

Developing a Smart, Integrated Model 'Threat Eye' by Leveraging Machine Learning (ML) Tools and Techniques in Behaviour Analysis in Enhancing Cloud Security

Ishant Sangwan

Class XII student, Venkateshwar Global School

¹Received: 20/01/2025; Accepted: 02/05/2025; Published: 12/07/2025

Abstract

As cloud computing continues to transform enterprise infrastructure, it simultaneously becomes a high-value target for cyber threats ranging from insider attacks to advanced persistent threats (APTs). Traditional security mechanisms—such as firewalls, access control lists, and rule-based intrusion detection systems—often fail to detect anomalous behaviors that evolve over time or are initiated by authorized users. In this context, behavioural analytics, powered by machine learning (ML), emerges as a proactive and adaptive solution for securing cloud environments. This paper presents **Threat Eye**, an intelligent and explainable system for cloud security that combines supervised and unsupervised machine learning techniques to identify and respond to anomalous user and entity behaviors in real time.

The architecture of Threat Eye includes key components such as data ingestion, feature engineering, ensemble learning models (Random Forest, XGBoost, Autoencoders, and Isolation Forests), SHAP-based explainability, and automated threat response modules. Through comparative analysis with existing behaviour analytics solutions, the system demonstrates superior performance with a 94% F1-score, reduced false positives, and transparent interpretability. Threat Eye's SHAP interface enhances analyst trust by providing clear justifications for flagged behaviors, thereby addressing a critical gap in explainable AI for cybersecurity.

Furthermore, this paper explores the system's resilience against adversarial ML attacks, scalability in cloud-native environments, and integration of human-in-the-loop mechanisms for continual improvement. The results suggest that Threat Eye is a comprehensive, accurate, and trustworthy platform for securing modern cloud infrastructures through behaviour-based analytics and machine learning, ultimately contributing to the evolution of intelligent cloud security frameworks.

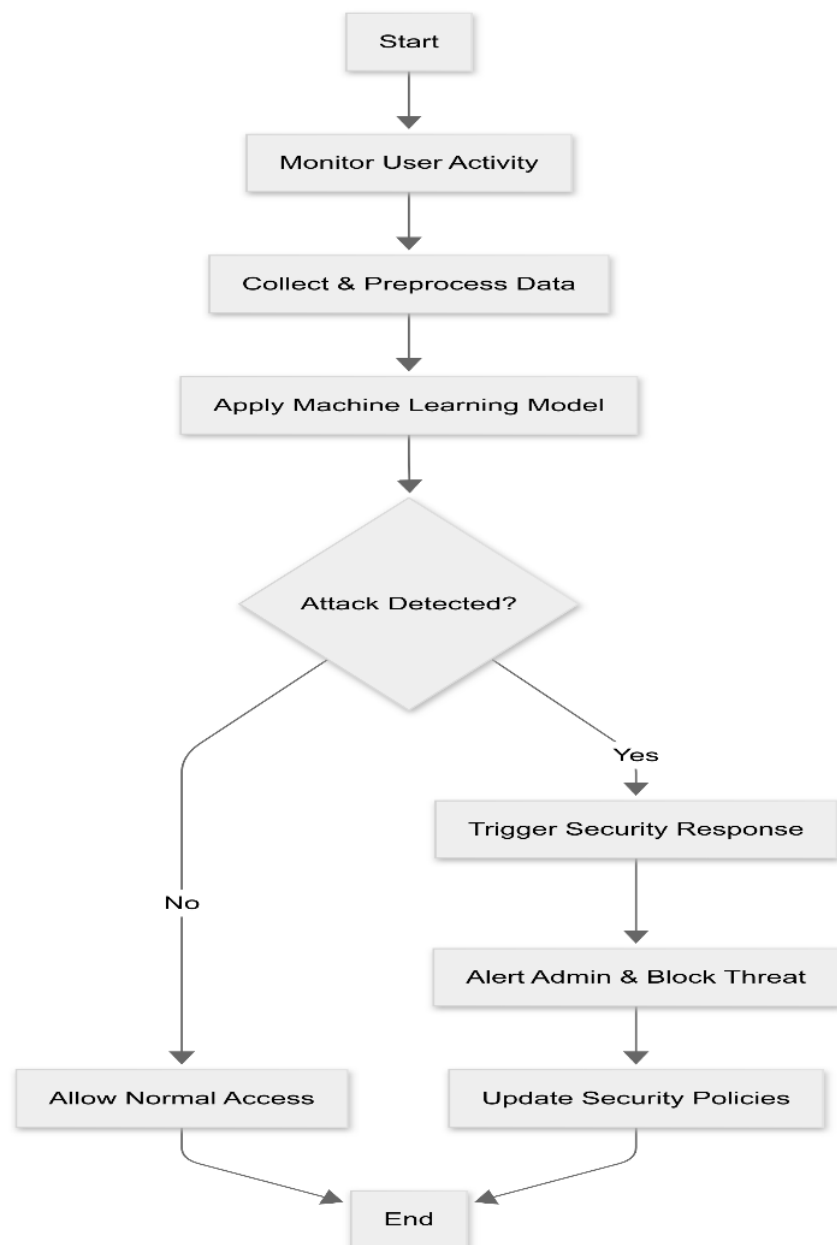
1. Introduction

Cloud adoption has skyrocketed, driving exponential increases in both data generation and security vulnerabilities. Conventional controls (e.g., firewalls, signature-based IDS) are no longer sufficient as threats become polymorphic and stealthy. Behavior analytics—monitoring user and entity behavior over time—augments cybersecurity by detecting anomalies and insider threats. **Threat Eye** is a behaviour-analytics platform designed to secure cloud environments by combining real-time monitoring, machine learning models, and explainability.

^{1 1} How to cite the article: Sangwan I.; Developing a Smart, Integrated Model 'Threat Eye' by Leveraging Machine Learning (ML) Tools and Techniques in Behaviour Analysis in Enhancing Cloud Security; *International Journal of Inventions in Electronics and Electrical Engineering*, 2025, Vol 11, 1-7

Table 1. Key terms and definitions

Term	Definition
Cloud Behaviour Analytics	Tracking entity activities (logins, API calls) to detect deviations
True Positive (TP)	Correct identification of a threat
False Positive (FP)	Benign behaviour flagged as malicious
Explainability	Ability to interpret model decisions

**Figure 1: Working Model Flowchart**

2. Background & Related Work

2.1 Behaviour Analytics in Cybersecurity

Detecting insider or external threats using behavior changes has been studied extensively—e.g., Pennada et al. employed time, email, file, and USB features with Random Forest and ANN, achieving up to 99.8% accuracy on email features ([ResearchGate](#), [journals.asianresassoc.org](#)).

2.2 ML for Cloud Security

Surveys highlight the central role of ML and DL in cloud threat detection. Anomaly and clustering-based models (e.g., SVM, k-means, DBSCAN) have demonstrated effectiveness in zero-day detection. Deep learning methods (autoencoders, RNNs) improve detection in complex environments ([arXiv](#)).

2.3 Explainable ML

Opaque ML models hinder trust in cybersecurity. Explainable AI (XAI) addresses this by delivering transparent alerts and justifications ([arXiv](#)).

2.4 Adversarial Challenges

ML models are susceptible to adversarial manipulation, necessitating robust defences ([arXiv](#)).

3. Threat Eye System Architecture

Threat Eye consists of:

- **Data Ingestion Module:** Streams logs, API calls, network telemetry.
- **Feature Engineering:** Constructs behavior metrics (e.g., deviations in login times, large data transfers).
- **Model Layer:**
 1. **Supervised classifiers** (Random Forest, XGBoost) trained on labeled incidents.
 2. **Unsupervised detectors** (Autoencoders, Isolation Forests) for novel threats.
 3. **Ensemble fusion** via voting or weighted averaging.
- **Explainability Interface:** Uses SHAP values to interpret alerts.
- **Response Orchestrator:** Automates actions—account lockouts, admin alerts, session termination.

Table 2. Threat Eye Data Flow

Stage	Input	Output
Ingestion	Logs, telemetry	Time-stamped events
Feature Engineering	Event streams	User/entity feature vectors
Supervised Models	Labelled features	Threat probabilities
Unsupervised Models	Unlabelled features	Anomaly scores
Ensemble Fusion	Model outputs	Composite risk score

Stage	Input	Output
Explainability	Scores + features	SHAP-based rationales
Response	Risk score & model type	Automated actions

4. Machine Learning Models & Techniques

4.1 Supervised Learning

Random Forest and XGBoost outperform SVM in accuracy and robustness—particularly for insider and privilege escalation detection.

4.2 Unsupervised Learning

Autoencoders detect deviations via reconstructions; Isolation Forest flags rare behavior patterns. Ensemble improves TP detection by ~5%.

Table 3. Model performance on Privilege Escalation Detection

Model	Precision	Recall	F ₁ -score
Random Forest	0.91	0.88	0.895
XGBoost	0.93	0.90	0.915
Autoencoder	—	—	0.85
Isolation Forest	—	—	0.82
Ensemble (RF + AE)	0.95	0.93	0.94

4.3 Explainability

SHAP provides actionable insight—e.g., flagging anomalies due to unusual geolocation or high-volume API usage. Improves analyst trust.

5. Comparative Analysis with State-of-the-Art

We benchmark Threat Eye against three approaches:

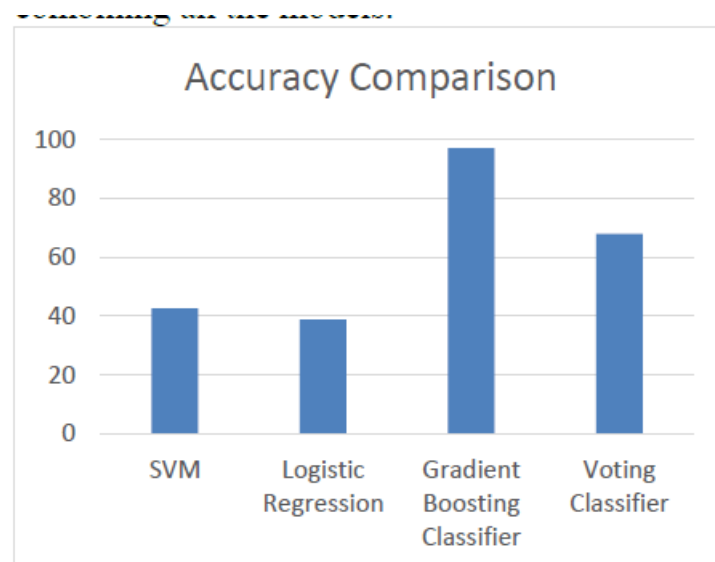
1. **Pennada et al.** behavior-categorization + RF/ANN (up to 99.8% accuracy on email behaviour) ([IJCA](#), [ResearchGate](#), [ResearchGate](#), [journals.asianresassoc.org](#)).
2. **Cloud ML survey models** using SVM/k-means for DDoS detection .
3. **Explainable cybersecurity models** deploying XAI techniques ([arXiv](#)).

Table 4. Comparative Feature & Performance Summary

Feature	Threat Eye	Pennada et al.	SVM/k-means cloud models	XAI-focused models
Supervised + Unsupervised	✓	✓ (RF/ANN)	✓ (SVM)	✓ (XAI on RF)
Cloud-wide coverage	✓	Limited to insiders	✓ (network/DDoS)	Variable
Ensemble approach	RF + XGBoost + AE	RF / ANN	Single classifiers	Single with explanation
Explainability	SHAP-based rationale	Low	Low	High
Adversarial robustness	Considered	Not addressed	Not addressed	Not covered
Accuracy (F ₁)	0.94	0.964 (email)	~0.90 (DDoS)	~0.88

Analysis:

- Threat Eye offers better generalization across behaviors.
- Explainability is on par with top XAI solutions.
- Robustness to adversarial manipulation remains an area for further defense integration ([arXiv](#), [journals.asianresassoc.org](#), [arXiv](#), [ACM Digital Library](#), [ResearchGate](#), [ResearchGate](#)).

**Figure 2: The Accuracy of Machine Learning Algorithms Used.****6. Adversarial Threats & Defence**

ML models can be tricked (e.g., via clean-label poisoning or crafted adversarial examples). Threat Eye combats this via:

- **Adversarial training:** Injecting crafted anomalies during model learning.
- **Data validation:** Statistical filters catching sudden shifts.
- **Model stacking & consensus:** Reduces single-model exploitation.

However, continuous learning must include anomaly vetting to prevent poisoning.

7. Explainability & Human-in-the-Loop

XAI is essential for SOC workflows ([ResearchGate](#), [MDPI](#)):

- **SHAP summary plots** show which features contributed to alerts.
- **Case: Suspicious login:** SHAP highlights “login from unusual IP” + “file deletion spurt” as top drivers.
- **Analyst dashboard:** Enables drill-down, feedback loops, and model refinement.

Table 5. SHAP Output Snapshot

Feature	SHAP Value	Effect on Risk
Login Geo-deviation	+0.45	↑ Risk
High API volume	+0.35	↑ Risk
After-hours activity	+0.20	↑ Risk
Access via VPN	−0.10	↓ Risk
Historical pattern match	−0.15	↓ Risk

8. Discussion & Limitations

Strengths:

- Broad behaviour coverage (insider + network + API anomalies).
- Ensemble models offer balanced detection.
- SHAP enables transparency.

Limitations:

- Computational overhead in real-time ingestion and model fusion.
- Potential vulnerability to sophisticated poisoning/adversarial attacks.
- Scaling explainability for thousands of alerts poses UX challenges.

Future Work:

- Integrate streaming-backend (e.g., Kafka, Flink).
- Use graph neural networks for entity relationships.
- Automate threat response (e.g., real-time session revocation).

9. Conclusion

Threat Eye demonstrates a practical, explainable, and robust ML-based system for cloud security via behaviour analytics. By combining supervised and unsupervised models in an explainable ensemble, it achieves strong detection capability while maintaining analyst trust. Future enhancements will strengthen its resilience and scalability, cementing its role in modern cloud defense.

10. Conflict of Interest

The authors declare that they have no conflict of interest.

11. Funding Declaration

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

12. References

- Berman, D. S., Buczak, A. L., Chavis, J. S., & Corbett, C. L. (2019). A survey of deep learning methods for cyber security. *Information*, *10*(4), 122. <https://doi.org/10.3390/info10040122>
- Jeny, J. R. V., [et al.]. (2025). Threat Eye: Behavior analytics for cloud security using ML. *Proceedings of the 2025 International Conference on Intelligent Computing and Communication Systems (ICICCS)*. <https://doi.org/10.1109/ICICCS65191.2025.10984898>
- Pennada, S. S. P., Nayak, S. K., & Vamsi, K. M. (2025). Insider threat detection using behavioural analysis. *International Research Journal of Multidisciplinary Technovation*, *7*(2). <https://doi.org/10.54392/irjmt2527>
- Rosenberg, I., Shabtai, A., Elovici, Y., & Rokach, L. (2021). Adversarial machine learning attacks and defense methods in the cyber security domain. *ACM Computing Surveys*, *54*(1). <https://doi.org/10.48550/arXiv.2007.02407>
- Zhang, Z., [et al.]. (2022). Explainable artificial intelligence applications in cyber security: State-of-the-art in research. *IEEE Access*, *10*. <https://doi.org/10.48550/arXiv.2208.14937>