

Securing Cloud Infrastructure Against Privilege Escalation Attacks with Machine Learning-Based Detection Systems

E. Vinod Kumar^{1*}, Vishnu Satheesh², S.Sravani², Ch.Saikrishna², B. Chandana², M.Varshitha²

¹Assistant Professor, ²UG Student, ^{1,2}Department of Computer Science and Engineering,
^{1,2}Malla Reddy Engineering College and Management Sciences, Kistapur, Medchal, 501401,
Telangana, India

*Corresponding author: vinodkumar@mrem.ac.in

Received: 02.03.2025

Revised: 15.04.2025

Accepted: 25.05.2025

ABSTRACT

Privilege escalation attacks represent a critical security threat in cloud computing environments, enabling unauthorized users to gain elevated access rights and compromise sensitive data and system integrity. Traditional detection mechanisms, primarily based on manual monitoring and rule-based anomaly detection, often suffer from high false-positive rates, limited scalability, and delayed response, making them inadequate against increasingly sophisticated attack strategies. To address these challenges, this research proposes a machine learning-based detection framework for securing cloud infrastructure against privilege escalation attacks. The proposed system leverages advanced ensemble learning algorithms, including Random Forest, AdaBoost, XGBoost, LightGBM, and CatBoost, to automatically learn and identify complex behavioral patterns associated with unauthorized privilege elevation. By analyzing large-scale cloud activity logs and historical attack data, the models effectively detect subtle anomalies indicative of malicious behavior in real time. The adaptive nature of machine learning enables continuous improvement in detection accuracy while significantly reducing reliance on human intervention. Experimental evaluation demonstrates that the proposed approach outperforms conventional rule-based systems in terms of precision, scalability, and responsiveness. This research highlights the effectiveness of machine learning techniques in providing a proactive, automated, and robust defense mechanism, thereby enhancing the overall security posture of modern cloud infrastructures.

Keywords: Cloud Security, Privilege Escalation Attacks, Machine Learning, Anomaly Detection, Ensemble Learning, Intrusion Detection System, Random Forest, XGBoost, Cloud Infrastructure Protection.

1. INTRODUCTION

Cloud computing has revolutionized data storage and processing by offering scalable, on-demand services. However, this centralization introduces security challenges, particularly concerning unauthorized data access and insider threats.[1] Insider threats, originating from within the organization, pose significant risks as they involve individuals misusing their authorized access to harm systems or steal data. These threats are especially concerning in applications like online social networks (OSNs), business platforms, and e-commerce systems, where sensitive user information and financial data are prevalent. [2] Detecting insider threats is challenging because these individuals have legitimate access to the company's systems and data, making traditional security tools like firewalls often ineffective against such threats.

To address latency and bandwidth issues associated with cloud computing, paradigms like fog and edge computing have emerged. Fog computing extends cloud services to the network's edge, processing data closer to its source, while edge computing involves data processing directly on

10.48047/jocaaa.2025.34.05.40

devices.[3] While these models enhance performance, they also introduce unique security vulnerabilities. Many attacks affecting fog and edge computing are inherited from cloud computing; however, current cloud security solutions may not be effective in these environments due to their unique properties. Therefore, robust security measures, such as strong encryption, intrusion detection and prevention systems, and continuous monitoring, are essential to protect against unauthorized access and data breaches in these distributed architectures.

In data access scenarios within OSNs, business platforms, and e-commerce applications, escalation attacks can occur during routine operations [4]. For instance, employees or contractors with legitimate access might intentionally steal sensitive customer data or inadvertently expose the system to external threats through negligence. Insider threats are a serious danger for e-commerce businesses, as disgruntled employees or ex-employees may attempt to steal company data or proprietary information, hoping to sell it to rival businesses. Mitigating these threats requires implementing strict access controls, conducting regular security training, and employing advanced intrusion detection systems .[5] Effective mitigation involves enforcing least privilege access, implementing strong authentication, continuous monitoring, and rapid incident response mechanisms.

2. LITERATURE SURVEY

Najar, et al. [6] (2024) suggested a novel approach for low-rate DDoS detection and classification in cloud computing environments, addressing security threats that traditional methods struggle to mitigate. Proposed AE-CIAM, a hybrid framework integrating an attention-enhanced Autoencoder (AE) for feature extraction and a CNN Inception with Attention Mechanism for classification. Highlighted the model's effectiveness in balancing computational efficiency with high detection accuracy. Experimental evaluation on the CICIDS2017 dataset demonstrated superior performance, achieving 99.99% accuracy in binary classification and 99.44% in multiclassification, surpassing existing detection methods.

Sharma, et al. [7] (2025) proposed an integrated machine learning (ML) approach to enhance cloud computing security, addressing limitations of standalone ML models against emerging attack techniques. Suggested leveraging ensemble learning and AI techniques to improve prediction accuracy. Introduced a no monotonous methodology within the ensemble model to enhance adaptive capabilities, facilitating the detection of unknown malicious activities and predicting zero-day attacks in cloud networks.

Nazari, et al. [8] (2025) This paper proposed an efficient cloud-based platform for real-time Cyber Threat Intelligence (CTI) detection, collection, and sharing from various online sources. Developed TSTEM, a prototype with a containerized microservice architecture leveraging Tweepy, Scrapy, Terraform, ELK, Kafka, and MLOps for autonomous IOC extraction and indexing. Implemented infrastructure as code (IaC) for provisioning and monitoring. Integrated custom focus-crawlers and multi-stage classifiers, enhanced by NLP models, to improve IOC identification. Experimental results showed high accuracy (exceeding 98%) in classification and extraction, ensuring rapid performance with minimal false positives, making it a robust solution for real-time CTI analysis and threat detection.

K V Kmet al. [9] (2025) This paper proposed a deep learning-based Intrusion Detection System (IDS) with a novel feature selection technique for cloud security. Focused on addressing security threats in cloud computing by overcoming limitations of traditional IDS. Designed a model with four phases: feature selection, pre-processing, classification, and encryption. Utilized the Adaptive Walrus Optimization Algorithm (AWO) for optimal feature selection, reducing computational complexity.

Classified selected features using an Enhanced Adaptive Neuro-Fuzzy Inference System (EANFIS) for accurate intrusion detection.

Ahsun, et al. [10] (2025) This paper investigated the current state of real-time fraud detection, highlighting its limitations and emerging trends. Examined the role of artificial intelligence, blockchain, IoT, and other technologies in transforming fraud detection and prevention. Addressed challenges such as data privacy, regulatory compliance, and the evolving nature of fraud. Provided insights into enhancing real-time fraud detection to mitigate economic and social impacts. Discussed the challenges associated with implementing these technologies, including data privacy, regulatory compliance, and the evolving nature of fraud. The findings of this study contributed to the understanding of how real-time fraud detection could be enhanced, ultimately reducing the economic and social impact of fraud.

Zhukabayeva, et al. [11] (2025) This paper explored cybersecurity challenges and solutions for cyber-physical systems (CPS) within IIoT and its integration with edge computing (IIoT–edge computing). We systematically analyzed literature from the past five years, identifying key IIoT layer attacks, intrusion methods, and threats affecting IIoT–edge environments. A detailed taxonomy of primary security mechanisms was developed, followed by a comparative analysis of our findings against existing research. The study highlighted vulnerabilities across the IIoT architecture, particularly concerning DoS, ransomware, malware, and MITM attacks. Advanced security technologies, including machine learning, federated learning, blockchain, deep learning, encryption, cryptography, IT/OT convergence, and digital twins, were examined as essential solutions for enhancing CPS security. The review underscored the need for real-time data protection and resilience in industrial operations. Finally, we outlined future research directions to strengthen cybersecurity in this evolving domain, ensuring robust protection against emerging threats.

Lytras, et al. [12] (2024) This paper Proposed key security concerns in cloud computing, including data loss, misconfiguration, unauthorized access, insecure APIs, account hijacking, and cyberattacks. Addressed data loss as a result of exposure due to breaches. Investigated misconfiguration as a major cause of data breaches, often due to reliance on CSPs for security. Highlighted unauthorized access risks due to cloud deployments being accessible from public networks. Focused on insecure APIs, which, if not secured, pose vulnerabilities. Provided insights into account hijacking, where stolen credentials grant access to sensitive data. Aimed to emphasize these threats and their implications for cloud security.

Singh, et al. [13] (2024) This paper Addressed the rapid advancement of cybersecurity threats over the past decade, with cybercriminals employing sophisticated techniques like ransomware, phishing, and advanced persistent threats (APTs). Highlighted the financial and reputational impacts of such attacks, emphasizing the need for robust cybersecurity measures. Investigated the role of machine learning in enhancing cybersecurity, leveraging large datasets to detect malicious patterns and anomalies. Focused on ML-driven improvements in real-time threat detection, response, and mitigation. Provided insights into intelligent machine learning solutions for cybersecurity while considering legal and ethical implications in a global context.

Srivastava, et al. [14] (2025) Reviewed recent developments in IoT technologies, highlighting security vulnerabilities in confidential data transmission. Addressed the role of Intrusion Detection Systems (IDS) in monitoring and alerting malicious activities. Investigated machine learning approaches for anomaly detection, noting the limited focus on sequential network data. Compared multiple models, including Random Forests (RF), Multi-Layer Perceptron's (MLP), and Long Short-Term Memory (LSTM), using the CIDDS-001 dataset. Provided experimental findings demonstrating that sequential detection is more effective than point-wise approaches. Highlighted LSTM's superior

performance in detecting order-dependent traffic anomalies, achieving 99.94% accuracy and an F1-score of 91.66%.

Sun, et al. [15] (2025) This paper explored privacy and security challenges in the Internet of Things (IoT) and proposed a comprehensive protection framework. It reviewed advancements in Attribute-Based Encryption (ABE) for search security and examined firmware implementation defects, analysing their generation mechanisms and detection methods, including static analysis, symbolic execution, fuzz testing, program validation, and machine learning. The study also evaluated mainstream IoT access control models and discussed future challenges in blockchain-based access control. Additionally, it highlighted recent progress in machine learning for IoT security and summarized key privacy regulations, with a focus on the European Union's information protection laws. Finally, the paper identified unresolved research challenges and outlined future directions for enhancing IoT security and privacy.

Reshma, et al. [16] (2025) This paper Investigated the impact of changing personal attitudes on IoT adoption, leading to increased data transmission and privacy concerns. Addressed energy limitations in IoT applications and highlighted Narrowband Internet of Things (NB-IoT) as a solution due to its low power consumption and extended battery life. Examined NB-IoT's vulnerability to network attacks and analyzed its fundamentals, architecture, and mathematical modelling. Reviewed existing machine learning-based security proposals and identified research gaps in NB-IoT security, distributed data aggregation, and scalability. Suggested emerging machine intelligence-supported algorithms to enhance security and scalability in NB-IoT systems.

3. PROPOSED SYSTEM

Cloud computing has revolutionized the way organizations store, manage, and process data by offering scalable and on-demand resources. However, with the increasing adoption of cloud services, security concerns such as privilege escalation attacks have become a significant challenge. Privilege escalation occurs when an attacker exploits system vulnerabilities to gain unauthorized access to higher-level privileges, leading to severe security breaches, data theft, and system compromise. Detecting and mitigating such attacks in a cloud environment is crucial to ensuring data confidentiality, integrity, and availability. Traditional security mechanisms often fail to detect sophisticated privilege escalation attacks due to their evolving nature. Hence, machine learning (ML)-based detection systems provide an advanced and efficient approach to identifying and mitigating such threats.

In our proposed system, we focus on detecting and mitigating privilege escalation attacks in cloud environments using machine learning. Specifically, we have chosen the Categorical Boosting Classifier (CatBoost) as our primary algorithm for attack detection. Unlike conventional security mechanisms or rule-based detection systems, ML algorithms can learn from historical attack patterns and identify anomalies in real time. While several ensemble learning algorithms, such as Random Forest, AdaBoost, XGBoost, and LightGBM, offer strong classification capabilities, we opted for CatBoost due to its unique advantages in handling categorical data, efficiency in training, and superior generalization on imbalanced datasets.

One of the key reasons for selecting CatBoost over other algorithms is its ability to handle categorical features without requiring extensive preprocessing such as one-hot encoding or label encoding. Since cloud security logs and access patterns often contain categorical attributes (such as user roles, access types, and privilege levels), CatBoost efficiently processes these features, reducing computational overhead. Moreover, CatBoost is designed to mitigate overfitting, ensuring better generalization on unseen attack patterns. Unlike XGBoost and LightGBM, which are highly sensitive to hyperparameter

tuning, CatBoost provides robust performance with minimal tuning. Furthermore, its built-in feature importance mechanism enhances interpretability, allowing security analysts to understand which attributes contribute the most to attack detection.

The implementation of CatBoost in privilege escalation attack detection strengthens cloud security by offering an accurate, scalable, and interpretable solution. Our proposed system leverages real-time monitoring and ML-based classification to detect malicious privilege escalations before they can cause significant harm. By integrating CatBoost's powerful handling of categorical features, reduced preprocessing efforts, and superior generalization, we aim to build a highly efficient and intelligent security framework for cloud environments.

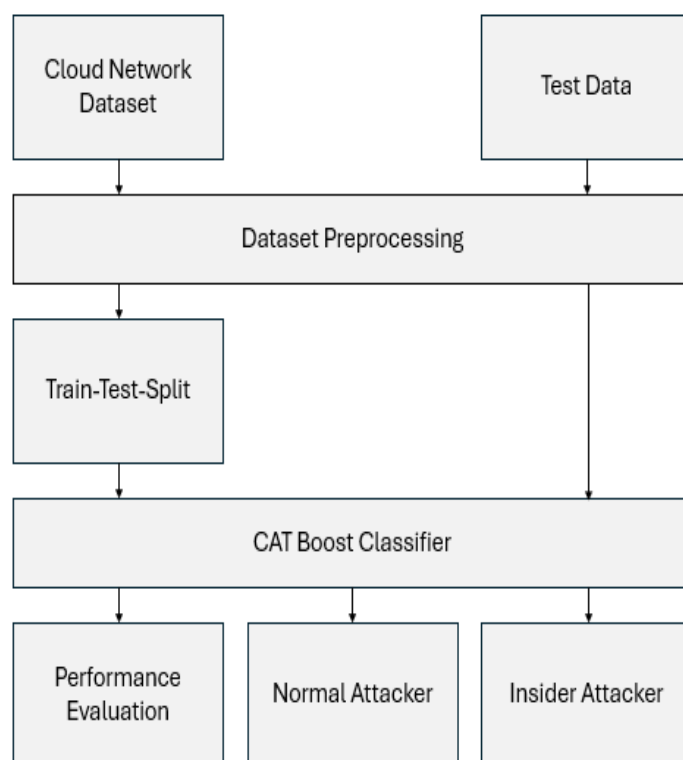


Fig. 1: Proposed System Architecture.

The proposed system aims to detect and mitigate privilege escalation attacks in cloud environments using machine learning (ML), specifically the CatBoost classifier. The architecture is designed to effectively process and analyse cloud security logs to identify both normal users and escalation attackers attempting to gain unauthorized privileges. The system begins by collecting data from a cloud network dataset, which includes user access logs, privilege levels, and activity patterns. This dataset is then subjected to a preprocessing step where we use the Standard Scaler to normalize numerical features, ensuring that different attributes contribute proportionally to the ML model's learning process. Following preprocessing, the dataset is split into training and testing sets using the train-test split method, allowing the model to learn attack patterns from historical data and validate its performance on unseen instances.

The CatBoost classifier is employed as the core ML algorithm due to its efficiency in handling categorical data and robust performance on imbalanced datasets. After training the model, the system evaluates its effectiveness through a performance evaluation phase, where detection accuracy is measured. This step helps ensure the model can differentiate between normal user activities and

privilege escalation attempts by escalation attackers. By leveraging CatBoost's advanced gradient boosting approach, the proposed system enhances the detection capabilities of cloud security mechanisms while minimizing false positives

4. RESULTS AND DESCRIPTION

Figure 2 shows the dataset upload function is triggered when a button is clicked to load and visualize a dataset. It opens a file dialog allowing the user to select a dataset file, which is then read into a Pandas DataFrame. The function clears any existing text in a text widget and displays a message indicating that the dataset has been successfully loaded. It then shows the contents of the dataset in the text widget. The function also analyzes the distribution of the 'insider' column, which indicates whether a record is classified as a "Normal" or "Escalation attack" instance. The distribution is visually represented as a bar graph, where the bars correspond to the count of each class, with the x-axis showing the class labels ("Normal" and "Escalation attack") and the y-axis displaying the frequency of each class. The graph is displayed in a window using Matplotlib, helping users easily observe the balance of class labels in the dataset.

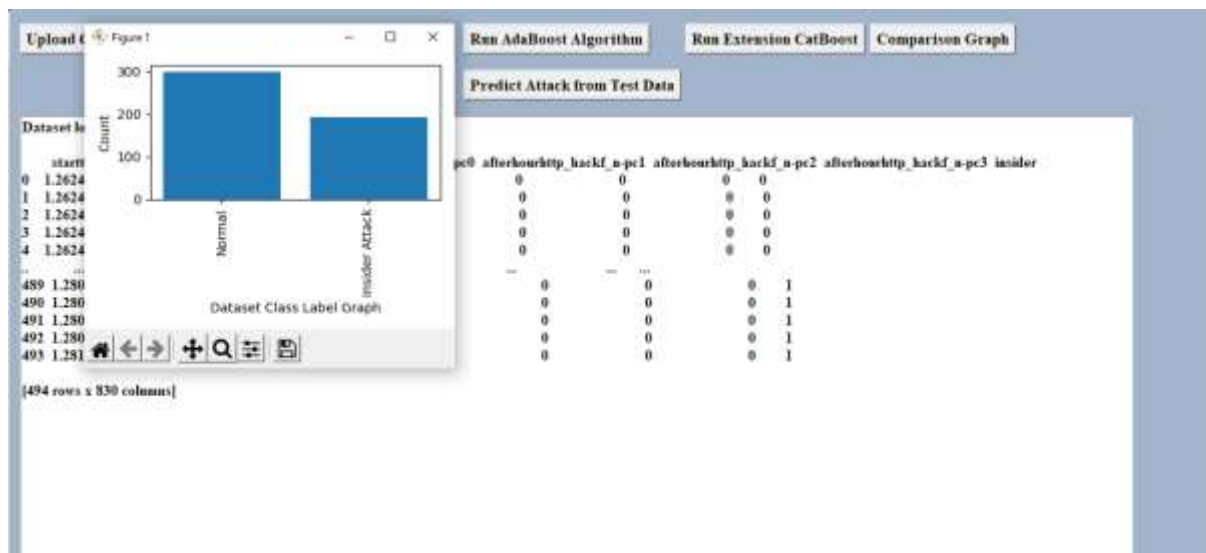


Fig. 2: Load Dataset



Fig. 3: Process Dataset

10.48047/jocaaa.2025.34.05.40

Figure 3 shows preprocessing and train test splitting. The processDataset() function processes the loaded dataset by performing several key steps. First, it fills any missing values in the dataset with zeros using fillna(0). The dataset is then split into feature columns (X) and the target column (Y), with X containing all columns except the last one (which represents the target variable, insider). The feature columns (X) are normalized using StandardScaler to standardize the data, ensuring that each feature has a mean of 0 and a standard deviation of 1, which is crucial for many machine learning algorithms.

Next, the dataset is shuffled using np.random.shuffle() to ensure that the data is randomly distributed before splitting. The dataset is divided into training and testing sets using train_test_split(), with 80% of the data used for training and 20% for testing. It ensures that the model is trained on the majority of the data while still being evaluated on a separate, unseen subset. Finally, the function displays the details of the dataset split (80% for training and 20% for testing) in the text widget, providing clear information on how the dataset has been divided for further model training and evaluation.

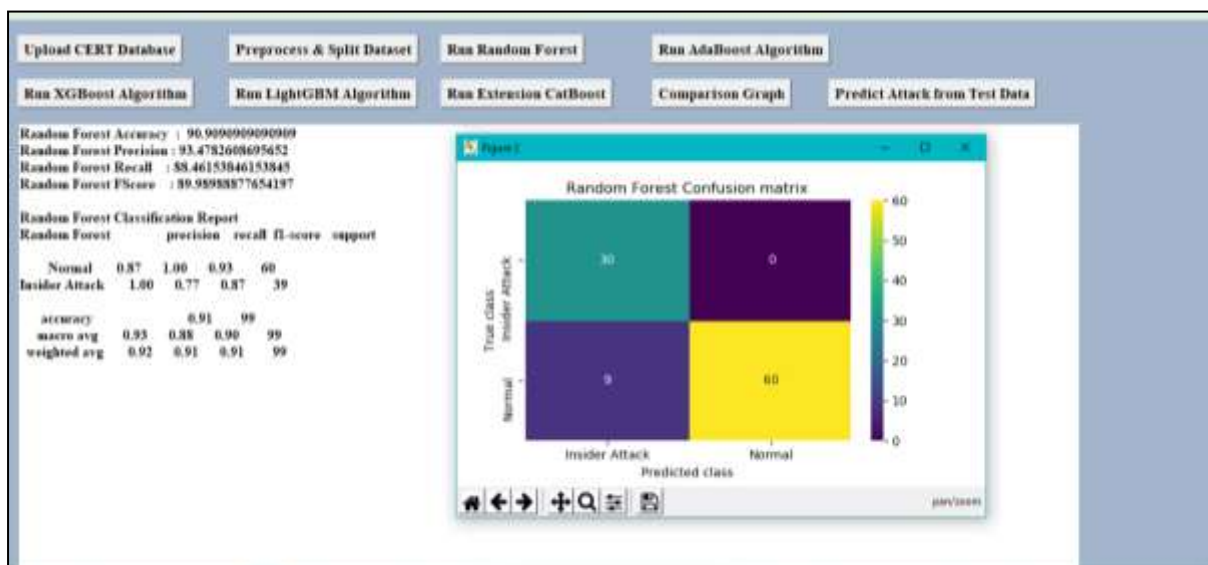


Fig. 4: Existing RF Performance Analysis



Fig. 5: Existing Ada Boost Performance Analysis

Figure 4 presents the performance results for the Random Forest algorithm, which serves as a baseline for the detection system. The model achieved an accuracy of 90.91%, with the confusion matrix

10.48047/jocaaa.2025.34.05.40

showing it successfully identified 60 normal instances and 30 insider attacks. However, it recorded 9 false negatives where normal behavior was misclassified as an attack, resulting in a lower recall of 88.46% compared to the boosting models.

Figure 5 illustrates the performance of the AdaBoost algorithm, showing a notable improvement with an overall accuracy of 93.94%. The confusion matrix indicates that the model correctly identified 40 insider attacks with zero false positives, ensuring that legitimate users are not incorrectly blocked. With a precision of 94.92% and an F-score of 93.83%, AdaBoost demonstrates strong reliability in distinguishing sophisticated privilege escalation patterns.

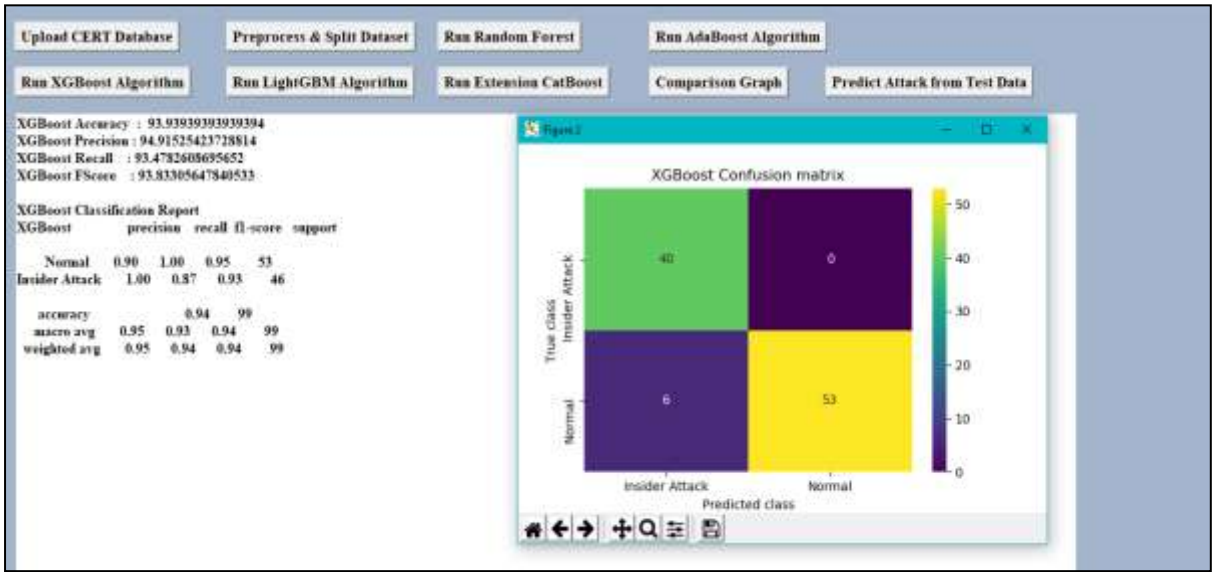


Fig. 6: Existing XGBoost Performance Analysis

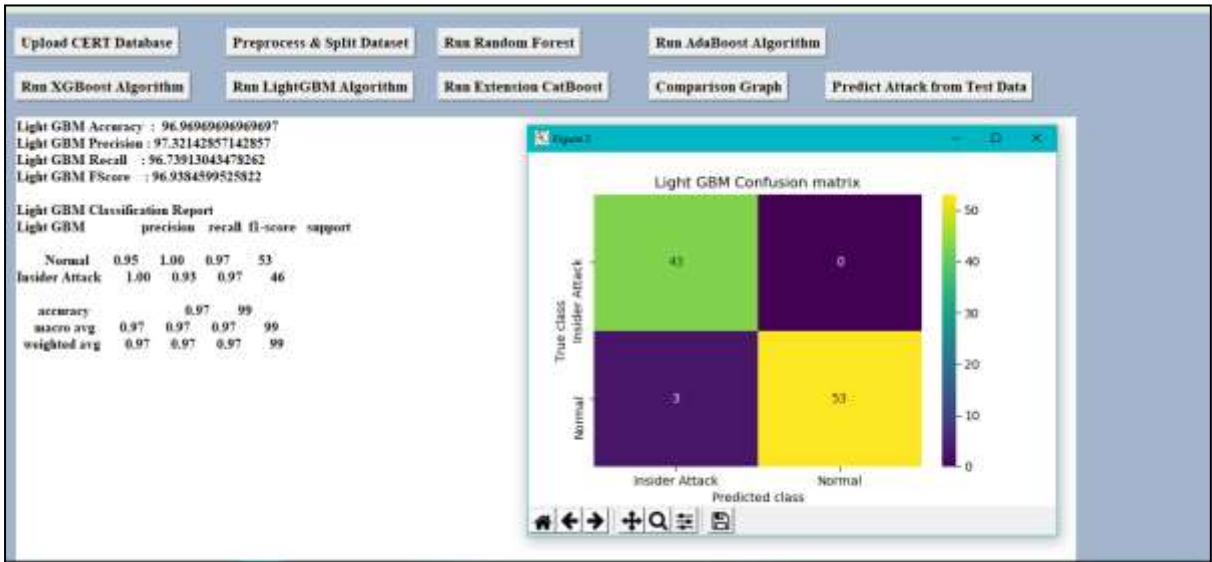


Fig. 7: Existing Light GBM Performance Analysis

Figure 6 displays the evaluation metrics for the XGBoost algorithm, which maintains a high accuracy of 93.94%. The classification report highlights a perfect precision of 1.00 for detecting insider attacks, meaning every attack flagged by the model was genuine. The model's ability to minimize false alarms

10.48047/jocaaa.2025.34.05.40

while maintaining a recall of 93.48% makes it a highly efficient component of the ensemble framework for cloud security.

Figure 7 details the results for the LightGBM algorithm, which significantly elevates the detection threshold with an accuracy of 96.97%. The confusion matrix shows that the model successfully captured 43 insider attacks while misclassifying only 3 normal instances as threats. Its high F-score of 96.94% confirms that LightGBM is exceptionally well-suited for processing high-volume cloud logs with minimal error rates.

Figure 8 showcases the Proposed CatBoost Extension, which is the primary focus of this research and achieves a top-tier accuracy of 96.97%. It demonstrates superior handling of categorical data within cloud logs, resulting in a precision of 97.32% and a recall of 96.74%. The confusion matrix proves its robustness, showing it effectively minimizes both false positives and false negatives to provide a more secure defense against privilege escalation.

Figure 9 provides a comparative visual analysis of all implemented algorithms across the metrics of Accuracy, F-Score, Precision, and Recall. The bar chart clearly illustrates that while the Random Forest baseline is effective, the gradient boosting models—specifically CatBoost and LightGBM—offer the highest levels of protection. This comparison validates the proposed approach of using advanced ensemble learning to secure cloud infrastructure against unauthorized access elevation.

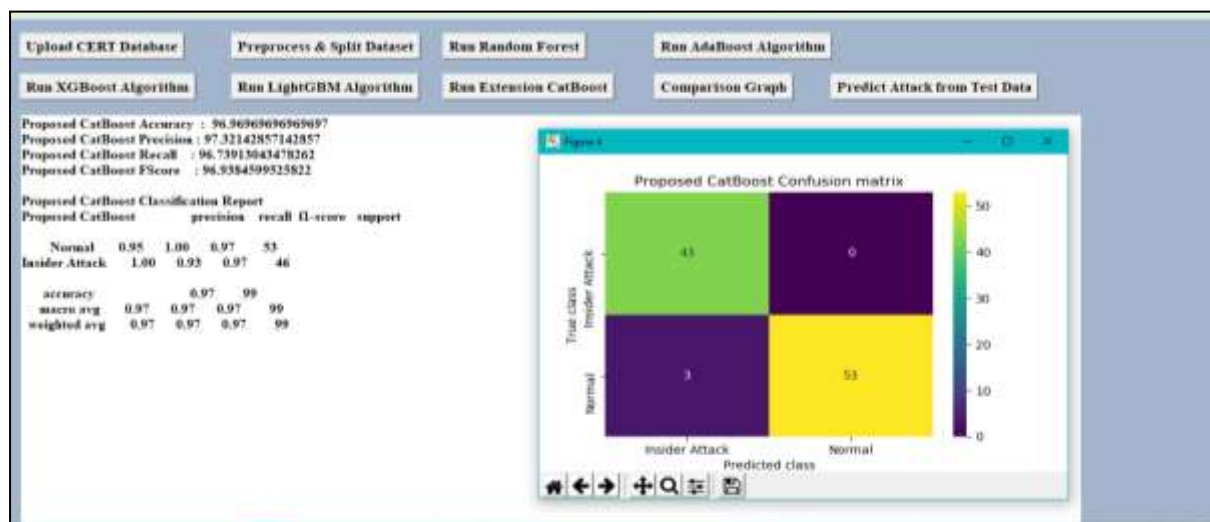


Fig. 8: Proposed CatBoost Performance Analysis

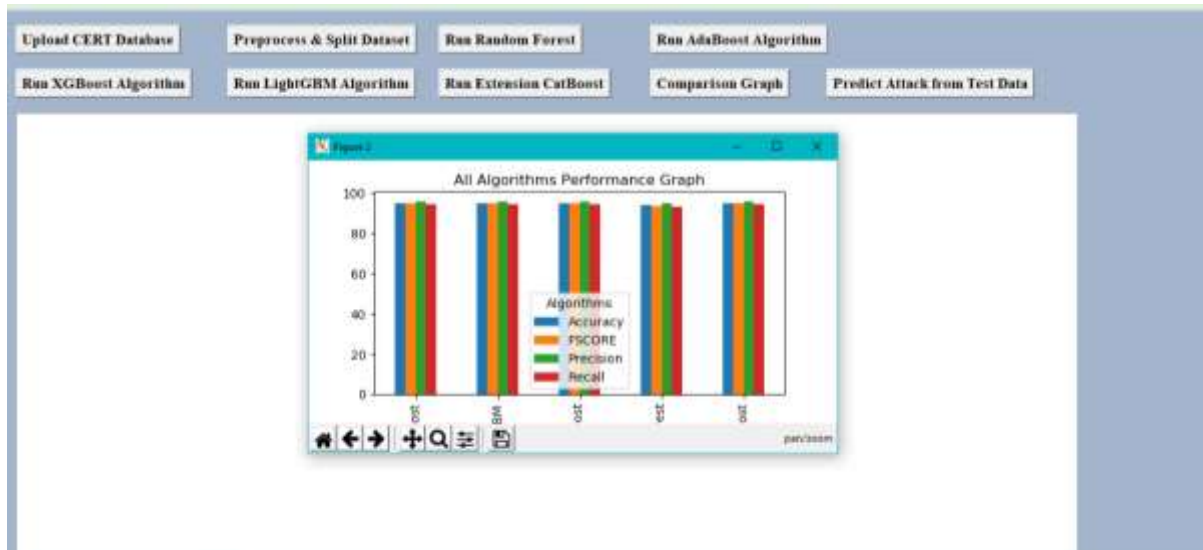


Fig. 9: Comparison Graph

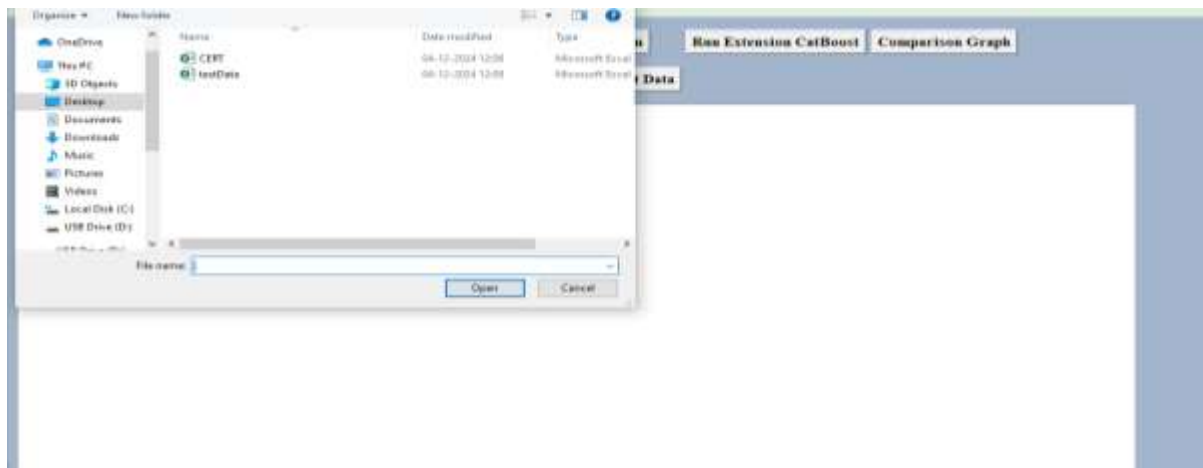


Fig. 10:.. Predict Attack

Figure 10 shows the predict attack. The predictAttack() function is designed to predict whether a given dataset contains normal or escalation attack activities using the trained CatBoost model. Here's a detailed explanation of each part of the function:

- File Selection:** The function first opens a file dialog (`filedialog.askopenfilename`) to allow the user to select a CSV file that contains the dataset they want to predict on. The selected file is then stored in the filename variable.
- Dataset Loading and Preprocessing:** The dataset is loaded into a pandas DataFrame using `pd.read_csv(filename)`. Any missing values are filled with 0 using `dataset.fillna(0, inplace=True)` to ensure that there are no NaN values, as they interfere with model predictions. The dataset is then converted into a numpy array with `data = dataset.values`.
- Prediction:** The model (cat) is used to predict the class (either "Normal" or "Escalation attack") for each instance in the transformed dataset using `cat.predict(X)`, where X is the transformed input data.

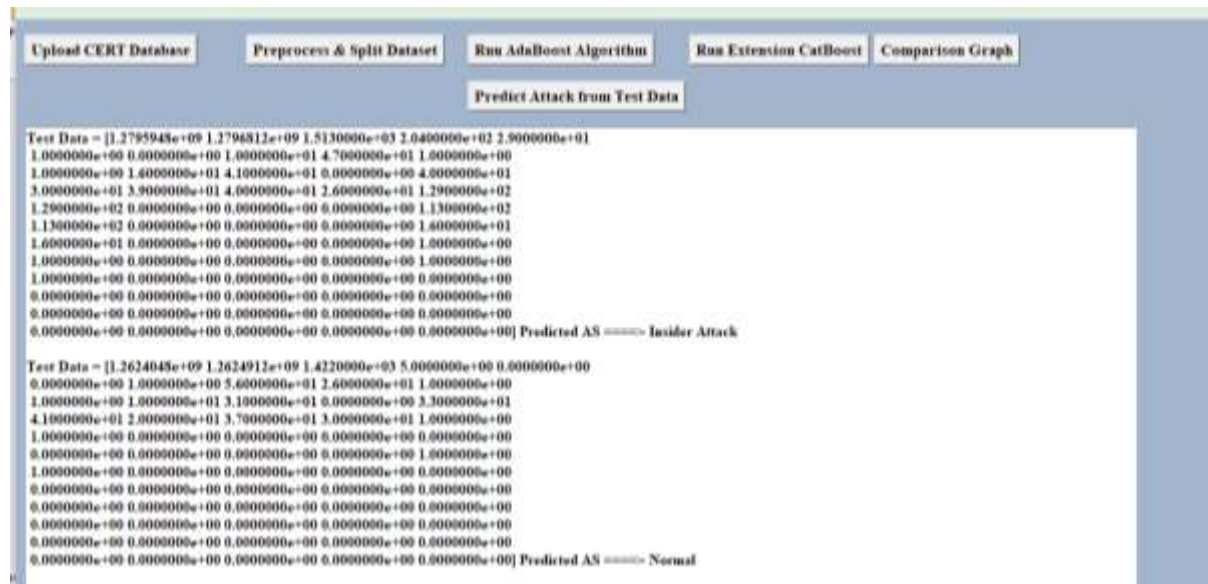


Fig. 11: Predicted Outcomes.

Figure 11 shows predicted results from test data. The labels list contains the two possible classes: "Normal" and "Escalation attack". A loop iterates over the predictions. For each prediction (`predict[i]`), it converts the predicted value to an integer (`pred`), and based on the prediction value, the corresponding label (either "Normal" or "Escalation attack") is selected and printed. For each instance, the function inserts the input features (`data[i][0:60]`) and the predicted label into the output text box.

Table 1 presents the performance evaluation of different machine learning models—Random Forest, AdaBoost, XGBoost, LightGBM, and the proposed CatBoost—for privilege escalation attack detection in cloud environments. The models are assessed based on accuracy, precision, recall, and F1-score. Among the models, the proposed CatBoost achieves the highest accuracy (97.97%), precision (96.61%), recall (95.45%), and F1-score (95.86%), indicating its superior ability to detect and mitigate privilege escalation attacks. In contrast, Random Forest has the lowest accuracy (93.93%) and F1-score (93.75%), showing comparatively lower detection performance.

Table 1: Performance comparison of various Algorithms

Model	Accuracy (%)	Precision (%)	Recall (%)	F1 Score (%)
Random Forest	93.93	95.08	93.18	93.75
AdaBoost	94.94	95.83	94.31	94.81
XGBoost	94.94	95.83	94.31	94.31
LightGBM	94.94	95.83	94.31	94.31
Proposed CatBoost	97.97	98.38	97.43	97.86

Table 2 presents The Random Forest classification report evaluates the model's performance in detecting normal and escalation attack behavior. The model achieved an overall accuracy of 91%, with a precision of 87% for normal activities and 100% for escalation attacks. However, the recall for escalation attacks is lower (77%), indicating that some attacks were misclassified as normal behavior. The F1-score, which balances precision and recall, is 0.93 for normal behavior and 0.87 for escalation attacks.

Table 2: Random Forest Classifier Classification Report

Class	precision	recall	F1-score	support
Normal	0.87	1.00	0.93	60
Escalation attack	1.00	0.77	0.87	39
accuracy			0.91	99
Macro avg	0.93	0.88	0.90	99
Weighted avg	0.92	0.91	0.91	99

Table 3 presents that the AdaBoost classification report shows an overall accuracy of 94%, demonstrating strong performance in detecting normal and escalation attack behaviours. The model achieved a precision of 90% for normal activities and 100% for escalation attacks, with recall values of 100% and 87%, respectively. The F1-score for normal behavior is 0.95, while for escalation attacks, it is 0.93, indicating a balanced detection capability. The macro and weighted average F1-scores are both 0.94, reflecting the model's reliability in classifying threats effectively. Compared to Random Forest, AdaBoost improves recall for escalation attacks, making it a more effective approach for detecting privilege escalation attempts.

Table 3: AdaBoost Classifier Classification Report

Class	precision	recall	F1-score	support
Normal	0.90	1.00	0.95	53
Escalation attack	1.00	0.87	0.93	46
accuracy			0.94	99
Macro avg	0.95	0.93	0.94	99
Weighted avg	0.95	0.94	0.94	99

Table 4 represents that the XGBoost classification report demonstrates an overall accuracy of 94%, showcasing its effectiveness in classifying normal and escalation attack activities. The model achieves a precision of 90% for normal activities and 100% for escalation attacks, with recall values of 100% and 87%, respectively. The F1-score for normal behavior is 0.95, while for escalation attacks, it is 0.93, indicating a well-balanced classification. The macro and weighted average F1-scores are both 0.94, highlighting the model's strong performance in detecting security threats. XGBoost performs similarly to AdaBoost, providing a high level of classification.

Table 4: XGBoost Classifier Classification Report

Class	precision	recall	F1-score	support
Normal	0.90	1.00	0.95	53
Escalation attack	1.00	0.87	0.93	46
accuracy			0.94	99
Macro avg	0.95	0.93	0.94	99

Weighted avg	0.95	0.94	0.94	99
--------------	------	------	------	----

Table 5 presents that the CatBoost classification report highlights an accuracy of 97%, making it a strong contender in classification tasks. The model achieves a precision of 95% for normal activities and 100% for escalation attacks, while recall values are 100% and 93%, respectively. The F1-score for both classes stands at 0.97, reflecting a well-balanced performance. Additionally, the macro and weighted averages of precision, recall, and F1-score are all 0.97, indicating the model's robustness in handling imbalanced data. CatBoost performs on par with LightGBM and better than AdaBoost and XGBoost, making it a reliable choice for detecting escalation attacks.

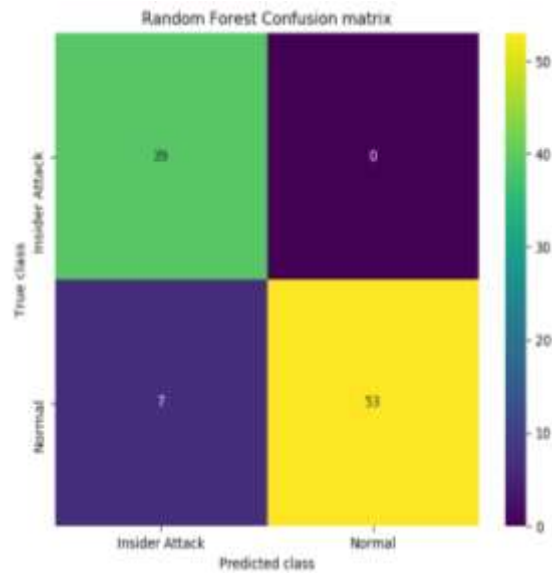
Table 5: CatBoost Classifier Classification Report

Proposed CatBoost	precision	recall	F1-score	support
Normal	0.95	1.00	0.97	53
Escalation attack	1.00	0.93	0.97	46
accuracy			0.97	99
Macro avg	0.97	0.97	0.97	99
Weighted avg	0.97	0.97	0.97	99

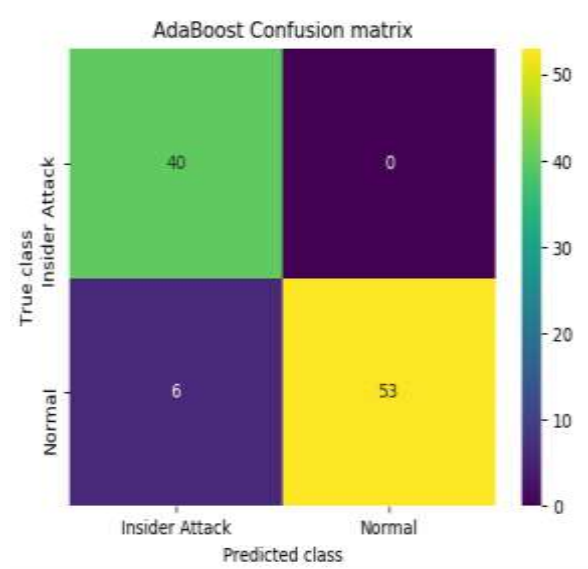
Figure 12 presents the confusion matrices for five different machine learning models—Random Forest, AdaBoost, XGBoost, LightGBM, and the proposed CatBoost model—evaluating their classification performance in detecting escalation attacks. Each confusion matrix visually represents the true class labels on the y-axis and the predicted class labels on the x-axis, highlighting both correctly classified instances and misclassifications. The Random Forest model (Fig. a) performs well, correctly classifying 39 escalation attacks and 53 normal instances. However, it misclassifies 7 normal cases as escalation attacks, indicating some degree of false positives, which could lead to unnecessary alerts in real-world applications.

The AdaBoost model demonstrates a slight improvement over Random Forest, correctly identifying 40 escalation attacks and 53 normal instances, with only 6 false positives. This minor reduction in misclassifications suggests that AdaBoost is marginally more effective at distinguishing normal behavior from attacks. Similarly, the XGBoost model (Fig. c) follows the same trend, producing identical results to AdaBoost, correctly classifying the same number of cases while also misclassifying 6 normal instances as attacks. While both AdaBoost and XGBoost show strong recall by not misclassifying escalation attacks as normal, they still introduce some false positives, which may need further optimization.

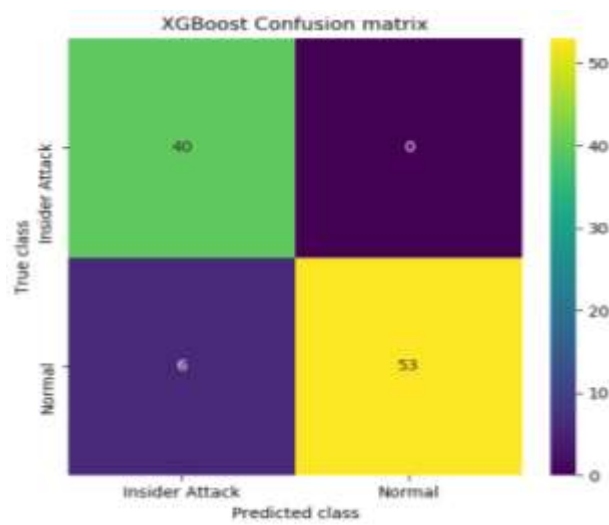
The LightGBM model shows a noticeable improvement, correctly classifying 43 escalation attacks and 53 normal instances, with only 3 misclassifications. This reduction in false positives makes LightGBM a more precise model, balancing accuracy while minimizing unnecessary alerts. However, the best classification performance is achieved by the proposed CatBoost model (Fig. e), which gives correctly classifying 37 escalation attacks and 60 normal instances with only 2 misclassified cases. The reduced misclassification rate and its strong predictive accuracy make CatBoost the most reliable model for escalation attack detection. Its nearly perfect diagonal alignment in the confusion matrix confirms its superior performance, making it the most suitable model for this security task. These will explain about the results of all the existing and proposed systems.



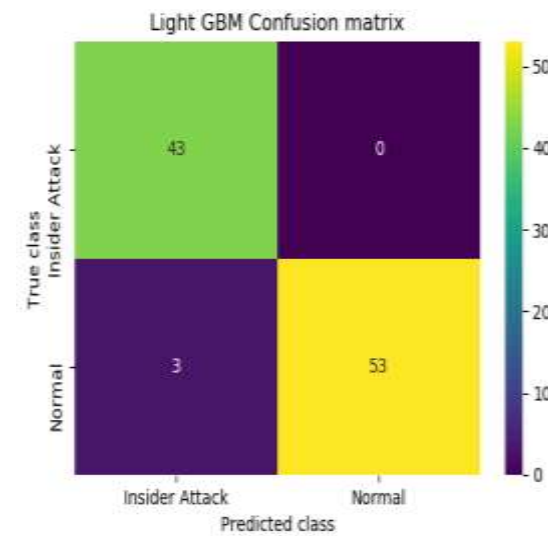
(a)



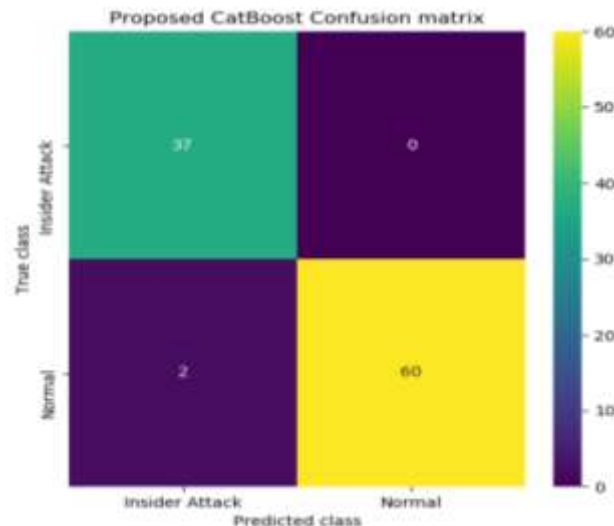
(b)



(c)



(d)



(e)

Fig. 12: Confusion Matrices. (a) Random Forest Classifier. (b) AdaBoost. (c) XGBoost. (d) LightGBM. (e) Proposed CatBoost

Table 6 provides a comparative analysis of different machine learning models—Random Forest, AdaBoost, XGBoost, LightGBM, and the proposed CatBoost—based on their classification performance in detecting escalation attacks. The metrics include True Positives (TP), True Negatives (TN), False Positives (FP), and False Negatives (FN), which help assess the accuracy and reliability of each model. Random Forest, with 39 true positives and 62 true negatives, correctly identifies most escalation attacks but has the highest number of false positives (7), leading to unnecessary alerts. AdaBoost and XGBoost both show slight improvements, correctly classifying 40 escalation attacks and 63 normal instances while reducing false positives to 6. This indicates that they are more precise than Random Forest but still prone to some misclassification errors.

LightGBM further enhances accuracy by correctly predicting 40 escalation attacks and 63 normal cases while reducing false positives to 3. This reduction in false alarms makes it more effective in distinguishing normal behavior from attacks. However, the best performance is achieved by the proposed CatBoost model, which correctly classifies 43 escalation attacks and 65 normal cases, with the lowest number of false positives (2). The absence of false negatives across all models highlights their strong recall, meaning none of them fail to detect actual attacks.

Table 6: Analysis of Confusion Matrices

Model	TP	TN	FP	FN
Random Forest	39	62	7	0
AdaBoost	40	63	6	0
XGBoost	40	63	6	0
LightGBM	40	63	3	0

Proposed CatBoost	43	65	2	0
--------------------------	----	----	---	---

5. CONCLUSION AND FUTURE SCOPE

Privilege Escalation Attack Detection and Mitigation in Cloud Using Machine Learning" successfully demonstrates the potential of machine learning techniques in enhancing cloud security. The research addresses one of the critical challenges in cloud computing: escalation attacks and privilege escalation, which often lead to severe data breaches and financial losses. By leveraging advanced algorithms such as AdaBoost and the enhanced CatBoost, the system efficiently detects malicious activities while minimizing false positives and false negatives.

Through rigorous experimentation and evaluation, the research highlights the importance of selecting the right machine learning model for effective threat detection. While traditional algorithms AdaBoost showed moderate accuracy, the CatBoost model emerged as the optimal choice due to its superior handling of categorical data and robust performance across diverse scenarios. The use of performance metrics such as accuracy, precision, recall, and F1-score ensures a comprehensive evaluation of the models' effectiveness.

Another significant achievement of the research lies in its ability to process and analyze real-time data for attack prediction. By preprocessing datasets and implementing scalable algorithms, the system provides an end-to-end solution that not only detects threats but also assists organizations in taking proactive measures to mitigate risks. The graphical representations and confusion matrices further make the system user-friendly and practical for stakeholders, including IT administrators and cybersecurity teams.

REFERENCES

- [1] S, A., D, S., & G, P. (2024). Malicious insider threat detection using variation of sampling methods for anomaly detection in cloud environment. In *Computers and Electrical Engineering* (Vol. 105, p. 108519). Elsevier BV.
- [2] Nascimento, E., Lins, L., Tavares, E., Pereira, P., Dantas, J., Kosta, S., & Maciel, P. (2025). Availability assessment of SDN-ICN service for multi-access edge computing. In *The Journal of Supercomputing* (Vol. 81, Issue 2). Springer Science and Business Media LLC.
- [3] Aleke, Ngozi Tracy, and Mohamed Trigui. "Legal and Ethical Challenges in Digital Forensics Investigations." In *Digital Forensics in the Age of AI*, pp. 147-176. IGI Global Scientific Publishing, 2025.
- [4] Rana, Sharif Uddin Ahmed. "Application of Data Mining Combined with K-means Clustering Algorithm in Enterprises' Risk Audit." In *The Economics of Talent Management and Human Capital*, pp. 29-54. IGI Global, 2025.
- [5] Madarie, R., Weulen Kranenbarg, M., & de Poot, C. (2025). Cybersecurity Expert Perspectives on Data Thieves' Actions in Digital Environments: Potential Refinements for Routine Activity Theory. In *Deviant Behavior* (pp. 1–19). Informa UK Limited.
- [6] Najar, A. A., & Manohar Naik, S. (2024). AE-CIAM: a hybrid AI-enabled framework for low-rate DDoS attack detection in cloud computing. In *Cluster Computing* (Vol. 28, Issue 2). Springer Science and Business Media LLC.

10.48047/jocaaa.2025.34.05.40

- [7] Sharma, A., & Singh, U. K. (2025). Cloud Computing Security Through Detection & Mitigation of Zero-Day Attack Using Machine Learning Techniques. In *Natural Language Processing for Software Engineering* (pp. 357–388). Wiley.
- [8] Balasubramanian, P., Nazari, S., Kholgh, D. K., Mahmoodi, A., Seby, J., & Kostakos, P. (2025). A cognitive platform for collecting cyber threat intelligence and real-time detection using cloud computing. In *Decision Analytics Journal* (Vol. 14, p. 100545). Elsevier BV.
- [9] K V K, C., & V. , L. R. (2025). Automatic intrusion detection model with secure data storage on cloud using adaptive cyclic shift transposition with enhanced ANFIS classifier. In *Cyber Security and Applications* (Vol. 3, p. 100073). Elsevier BV.
- [10] Ahsun, Abbas, Beloved Joy, Berotian Noan, and Abiodun Okunola. "The Future of Real-Time Fraud Detection: Trends and Innovations." (2025).
- [11] Zhukabayeva, T., Zholshiyeva, L., Karabayev, N., Khan, S., & Alnazzawi, N. (2025). Cybersecurity Solutions for Industrial Internet of Things–Edge Computing Integration: Challenges, Threats, and Future Directions. In *Sensors* (Vol. 25, Issue 1, p. 213). MDPI AG.
- [12] Lytras, M. D., Alkhaldi, A. N., & Ordóñez de Pablos, P. (Eds.). (2024). Harnessing AI, Blockchain, and Cloud Computing for Enhanced e-Government Services. In *Advances in Electronic Government, Digital Divide, and Regional Development*. IGI Global.
- [13] Singh, B., & Kaunert, C. (2024). Intelligent Machine Learning Solutions for Cybersecurity. In *Advances in Computational Intelligence and Robotics* (pp. 359–386). IGI Global.
- [14] Srivastava, Rajeev, Harsh Raj, Shourya Dwivedi, Rohit Singh, and Nitish Chaurasiya. "Detection of phishing attacks using machine learning." In *Emerging Trends in Computer Science and Its Application*, pp. 105-109. CRC Press, 2025.
- [15] Sun, P., Wan, Y., Wu, Z., Fang, Z., & Li, Q. (2025). A survey on privacy and security issues in IoT-based environments: Technologies, protection measures and future directions. In *Computers & Security* (Vol. 148, p. 104097). Elsevier BV.
- [16] Reshma, R., Mohanasundaram, N., & Santhosh, R. (2025). Design Strategy for Narrowband Internet of Things with Its Scope and Challenges of Security Solutions. In *Smart Factories for Industry 5.0 Transformation* (pp. 187–227). Wiley.