, the RF Model's curve which is another blue solid line shows the model's performance which is

superior to random guessing as the curve curves towards the ideal model, demonstrating that it ranks instances of truth consequently better than random guessing. The Diagnostic Accuracy of the RF Model Relevants is being measured by the Area Under the Curve Calculation (AUC CAP) works very well, which has a value close to 1, and this one is far better than a random approach. The R^2 Model exhibits a pattern of improvement against unintended chance and is approaching the most desired alternative.

i) RF Distribution of Total Reports

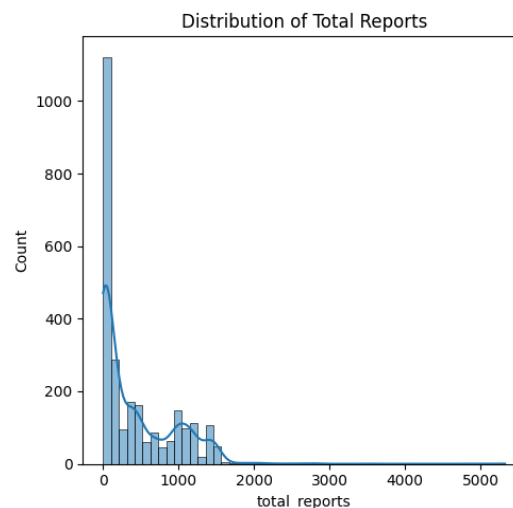


Figure 9: Distribution of Total Reports

Figure 9 shows the histogram together with a kernel density estimate (KDE) in which the distribution of the value "total_reports" is plotted. The y-axis represents values count or items are divided into bins. On the y-axis, data is presented in the form of how often these words are in the analyzed texts. The height of every bar you can see indicates the number of occurrences of the analyzed word in every bin range. The graph illustrates the right-tailed pattern to indicate over the x-axis mark, the greatest number of data exhibits low "total_reports"

j) AWS WAF IP-Deny List for Blocked IP Address

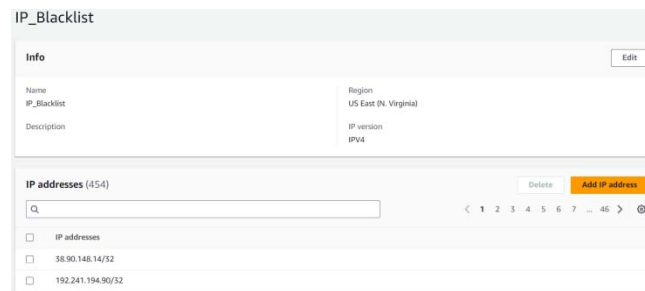


Figure 10: Blacklist over AWS WAF

Figure 10 shows that an automated IP was blacklisted by The IP Reputation validation system accurately minimizing False Positive IP blocking to allow legitimate services not been getting blocked by the AWS WAF in the corresponding IP List section. Also, this solution successfully blocks these kinds of bad attacks

count, and less and less as we move toward the end of the axis. The KDE line shows the distribution smoothed by connecting the points representing the peaks on the left and the end of the tails indicating some extreme cases with high report counts. This smoothed image indicates the presence of the skew right value with most reports around the centre and the tail on the top side and a single point or few isolated points on the bottom curve.

k) AWS WAF WAF Blocked and Allowed IP Percentage

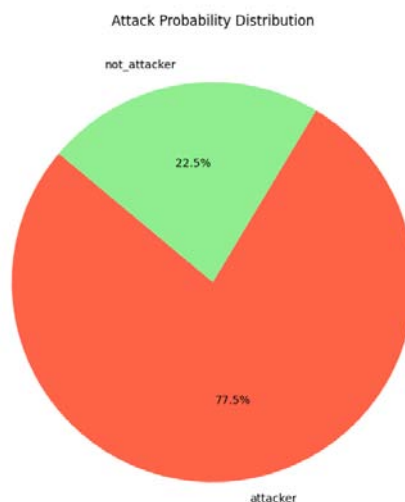


Figure 11: WAF Blocked and Allowed Illustration

Figure 11 shows both blocked and allowed IP address percentages based on a defined period in the WAF IP-Blacklist updated module based on ML prediction auto-generated pie charts as shown above.

IX. DISCUSSION

This is an enhanced solution that the showcase introduces, with strong application value in cybersecurity

by bad actors and automatically blacklines all relevant addresses based on machine learning-in effectively, those related to checking ML-Driven signatures verification process while the solution uses the IP-List section of AWS WAF to automatically blacklist attacks from bad IP addresses.

operations. It can greatly simplify many other aspects of the trust verification process, one prerequisite for network security has to do with IPR. It connects to AbuseIPDB and OpenAI Analytics Engine. This entails checking a database of abuse reports, so the assessments are accurate and topical. This approach quickly selects, verifies, and classifies response data for analysis or combination with other systems. Using its

scalability and automation capabilities, it can track down threats from virtually any IP. However, weaknesses exist such that the use of handle API rate limits is convenient but results in delays when network activity is high that can become a bottleneck. Misses from a lack of sophistication in error handling threaten detection. AbuseIPDB goes down while continuously validated by Trane ML dynamic signatures until API is back to normal if any case has occurred for the API fetch process., and the whole network is exposed. The solution also needs to have resilience (rate limit checks and error logging), but it cannot cover all eventualities.

X. FUTURE WORK

Looking at the upcoming future, a hybrid system will become a key factor in evolving the solution which will be accomplished by the application of the blend of the machine learning techniques Random Forest in collaboration with deep learning ML. [39] With the joint compilation of many works, the power of future predictions can be sharpened, leading to a level of accuracy even with a chance of less than 1%. Using these advanced computations, we can both have highly accurate results in this regard as well as the aspect of a considerable decrease in the time both in the process of detention and in the prevention of any event.

XI. CONCLUSION

During that thorough research, the proposed option attempts to prove the soundness of reputation validation for cloud firewalls with the help of modern ML-driven technologies. The foundational objective of this research was to find a solution to the shortcomings in the existing cloud-based firewall security mechanisms that usually fail to discriminate between the harmful and innocent firewall rules. The given study suggests a solution by applying a combination of RF algorithm (ML) and deep learning (DL) methods which have not been seen before.

Also, this combination was specifically chosen to leverage the strengths of both methodologies such that ML supplies a provisional predictive precision, whereas DL elevates the model's ability to analyze and distinguish complicated data structures. This technical improvement, however, is of the highest accuracy ever at more than 99%. This great level of accuracy is because of the design of triple filtering architecture into the AWS cloud firewall. This function brings different measures such as several tests and balancing into a system for the IP to check and report on the IP accuracy. This mechanism offers a new propitious approach that separates harmful traffic while leaving those who legitimately want to use the web applications unharmed. The paper shows that the given method can easily be transferred and used for different types of web applications and threats. This adaptability is of utmost

importance since the industry always faces the challenge of continually dealing with new cybersecurity threats. They highlight that the systematical method that they have developed is not only a static solution.

Also, this suggests that the dominant role of ML and artificial intelligence (AI) in the creation and implementation of such security programs should be highlighted as well. These technologies provide the groundwork for the expansion and refinement of existing cyber defence capabilities in the face of potential complex cyber-attacks from what this paper has shown us it can be inferred that learning algorithms and neural networks yield a complete turnover of cloud firewall security systems. The reputation validation for the IP of web applications that use clouds is very accurate with this approach, therefore, it lays a solid foundation for protecting against malicious threats for online cloud web applications as well.

ACKNOWLEDGEMENT

The authors equally contributed to the present research, at all stages from the formulation of the problem to the final findings and solution.

REFERENCES RÉFÉRENCES REFERENCIAS

1. P. L. S. Jayalaxmi, R. Saha, G. Kumar, M. Conti, and T. H. Kim, "Machine and Deep Learning Solutions for Intrusion Detection and Prevention in IoTs: A Survey," *IEEE Access*, vol. 10, no. October, pp. 121173–121192, 2022, doi: 10.1109/ACCESS.2022.3220622.
2. R. Vinayakumar, M. Alazab, S. Srinivasan, Q. V. Pham, S. K. Padannayil, and K. Simran, "A Visualized Botnet Detection System Based Deep Learning for the Internet of Things Networks of Smart Cities," *IEEE Trans. Ind. Appl.*, vol. 56, no. 4, pp. 4436–4456, 2020, doi: 10.1109/TIA.2020.2971952.
3. A. Georgiadou, S. Mouzakitis, and D. Askounis, "Assessing mitre att&ck risk using a cyber-security culture framework," *Sensors*, vol. 21, no. 9, 2021, doi: 10.3390/s21093267.
4. N. A. Sankar and K. A. Fasila, "Implementation of SOC using ELK with Integration of Wazuh and Dedicated File Integrity Monitoring," *2023 9th Int. Conf. Smart Comput. Commun.*, pp. 350–354, 2023, doi: 10.1109/ICSCC59169.2023.10334992.
5. N. Usman *et al.*, "Intelligent Dynamic Malware Detection using Machine Learning in IP Reputation for Forensics Data Analytics," *Futur. Gener. Comput. Syst.*, vol. 118, pp. 124–141, 2021, doi: 10.1016/j.future.2021.01.004.
6. H. Manocha, A. Srivastava, C. Verma, R. Gupta, and B. Bansal, "Security Assessment Rating Framework for Enterprises using MITRE ATT&CK Matrix," 2021, [Online]. Available: <http://arxiv.org/abs/2108.06559>

7. W. Fang, C. Zhang, Z. Shi, Q. Zhao, and L. Shan, "BTRES: Beta-based Trust and Reputation Evaluation System for wireless sensor networks," *J. Netw. Comput. Appl.*, vol. 59, pp. 88–94, 2016, doi: 10.1016/j.jnca.2015.06.013.
8. E. S. Sagatov, D. A. Shkirdov, and A. M. Sukhov, "Analysis of network threats based on data from server-traps," *2019 10th IFIP Int. Conf. New Technol. Mobil. Secur. NTMS 2019 - Proc. Work.*, pp. 1–5, 2019, doi: 10.1109/NTMS.2019.8763847.
9. S. Benedict, "EA-POT: An Explainable AI Assisted Blockchain Framework for Honeypot IP Predictions," *Acta Cybern.*, vol. 26, no. 2, pp. 149–173, 2023, doi: 10.14232/actacyb.293319.
10. L. Deri and F. Fusco, "Evaluating IP Blacklists Effectiveness," pp. 1–8, 2023.
11. H. S. Sikandar, U. Sikander, A. Anjum, and M. A. Khan, "An Adversarial Approach: Comparing Windows and Linux Security Hardness Using Mitre ATT&CK Framework for Offensive Security," *IEEE 19th Int. Conf. Smart Communities Improv. Qual. Life Using ICT, IoT AI, HONET 2022*, pp. 22–27, 2022, doi: 10.1109/HONET56683.2022.10018981.
12. A. Kuppa, L. Aouad, and N. A. Le-Khac, "Linking CVE's to MITRE ATT and CK Techniques," *ACM Int. Conf. Proceeding Ser.*, 2021, doi: 10.1145/3465481.3465758.
13. J. L. Lewis, G. F. Tambaliuc, H. S. Narman, and W. S. Yoo, "IP Reputation Analysis of Public Databases and Machine Learning Techniques," *2020 Int. Conf. Comput. Netw. Commun. ICNC 2020*, pp. 181–186, 2020, doi: 10.1109/ICNC47757.2020.9049760.
14. S. Goel and S. Kumar, "An improved method of detecting spoofed attack in wireless LAN," *1st Int. Conf. Networks Commun. NetCoM 2009*, pp. 104–108, 2009, doi: 10.1109/NetCoM.2009.75.
15. A. Tosun, M. De Donno, N. Dragoni, and X. Fafoutis, "RESIP Host Detection: Identification of Malicious Residential IP Proxy Flows," *Dig. Tech. Pap. - IEEE Int. Conf. Consum. Electron.*, vol. 2021-Janua, 2021, doi: 10.1109/ICCE50685.2021.9427688.
16. E. Chiapponi, M. Dacier, and O. Thonnard, "Inside Residential IP Proxies: Lessons Learned from Large Measurement Campaigns," *Proc. - 8th IEEE Eur. Symp. Secur. Priv. Work. Euro S PW 2023*, pp. 501–512, 2023, doi: 10.1109/EuroSPW59978.2023.00062.
17. X. Mi et al., "Resident evil: Understanding residential IP proxy as a dark service," *Proc. - IEEE Symp. Secur. Priv.*, vol. 2019-May, pp. 1185–1201, 2019, doi: 10.1109/SP.2019.00011.
18. E. Chiapponi, M. Dacier, and O. Thonnard, "Poster: The Impact of the Client Environment on Residential IP Proxies Detection," *Proc. ACM SIGCOMM Internet Meas. Conf. IMC*, pp. 712–713, 2023, doi: 10.1145/3618257.3624993.
19. Y. Huang et al., "Detect Malicious IP Addresses using Cross-Protocol Analysis," *2019 IEEE Symp. Ser. Comput. Intell. SSCI 2019*, pp. 664–672, 2019, doi: 10.1109/SSCI44817.2019.9003003.
20. R. Maurya, "Analyzing the Role of AI in Cyber Security Threat Detection & Prevention," *Int. J. Res. Appl. Sci. Eng. Technol.*, vol. 11, no. 11, pp. 514–519, 2023, doi: 10.22214/ijraset.2023.56510.
21. A. Z. Faridee and V. P. Janeja, "Measuring Peer Mentoring Effectiveness," *Am. J. o*, vol. 15, no. 2, pp. 7–22, 2020.
22. R. Ganeshan, C. S. Kolli, C. M. Kumar, and T. Daniya, "A Systematic Review on Anomaly Based Intrusion Detection System," *IOP Conf. Ser. Mater. Sci. Eng.*, vol. 981, no. 2, 2020, doi: 10.1088/1757-899X/981/2/022010.
23. A. Aldweesh, A. Derhab, and A. Z. Emam, "Deep learning approaches for anomaly-based intrusion detection systems: A survey, taxonomy, and open issues," *Knowledge-Based Syst.*, vol. 189, p. 105124, 2020, doi: 10.1016/j.knosys.2019.105124.
24. M. Zolanvari, M. A. Teixeira, L. Gupta, K. M. Khan, and R. Jain, "Machine Learning-Based Network Vulnerability Analysis of Industrial Internet of Things," *IEEE Internet Things J.*, vol. 6, no. 4, pp. 6822–6834, 2019, doi: 10.1109/JIOT.2019.2912022.
25. Z. K. Maseer, R. Yusof, N. Bahaman, S. A. Mostafa, and C. F. M. Foozy, "Benchmarking of Machine Learning for Anomaly Based Intrusion Detection Systems in the CICIDS2017 Dataset," *IEEE Access*, vol. 9, pp. 22351–22370, 2021, doi: 10.1109/ACCESS.2021.3056614.
26. D. Jeon and B. Tak, "BlackEye: automatic IP blacklisting using machine learning from security logs," *Wirel. Networks*, vol. 28, no. 2, pp. 937–948, 2022, doi: 10.1007/s11276-019-02201-5.
27. R. Al-Shaer, J. M. Spring, and E. Christou, "Learning the Associations of MITRE ATT CK Adversarial Techniques," *2020 IEEE Conf. Commun. Netw. Secur. CNS 2020*, vol. 1345, 2020, doi: 10.1109/CNS48642.2020.9162207.
28. D. Jeon and B. Tak, "automatic IP blacklisting using machine learning," *Wirel. Networks*, vol. 28, no. 2, pp. 937–948, 2022, doi: 10.1007/s11276-019-02201-5.
29. N. Usman et al., "Intelligent Dynamic Malware Detection using Machine Learning in IP Reputation for Forensics Data Analytics," *Futur. Gener. Comput. Syst.*, vol. 118, pp. 124–141, 2021, doi: 10.1016/j.future.2021.01.004.
30. S. Shaw and P. Choudhury, "MAC address spoofing," *Conf. Proceeding - 2015 Int. Conf. Adv. Comput. Eng. Appl. ICACEA 2015*, pp. 347–350, 2015, doi: 10.1109/ICACEA.2015.7164728.
31. Y. Huang et al., "Graph neural networks and cross-protocol analysis for detecting malicious IP addresses," *Complex Intell. Syst.*, vol. 9, no. 4, pp.



- 3857–3869, 2023, doi: 10.1007/s40747-022-00838-y.
32. D. Ocampo, F. B. C, D. Castillo, T. M. L, and M. A. N, "A New Local Area Network Attack through IP and M," pp. 198–205, 2013.
33. "AbuseIPDB - IP address abuse reports." <https://www.abuseipdb.com/> (accessed Mar. 06, 2024).
34. "SANS Institute." <https://www.sans.org/cyber-security-training-overview/?msc=main-nav> (accessed Mar. 06, 2024).
35. P. Kumar and S. Rana, "Development of modified AES algorithm for data security," *Optik (Stuttg)*, vol. 127, pp. 2341–2345, 2016, doi: 10.1016/J.IJLEO.2015.11.188.
36. A. M. Carrington *et al.*, "Deep ROC Analysis and AUC as Balanced Average Accuracy, for Improved Classifier Selection, Audit and Explanation," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 45, pp. 329–341, 2021, doi: 10.1109/TPAMI.2022.3145392.
37. Z. Wen, R. Zhang, and K. Ramamohanarao, "Enabling Precision/Recall Preferences for Semi-supervised SVM Training," *Proc. 23rd ACM Int. Conf. Conf. Inf. Knowl. Manag.*, 2014, doi: 10.1145/2661829.2661977.
38. C. K. I. Williams, "The Effect of Class Imbalance on Precision-Recall Curves," *Neural Comput.*, pp. 1–5, 2020, doi: 10.1162/neco_a_01362.
39. J. Zhang and S. Li, "A Review of Machine Learning Based Species' Distribution Modelling," *Proc. - 2017 Int. Conf. Ind. Informatics - Comput. Technol. Intell. Technol. Ind. Inf. Integr. ICIIICII 2017*, vol. 2017-Decem, pp. 199–206, 2017, doi: 10.1109/ICIIICII.2017.76.