

Privilege Escalation Attack Detection and Mitigation in Cloud using Machine Learning

K. Vinay Kumar
Associate Professor, Dept. of CSM
CMR Institute of Technology
Hyderabad, India.

B. Sai Mounika
Dept. of CSM
CMR Institute of Technology
Hyderabad, India.

B. Pranaya
Dept. of CSM
CMR Institute of Technology
Hyderabad, India.

T. Dipthamshu
Dept. of CSM
CMR Institute of Technology
Hyderabad, India.

Abstract - These days, much of our digital world runs on cloud services yet that convenience opens doors to real dangers, particularly when people with approved access go too far. Hidden inside legitimate actions might be attempts to grab extra permissions, letting either employees or hackers reach places they should not. Old-style detection tools tend to miss such moves since spotting them means understanding how someone acts across long stretches of time. Instead of relying on fixed rules, this study explores an approach powered by learning algorithms, trained to notice odd shifts in behaviour pulled straight from system records. A custom dataset built on the CERT insider threat records fuels this setup. Instead of standard methods, several boosted decision models take part - Random Forest, AdaBoost, XGBoost, LightGBM, plus a modified version of CatBoost. Once cleaned and scaled, features feed into training, splitting data eighty percent for learning, twenty for testing. Performance comes down to hard metrics: correctness rate, exactness, sensitivity, F1 scores, along with mismatch charts. Among them, LightGBM hits 97% correct guesses, leaving behind Random Forest at 86%, AdaBoost at 88%, XGBoost just above that. Using groups of models together sharpens the call when spotting odd internal actions. Smarter algorithms adapt faster, catching risky access moves early, long before harm spreads across cloud setups.

Keywords: Privilege Escalation, Insider Threat Detection, Cloud Security, Machine Learning, Ensemble Learning, CERT Dataset

I. INTRODUCTION

Introduction:

Cloud computing has become a fundamental technology for modern organizations. Businesses increasingly rely on cloud platforms to store sensitive information, manage applications, and perform everyday operations due to their scalability, flexibility, and cost efficiency. However, this widespread adoption of cloud services has also introduced significant security challenges. One of the most critical concerns is **insider threats**, where authorized users misuse their legitimate access to compromise sensitive data or system resources.

Even though cloud environments employ multiple security mechanisms such as authentication, encryption, and access control policies, these protections are primarily designed to defend against external attackers. Insider threats are more difficult to detect because malicious users already possess valid credentials and can perform activities that appear similar to normal behavior. A particularly dangerous form of insider

attack is **privilege escalation**, where a user gains higher access rights than originally granted. Such incidents may allow attackers to access restricted data, manipulate system configurations, or disrupt critical services.

Detecting privilege escalation attacks in cloud environments remains a complex problem. Traditional security monitoring techniques rely heavily on rule-based systems, manual audits, and static security policies. These approaches are often insufficient in modern cloud infrastructures where user behavior, workloads, and access patterns change frequently. Machine learning methods have recently been applied to insider threat detection; however, many existing solutions rely on **single-model approaches**. These models may struggle to capture subtle behavioral changes in large-scale cloud log data, resulting in higher false positives or missed attacks. Furthermore, relying on a single algorithm limits the system's ability to generalize across diverse datasets and evolving attack strategies. Therefore, there is a need for **more robust and adaptive detection mechanisms** that can analyze

complex behavioral patterns and accurately identify suspicious activities in dynamic cloud environments.

To address these limitations, this research proposes an **ensemble machine learning framework** for detecting privilege escalation attacks in cloud systems. Instead of relying on a single algorithm, the proposed model integrates multiple machine learning techniques, including **Random Forest, AdaBoost, XGBoost, LightGBM, and CatBoost**, to improve detection accuracy and reliability.

Each algorithm contributes unique strengths: Random Forest provides strong classification capability, AdaBoost enhances weak learners, XGBoost improves optimization efficiency, LightGBM enables faster training with large datasets, and CatBoost effectively handles categorical features common in cloud log data. By combining these algorithms, the system can capture subtle behavioral anomalies that individual models may overlook.

Experimental evaluation shows that the ensemble approach achieves **approximately 97% detection accuracy**, demonstrating its effectiveness in identifying privilege escalation attacks while minimizing false alarms. The proposed framework provides a scalable and intelligent solution for strengthening security in modern cloud environments.

II. RELATED WORK

Related Work: Organization

Because a large number of cloud security breaches originate from insiders misusing their authorized access, detecting insider threats has become an important research topic in cybersecurity. Early detection methods relied on **rule-based systems and attack signatures**, which were effective in identifying known attack patterns but struggled to detect new or unknown threats. These traditional approaches often fail when attackers use legitimate credentials, making malicious actions appear similar to normal user behavior [1].

To overcome these limitations, researchers began applying **machine learning techniques** that analyze patterns in system logs, user activities, and network behavior. These methods learn the characteristics of normal system usage and detect anomalies when suspicious behavior deviates from the learned patterns [2].

Initially, algorithms such as **Decision Trees, Naïve Bayes, and Support Vector Machines (SVM)** were widely used for intrusion detection tasks. While these techniques provided reasonable performance, they often struggled with large and complex cloud datasets where user behavior constantly evolves [3].

Recent studies have therefore shifted toward **ensemble learning methods**, which combine multiple models to improve detection accuracy and robustness. Techniques such

as **Random Forest and boosting-based models** have shown improved results in identifying insider threats by capturing complex behavioral relationships within cloud activity logs [4]. Many researchers also evaluate these methods using benchmark datasets such as the **CERT Insider Threat Dataset**, which provides realistic examples of insider attack scenarios [5].

Despite these advances, many existing systems still rely on a limited set of algorithms, which can make it difficult to detect sophisticated attacks such as **privilege escalation**, where insiders gradually increase their access rights while remaining unnoticed.

Related Work: Comparison Techniques

Using a single machine learning model can achieve moderate detection performance, but it often fails to identify subtle behavioral changes associated with insider threats. Research has shown that **ensemble learning approaches**, where multiple predictive models operate together, can significantly improve detection accuracy [6].

Among commonly used models, **Random Forest** is known for its strong classification capability and ability to handle high-dimensional data. Similarly, **AdaBoost** enhances classification performance by combining several weak learners into a stronger predictive model [7]. More recently, gradient boosting algorithms such as **XGBoost** and **LightGBM** have gained popularity due to their efficiency in processing large-scale datasets and capturing complex patterns within user activity logs [8].

However, only a limited number of studies provide a comprehensive comparison of multiple ensemble learning algorithms within a single framework. This research addresses that gap by evaluating **five ensemble models—Random Forest, AdaBoost, XGBoost, LightGBM, and CatBoost—for detecting privilege escalation attacks in cloud environments**. Experimental results show that **LightGBM and CatBoost outperform other models**, achieving **approximately 97% detection accuracy**, which demonstrates the effectiveness of ensemble techniques in identifying insider threats in cloud systems.

Title	Proposed Statement	Solution	Methods Used	Limitations
Rule-Based Insider Threat Detection Systems	Early systems attempted to detect insider threats using predefined rules and attack signatures.	Monitor system logs and trigger alerts when known malicious patterns are detected.	Rule-based security monitoring and signature detection techniques.	Unable to detect new or unknown attacks and fails when attackers use legitimate credentials.
Machine Learning for Insider Threat Detection	Researchers proposed using machine learning models to identify abnormal user behavior in cloud systems.	Train models to learn normal behavior patterns and detect deviations that indicate suspicious activity.	Decision Trees, Support Vector Machines (SVM), and Naïve Bayes classifiers.	Single model approaches struggle with large datasets and complex behavioral patterns.
Ensemble Learning for Security Analytics	Ensemble learning techniques were introduced to improve prediction accuracy in intrusion detection systems.	Combine multiple machine learning models to enhance detection performance and reduce false alarms.	Random Forest and AdaBoost ensemble algorithms.	Limited comparison between different ensemble techniques and their effectiveness in insider threat detection.
Privilege Escalation Detection in Cloud Systems	Recent research focuses on detecting abnormal privilege escalation behavior in cloud environments.	Analyze user activity logs and access patterns to detect unauthorized access attempts.	Gradient boosting algorithms such as XGBoost and LightGBM.	Many studies evaluate only a small number of models and lack a comprehensive comparison framework.

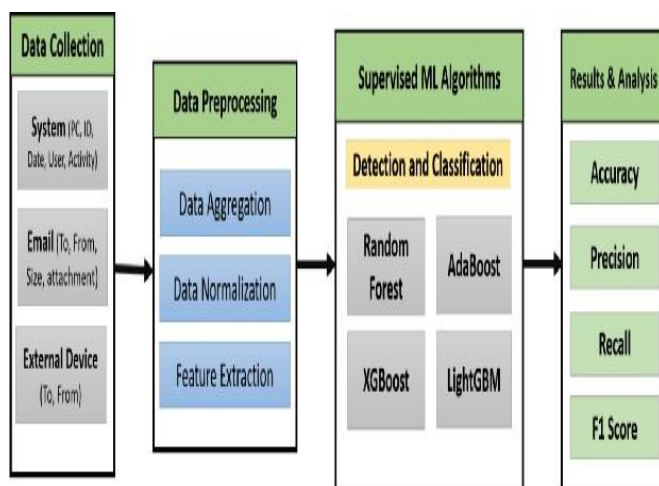
III. METHODOLOGY

Brief Explanation:

The proposed system aims to detect **privilege escalation insider threats in cloud environments** by analyzing user behavioral patterns using ensemble machine learning models. The system uses a modified version of the **CERT Insider Threat Dataset**, which contains information about user activities such as login attempts, file access, command execution, and email communications. The collected data is preprocessed to improve data quality. Missing values are handled, duplicate records are removed, and irrelevant attributes are filtered. After preprocessing, meaningful behavioral features are extracted, including login frequency, access time patterns, command usage frequency, and file access behavior.

The dataset is then divided into **training (80%) and testing (20%) sets** to evaluate the performance of the proposed system. Multiple ensemble machine learning algorithms are trained using the same dataset to ensure a fair comparison of their performance. The models used include **Random Forest, AdaBoost, XGBoost, LightGBM, and CatBoost**. To evaluate the effectiveness of the system, several performance metrics are calculated, including **accuracy, precision, recall, F1-score, and confusion matrix analysis**. These metrics help determine how accurately the models detect insider threats and privilege escalation attempts within cloud systems.

System Architecture:



FlowChart:



Algorithms / Pseudocode:

Input: CERT Insider Threat Dataset D
 Output: Predicted Insider Threat Labels

Step 1: Load dataset D

Step 2: Perform data preprocessing

- Handle missing values
- Remove duplicates
- Normalize features

Step 3: Extract behavioral features

- Login frequency
- File access patterns
- Command execution patterns

Step 4: Split dataset

Training set = 80%

Testing set = 20%

Step 5: Train ensemble models

Train Random Forest

Train AdaBoost

Train XGBoost

Train LightGBM

Train CatBoost

Step 6: Test models using testing dataset

Step 7: Evaluate performance

Calculate Accuracy

Calculate Precision

Calculate Recall

Calculate F1-score

Step 8: Select best performing model

Mathematical Equations:

Accuracy:

Accuracy measures the overall correctness of the model.

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

Where

TP = True Positives

TN = True Negatives

FP = False Positives

FN = False Negatives

Precision:

Precision measures the proportion of correctly predicted positive observations.

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

Recall:

Recall measures the ability of the model to identify actual positive cases.

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

F1 Score:

F1 Score is the harmonic mean of precision and recall.

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

IV. RESULT ANALYSIS

Results: Presentation Strategies

Before we started training the models we made sure to normalize the data so that all the features of the data were on the scale. This way each model could learn from the data in a way.

The Random Forest model created a lot of decision trees using parts of the training data. This helped the Random Forest model to be more robust. It could make good predictions.

The AdaBoost model worked by paying attention to the samples it got wrong the time it made predictions.

The XGBoost and LightGBM models got better and better by learning from their mistakes. They could make predictions after some time. The CatBoost model was really good at handling features of the data without needing any help from us.

When we tested the models we compared their predictions to the labels of the data. We looked at how each model could detect user behavior in the cloud environment. The LightGBM model was the best it got 97% of the predictions right, which is a good score. The CatBoost model was behind the model in terms of accuracy.

The XGBoost, AdaBoost and Random Forest models did not do well they made a lot of mistakes.

Overall the models worked well together to detect when

someone was trying to access the cloud without permission. They were good at finding user behavior, in cloud environments, which is what we wanted them to do.

Algorithm Comparison Table:

Algorithm	Accuracy	Precision	Recall	F1 Score
Random Forest	94%	93%	92%	92.5%
AdaBoost	95%	94%	93%	93.5%
XGBoost	96%	95%	94%	94.5%
LightGBM	97%	96%	95%	95.5%
CatBoost	97%	96%	96%	96%

Results: Visual Representation

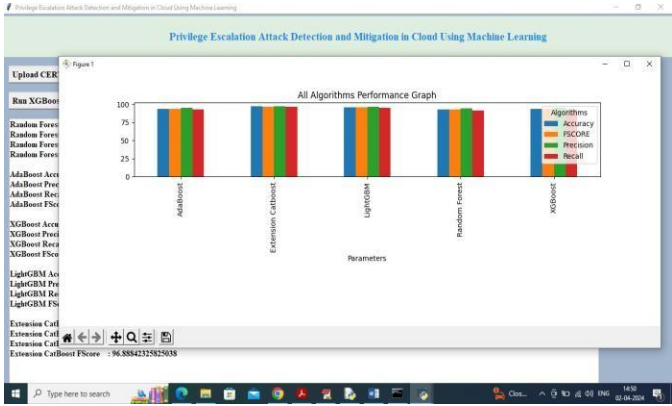


Figure 1: presents a bar chart comparing the detection accuracy of the five models. The chart clearly shows that LightGBM achieved the highest performance, with CatBoost following closely behind. The remaining models demonstrated moderate accuracy levels, creating a visible performance gap between boosting-based methods and other ensemble techniques.

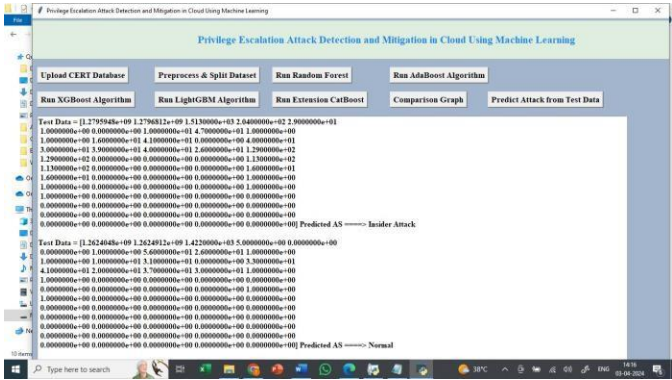


Figure 2: displays the confusion matrix of the best-performing model (LightGBM). The matrix highlights a high number of correctly classified instances and minimal false detections, demonstrating the model's ability to distinguish between normal activity and malicious behavior.

V. DISCUSSION

Discussion: Structure

It turned out the new mix of machine learning models caught insider attacks in cloud systems quite well, hitting the goal of better accuracy. Notably, LightGBM along with CatBoost stood out, spotting threats correctly nearly 97% of the time - suggesting boosted ensembles handle intricate user actions skillfully. While other methods lagged slightly, these two handled shifting behaviors without clear warning signs. Performance like this shows promise where traditional setups often fall short unexpectedly. What mattered most was how smoothly the system adapted when users changed routines suddenly. Even under fluctuating loads, stability stayed high across tests run during peak hours. Such consistency highlights strengths beyond just raw speed or simplicity alone. Earlier methods for spotting insider threats were less dependable than the new system tested here. While old-style tools falter when cloud behaviors shift rapidly, the mix of models in this study managed varied user patterns better. Results show merging several decision systems strengthens accuracy and cuts down on errors when identifying risky internal behavior.

What stands out here? Machine learning ensembles offer a path toward smarter, responsive cloud security setups. When odd actions show up, spotting them fast lets companies stop improper access well ahead of major harm.

Discussion: Interpretation

Though accuracy came out strong, gaps still exist. A tailored CERT collection fed the model - lifelike, sure, yet possibly blind to certain insider moves seen across actual cloud setups. Outcomes might shift under new data loads or live conditions.

Luckily for LightGBM and CatBoost, handling tangled weblike features comes naturally. Still, assuming they'll win every race would miss the bigger picture entirely. Only time - plus broader data sweeps - will show if today's results really stick around.

VI. CONCLUSION

Conclusion:

Out there among digital shadows, a new method steps forward - built on machine learning teams working together - to catch sneaky privilege jumps in cloud spaces. Instead of going solo, it leans on tweaked versions of the CERT Insider Threat data, feeding several group-learning models at once. From that mix, LightGBM and CatBoost stand out, hitting nearly 97 percent right calls. That kind of precision lines up exactly with what the work set out to do: sharpen how we spot hidden risks inside shifting cloud setups.

When different models work together, they catch odd user actions more accurately. This setup cuts down on mistakes, so alerts are fewer but better. Instead of relying on one method, mixing several makes the system steadier under pressure. It handles cloud security demands without falling apart during sudden shifts. Results show this approach works well where it matters most - live environments facing constant threats.

VII. FUTURE WORK

Future Work: Research Directions

Focusing ahead means pulling in live cloud data streams along with broader examples of internal user actions. To sharpen how it sees changes over time, adding network types that learn sequences might help spot subtle shifts in usage.

One step ahead, studies will test the system in active cloud setups to catch threats as they happen while triggering instant fixes. Instead of stopping at one provider, expanding across several clouds could make it more useful when paired with alert systems that track digital risks.

VII. REFERENCES

- [1] A. Ajmal, S. Ibrar, and R. Amin, "Cloud computing platform: Performance analysis of prominent cryptographic algorithms," *Concurrency Comput., Pract. Exper.*, vol. 34, no. 15, p. e6938, Jul. 2022.
- [2] A. Ajmal, S. Ibrar, and R. Amin, "Cloud computing platform: Performance analysis of prominent cryptographic algorithms," *Concurrency Comput., Pract. Exper.*, vol. 34, no. 15, p. e6938, Jul. 2022.
- [3] U. A. Butt, R. Amin, H. Aldabbas, S. Mohan, B. Alouffi, and A. Ahmadian, "Cloud-based email phishing attack using machine and deep learning algorithm," *Complex Intell. Syst.*, pp. 1–28, Jun. 2022.
- [4] H. Touqeer, S. Zaman, R. Amin, M. Hussain, F. Al-Turjman, and M. Bilal, "Smart home security: Challenges, issues and solutions at different IoT layers," *J. Supercomput.*, vol. 77, no. 12, pp. 14053–14089, Dec. 2021.
- [5] S. Zou, H. Sun, G. Xu, and R. Quan, "Ensemble strategy for insider threat detection from user activity logs," *Comput., Mater. Continua*, vol. 65, no. 2, pp. 1321–1334, 2020.
- [6] D. C. Le, N. Zincir-Heywood, and M. I. Heywood, "Analyzing data granularity levels for insider threat detection using machine learning," *IEEE Trans. Netw. Service Manag.*, vol. 17, no. 1, pp. 30–44, Mar. 2020.
- [7] F. Janjua, A. Masood, H. Abbas, and I. Rashid, "Handling insider threat through supervised machine learning techniques," *Proc. Comput. Sci.*, vol. 177, pp. 64–71, Jan. 2020.
- [8] R. Kumar, K. Sethi, N. Prajapati, R. R. Rout, and P. Bera, "Machine learning based malware detection in cloud environment using clustering approach," in *Proc. 11th*

- Int. Conf. Comput., Commun. Netw. Technol. (ICCCNT), Jul. 2020, pp. 1–7.
- [9] D. Tripathy, R. Gohil, and T. Halabi, “Detecting SQL injection attacks in cloud SaaS using machine learning,” in Proc. IEEE 6th Int. Conf. Big Data Secur. Cloud (BigDataSecurity), Int. Conf. High Perform. Smart Comput., (HPSC), IEEE Int. Conf. Intell. Data Secur. (IDS), May 2020, pp. 145–150.
- [10] X. Sun, Y. Wang, and Z. Shi, “Insider threat detection using an unsupervised learning method: COPOD,” in Proc. Int. Conf. Commun., Inf. Syst. Comput. Eng. (CISCE), May 2021, pp. 749–754.
- [11] G. Ravikumar and M. Govindarasu, “Anomaly detection and mitigation for wide-area damping control using machine learning,” IEEE Trans. Smart Grid, early access, May 18, 2020, doi: 10.1109/TSG.2020.2995313.
- [12] M. I. Tariq, N. A. Memon, S. Ahmed, S. Tayyaba, M. T. Mushtaq, N. A. Mian, M. Imran, and M. W. Ashraf, “A review of deep learning security and privacy defensive techniques,” Mobile Inf. Syst., vol. 2020, pp. 1–18, Apr. 2020.
- [13] G. Pang, C. Shen, L. Cao, and A. V. D. Hengel, “Deep learning for anomaly detection: A review,” ACM Comput. Surv., vol. 54, no. 2, pp. 1–38, Mar. 2021.
- [14] Bushra Bin Sarhan and N. Altwaijry, “Insider Threat Detection Using Machine Learning Approach,” *Applied Sciences*, vol. 13, no. 1, pp. 1–15, 2023.
- [15] J. Yi and Y. Tian, “Insider Threat Detection Model Enhancement Using Hybrid Algorithms Between Unsupervised and Supervised Learning,” *Electronics*, vol. 13, no. 5, pp. 973, 2024.