
RESEARCH ARTICLE**Enhancing Data Privacy and Governance in Cloud-Deployed Large Language Models Using Deep Learning-Based Risk Mitigation****Naresh Alapati¹, Ramachander Rao Thallada² and Koteswararao Nallabothu³**¹ *Principal Cloud Architect, Walmart, Bentonville, USA*² *GRC Executive, Manulife, Canada*³ *Lead Software Engineer, Lowes, Charlotte, USA***Corresponding Author:** Naresh Alapati, **E-mail:** narencachy@gmail.com

ABSTRACT

The growing number of problems related to the privacy and security of data in cloud computing requires the need for efficient protection schemes. In this paper, we present an innovative solution for ensuring the privacy and management of data in large language models deployed in clouds by applying risk mitigation techniques based on deep learning. The technique involves risk mitigation, secure data management, and model performance enhancement. First, data input is acquired from a synthesized cybersecurity logs dataset and filtered using data preprocessing to eliminate any form of inconsistencies. Next, data preprocessing enables the input to be directed to the feature extraction phase utilizing the Dense Channel Spatial Semantic Guidance Attention UNet (DCSGA-UNet), which allows for the extraction of relevant features, such as Log_ID, IP_Address, Request_Type, and Response_Time_ms. After that, these features undergo processing through a Bayesian Constitutive Artificial Neural Network (BCANN) classifier that helps predict the probability of a threat or non-threat occurrence. The technique is designed using Python and evaluated against current methods. The experiments show that our innovative solution exhibits a higher accuracy of 98% in a lesser amount of time of 1.07 seconds.

KEYWORDS

Cloud security, data privacy, risk mitigation, deep learning, governance

ARTICLE INFORMATION**ACCEPTED:** 02 June 2026**PUBLISHED:** 18 August 2026**DOI:** 10.32996/fcsai.2026.5.9.18

1. INTRODUCTION

The rapid adoption of cloud-hosted large language models (LLMs) has significantly transformed data-driven applications across domains such as healthcare, finance, and intelligent networking (Pendy , 2025) However, the large-scale data processing requirements of these models introduce critical concerns related to data governance, including privacy preservation, regulatory compliance, and secure data handling (Yan et al., 2023) Ensuring robust governance mechanisms is essential to maintain trust, protect sensitive information, and support ethical AI deployment. In this context, integrating deep learning techniques for intelligent risk analysis and automated governance enforcement becomes highly important (Vankayalapati , 2025) Such an approach enables proactive identification of anomalies, enhances data integrity, and ensures that LLM systems operate within defined legal and ethical boundaries while maintaining scalability and performance in cloud environments.

Despite these advantages, several challenges persist in cloud-based LLM systems, primarily due to the opaque nature of deep learning models, heterogeneous data sources, and dynamic threat landscapes (Agarwal et al., 2024) . Issues such as data leakage, bias propagation, lack of transparency, and vulnerabilities in cloud infrastructures can significantly impact system reliability and compliance. Moreover, traditional governance frameworks often fail to adapt to real-time risks and complex data flows (Tsai et al., 2026) To overcome these limitations, this work proposes a comprehensive governance-aware framework that incorporates secure data preprocessing, encryption-based storage, and deep learning-driven risk assessment mechanisms. Additionally, explainable AI

techniques and continuous monitoring strategies are employed to improve model transparency and detect anomalies in real time, thereby ensuring a resilient, compliant, and trustworthy cloud-based LLM ecosystem.

2. LITERATURE SURVEY

Several research works presented in literature were primarily concerned with Cloud-Deployed Large Language Models a few of them were reviewed here

Vincent et al., (2025) A Security framework has been created to ensure that there is no data leakages or injection attacks in the large language models. The model will become more secure while analyzing the users' input information. This can be done through adversarial training, validation of the inputs, and privacy preservation techniques. However, the technique may increase the computation cost, and still not be completely immune to attacks.

Tsai et al., (2026) has designed a cloud-native, LLM-driven dynamic honeypot framework that is governed by stringent policies and state checks to provide secure and authentic engagements. The framework made use of microservice-based deployment, session brokering, deterministic emulation, constrained LLM generation, and governance techniques. Nevertheless, the methodology had limitations such as increased complexity, scalability issues, and inability to counter highly adaptive or zero-day attacks.

Naik, (2023) have managed to develop a security analysis on cloud computing regarding data and network security. The aim here was to make sure that confidentiality, integrity, and availability were maintained within the cloud environment. Security frameworks and risk management principles were used in this case, but some difficulties have emerged.

Salako et al., (2024) have created a dynamic information governance framework for AI-based cloud systems to improve security and compliance. It aimed at filling the gaps in privacy, incident management, and governance efficiency. For this research, PETs, ethics committee, and SEM, Cox proportional hazard models, and DiD methodology were applied. However, potential difficulties lie in sectoral weaknesses and complexity of governance frameworks implementation.

Hareram, (2025) have developed an architecture for protecting privacy in a cloud computing environment for large enterprises considering issues of IAM, DLP, and governance structures. This research aimed to protect and manage the data lifecycle within cloud computing in a manner that ensures regulatory compliance. It made use of various models based on combining cloud computing with regulatory requirements such as GDPR and HIPAA.

Table 1: Summary of existing literature

Ref	Author	Contribution (Key Idea)	Techniques Used	Limitations
(2025)	Vincent et al.,	Security framework for LLMs to prevent data leakage and injection attacks	Adversarial training, input validation, privacy preservation	High computational cost; not fully attack-proof
(2026)	Tsai et al.,	Cloud-native LLM-based dynamic honeypot for secure interaction	Microservices, session brokering, constrained LLM generation	High complexity; scalability issues; weak against zero-day attacks
(2023)	Naik et al.,	Security analysis for cloud ensuring CIA (Confidentiality, Integrity, Availability)	Risk management, security frameworks	Implementation challenges; limited adaptability
(2024)	Salako et al.,	Dynamic governance framework for AI-cloud systems	PETs, SEM, Cox model, DiD methodology	Complex implementation; sector-specific limitations
(2025)	Hareram et al.,	Privacy-preserving cloud architecture for enterprises	IAM, DLP, GDPR, HIPAA compliance models	Complex integration; regulatory overhead

Research Gap

While considerable progress has been made toward ensuring the security of cloud computing and artificial intelligence technologies, existing literature identifies a number of research gaps related to the development of a single, scalable, and adaptive security framework. (<https://www.kaggle.com/datasets/sushrutha123/synthetic-cybersecurity-logs-dataset/data>) Indeed, most of the current research efforts consider individual issues of cloud security and artificial intelligence separately, ranging from LLM security, detection of threats through honeypots, cloud data security, and governance models. While each of these areas represents an interesting subject of study, existing frameworks suffer from a number of limitations, including poor scalability, computational complexity, and lack of robustness against zero-day attacks. Governance models, in turn, may prove difficult to implement due to the nature of sector-specific vulnerabilities and organizational support needs. Finally, privacy mechanisms, which have shown considerable success, fail to provide an integrated framework for both privacy protection and threat detection.

The major contribution of this work is below ,

- Introduced a new deep learning-based system for improving data privacy and governance in cloud-based large language models through risk management.
- Designed a hybrid model combining DCSGA-UNet and BCANN, where DCSGA-UNet successfully extracts features and BCANN precisely estimates privacy risks and categorizes threats.

- Demonstrated outstanding results with an accuracy rate of 98% and processing time of 1.07 sec, surpassing other systems based on machine learning, artificial neural network, and deep learning techniques.

Novelty

This study suggests an innovative model of hybrid deep learning system, which consists of DCSGA-UNet and BCANN, providing comprehensive security against risks in cloud-based language models. What sets it apart from the traditional approaches is its ability to combine attention-based feature extraction with probabilistic risk prediction techniques, ensuring better performance and speed (accuracy of 98%, time of 1.07 seconds).

3. RESEARCH METHODOLOGY

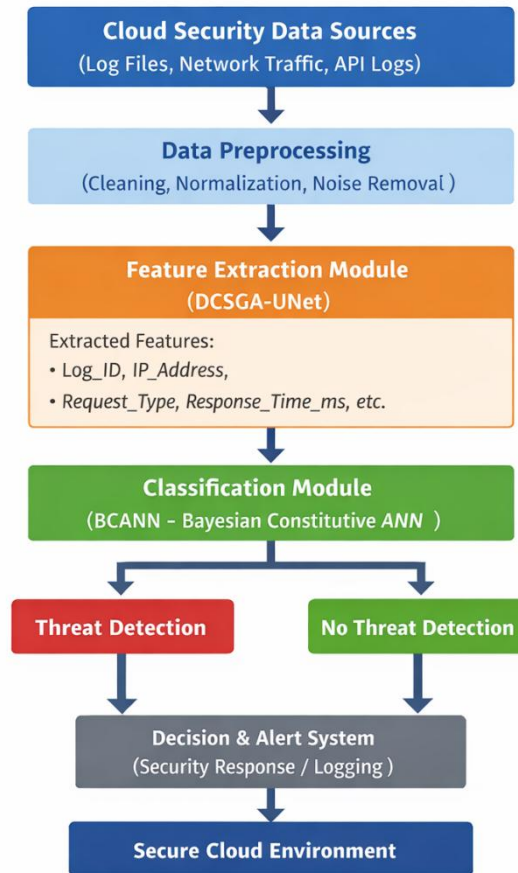


Figure 1 : Block Diagram of the proposed methodology

The proposed based on deep learning for large language models in the cloud to increase data privacy and governance. In the methodology, first, the data are pre-processed, followed by the feature extraction step using DCSGA-UNet to extract the spatial and semantic features of the data. Then, using BCANN, the probability risk assessment of the cyber attacks is performed to find any threats. Figure 1 shows the Block Diagram of proposed methodology

3.1 Data Description

In this research, the Synthetic Cyber security Logs Dataset on Kaggle is applied to generate an imitation cloud-based setting to identify possible data governance risks within the LLM architecture. (Hussain et al., 2025; Linka et al., 2025) The data used in this study includes synthetic logs related to users' and systems' actions, which include logging in actions and access patterns, categorized into Threat and No Threat types. In the current study, this data helps test whether the proposed solution can effectively identify anomalous activities, malicious logins, or any other activities that pose risks to LLMs' data governance compliance. Despite being synthetic, the data successfully enables the risk identification processes in cloud-based infrastructure and governance of machine learning solutions.

3.2 Data Pre-processing

In this section, data pre-processing was conducted to make sure that quality input data was provided for the secure analysis in AI-based cloud computing systems. First of all, data cleaning was conducted to eliminate noise, discrepancies, missing values, and

duplicate records in the data originating from various cloud sources. Next, normalization was done to scale the data within a specific range, thereby optimizing the effectiveness of the models being designed. Moreover, private data was managed prudently during data pre-processing.

3.3 Feature extraction using Dense Channel Spatial Semantic Guidance Attention UNet (DCSGA-UNet)

This section discusses the process of feature extraction utilizing the Dense Channel Spatial Semantic Guidance Attention U-Net (DCSGA-UNet). The approach has been chosen owing to its effectiveness in capturing dependencies among channels and space. This allows more effective learning of features from cloud security logs that have high-level and spatial representations. DCSGA-UNet provides an enhanced representation capability that facilitates the extraction of useful features including Log_ID, IP_Address, Request_Type, and Response_Time_ms through minimizing the inclusion of noise. As compared to other traditional techniques, the DCSGA-UNet has improved feature discrimination capabilities and enhanced learning capability, thus ensuring accuracy.

$$M_c(F_c) = \sigma[MLP(F_{cMax}) \oplus MLP(F_{cAvg})] \quad (1)$$

Where, F_c which are then concatenated using $\oplus$ for further processing. The activation function σ refines the result for the next detection steps. The refined spatial features F_{sr} are computed by multiplying the spatial attention weights. The model extracts structural features from cloud security log data, helping to identify patterns associated with malicious activity, intrusion behavior, and system anomalies, is expressed in Equation (2)

$$F_{cr} = M_c(F_c) \otimes F_c \quad (2)$$

Where, F_c represent the input feature map after convolutional layers, $M_c(F_c)$ represent a channel attention map generated from F_c , F_{cr} was the refined feature map after applying channel attention. DCSGA-UNet it can be seen that feature extraction based on DCSGA-UNet results in better representations and contributes to the detection of security threats in cloud computing systems.

3.4 Classification using Bayesian constitutive Artificial Neural Network (BCANN)

This part highlights the use of classification through the Bayesian Constitutive Artificial Neural Network (BCANN) to classify cloud security events into two categories: Threat and No Threat, depending on the extracted features. The reason why this technique was selected is that it performs exceedingly well in dealing with uncertainties, probability theory, and nonlinearity associated with cybersecurity information. BCANN involves the incorporation of Bayesian analysis and artificial neural networks together for efficient reliability and robustness when presented with noisy input data, among others.

$$q(\theta) = \prod_{j=1}^m q_i(N_j, N_j^*) \quad (3)$$

Where, N_j is a j -th input feature, N_j^* is a predicted output for the j -th input feature, m is a number of elements in the neural network. The posterior distribution is updated based on the observed data and is computed, which was expressed in equation (4).

$$G = \nabla_y \phi \quad (4)$$

Where, G is the Gradient of the decision function with respect to the output class, ∇_y is the Gradient operator with respect to the output class, finally the BCANN classifier proves to be quite reliable and precise when detecting cloud security threats through a proper distinction between threat and no-threat cloud behaviors in uncertain and complicated conditions.

4.RESULT AND DISCUSSION

The efficacy of the suggested deep learning-enabled framework for improved privacy protection and governance in LLMs deployed on cloud systems was assessed based on its implementation via Python programming on an Intel® Xeon® CPU (2.2 to 2.3 GHz) with 12.7 GB RAM and Ubuntu 18.04/20.04. The performance of the method was compared with traditional techniques such as ML, ANN, and DL. The results revealed that the proposed technique offered improved performance.

Figure 2: Analysis the Accuracy

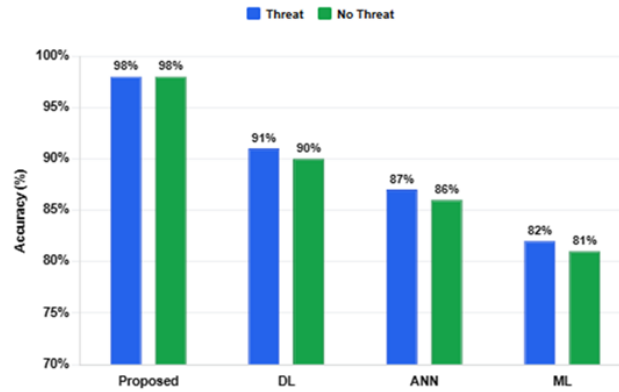


Figure 2: Analysis the Accuracy

The accuracies of the proposed framework as compared with those of other approaches have been demonstrated in Figure 2 when there is both a threat and absence of a threat to information systems. The proposed framework shows an accuracy rate of 98% in all cases, which is higher than those shown by DL (91%, 90%), ANN (87%, 86%), and ML (82%, 81%). It is apparent from the results that the proposed framework performs better.



Figure 3: Analysis the precision

Figure 3 shows the comparison of the precision of the proposed model and other models in presence and absence of threats. It can be observed that the proposed method provides the best precision at 97% in both scenarios, while DL, ANN, and ML provide 90%, 89%; 86%, 85%; and 81%, 80% precision respectively.

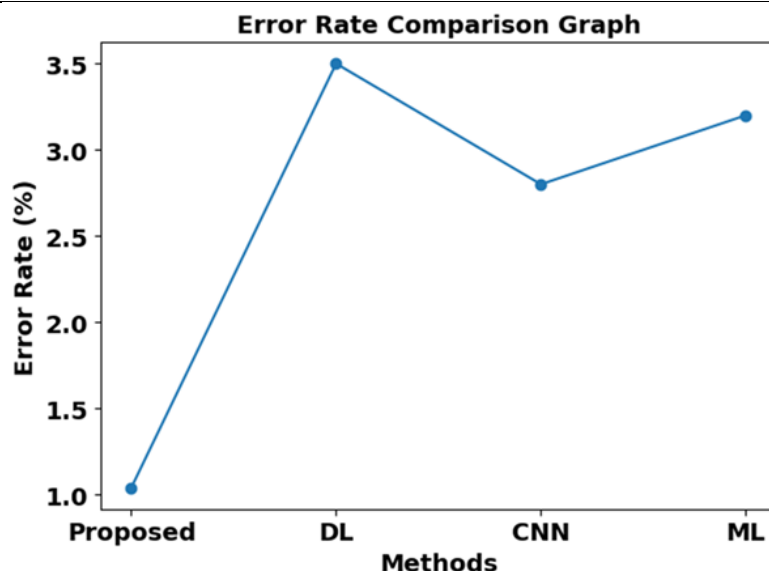


Figure 4: Analysis the Error Rate

Error rates for various approaches have been illustrated in Figure 4. The proposed model has produced an error rate of around 1.0% which is quite better than the error rates provided by DL, CNN, and ML, i.e., 3.5%, 2.8%, and 3.2%, respectively.

Table 2: Comparison of performance proposed and exiting approach

Methods	RMSE (%)	Computational Time (sec)
Proposed	1.37	1.07
DL	3.50	2.80
ANN	2.80	2.10
ML	2.10	1.75

According to the findings provided in Table, the proposed approach obtains the minimum RMSE of 1.37% and least computing time of 1.07 seconds. On the other hand, DL, ANN, and ML have larger RMSEs varying between 2.10% and 3.50%, as well as computing time within the range of 1.75 to 2.80 seconds. This clearly suggests that the proposed methodology offers more accurate estimates with better computation efficiency.

5. CONCLUSION

In summary, the framework presented above provides an effective way to protect data privacy and management in large language models used in cloud computing through threat recognition and classification. The incorporation of the DCSGA-UNet and BCANN model in the framework ensures improved features extraction and risk prediction, giving an accuracy of 98% and minimal RMSE of 1.37% as compared to the other techniques available. The effectiveness of the framework in securing cloud computing shows its applicability in various fields including cybersecurity, health care, enterprise applications, and financial security. Further research on the framework will include real-time implementation among others.

Limitation

The first limitation of this study is the evaluation of the model through a synthesized cybersecurity logs dataset, which might not be able to reflect all the complexities and uncertainties of actual cloud systems. The performance and reliability of the proposed architecture in dealing with real-world challenges still remains unknown. Furthermore, the implementation of DCSGA-UNet along with BCANN results in increased computational complexity, making its deployment more difficult.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers.

REFERENCES

- [1]. Agarwal, G., Katiyar, V., & Tiwari, A. (2026). Cloud-Based AI Development Tools: Making Machine Learning is Accessible to Everyone. *Lecture Notes in Networks and Systems*, 15–24. https://doi.org/10.1007/978-981-95-3309-1_2
- [2]. Sruthi Erra Hareram. (2025). Data Privacy in Cloud-Based Pipelines: IAM, DLP, and Governance. *Journal of Computer Science and Technology Studies*, 7(11), 23–29. <https://doi.org/10.32996/jcsts.2025.7.11.4>
- [3]. Hussain, T., Shouno, H., Mohammed, M. A., Marhoon, H. A., & Alam, T. (2025). DCSSGA-UNet: Biomedical image segmentation with DenseNet channel spatial and Semantic Guidance Attention. *Knowledge-Based Systems*, 314, 113233. <https://doi.org/10.1016/j.knosys.2025.113233>
- [4]. Linka, K., Holzapfel, G. A., & Kuhl, E. (2025). Discovering uncertainty: Bayesian constitutive artificial neural networks. *Computer Methods in Applied Mechanics and Engineering*, 433, 117517. <https://doi.org/10.1016/j.cma.2024.117517>
- [5]. Naik, S. (2023). Cloud-Based Data Governance: Ensuring Security, Compliance, and Privacy. *The Eastasouth Journal of Information System and Computer Science*, 1(01), 69–87. <https://doi.org/10.58812/esiscs.v1i01.452>
- [6]. Bimalendu Pandy. (2025). ChatGPT, Machine Learning, and Cloud AI Technologies for Smarter Healthcare and Enhanced Cybersecurity. *Global Research Repo*, 1(3), 37–59. <https://globalfoodresearch.com/index.php/grr/article/view/32>
- [7]. Salako, A. O., Fabuyi, Jumai Adedoja, Taiwo, A. N., Selesi-Aina, O., Louisa, D.-O. D., & Olaniyi, Oluwaseun Oladeji. (2024). Advancing Information Governance in AI-Driven Cloud Ecosystem: Strategies for Enhancing Data Security and Meeting Regulatory Compliance - ScienceOpen Library. *Submanuscript.com*. <http://journal.submanuscript.com/id/eprint/2644/1/Salako17122024AJRCOS127573.pdf>
- [8]. Tsai, S. E., Tsai, T. T., & Aun, M. H. (2026). *LLM-Enabled Cloud-Native Dynamic Honeypot Systems: Architecture, Ethical Governance, and Empirical Evaluation*. <https://doi.org/10.21203/rs.3.rs-9003091/v1>
- [9]. Vankayalapati, R. K. (2025). Zero-Trust Security Models for Cloud Data Analytics: Enhancing Privacy in Distributed Systems. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.5121185>
- [10]. Alapati, N., Thallada, R. R., & Nallabothu, K. (2026). *Enhancing Data Privacy and Governance in Cloud-Deployed Large Language Models Using Deep Learning-Based Risk Mitigation*. <https://doi.org/10.21203/rs.3.rs-9420186/v1>