

A Comparative Analysis of Supervised Machine Learning Models for Threat Detection in Domain Name System

Amir Haris Ahmad¹, Hadhrami Ab Ghani¹, Muhammad Muzzammal Mirza¹,
Muhammad Anwar^{1,2*}

¹Faculty of Data Science and Computing, Universiti Malaysia Kelantan, Kota Bahru 16100, Kelantan, Malaysia

²Department of Information Sciences, Division of Science and Technology, University of Education, Lahore 54000, Pakistan

ARTICLE INFO

Article history:

Received 7 January 2026

Revised 20 March 2026

Accepted 2 April 2026

Online first

Published 30 April 2026

Keywords:

DSN security
supervised machine learning
threat detection
comparative analysis

DOI:

10.24191/mij.v7i1.11939

ABSTRACT

The Domain Name System (DNS) is an important element of the Internet. It serves as the standard method for translating human-readable domain names into machine-readable IP addresses. Despite its importance, DNS has persistently challenged by various threats that compromise its security and functionality. Various machine learning models have been proposed to detect and classify DNS attacks. This work aims to provide a comparative analysis of three supervised machine learning models: Random Forest, Vector-Based, and XGBoost. These models were trained and tested for the classification of DNS threats in 11 different categories. Among all three models, XGBoost consistently outperforms Random Forest and the Vector-Based model in terms of accuracy, speed, confidence, precision, recall, and F1-score. It provides the highest detection accuracy (56.4%), the fastest processing speed (10,263 domains/sec), and the lowest false alarm rate (4.3%), making it the most reliable choice for malicious domain detection. In addition to examining the DNS threat landscape and existing challenges, this research also highlights the strengths and limitations of the existing state of the art and provides future research directions for researchers to understand DNS vulnerabilities and current gaps in their mitigation techniques.

1. INTRODUCTION

The expansion of the Internet has been remarkable, characterized by the surge in connected hosts and the development of high-speed access (Choi & Yi, 2018). The Internet uses TCP/IP (Transmission Control Protocol/Internet Protocol) for communication among network devices (Nath & Uddin, 2015). Every device

^{1,2*}Corresponding author. E-mail address: anwar.muhammad@ue.edu.pk

should have a unique IP address to communicate over the Internet. It is difficult for humans to remember the IP addresses of websites. However, website names can easily be remembered. For this reason, the Domain Name System (DNS) was introduced. DNS translates domain names into IP addresses and vice versa (Bonastre & Veal, 2019).

Due to the importance of DNS, it has been targeted by attackers (Schmid, 2021). Several works have been proposed in the literature to secure DNS servers. However, conventional methods lack the ability to provide comprehensive security (Al-Mashhadi & Manickam, 2020). The recommended solution is to adapt artificial intelligence and machine learning based methods to defend against increasingly powerful attacks (Marques et al., 2021). This study aims to provide a comprehensive analysis of supervised machine learning models for DNS threat detection and classification. This paper focuses on exploring three supervised machine learning models namely Random Forest, Vector-Based, and XGBoost. The dataset contained 11 threat categories (e.g., abuse, adware, cryptojacking, drugs, fraud, gambling, malware, phishing, piracy, ransomware, and scam). For each category, five representative domains were collected for testing. The performance of these models was assessed in terms of precision, recall, F1-score, and False Alarm Ratio (FAR). A confusion Matrix was generated to evaluate misclassification rates across categories.

The key contributions of this research are as follows:

- (i) The design of a supervised ML-based DNS threat detection framework.
- (ii) Comparative evaluation of multiple supervised models.
- (iii) Insights into model performance and computational trade-offs.

The structure of this paper is as follows: Section 2 reviews relevant literature, Sections 3 and 4 describe the research methodology and experimental findings, and Section 5 presents the conclusion.

2. RELATED WORK

Several supervised machine learning models have been proposed in the literature for DNS security. The study by Barradas et al. (2021) introduced FlowLens, which is a proposed framework that demonstrates how programmable network switches can be employed to mine linearly independent flow features directly from network traffic, allowing parallel and fast machine learning based classification. It focuses on important security cases of covert channel detection, website fingerprinting, and botnet identification using supervised models such as XGBoost, Random Forest, and Naïve Bayes, respectively. FlowLens lowers the cost and enhances the scalability of low-latency feature extraction by offloading it to the data plane so that real-time analytics can be achieved at high throughput. The paper demonstrates the effectiveness of the system across several datasets and shows that it can achieve high accuracy while adding minimal latency on average. While the findings are positive, dependence on specialized switch capabilities and labeled datasets may limit generalizability in diverse or evolving network environments.

The study conducted by Singh and Roy (2022) investigates the ability of ensemble machine learning to distinguish between malicious and benign DNS over HTTPS traffic. Many models, such as Random Forest, KNN, and Logistic Regression, are tested using the CIRA-CIC-DoHBrw-2020 dataset and thirty selected characteristics. Since ensemble methods are employed, the ensemble trials yield an accuracy of 100 percent, thus proving the effectiveness of ensemble methods. Given the focus on offline static data and an the omission of real-world encrypted traffic conditions, the research also demands further study. Moreover, the consideration of the evolving attacker behaviours is excluded, reducing its advantage in preventing emerging threats.

The study by Jerabek et al. (2023) used flow-level statistical features extracted from NetFlow and IPFIX records and applied several supervised machine learning algorithms, such as Random Forest, Decision Tree, and Gradient Boosting. The study implements this approach on real-world datasets and demonstrates strong generalization for attack detection in the differential classification of DoH traffic from regular HTTPS traffic. In this work, the authors propose a privacy-preserving detection scheme that does not depend on payloads and hence can be applied in encrypted domains. Nonetheless, the methodology might be limited in its adaptability to adversaries that are dynamically changing their traffic patterns to bypass statistical detection, calling for frequent model retraining.

The study by Calle et al. (2024) addresses large scale DNS tampering, especially by authoritarian states such as China's Great Firewall. The study uses both supervised and unsupervised models on OONI's global DNS data, as shown in Fig. 2, achieving high precision to learn tampering heuristics broadly and detect previously unknown manipulations. In addition, the study analyzes performance under different time windows (1–24 months) to give insight into how censorship changes over time. It demonstrates solid scalability and cross-regional deployability but lacks explicit interpretability or consideration of false-positive costs in benign anomalies (e.g., load balancing). More work is necessary to accommodate adversarial resilience and active response mechanisms.

El Attar et al. (2024) examine various machine learning models (i.e. XGBoost, MLP, and SVM) by focusing on DNS-flooding DDoS attacks using the CIC-DDoS2019 dataset. This study improves upon previous limitations through model feature selection (filter, wrapper, and embedded) strategies. Even though the method is proven in emulated attack scenarios, this work does not cover zero-day or polymorphic attacks and does not discuss integration with real-world DNS infrastructure. It also fails to investigate model robustness under attack or explore the evolution of detection performance over time.

Schummer et al. (2024) introduce a complete approach to the detection of network anomalies as well as DNS-specific irregularities using different supervised and unsupervised machine learning mechanisms, clustering, change-point detection, and classification models. This is not limited to DNS attacks; however, the system's ability to monitor packet loss, congestion, and behavioral anomalies can apply to DNS anomaly settings. By allowing the interpretation of model confidence, the paper addresses one of the principal drawbacks in ML-based cybersecurity systems: the interpretability of such models. The paper is not comprehensive with respect to considerable DNS protocol nuances, such as tunneling, DoH, or DGA-based evasion that are abstracted away in this study and therefore restrict its direct applicability for deployment in specialized DNS security systems. Without tuning for the DNS-specific deployment, this approach remains suitable as a general-purpose anomaly detector.

3. METHODOLOGY

To evaluate the effectiveness of different supervised machine learning approaches for malicious domain detection, three models were developed and tested namely Random Forest, Vector-Based, and XGBoost, as shown in Fig.1. These models were trained and tested for classification of DNS threats in 11 threat categories, as shown in Table 1. For each category, five representative domains were collected for testing. Synthetic domains were also used for training purposes to ensure balanced learning across all classes.

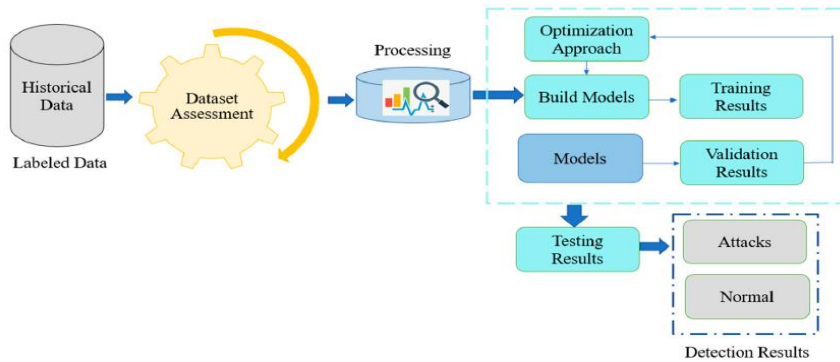


Fig. 1. Workflow of supervised machine learning model

4. EXPERIMENTAL RESULTS

The assessment of each model was conducted based on precision, recall, F1-score, false alarm rate, and the confusion matrix. The precision and False Alarm Rate (FAR) tests were conducted by determining the True Positive, False Positive, False Negative, and True Negative values. Eq. (1) to Eq. (4) used to calculate precision, recall value, F1-score, and FAR. Table 2, Table 3, and Table 4 present the results of precision, recall, F1-score, and FAR for Random Forest, Vector-Based, and XGBoost, respectively.

$$\text{Precision} = \text{TP} / (\text{TP} + \text{FP}) \quad (1)$$

$$\text{Recall} = \text{TP} / (\text{TP} + \text{FN}) \quad (2)$$

$$\text{F1-Score} = 2 * (\text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall}) \quad (3)$$

$$\text{FAR} = \text{FP} / (\text{FP} + \text{TN}) \quad (4)$$

Table 1. Categories of DNS traffic in dataset

Category	Domains
Abuse	casino-engine.ru
	00hidemyprivacy.ws
	00sexus.com
	012469af389a1d1246d.com
	01apple.com
Adware	10130.engine.mobileappttracking.com
	www.iuqerfsodp9ifjaposdfjhgosurijfaewrwergwea.com
	www.ivideopiuvisti.com
	www.iwanttodeliver.com
	www.iyogi.com
Cryptojacking	ugfxrrqz.bid
	sncx.pool.mn
	sng.ehashcoins.org
	snowgem.minerspool.cc
	social.moneroworld.com
Drugs	epd.bz
	epharmacy.com.au
	vodrx.com
	volpills.com
	vpillsal.org
Fraud	www.instagramconfirmation.info
	9cb4afde731-google.com
	islamicmantraforlovemarriage.com
	islanded.us
	ismartsupport.com
Gambling	macau.williamhill.com
	folkeautomaten.com
	casinoland.com
	www.palladiumcasino.eu
	www.instantactionsports.com
Malware	020zz.com
	024dna.cn
	028happy.com
	02ip.ru
	0315st.com
Phishing	904facebook.com
	www.secure-infos-appleid.com
	www.secure-link-paypal.com
	www.secure-login-appleid-verification-account-mailconfirmations.cloud
	www.secure-paypal-protection.info
Piracy	kickasstorrents.to
	3dmgame.com
	4allprograms.net
	a2zapk.com
	adobecccrack2017formacandwindows.yolasite.com
Ransomware	vyohacxzoue32vvk.34o9h1.bid
	3qbyaohkcqkzrz6.tordoor.li
	4kqd3hmqgptupi3p.0vgu64.top
	4kqd3hmqgptupi3p.143h2a.top
	4kqd3hmqgptupi3p.1tvjk1.to
Scam	auxiliarymine.com
	binancetrade.live
	catchemboxe.com
	directtrade24.com
	hqoptions.com

Table 2. Precision, Recall, F1-score, and FAR of Random Forest

Category	Precision	Recall	F1-Score	False Alarm
Abuse	0.0000	0.0000	0.0000	0.0200
Adware	0.0000	0.0000	0.0000	0.0800
Crypto	1.0000	0.2000	0.3333	0.0000
Drugs	0.2500	0.8000	0.3810	0.2400
Fraud	0.0000	0.0000	0.0000	0.0400
Gambling	0.0000	0.0000	0.0000	0.0600
Malware	0.0000	0.0000	0.0000	0.1200
Phishing	0.2500	0.2000	0.2222	0.0600
Piracy	0.0000	0.0000	0.0000	0.0000
Ransomware	0.0909	0.2000	0.1250	0.2000
Scam	0.2857	0.4000	0.3333	0.1000
Overall	0.1706	0.1636	0.1268	0.0836

Table 3. Precision, Recall, F1-score, and FAR of Vector-Based

Category	Precision	Recall	F1-Score	False Alarm
Abuse	0.3333	0.4000	0.3636	0.0800
Adware	0.2857	0.4000	0.3333	0.1000
Crypto	1.0000	0.6000	0.7500	0.0000
Drugs	0.3333	0.6000	0.4286	0.1200
Fraud	0.2857	0.4000	0.3333	0.1000
Gambling	0.4000	0.4000	0.4000	0.0600
Malware	0.0000	0.0000	0.0000	0.4000
Phishing	0.4444	0.8000	0.5714	0.1000
Piracy	0.5000	0.4000	0.4444	0.0400
Ransomware	0.0000	0.0000	0.0000	0.0400
Scam	0.0000	0.0000	0.0000	0.0200
Overall	0.3257	0.3636	0.3295	0.0636

Table 4. Precision, Recall, F1-score, and FAR of XGBoost

Category	Precision	Recall	F1-Score	False Alarm
Abuse	0.5000	0.4000	0.4444	0.0400
Adware	0.5000	0.6000	0.5455	0.0600
Crypto	1.0000	0.8000	0.8889	0.0000
Drugs	1.0000	1.0000	1.0000	0.0000
Fraud	1.0000	0.6000	0.7500	0.0000
Gambling	0.6667	0.8000	0.7273	0.0400
Malware	0.0000	0.0000	0.0000	0.0200
Phishing	0.5000	0.4000	0.4444	0.0400
Piracy	1.0000	0.6000	0.7500	0.0000
Ransomware	0.3125	1.0000	0.4762	0.2200
Scam	0.0000	0.0000	0.0000	0.0600
Overall	0.5890	0.5479	0.5479	0.0436

The confusion matrix test was created by comparing the true labels for the dataset and the prediction made by the machine learning model. In this experiment, five domains for each of the 11 categories were tested, and the result obtained from the test was converted into a confusion matrix. Figs. 2 to 4 represent the confusion matrix of Random Forest, Vector-Based, and XGBoost, respectively.

Categories	Abuse	Adware	Crypto	Drug	Fraud	Gambling	Malware	Phishing	Piracy	Ransomware	Scam
Abuse				1		1					3
Adware							3			2	
Crypto	1		1	2	1						
Drugs				4		1					
Fraud				1				1		2	1
Gambling				1	1					1	2
Malware				5							
Phishing		1						1		3	
Piracy		1		1		1				1	1
Ransomware		1		1			1	1		1	
Scam		1					1	1			2

Fig. 2. Confusion matrix of Random Forest

Categories	Abuse	Adware	Crypto	Drug	Fraud	Gambling	Malware	Phishing	Piracy	Ransomware	Scam
Abuse	2					1		1	1		
Adware	1	2					1	1			
Crypto	1		3								
Drugs	1			3	1						
Fraud					2		1	1		1	
Gambling					1	2		2			
Malware				1	3						
Phishing					1						
Piracy		1		1		1			2		
Ransomware	1			4							
Scam		4								1	

Fig. 3. Confusion matrix of Vector-Based

The computational cost or training time of the Random Forest, Vector-Based (SVM), and XGBoost models are presented in Table 5.

Table 5. Computational cost or training time

Model	Approx. Training Time	Computational Cost	Key Observation
Random Forest	18-25 seconds	Medium	Balanced cost and accuracy
Vector-Based (SVM)	40-60 seconds	High	High accuracy but slow for large data
XGBoost	25-35 seconds	Medium-High	Best accuracy with efficient cost

Categories	Abuse	Adware	Crypto	Drug	Fraud	Gambling	Malware	Phishing	Piracy	Ransomware	Scam
Abuse	2					1	1	1			
Adware	1	3								1	
Crypto			4							1	
Drugs				5							
Fraud					3					1	1
Gambling		1				4					
Malware										4	1
Phishing	1							1		2	
Piracy									3	1	1
Ransomware										5	
Scam		2				1		1		1	

Fig. 4. Confusion matrix of XGBoost

5. CONCLUSION

The research focuses on analyzing event data generated through the Domain Name System (DNS), a fundamental internet protocol, for threat identification and classification. Supervised machine learning models, namely Random Forest, Vector Based, and XGBoost, were explored. Among all three models, XGBoost consistently outperforms Random Forest and the Vector-Based model in terms of accuracy, speed, confidence, precision, recall, and F1-score. It provides the highest detection accuracy (56.4%), the fastest processing speed (10,263 domains/sec), and the lowest false alarm rate (4.3%), making it the most reliable choice for malicious domain detection. While Random Forest and Vector-Based models show some effectiveness, they suffer from low accuracy and higher error rates. In contrast, XGBoost demonstrates strong generalization, efficiency, and robustness, making it the preferred model for deployment in real-world cybersecurity systems.

6. ACKNOWLEDGEMENT/FUNDING

The authors would like to acknowledge the support of Universiti Malaysia Kelantan (UMK) for providing the facilities and financial support for this research.

7. CONFLICT OF INTEREST STATEMENT

The authors agree that this research was conducted in the absence of any self-benefits, commercial or financial conflicts and declare the absence of conflicting interests with the funders.

8. AUTHORS' CONTRIBUTIONS

Amir Haris Bin Ahmad: Conceptualization, investigation, methodology, writing – original draft, formal analysis; Hadhrami Bin Ab Ghani: Supervision, methodology, project administration; Muhammad Muzzammal Mirza: Visualization, data curation, writing – review & editing; Muhammad Anwar: Writing – review & editing, validation, submission management.

9. DECLARATION OF GENERATIVE AI IN THE WRITING PROCESS

The authors declare that generative artificial intelligence (AI) tools may have been used solely to support language refinement during the preparation of this manuscript. Such tools were used only to improve grammar, sentence structure, clarity, and overall readability. All aspects of the research, including the study design, data collection, data extraction, analysis, interpretation of findings, and final academic decisions, were conducted, reviewed, and approved by the authors. The authors take full responsibility for the accuracy, integrity, originality, and scholarly content of this manuscript.

10. DATA AVAILABILITY/SUPPLEMENTARY MATERIALS

The data and/or supplementary materials supporting the findings of this study are available from the corresponding author upon reasonable request, where applicable.

11. ETHICS STATEMENT

The authors declare that this study did not involve human participants or animal subjects. Where applicable, all procedures were conducted in accordance with relevant institutional guidelines, safety requirements, and ethical standards.

REFERENCES

- Al-Mashhadi, S., & Manickam, S. (2020). A brief review of blockchain-based DNS systems. *International Journal of Internet Technology and Secured Transactions*, 10(4), 420-432. <https://doi.org/10.1504/ijitst.2020.108134>
- Barradas, D., Santos, N., Rodrigues, L., Signorello, S., Ramos, F. M. V., & Madeira, A. (2021). FlowLens: Enabling efficient flow classification for ML-based network security applications. In *Proceedings of the Network and Distributed System Security Symposium (NDSS 2021)*. <https://doi.org/10.14722/ndss.2021.24067>
- Bonastre, O. M., & Veà, A. (2019). Origins of the domain name system. *IEEE Annals of the History of Computing*, 41(2), 48–60. <https://doi.org/10.1109/MAHC.2019.2913116>
- Calle, P., Savitsky, L., Bhagoji, A. N., Hoang, N. P., & Cho, S. (2024). Toward automated DNS tampering detection using machine learning. In *Free and Open Communications on the Internet 2024 (FOCI '24)*.
- Choi, C., & Yi, M. H. (2018). The Internet, R&D expenditure and economic growth. *Applied Economics Letters*, 25(4), 264–267. <https://doi.org/10.1080/13504851.2017.1316819>
- El Attar, A., Khatoun, R., Chbib, F., Fadlallah, A., & Serhrouchni, A. (2024). DNS flooding attack detection scheme through machine learning. In *2024 International Wireless Communications and Mobile*

Computing (IWCMC) (pp. 132–137). IEEE.
<https://doi.org/10.1109/IWCMC61514.2024.10592588>

Jeřábek, K., Hynek, K., Ryšavý, O., & Burgetová, I. (2023). DNS over HTTPS detection using standard flow telemetry. *IEEE Access*, 11, 50000–50012. <https://doi.org/10.1109/ACCESS.2023.3275744>

Marques, C., Malta, S., & Magalhães, J. (2021). DNS firewall based on machine learning. *Future Internet*, 13(12), Article 309. <https://doi.org/10.3390/fi13120309>

Nath, P. B., & Uddin, M. M. (2015). TCP/IP model in data communication and networking. *American Journal of Engineering Research*, 4(10), 102–107.

Schmid, G. (2021). Thirty years of DNS insecurity: Current issues and perspectives. *IEEE Communications Surveys & Tutorials*, 23(4), 2429–2459. <https://doi.org/10.1109/COMST.2021.3105741>

Schummer, P., del Rio, A., Serrano, J., Jimenez, D., Sánchez, G., & Llorente, Á. (2024). Machine learning-based network anomaly detection: Design, implementation, and evaluation. *AI*, 5(4), 2967–2983. <https://doi.org/10.3390/ai5040143>

Singh, S. K., & Roy, P. K. (2022). Malicious traffic detection of DNS over HTTPS using ensemble machine learning. *International Journal of Computing and Digital Systems*, 11(1), 1061–1069. <https://doi.org/10.12785/ijcds/110185>



© 2026 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY-NC-ND) license (<http://creativecommons.org/licenses/by/4.0/>).