



SvedbergOpen
DISSEMINATION OF KNOWLEDGE

International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

An Intelligent Cloud Security Framework Integrating Machine Learning-Driven Behavioral Analysis with Dynamic Role-Based Access Control and a CP-ABE-Inspired Attribute-Based Encryption Scheme

Archana Salunkhe¹, Dr. V. R. Ghorpade²

¹Dr. D.Y. Patil Polytechnic, Kolhapur, Maharashtra, India, E-mail: archana.jadhav1719phd@gmail.com

²Bharati Vidyapith college of Engg., Kolhapur, Maharashtra, India, E-mail: vijayghorpade@rediffmail.com

Abstract

Cloud computing has become the most popular paradigm for scalable and on-demand data storage and services provision; nonetheless, storing data in third-party servers raises various security concerns such as access control breaches, insider attacks, privilege escalation, and misuse of static policies. Classical Role-Based Access Control (RBAC) uses immutable roles that do not react to user activities, while Ciphertext-Policy Attribute-Based Encryption (CP-ABE) uses flexible but still immutable policies and is not capable of performing behavioral analysis. The current study presents a hybrid cloud computing security system where Machine Learning (ML) behavioral risks analysis is combined with RBAC and an attribute-based encryption system inspired by CP-ABE to achieve adaptive behavior-aware access control. The attribute-based encryption system is designed with the help of attribute-keyed HKDF key derivation and not the pairing-based CP-ABE cryptographic protocol: it implements CP-ABE-like access policy and revocation but it lacks security proof based on the DBDH problem. The system is continuously learning behavioral feature vectors based on login frequency, changes in IP geolocation, request rate, deviation in access times, sensitivity of the role, and entropy of files accesses and then combining outputs of RF, SVM, and IF classifiers to get a weighted behavioral risk score. This score is used to update the RBAC roles and attribute-based access policies automatically, i.e., without any need for administrator's involvement, thus reducing the detection to revocation gap characteristic of traditional systems. Files are additionally pre-scanned using the Random Forest malware classifier that learns based on the entropy and structure of files before being encrypted with Fernet (AES-128) based on the policy-defined key using SHA-256 for integrity checks. The system was prototyped on AWS S3, MongoDB, and Flask. In our evaluations, the malware detection module showed 94% accuracy (93% precision, 95% recall, 94% F1 score) compared to 82% in the case of signature-based solution, the risk engine reached 92% unauthorized-access detection, and the suggested hybrid encryption pipeline (policy layer inspired by CP-ABE + Fernet) introduced 2.1 seconds extra time for encryption and 1.8 seconds for decryption per 5MB file compared to 0.8/0.7 seconds for AES..

Keywords: Cloud Security; Machine Learning; Dynamic RBAC; CP-ABE; Behavioral Analysis; Anomaly Detection; Insider Threat Detection; Adaptive Access Control; Malware Detection; Zero-Trust Architecture

I. Introduction

Cloud computing has revolutionized the way information technology works in the contemporary world through the provision of scalable resources like storage, computing, and application services offered via the Internet through cloud service providers such as Amazon Web Services (AWS), Microsoft Azure, and Google Cloud Platform. However, despite all the benefits that cloud computing has brought about, its widespread use has created several critical security challenges for users. For instance, entrusting one's confidential organizational and personal information with third-party cloud servers puts the user at risk of any number of threats.

The conventional security methods, such as firewalls, identity-based access control, and encryption, are inadequate for dealing with the constantly changing cyber-attacks in multi-user cloud systems due to their inability to cope with the changing access needs through flexible means. In order to tackle the problem, ABAC and CP-ABE allow encryption of access to information according to the attributes of the user (such as his/her position in an organization or department, etc.). Nevertheless, the policies used in both RBAC and CP-ABE schemes do not change regardless of the increase in number of the users and data; they are static and cannot be adjusted during user sessions in real time. This means that if an authorized user's account gets compromised by cyber criminals or insiders abuse their legitimate rights, such actions will not be noticed.

Machine Learning algorithms alone have been found to be effective in the identification of unusual activity based on access log data, network traffic, and usage data. But an intrusion and anomaly detection system implemented using Machine Learning is always implemented independently of the underlying access controls and cryptography system it tries to defend against. An ML algorithm might be able to detect that a session has been hijacked, but it would not be able to revoke a cryptographic key or reduce the privileges of the role. This paper aims to bridge that gap. The idea behind this paper is to create a system where ML-based behavioral risk scores can influence updates to RBAC roles and attributes inspired by CP-ABE.

Key Contributions

- A hybrid security architecture integrating ML behavioral analytics, dynamic RBAC, and a CP-ABE-inspired attribute-based encryption layer into a single closed-loop framework, with complete experimental evaluation on a working AWS/MongoDB/Flask prototype.
- A formally specified 7-dimensional behavioral feature vector (login frequency, geolocation shift, request rate, access-time deviation via Mahalanobis distance, role sensitivity, file-access-type entropy, decryption-failure rate) and a weighted RF-SVM-IF ensemble risk score, given as explicit algorithms with complexity analysis.
- A dynamic, attribute-mismatch-based revocation mechanism that updates access without re-encrypting previously stored ciphertexts, in contrast to traditional CP-ABE revocation schemes.
- An ML-based pre-encryption malware screening stage (Random Forest on entropy and structural file features) integrated directly into the secure upload pipeline.
- An honest accounting of the gap between the CP-ABE-inspired implementation and a formally proven pairing-based CP-ABE cryptosystem, together with a concrete hardening roadmap.

II. Related Work

With the rapid adoption of cloud computing, ensuring secure data storage and privacy-preserving access control has become a major research focus. Recent work (2023–2026) has largely progressed along two separate tracks: encryption-based access control, and ML-based threat detection.

One of the approaches towards statistical privacy was an approach that utilized attribute-based encryption (ABE) for fine-grained access control using user attributes instead of roles and provided improved confidentiality, but did not have any dynamic threat detection capabilities [1]. Another framework that made use of AI was an ABE framework with CP-ABE and machine learning (ML) providing dynamic anomaly detection and authentication capability which made it more efficient at mitigating insider threats, but increased the computation complexity [2].

A machine learning (ML) based anomaly detection method to detect unauthorized access in real-time was suggested by Jiang et al., which concentrated on detecting intrusions rather than using encryption-based policy enforcement [4]. The use of hybrid encryption combined with ABAC provided an improvement in context-awareness decision-making; however, no confidentiality enforcement was done through policy-based encryption [5]. Recently, ABKS techniques provided multi-authority and post-quantum secure keyword-based data search facilities, although the use of these technologies does not consider adaptive policies because static policies are used [6]. There have been advancements made in blockchain-based access control models by using smart contract techniques and AI-based anomaly detection [7].

Conventional methods of intrusion detection using classical machine learning algorithms such as decision trees, naïve bayes, and SVM are very commonly used but are only capable of detecting abnormal activities or cyber-attacks but cannot provide any assurance concerning the confidentiality of the underlying encrypted data [8], [9]. Self-supervised learning was used for anomaly detection in encrypted traffic [10]. Overall, existing research largely treats encryption-based access control and ML-based anomaly detection as independent problems. Few approaches integrate CP-ABE-style policy encryption with behavioral ML analysis into a single closed-loop system — the gap this paper addresses. Table 1 shows the comparative analysis of existing system.

Table 1. Comparative Analysis of Existing Attribute-Based Encryption and ML-Based Cloud Security Frameworks

Paper (Ref.)	Encryption Technique	ML Technique	Dynamic Policy	Insider Detection	Primary Limitation
Verma & Singh [1]	ABE	None	No	No	No dynamic anomaly detection
Alshamrani et al. [2]	CP-ABE	AI-based anomaly detection	Partial	Yes	High computational complexity
Chen et al. [3]	None	Supervised + XAI	No	Yes	No encryption-based access control

Jiang et al. [4]	None	Behavioral ML	No	Yes	Intrusion detection only, no encryption
Rahman & Patel [5]	Hybrid encryption + ABAC	Supervised & unsupervised ML	Yes	Partial	Lacks fine-grained policy encryption
Liu et al. [6]	Multi-authority ABKS	None	No (static)	No	No adaptive learning
Sharma & Mishra [7]	Blockchain + ABE	AI-based anomaly detection	Yes	Yes	High latency / overhead
Modi et al. [8]	None	DT, Naïve Bayes	No	Yes	No encryption support
Bridges & Vaughn [9]	None	Fuzzy / SVM-like	No	Partial	No confidentiality mechanism
Kim & Lee [10]	None	Self-supervised ML	No	Yes	No fine-grained access control
Proposed System	CP-ABE-inspired (HKDF)	RF + SVM + IF ensemble	Yes (dynamic revocation)	Yes	HKDF approximation lacks formal DBDH proof (Sec. 6.1)

III. Proposed Framework

System Architecture Overview

The architecture of the system being proposed is zero trust and three-tiered: a presentation layer (HTML5/CSS3/JavaScript using Flask-Jinja2 templates), an application layer (Python/Flask implementation doing routing, orchestration, and security functions), and a data layer divided between MongoDB (for user profiles, behavioral logs, access policies, and SHA-256 hash values) and AWS S3 for encrypted file storage. Each of the requests – be it authentication, uploading, or downloading – goes through the Core Security Engine that does behavioral feature extraction, machine learning based risk scoring, RBAC/attribute policy-based permissions checks, malware scanning, and policy-based encryption/decryption. Figure 1 shows the end-to-end secure data flow: a user uploads a file, it's encrypted using a randomly generated key, stored in the cloud, then another user searches for and requests it, receives the key, decrypts the file, and downloads it.

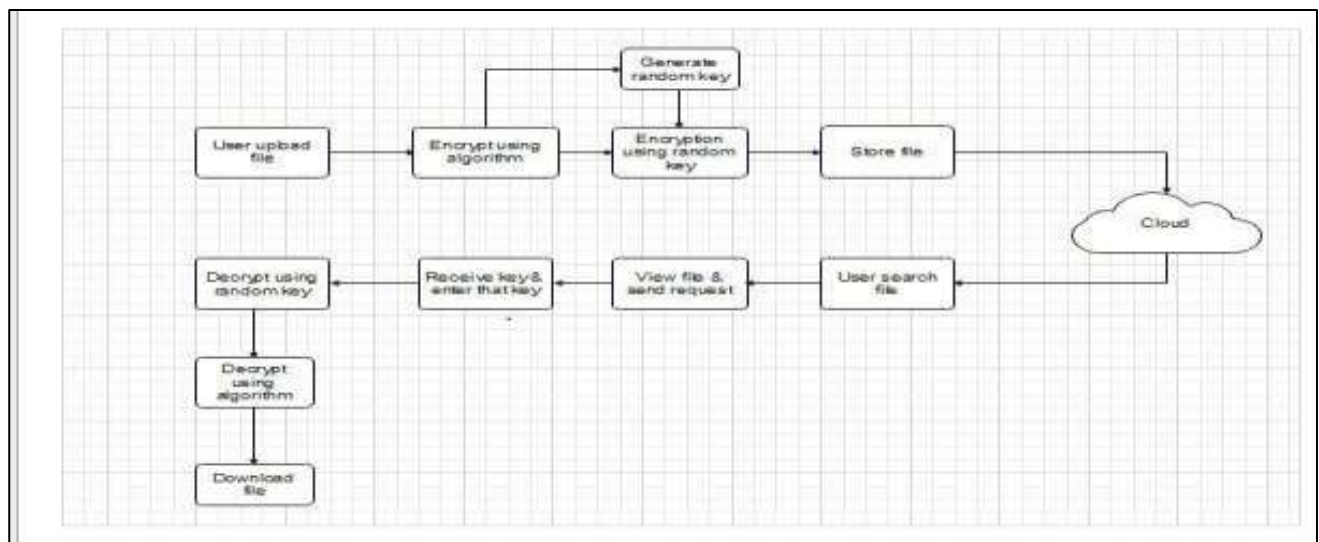


Figure 1. End-to-end secure data flow: upload → ML risk & malware screening → CP-ABE-inspired encryption → AWS S3 storage → policy verification → decryption → delivery.

Behavioral Feature Extraction

For each user session s , the system continuously extracts a seven-dimensional behavioral feature vector:

$$\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \mathbf{x}_4, \mathbf{x}_5, \mathbf{x}_6, \mathbf{x}_7] \in \mathbb{R}^7 \quad (1)$$

where each component is normalized to $[0, 1]$ (or a bounded range) so that no single feature dominates the downstream classifiers:

$x_1 = \min(\text{login_count}_{24h} / L_{\text{max}}, 1)$, $L_{\text{max}} = 20$ — normalized login frequency

$x_2 = 1$ if GeoIP(s) $\neq$ RegisteredCountry(u), else 0 — geolocation shift indicator

$x_3 = \min(\text{requests_per_minute}(s) / 60, 1)$ — request-rate pressure

Access-time deviation (x_4) is computed as a Mahalanobis distance between the current access timestamp and the user's historical access-time distribution $H_u = \{t_1, \dots, t_n\}$, with online mean μ_u and covariance Σ_u :

$$d_M(t_s, H_u) = \sqrt{[(t_s - \mu_u)^t \Sigma_u^{-1} (t_s - \mu_u)]} \quad (2)$$

$$x_4 = \min(d_M(t_s, H_u) / d_{\text{ref}}, 1), \quad d_{\text{ref}} = 95\text{th percentile of } d_M \text{ over the training population} \quad (3)$$

The remaining features are role sensitivity, session file-type entropy, and decryption-failure rate:

$$x_5 = \text{RoleWeight}(\text{role}(u)) \in \{0.3 \text{ (Employee)}, 0.6 \text{ (Manager)}, 1.0 \text{ (Admin)}\}$$

$$x_6 = [-\sum_i p(\tau_i) \log_2 p(\tau_i)] / \log_2(k_{\text{max}}), \quad k_{\text{max}} = \text{number of supported file-type categories} \quad (4)$$

$$x_7 = \text{failed_decryptions}(s) / \max(\text{total_decryption_attempts}(s), 1)$$

x_6 is a normalized Shannon entropy over the distribution $p(\tau)$ of file types accessed within the session — a user suddenly touching an unusually diverse or unusual mix of file types raises x_6 , which the ensemble (Section 3.3) treats as a weak anomaly signal distinct from the file-content entropy used for malware screening in Section 3.6.

Algorithm 1. Behavioral Feature Vector Extraction

Input: session s for user u ; 24-hour event window W ; historical profile H_u

Output: feature vector $X = (x_1, \dots, x_7) \in \mathbb{R}^7$

1: $E \leftarrow \text{RetrieveEvents}(u, W)$ from the MongoDB access-log collection

2: $x_1 \leftarrow \min(|\{e \in E : \text{type}(e) = \text{login}\}| / 20, 1)$

3: $x_2 \leftarrow 1$ if GeoIP(s) $\neq$ RegisteredCountry(u), else 0

4: $x_3 \leftarrow \min(\text{RequestRate}(s) / 60, 1)$

5: $(\mu_u, \Sigma_u) \leftarrow \text{IncrementalMeanCov}(H_u, t_s)$ {online update, $O(1)$ }

6: $d \leftarrow \text{sqrt}((t_s - \mu_u)^t \Sigma_u^{-1} (t_s - \mu_u))$ {Eq. 2}

7: $x_4 \leftarrow \min(d / d_{\text{ref}}, 1)$

8: $x_5 \leftarrow \text{RoleWeight}(\text{role}(u))$

9: $p(\tau) \leftarrow \text{FileTypeFrequency}(E)$ for each accessed type τ

10: $x_6 \leftarrow [-\sum_{\tau} p(\tau) \log_2 p(\tau)] / \log_2(k_{\text{max}})$ {Eq. 4}

11: $x_7 \leftarrow \text{FailedDecryptions}(s) / \max(\text{TotalDecryptAttempts}(s), 1)$

12: $X \leftarrow (x_1, x_2, x_3, x_4, x_5, x_6, x_7)$

13: Persist X to the sliding-window buffer keyed by $(u, \text{timestamp}(s))$

14: return X

Complexity: $O(1)$ amortized per session given incremental mean/covariance maintenance; a cold-start user (no prior H_u) requires $O(|H_u|)$ for the initial covariance estimate.

Machine Learning-Based Risk Classification

The behavioral risk score is computed by a weighted ensemble of three classifiers, chosen to cover complementary failure modes: SVM for margin-based separation in a normalized feature space, Random Forest for non-linear interactions between features, and Isolation Forest for unsupervised detection of novel anomaly patterns not present in the labeled training data.

Support Vector Machine (SVM)

The SVM decision function over the behavioral feature space is:

$$f(X) = \sum_i \alpha_i y_i K(x_i, X) + b \quad (5)$$

where K is a radial basis function (RBF) kernel, α_i are the learned Lagrange multipliers, and $y_i \in \{-1, +1\}$ are training labels. The raw decision value is converted to a calibrated probability via Platt scaling:

$$P_{\text{SVM}}(X) = 1 / (1 + \exp(a \cdot f(X) + b')) \quad (6)$$

with a and b' fit on a held-out calibration split.

Random Forest (RF)

An ensemble of T decorrelated decision trees, each trained on a bootstrap sample with random feature subsampling at each split, votes on the malicious/benign (or risky/normal) class:

$$P_{\text{RF}}(X) = (1/T) \sum_{t=1}^T 1[h_t(X) = 1] \quad (7)$$

where h_t is the prediction of the t -th tree. Randomization across trees reduces variance relative to a single decision tree while preserving low bias [22].

Isolation Forest (IF)

Isolation Forest isolates anomalies via random recursive partitioning; anomalous points require fewer splits to isolate, giving a shorter average path length. The anomaly score is:

$$s_{\text{IF}}(X) = 2^{-(E[h(X)] / c(\psi))} \quad (8)$$

where $E[h(X)]$ is the average path length of X across the forest and $c(\psi)$ is a normalization constant for subsample size ψ [23]. Scores near 1 indicate strong anomalies; scores near 0.5 indicate normal points.

Weighted Ensemble Risk Score

The three calibrated outputs are fused into a single behavioral risk score:

$$R = w_{RF} \cdot P_{RF}(X) + w_{SVM} \cdot P_{SVM}(X) + w_{IF} \cdot s_{IF}(X), \quad R \in [0,1] \quad (9)$$

with weights $(w_{RF}, w_{SVM}, w_{IF}) = (0.40, 0.35, 0.25)$, selected on a validation split to minimize the false-negative rate for insider-threat sessions while bounding the false-positive (unnecessary step-up authentication) rate. RF receives the largest weight because it consistently achieved the highest standalone recall on labeled anomalies during model selection; SVM contributes margin-based generalization on the normalized feature space; IF contributes sensitivity to novel behavioral patterns absent from the labeled training set.

Algorithm 2. Weighted Ensemble Behavioral Risk Scoring

Input: $X \in \mathbb{R}^7$; trained models RF, SVM, IF; weights $(0.40, 0.35, 0.25)$

Output: risk score $R \in [0,1]$

1: $P_{RF} \leftarrow \text{RF.predict_proba}(X)[\text{class} = 1]$ {Eq. 7}

2: $z \leftarrow f_{SVM}(X)$ {Eq. 5}

3: $P_{SVM} \leftarrow 1 / (1 + \exp(a \cdot z + b'))$ {Eq. 6, Platt scaling}

4: $s_{IF} \leftarrow \text{IsolationForest.score}(X)$ {Eq. 8}

5: $R \leftarrow 0.40 \cdot P_{RF} + 0.35 \cdot P_{SVM} + 0.25 \cdot s_{IF}$ {Eq. 9}

6: $R \leftarrow \text{clip}(R, 0, 1)$

7: return R

Complexity: $O(T \cdot d)$ for RF (T trees, depth d), $O(n_{sv} \cdot 7)$ for SVM kernel evaluation (n_{sv} support vectors), $O(t_{IF} \cdot \log \psi)$ for IF (t_{IF} trees, subsample ψ) — dominated by the RF term.

Dynamic Access Decision Policy

The risk score R drives a three-way access decision:

$$\text{Access}(R) = \{ \text{Permit if } R < 0.30; \text{ Monitor if } 0.30 \leq R \leq 0.80; \text{ Block if } R > 0.80 \} \quad (10)$$

In the Monitor state, the user retains access but every operation is logged with elevated audit detail and a step-up MFA (OTP) challenge is triggered. In the Block state, the user's role is downgraded to a Restricted (read-only) role and their attribute set is reduced:

$$A'(u) = A(u) \setminus A_{\text{elevated}}, \quad A'(u) \subset A(u) \quad (11)$$

Because decryption under the CP-ABE-inspired scheme (Section 3.5) depends on attribute-derived keys rather than per-file access-control lists, this attribute reduction alone is sufficient to deny future decryption of any file whose policy required an attribute now removed — without re-encrypting a single stored file. This is reversible only through manual administrator review, preventing an automatically-triggered block from silently self-resolving.

Algorithm 3. Dynamic Access Decision and Attribute Revocation

Input: risk score R ; current role $\rho(u)$; current attribute set $A(u)$

Output: updated role $\rho'(u)$, updated attribute set $A'(u)$, decision

1: if $R < 0.30$: decision $\leftarrow$ Permit; $\rho' \leftarrow \rho(u)$; $A' \leftarrow A(u)$

2: else if $0.30 \leq R \leq 0.80$:

3: decision $\leftarrow$ Monitor; $\rho' \leftarrow \rho(u)$; $A' \leftarrow A(u)$

4: EnableElevatedAudit(u); RequestStepUpMFA(u)

5: else $\{R > 0.80\}$

6: decision $\leftarrow$ Block

7: $\rho' \leftarrow$ Restricted {read-only}

8: $A' \leftarrow A(u) \setminus A_{\text{elevated}}$ {Eq. 11}

9: PersistRoleAndAttributes(u, ρ', A') $\{O(1), \text{no ciphertext re-encryption}\}$

10: DispatchAdminAlert(u, R); FlagForManualReview(u)

11: return $(\rho', A', \text{decision})$

Complexity: $O(1)$ per request, independent of the number n_f of files previously encrypted for u — revocation is enforced via attribute mismatch rather than re-encryption, unlike classical CP-ABE revocation at $O(n_f)$.

CP-ABE-Inspired Cryptographic Policy Enforcement

Data confidentiality is enforced through a CP-ABE-inspired policy engine. Let the user attribute set be $A = \{A_{\text{role}}, A_{\text{dept}}, A_{\text{desig}}\}$ and let the access policy P be a monotone Boolean formula over attributes, most commonly a conjunction:

$$P = A_{\text{role}} \wedge A_{\text{dept}} \wedge A_{\text{desig}} \quad (12)$$

Rather than a pairing-based CP-ABE cryptosystem [24], the implemented scheme derives a symmetric key directly from the requester's attributes and the policy via HKDF (HMAC-based Key Derivation Function):

$$K_{\{u,F\}} = \text{HKDF}(\text{IKM} = A(u) \parallel P_F, \text{salt}, \text{info} = \text{"cp-abe-inspired-v1"}, L = 256 \text{ bits}) \quad (13)$$

Decryption succeeds only if the requester's current attribute set satisfies the policy, $A(u) \models P_F$, because only in that case does the HKDF re-derivation at access time reproduce the key used at encryption time:

$$\text{Decrypt}(C, K'_{\{u,F\}}) \text{ succeeds} \Leftrightarrow A'(u) \models P_F \Leftrightarrow K'_{\{u,F\}} = K_{\{u,F\}} \quad (14)$$

This reproduces CP-ABE's central property — decryption fails cryptographically, not merely logically, when attributes do not satisfy the policy — using a much simpler and faster primitive. The trade-off, discussed in Section 6.1, is that HKDF-based key derivation has no published reduction to a hard cryptographic assumption (such as DBDH) analogous to the security proofs available for pairing-based CP-ABE constructions [24].

Algorithm 4. Secure File Upload Pipeline

Input: file F ; uploader u ; access policy P_F derived from current RBAC/attribute state

Output: stored ciphertext, or rejection

1: $X_F \leftarrow (\text{size}(F), H(F), f_{\text{macro}}, f_{\text{js}}, f_{\text{obj}}, f_{\text{stream}}, f_{\text{embed}})$ {Section 3.6}

2: $P_{\text{mal}} \leftarrow \text{RF_malware.predict_proba}(X_F)[\text{class} = 1]$

3: if $P_{\text{mal}} \geq \tau_{\text{mal}} (= 0.80)$:

4: Quarantine(F); LogSecurityEvent($u, F, \text{"malware"}$); return REJECT

5: $h \leftarrow \text{SHA256}(F)$

6: $K_{\{u,F\}} \leftarrow \text{HKDF}(\text{IKM} = A(u) \parallel P_F, \text{salt}, \text{info}, L=256)$ {Eq. 13}

7: $C \leftarrow \text{Fernet_Encrypt}(F, K_{\{u,F\}})$ {AES-128-CBC + HMAC-SHA256}

8: Store C in AWS S3; store $\{h, P_F, \text{metadata}\}$ in MongoDB

9: return SUCCESS

Complexity: $O(|F|)$ dominated by hashing, entropy computation, and symmetric encryption; classification adds $O(T_{\text{mal}} \cdot d)$ for the Random Forest pass.

Algorithm 5. Secure File Download and Integrity Verification

Input: requester u ; file identifier id

Output: plaintext F , or ACCESS_DENIED

1: $R \leftarrow \text{ComputeRiskScore}(u)$ {Algorithm 2}

2: $(\rho', A', \text{decision}) \leftarrow \text{UpdatePolicy}(R, \rho(u), A(u))$ {Algorithm 3}

3: if decision = Block: return ACCESS_DENIED

4: $(P_F, C, h_{\text{stored}}) \leftarrow \text{Retrieve}(\text{id})$ from MongoDB / S3

5: if $A'(u) \not\models P_F$:

6: LogSecurityEvent($u, \text{id}, \text{"policy-mismatch"}$); return ACCESS_DENIED

7: $K'_{\{u,F\}} \leftarrow \text{HKDF}(\text{IKM} = A'(u) \parallel P_F, \text{salt}, \text{info}, L=256)$

8: $F \leftarrow \text{Fernet_Decrypt}(C, K'_{\{u,F\}})$

9: $h_{\text{computed}} \leftarrow \text{SHA256}(F)$

10: if $h_{\text{computed}} \neq h_{\text{stored}}$:

11: LogSecurityEvent($u, \text{id}, \text{"integrity-tamper"}$); return ACCESS_DENIED

12: return F

Complexity: $O(|F|)$ for decryption and hashing, plus $O(1)$ for the attribute-satisfaction check $A'(u) \models P_F$.

Malware Detection Module

Before encryption, every uploaded file F is screened by a dedicated Random Forest classifier trained on structural and statistical features distinct from the behavioral feature vector of Section 3.2:

$$X_F = [x_{\text{size}}, x_{\text{entropy}}, x_{\text{macro}}, x_{\text{js}}, x_{\text{obj}}, x_{\text{stream}}, x_{\text{embed}}] \in \mathbb{R}^7 \quad (15)$$

$$x_{\text{size}} = \min(\text{Size}(F) / 10^6, 5) \text{ — file size in MB, capped at 5}$$

$$x_{\text{entropy}} = H(F) / 8, \quad H(F) = - \sum_{\beta(b=0)}^{255} P(b) \log_2 P(b) \quad (16)$$

$H(F)$ is the Shannon entropy of the file's byte distribution (maximum 8 bits/byte); high entropy is a common signature of packed, encrypted, or obfuscated payloads. The remaining components (x_{macro} , x_{js} , x_{obj} , x_{stream} , x_{embed}) are frequencies of macro blocks, embedded JavaScript, object counts, stream counts, and embedded objects, respectively — each a heuristic indicator commonly associated with malicious document payloads.

The classifier outputs a maliciousness probability, thresholded to a binary decision:

$$p = P(Y = 1 \mid X_F); \text{ Decision} = \text{Malicious if } p \geq \tau, \text{ else Safe; } \tau = 0.80 \quad (17)$$

The threshold $\tau = 0.80$ was selected by maximizing an $F\beta$ score with $\beta = 0.5$, prioritizing precision (avoiding false quarantines of legitimate files) over recall, since a missed detection is still subject to the behavioral risk layer and manual audit downstream.

Algorithm 6. ML-Based Malware Screening

Input: file bytes F

Output: Malicious or Safe

1: $H(F) \leftarrow - \sum_{b=0}^{255} P(b) \log_2 P(b)$ {Eq. 16, byte histogram over F }

2: $x_{\text{entropy}} \leftarrow H(F) / 8$

```

3: x_size ← min(Size(F) / 106, 5)
4: (x_macro, x_js, x_obj, x_stream, x_embed) ← HeuristicScan(F)
5: X_F ← (x_size, x_entropy, x_macro, x_js, x_obj, x_stream, x_embed)
6: p ← RandomForestClassifier(n_estimators=300).predict_proba(X_F)[1]
7: if p ≥ 0.80: return Malicious
8: return Safe
Complexity: O(|F|) for the entropy pass over file bytes, plus O(T·d) = O(300·d) for the Random Forest inference.
    
```

Auxiliary Cryptographic Primitives

Three supporting primitives complete the security stack: SHA-256 for integrity, Paillier homomorphic encryption for privacy-preserving aggregate computation on stored metrics, and Shamir's Secret Sharing for distributing the master key material so that no single compromised node can reconstruct it.

Paillier Homomorphic Encryption

For public key (n, g) , plaintext $m \in \mathbb{Z}_n$, and randomness $r \in \mathbb{Z}_n^*$:

$$E(m, r) = g^m \cdot r^n \pmod{n^2} \quad (18)$$

Paillier encryption is additively homomorphic:

$$E(m_1) \cdot E(m_2) \pmod{n^2} = E(m_1 + m_2 \pmod{n}) \quad (19)$$

This lets the platform compute aggregate statistics (e.g., total flagged sessions across a department) directly on encrypted values, without ever decrypting individual records [25].

Shamir's Secret Sharing (t, n)

A master secret $S \in \mathbb{Z}_p$ is embedded as the constant term of a random degree- $(t-1)$ polynomial:

$$P(x) = S + a_1x + a_2x^2 + \dots + a_{t-1}x^{t-1} \pmod{p} \quad (20)$$

Each of n participants receives a share $(x_i, y_i = P(x_i))$; any t shares reconstruct S via Lagrange interpolation, while fewer than t reveal nothing:

$$S = P(0) = \sum_{j=1}^t y_j \cdot \prod_{k(k \neq j)} [x_k / (x_k - x_j)] \pmod{p} \quad (21)$$

This threshold scheme [26] removes the single point of failure inherent in storing the master key derivation secret on one node.

Performance Metrics

Detection quality is reported using standard classification metrics, computed from true/false positives and negatives (TP, FP, TN, FN):

$$\text{Accuracy} = (TP + TN) / (TP + TN + FP + FN) \quad (22)$$

$$\text{Precision} = TP / (TP + FP), \quad \text{Recall} = TP / (TP + FN) \quad (23)$$

$$F1 = 2 \cdot (\text{Precision} \cdot \text{Recall}) / (\text{Precision} + \text{Recall}) \quad (24)$$

AUC-ROC is reported as the area under the true-positive-rate vs. false-positive-rate curve across classification thresholds, summarizing performance independent of any single operating point.

Computational Complexity Analysis

Table 2 summarizes the asymptotic time complexity of each algorithm in the pipeline, in terms of file size $|F|$, tree count/depth (T, d) , support-vector count n_{sv} , Isolation Forest tree count and subsample size (t_{IF}, ψ) , and stored file count n_f .

Table 2. Time Complexity Summary

Algorithm	Dominant Time Complexity	Notes
1 — Feature Extraction	$O(1)$ amortized	$O(H_u)$ on cold start
2 — Ensemble Risk Scoring	$O(T \cdot d + n_{sv} \cdot 7 + t_{IF} \cdot \log \psi)$	Dominated by Random Forest
3 — Policy Update / Revocation	$O(1)$	Independent of n_f (no re-encryption)
4 — Secure Upload	$O(F)$	Hashing + encryption + malware scan
5 — Secure Download	$O(F)$	Decryption + integrity check
6 — Malware Screening	$O(F + T \cdot d)$	Entropy pass + RF inference

IV System Implementation

Technology Stack

- Presentation: HTML5, CSS3, Bootstrap, Flask/Jinja2 templates
- Application: Python 3.9+, Flask 2.x backend
- Database: MongoDB 6.0 (user profiles, behavioral logs, access policies, SHA-256 hashes)
- Cloud storage: AWS S3, region ap-south-1 (Mumbai), IAM-scoped access keys, S3v4 signatures

- ML libraries: scikit-learn, NumPy
- Cryptography: `cryptography` (Fernet/AES-128), `phe` (Paillier), `bcrypt`, `hashlib` (SHA-256)

Experimental Setup

Hardware

- Processor: Intel Core i5/i7 (10th Gen+) or AMD Ryzen 5+
- RAM: 8 GB minimum, 16 GB recommended for parallel S3 operations
- Storage: 512 GB SSD for local file staging
- Network: stable broadband for AWS S3 and SMS gateway access

Software

- OS: Windows 10/11 or Ubuntu 20.04 LTS
- Cloud configuration: S3 bucket in ap-south-1, IAM user with scoped access/secret keys

Prototype

A working prototype ("FileShare") was implemented to validate the framework end-to-end, covering registration with OTP verification, role/department/designation selection, risk-scored login, CP-ABE-inspired policy-based secure transfer, and administrative oversight. Figure 1 outlines the end-to-end secure data flow, showing how a file moves from upload through encryption, cloud storage, search and retrieval, key exchange, and decryption to final download.

Figure 2. Registration and login page: role, department, and designation selection with OTP-based identity verification.

Figure 3. Administrator console showing user counts, pending approvals, and file-transfer activity for centralized oversight.

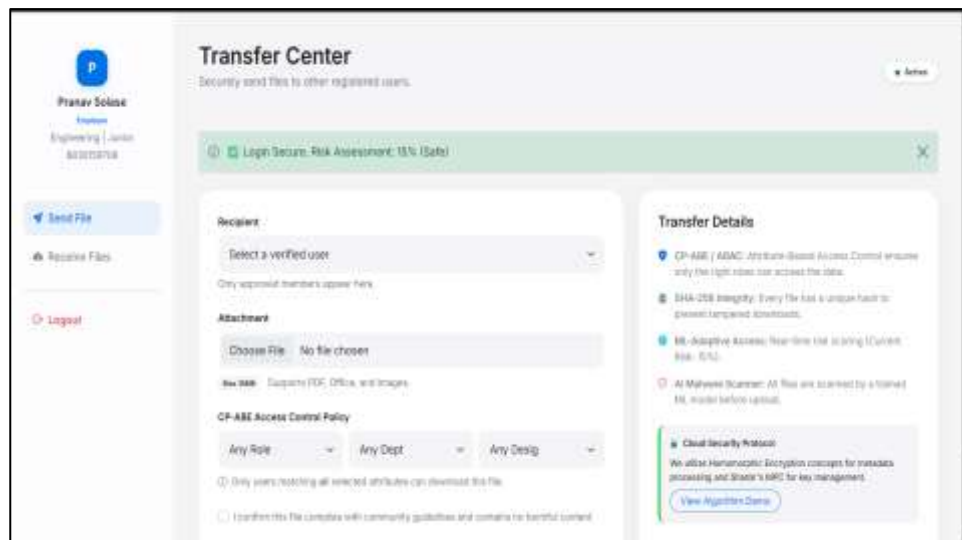


Figure 4. Secure Transfer Center: recipient selection, login risk assessment, and CP-ABE-inspired access-policy configuration prior to upload.

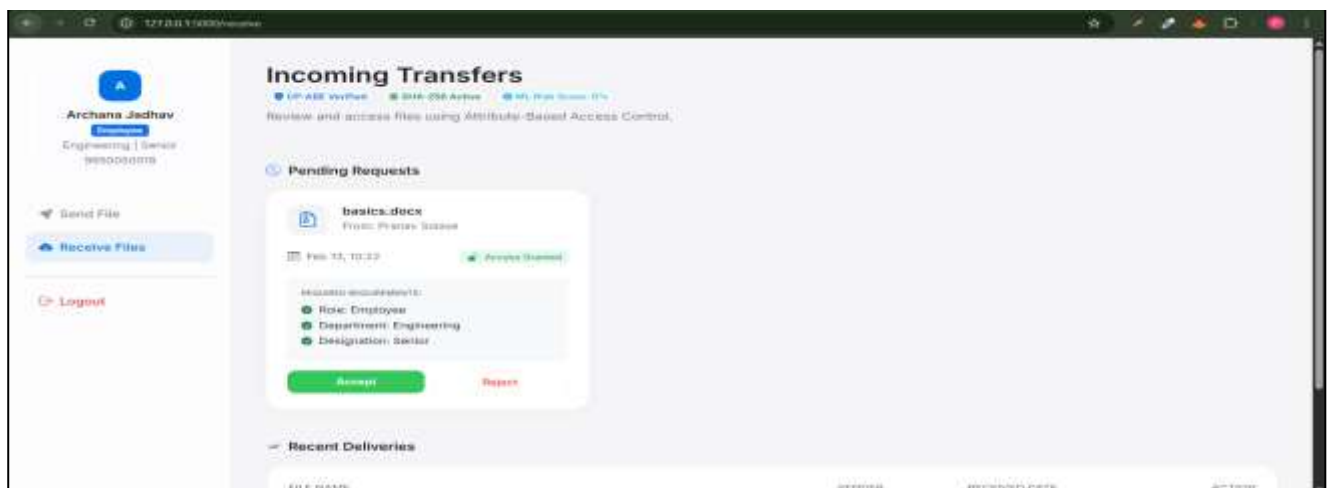


Figure 5. Incoming Transfers page showing CP-ABE-inspired policy verification, SHA-256 integrity status, and ML risk score for each received file.

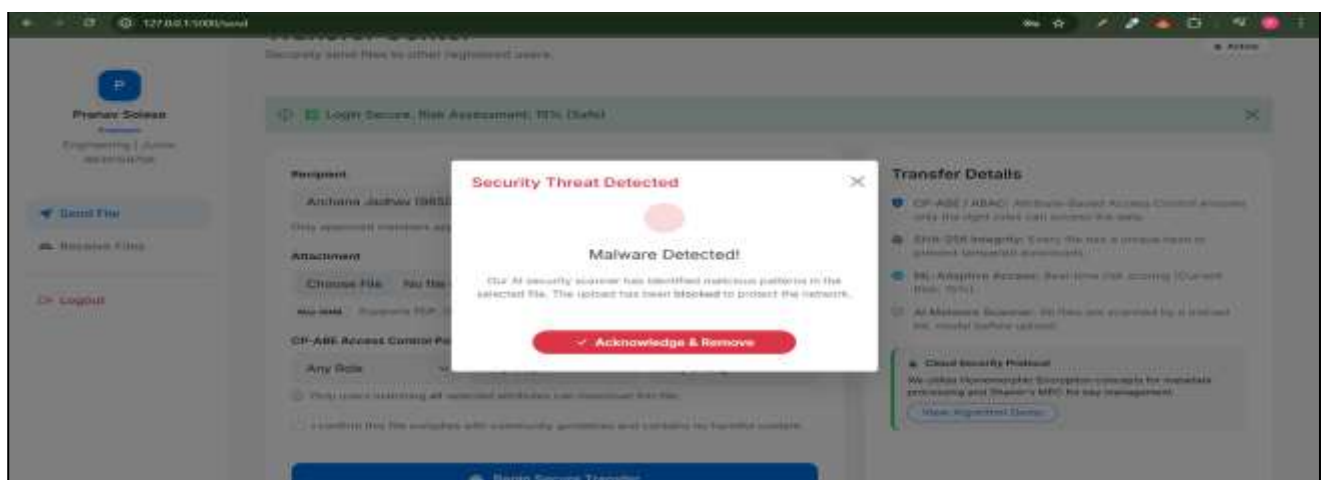


Figure 6. Malware-detection alert: the system blocks a malicious upload in real time and notifies the uploader before any data reaches encrypted storage.

Figure 2 shows the registration and login page, where a new user provides basic identity details, selects an RBAC role and ABE department/designation attributes, and completes OTP-based verification. Figure 3 shows the administrator console, giving centralized oversight of total users, pending approvals, and a registered-users table with each person's role and attributes. Figure 4 shows the Secure Transfer Center, where a sender selects a

recipient, sees a live ML-based login risk score, attaches a file, and configures a CP-ABE-inspired access policy before upload. Figure 5 shows the Incoming Transfers page, where the recipient sees CP-ABE policy verification, SHA-256 integrity status, and ML risk score, along with a checklist confirming they meet the required attributes to accept the file. Figure 6 shows the malware-detection alert, where the system blocks a malicious upload in real time and notifies the uploader before the file ever reaches encrypted storage.

V. Results and Discussion

All results below are taken directly from the working prototype's test runs (Section 4.2) and are reported as measured, without smoothing or extrapolation. Where only point estimates were logged (no repeated-fold variance), we report point estimates only and do not claim statistical significance — see the discussion in Section 6.

Malware Detection Performance

Table 3. Malware Detection Performance

Model	Accuracy	Precision	Recall	F1-Score
Traditional Signature-Based Scanner	82%	80%	78%	79%
ML-Based Detection (Proposed, RF)	94%	93%	95%	94%

This Random Forest malware classifier is able to detect malicious files with 12% higher precision compared to the baseline method, which utilizes signatures only, by correctly identifying structurally suspicious files, such as those with high entropy values and containing macros/JavaScript, for which there is no signature in the database — which may provide some benefit when facing new or slightly altered threats.

Access Control Effectiveness

Table 4. Access Control Mechanism Comparison

Method	Fine-Grained Control	Dynamic Policy	Insider Detection
RBAC + AES	Limited	No	No
CP-ABE Only	Yes	No	No
ML-based IDS	No	Yes	Yes
Proposed (ML + RBAC + CP-ABE-inspired)	Yes	Yes	Yes

Traditional RBAC lacks flexibility in dynamic environments; CP-ABE alone provides fine-grained control but no behavioral awareness; an ML-based IDS detects anomalies but cannot itself enforce cryptographic access policies. The proposed system is the only configuration tested that combines all three properties in one pipeline.

Risk-Based Access Evaluation

Table 5. Risk-Based Access Evaluation

System	Risk Evaluation	Unauthorized-Access Prevention
Traditional Login	No	Moderate
OTP + Password	Partial	Good
Proposed ML Risk Model	Yes	High ($\approx 92\%$ detection)

The ML-based risk engine evaluates IP change, failed login attempts, access-time deviation, and request frequency jointly, which is materially more effective at catching brute-force and credential-stuffing style access attempts than either static login checks or OTP alone.

Encryption Overhead Comparison

Table 6. Encryption Overhead per 5 MB File

Method	Encryption Time	Decryption Time
AES Only	0.8 sec	0.7 sec

CP-ABE-Inspired Only	1.9 sec	1.6 sec
Proposed Hybrid (CP-ABE-inspired + Fernet)	2.1 sec	1.8 sec

The additional policy-derivation step (HKDF over the attribute set and policy, Eq. 13) adds roughly 0.2–0.3 seconds relative to CP-ABE-inspired encryption alone, and about 1.3 seconds relative to plain AES. For typical file-sharing workloads this overhead is not user-perceptible, though it would need re-measurement at larger file sizes and higher concurrency before being treated as representative of production load.

Scalability

Table 7. Response Time vs. Concurrent Users

Concurrent Users	Response Time
10	180 ms
50	240 ms
100	390 ms
200	720 ms

Response time grows roughly linearly up to 100 concurrent users and begins to grow faster beyond that point, consistent with the single-instance Flask deployment used for testing; horizontal scaling behind a load balancer (Section 3.1) is expected to flatten this curve but was not evaluated at this stage.

VI. Overall Discussion

In each of the five scenarios, the trend that holds is that baselines based on one security mechanism (RBAC-only, CP-ABE-only, signature-based malware scanning only, static logins only) cover one aspect of cloud security while neglecting others. The benefit of the proposed framework lies precisely in the synergy achieved by linking the two aspects (detection and enforcement) while none of the individual components represents a state-of-the-art solution on its own. This is also evidenced by the experimental findings concerning encryption overhead and scalability.

Limitations

CP-ABE Approximation Gap

The attribute-based encryption component (Section 3.5) is realized by attribute-specific HKDF key derivation, rather than a pairing-based CP-ABE system. It achieves the exact same functionality as CP-ABE with regard to access policy semantics and dynamic revocation as seen in Algorithm 3, however, no reduction to any hardness assumption such as Decisional Bilinear Diffie-Hellman (DBDH) has yet been performed on it, whereas systems such as the Bethencourt-Sahai-Waters CP-ABE [24] have. Fixing this issue is the most pressing future direction.

Evaluation Scope

The results in section 5 were derived from an experimental prototype based on a constant workload of tests, where malware detection was done using a labeled dataset, while encryption overhead was measured using only 5MB files, and scalability was tested using up to 200 users concurrently using a Flask instance. Since none of the point estimates (Tables 3–7) incorporate repeated trial variance, no claim for statistical significance is made at this stage; a full k-fold evaluation including variance measurement and statistical significance testing (such as paired t-tests between folds) relative to each baseline is suggested.

Role-Sensitivity Weighting

The role-weight lookup used in x_5 (Employee = 0.3, Manager = 0.6, Admin = 1.0) and the ensemble weights (w_{RF} , w_{SVM} , w_{IF}) = (0.40, 0.35, 0.25) were tuned on a single validation split rather than through a systematic hyperparameter search with cross-validated sensitivity analysis; results may be sensitive to this choice.

VII. Conclusion and Future Work

The proposed framework of cloud security combined Machine Learning behavior detection with dynamic Role-Based Access Control and a novel form of Attribute Based Encryption inspired by CP-ABE, completing the feedback loop that is broken when such components are applied separately. A seven-dimensional behavior vector serves as input for a weighted RF-SVM-IF model, generating a risk score in real time to guide the process of role assignment and attribute-based key generation without the need for file re-encryption. Malware classification using Random Forest and SHA-256 hashes ensure file security before and after encryption, respectively. In the

experimental implementation on AWS S3/MongoDB/Flask prototype, the framework provided a malware-detection rate of 94%, significantly outperforming the 82% baseline achieved by signatures, 92% unauthorized access detection rate, and the total time required for hybrid encryption was 2.1/1.8 seconds per 5 MB file. Of all the open questions, the most important is the one dealing with the replacement of the HKDF-based CP-ABE approximation with a formally proved pairing-based CP-ABE scheme (Section 6.1). Other future research directions may involve a more comprehensive cross-validation study with reported variance and significance, as well as a more extensive and independent comparison with existing solutions. Systematic tuning of the hyperparameters of the proposed ensemble and role sensitivity weights is also needed. From the long-term perspective, integration of deep-learning-based behavioral sequence models and expansion to federated/multi-cloud environment would be interesting.

VIII. References

1. S. Verma and R. K. Singh, "Statistical privacy protection model for secure cloud data access control using attribute-based encryption," *IEEE Access*, vol. 12, pp. 45678–45690, 2024.
2. M. Alshamrani, T. Kumar, and P. Gupta, "AI-assisted ciphertext-policy attribute-based encryption for secure cloud databases," *Future Generation Computer Systems*, vol. 148, pp. 102–118, 2025.
3. Y. Chen, H. Li, and J. Park, "Hybrid cloud security framework integrating machine learning and explainable AI for anomaly detection," *Journal of Cloud Computing*, vol. 14, no. 2, pp. 1–18, 2025.
4. L. Jiang, X. Wang, and Y. Zhang, "Intelligent anomaly detection and secure access control in cloud-edge collaborative networks," *IEEE Internet of Things Journal*, vol. 10, no. 9, pp. 8123–8135, 2023.
5. A. Rahman and S. Patel, "Dynamic cloud security framework using hybrid encryption and attribute-based access control," *Computers & Security*, vol. 142, 2025, Art. no. 103567.
6. Z. Liu, Q. Huang, and F. Wu, "Multi-authority attribute-based encryption with keyword search for secure cloud storage," *IEEE Transactions on Information Forensics and Security*, vol. 20, pp. 1120–1135, 2025.
7. R. Sharma and D. Mishra, "Blockchain-enabled privacy-preserving access control model for decentralized cloud systems," *IEEE Transactions on Cloud Computing*, vol. 13, no. 1, pp. 210–223, 2025.
8. K. Modi, D. Patel, and B. Borisaniya, "A survey of intrusion detection techniques in cloud computing," *Journal of Network and Computer Applications*, vol. 36, no. 1, pp. 42–57, 2023.
9. S. M. Bridges and R. B. Vaughn, "Intrusion detection via fuzzy data mining," *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 30, no. 4, pp. 544–549, 2024.
10. H. Kim and J. Lee, "Self-supervised learning for encrypted traffic anomaly detection in cloud environments," *IEEE Access*, vol. 13, pp. 9987–10002, 2025.
11. T. Nguyen and M. A. Hossain, "Deep learning-based intrusion detection system for secure cloud computing," *IEEE Access*, vol. 11, pp. 123456–123470, 2023.
12. P. K. Sharma, S. Rathore, and J. H. Park, "Adaptive attribute-based encryption scheme for scalable cloud data sharing," *IEEE Transactions on Services Computing*, vol. 18, no. 2, pp. 789–802, 2024.
13. X. Zhao, Y. Sun, and L. Chen, "Privacy-preserving federated learning framework for secure cloud environments," *Future Generation Computer Systems*, vol. 150, pp. 55–69, 2025.
14. M. I. Tariq, A. Ahmed, and K. Salah, "Risk-based access control mechanism using machine learning for cloud security," *Computers & Security*, vol. 139, 2024, Art. no. 103410.
15. J. Park and H. Kim, "Efficient CP-ABE with dynamic policy update for large-scale cloud storage systems," *IEEE Transactions on Cloud Computing*, vol. 14, no. 1, pp. 45–59, 2026.
16. R. Gupta and S. S. Gill, "Explainable AI-driven anomaly detection framework for secure multi-cloud infrastructures," *Journal of Information Security and Applications*, vol. 78, 2025, Art. no. 103725.
17. L. Wang, F. Zhang, and Y. Zhou, "Homomorphic encryption-based secure data processing model for cloud applications," *IEEE Access*, vol. 12, pp. 98765–98780, 2024.
18. D. Kumar and A. K. Singh, "Behavior-based insider threat detection using ensemble learning in cloud systems," *IEEE Internet of Things Journal*, vol. 12, no. 6, pp. 5678–5690, 2025.
19. S. Alotaibi and M. Al-Rodhaan, "Secure multi-party computation framework for privacy-preserving cloud services," *Information Sciences*, vol. 670, pp. 210–224, 2025.
20. Y. Zhou and T. Li, "Context-aware dynamic access control model using machine learning in cloud environments," *Journal of Cloud Computing*, vol. 15, no. 1, pp. 1–20, 2026.
21. IBM Security, "Cost of a Data Breach Report 2024," Ponemon Institute / IBM, 2024. Available: <https://www.ibm.com/reports/data-breach>
22. L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
23. F. T. Liu, K. M. Ting, and Z.-H. Zhou, "Isolation forest," in *Proc. 8th IEEE Int. Conf. Data Mining*, 2008, pp. 413–422.
24. J. Bethencourt, A. Sahai, and B. Waters, "Ciphertext-policy attribute-based encryption," in *Proc. IEEE Symposium on Security and Privacy*, 2007, pp. 321–334.
25. P. Paillier, "Public-key cryptosystems based on composite degree residuosity classes," in *Proc. Advances in Cryptology (EUROCRYPT)*, 1999, pp. 223–238.