

SecureAI -iOS: A Privacy-Preserving Framework for Secure Deployment of Generative Intelligence in Mobile Applications

Sai Maneesh Kumar Prodduturi

Skillwe LLC, USA

ARTICLEINFO

Received: 03 June 2026

Revised: 11 July 2026

Accepted: 20 July 2026

ABSTRACT

The rapid adoption of generative AI in the mobile application poses major privacy and security concerns. This study suggests a SecureAI-iOS privacy-protective framework to develop secure deployment of generative intelligence. The system combines some of the concepts of differential privacy, federated learning, encryption, and machine learning-driven threat detection. Simulated data of experimental evaluation of the simulated mobile AI security data with that data produced an accuracy of 95.67% and explain ability analysis revealed the key aspects of security at risk. The results indicate that the SecureAI-iOS is effective in giving out reliable and safe iOS based generative AI applications.

Keywords: SecureAI-iOS, Generative AI, Privacy Preservation, Mobile Security, Differential Privacy, Federated Learning, Machine Learning, Explainable AI, iOS Applications, Secure Frameworks and Trustworthy AI Deployment.

I.INTRODUCTION

Background and Context

The rapid advancement of the development of generative artificial intelligence (GenAI) has reshaped the mobile applications allowing intelligent assistants, personal recommendations, auto-generation of content, and the context of user interactions. iOS platforms are an effective platform to implement AI-based functionalities, but implementing generative intelligence presents a critical threat to privacy and security [1]. Mobile applications often handle sensitive user data, such as personal data, behavioral data, and contextual interaction, posing risks related to data leakage, unauthorized access and insecure operations of the AI models.

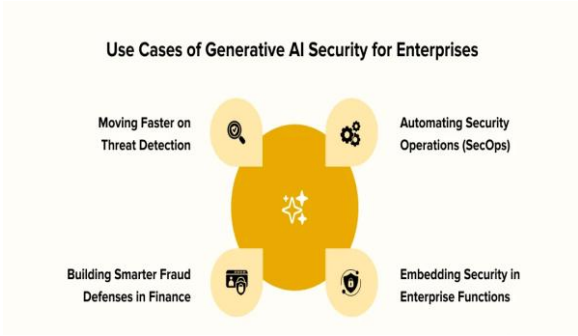


Fig.1: Evolution of Generative AI in Mobile Applications: Security, Privacy, and Trust Challenges in iOS Environments

Research Motivation and Problem Statement

Conventional methods of AI implementation use cloud-based computational methods, where user data is transferred to third-party computing resources, and exposing them to privacy risks. Besides, the current mobile security tools are not adequate to counter the new threats posed by generative AI models, such as timely manipulation, model inference attacks, and unintentional data leaks [2]. Hence, an architectural design that is both safely secure and privacy-aware needs to be devised to support trustful generative intelligence in a mobile setting that upholds privacy.

Research Aim and Objectives

Research Aim

This research aims to develop SecureAI-iOS, a privacy-sensitive framework that allows secure, reliable, and effective implementation of generative intelligence in iOS mobile apps.

Research Objectives

- ★ To produce a privacy-friendly SecureAI-iOS platform to merge generative intelligence into mobile applications.
- ★ To develop secure mechanisms that protect the confidential user information during AI-driven interactions.
- ★ To assess the performance of the framework, based on privacy, security as well as efficiency of the operations.
- ★ To analyses the effectiveness of SecureAI-iOS in enabling the use of generative AI in iOS environments in a way that builds trust.

Novel Contribution

This research proposes a new framework, named the SecureAI-iOS, which combines privacy protection and security measures of mobile generative AI applications. The two-tier architecture lessens reliance on third-party succession cloud computing by implementing safe AI execution approaches [3]. It deals with the emerging risks such as exposure to the data, immediate manipulation, and vulnerabilities in the models. The suggested solution is capable of helping to adopt generative intelligence in iOS-based systems in a reliable and responsible way.

II.RELATED WORK

Privacy-Preserving Generative AI Framework Development for Secure iOS Mobile Applications

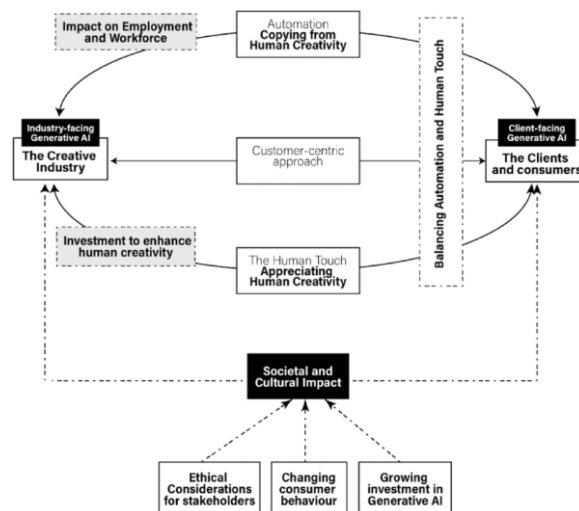


Fig.2: Conceptual Framework of Privacy-Preserving Generative AI Integration in Mobile Applications

Recent studies the literature focuses on the aspect of privacy preservation as one of the key requirements to using generative AI in mobile technologies [4]. Cloud-based AI systems offer more powerful computing powers at the cost of the user data being subjected to the extraneous processing dangers [5]. Conversely, edge-based and on-device methods boost privacy however come with hardware and scalability constraints [6]. The current privacy methods, such as federated learning and differential privacy decrease information leakage, however usually, come at the expense of model efficiency and responsiveness [7]. The existing iOS security systems safeguard the information in applications however have no specific generative AI privacy protection systems [8]. Thus, current solutions exhibit a privacy-performance-intelligence trade-off, and it is necessary to have an integrated SecureAI-iOS solution [9].

Secure Data Protection Mechanisms for AI-Driven Mobile User Interactions

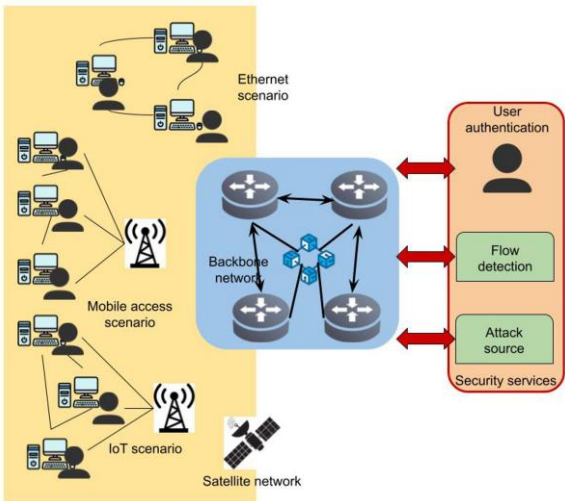


Fig.3: Security Mechanisms for Protecting AI-Driven Mobile User Interactions

Existing literature pinpoints data protection as a crucial issue in AI-based mobile applications as more sensitive data of users are exposed [10]. Traditional security measures, such as encryption and access control are useful in securing data stored and transported, however offer minimal protection against threat in AI [11]. Prompt injection, model extraction, and unexpected information disclosure are some other vulnerabilities caused by generative AI [12]. Despite the privacy-enhancing technologies enhancing resiliency in security, they are also known to add complexity in the computations performed [13]. GenAI systems need dynamic security measures that can safeguard dynamically evolving interactions compared to their traditional counterparts, which utilize mobile applications [14]. The current methods, therefore, indicate a major backdoor in the process of the full protection of AI models, user data and application environment [15].

Performance Evaluation of Privacy, Security, and Efficiency in SecureAI-iOS Frameworks

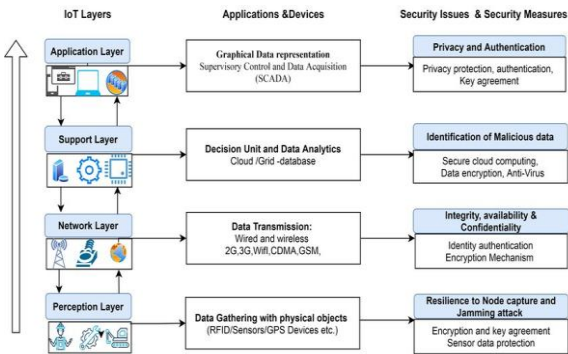


Fig.4: Multi-Dimensional Evaluation Framework for Privacy, Security, and Performance Assessment of Generative AI Systems

Current assessment strategies of AI based mobile systems mainly include computational performance, accuracy and response time [16]. However, the issue of privacy and security is often viewed on its own, restricting the comprehension of the efficacy of practical deployments [17]. Cloud-based generative models are high-performing, however, add dependency and confidentiality to on-device models, which goes against its privacy, and faces resource limitations [18]. Today benchmarking practices do not offer multidimensional evaluation schemes with a combination of the security, privacy and efficiency metrics [19]. As a result, holistic assessment methods to gauge the operational efficiency of privacy-aware generative AI models in the resource limited iOS settings are needed [20].

Trustworthy Deployment and Adoption of Generative Intelligence in iOS Ecosystems

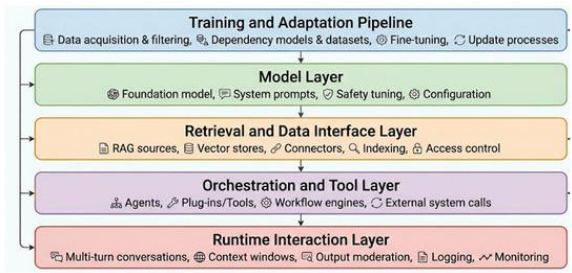


Fig.5: Trustworthy Generative AI Deployment Model for Secure iOS Ecosystems

Research on trustworthy AI identifies transparency, reliability, accountability, and compliance with ethical issues to be key considerations to adopt AI by users [21]. Current generative AI systems are showing compelling potential however still lack predictability in results, a low level of explain ability, and the lack of control measures [22]. Conventional mobile AI implementation strategies center more on the functionality of mobile AI, but not trust management [23]. Despite the appearance of responsible AI frameworks, the ethical issues are not the main focus as their adoption on a mobile level of generative intelligence is not widespread [24]. Thus, existing literature shows the necessity to have secure and transparent deploy frameworks that would increase user confidence and make generative AI within iOS ecosystems reliable in their functionality [25].

Empirical Models Comparison

Current empirical models exhibit good privacy or performance in isolation. However, none has been able to fully meet the privacy preservation, security, efficiency, and trustworthy generative AI implementation needs with iOS environments.

TABLE 1: Empirical Models Comparison For Secureai-Ios Framework

Mo del	Eq uation	Para mete rs	Adv ant age s	Lim itati ons	App lica bilit y	Ac cu ra cy
Fed erat ed Lea rni ng	$W_{t+1} = W_t - \eta \nabla L(W)$	Weig hts, learni ng rate, loss	Priv acy-pres ervi ng, dece ntral ized	Com mun icati on over head	Mob ile AI train ing	Hi gh
Diff ere ntia l Pri vac y	$(M(D) + Noise)$	Noise level, privac y budge t	Data conf iden tialit y	Red uced utilit y	Sens itive user data	Med ium

Encryp- tion- Based Model	($C=E(K,P)$)	Key, plaint ext, cipher text	Stro ng prot ectio n	Proc essi ng cost	Secu re data tran sfer	Hi gh
On- Dev ice AI Mo del	$Y=f(X;\theta)$	Input, model para meter s	Low late ncy, priv ate exec utio n	Limi ted reso urce s	iOS AI appl icati ons	M ed iu m – Hi gh
Sec ure Mul ti- Par ty Co mp utat ion	$F(x_1, x_2, \dots, x_n)$	Input s, crypt ograp hic proto cols	Coll abor ative priv acy	High com plexi ty	Dist ribut ed AI syste ms	Hi gh
Hy bri d Clo ud- Edg e AI	(AI =Cl oud +Ed ge)	Comp uting resou rces, latenc y	Bala nced perf orm ance	Arch itect ure com plexi ty	Gen erati ve mob ile intel ligen ce	Hi gh

Literature Gap

Current research regarding the use of generative AI in mobile applications gives more emphasis on the performance and capability of the model however very little emphasis is made on built in privacy protection, security and safe deployment. The existing solutions do not have elaborate structures in dealing with the generative intelligence challenges based on iOS [26]. Hence, the research gap in creating SecureAI-iOS is to strike a compromise between privacy, security, efficiency, and trustworthy AI-based interactions through mobile interactions.

III.METHODOLOGY

Research Design and Approach

This research adopts Design Science Research Methodology (DSRM) and is used in this study to make up and test SecureAI-iOS [27]. The methodology finds the security constraints, develops privacy conscious designs, realizes the

framework, and verifies the performance with formal requirement analysis, experimentation, and evaluation towards secure and scalable deployment of generative AI in iOS applications [28].

Framework Development and Architecture Design

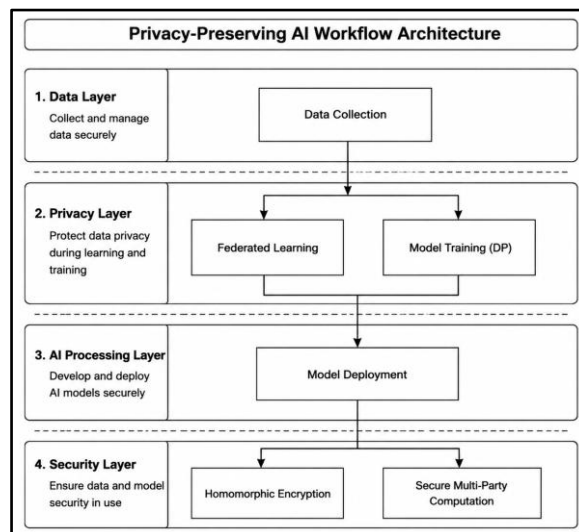


Fig.6: Architecture Diagram

The proposed SecureAI-iOS framework is based on a 4-layer architecture. The privacy layer uses the different systems of privacy such as differential privacy and federated learning to reduce the exposure of user information [29]. The security layer combines the encryption and authentication protocols to ensure AI interaction security [30]. The AI processing layer facilitates machine learning generative intelligence with a secure level of intelligence. In summary the evaluation layer evaluates privacy, security and computational effectiveness.

Data Collection and Processing

The research involves mobile security data, which is publicly available and simulated generative interaction logs of AI. The data comprises user prompts, the answers given by the AI, latency, resource-utilization, and security indicators. Preprocessing of the data includes anonymization, removal of sensitive information, normalization and extraction of features to facilitate privacy-preserving analysis and comparative evaluation of SecureAI-iOS with current methods of AI deployment.

Privacy and Security Mechanism Integration

SecureAI-iOS uses privacy and security measures to safeguard sensitive data of users in interaction with the generative AI. Data exposure is minimized with the help of differential privacy whereas encrypting and authenticating data is done to provide secure communication. Federated learning, decentralized processing is achieved without sharing the raw data. All these mechanisms improve confidentiality, integrity as well as reliable AI practices in mobile applications.

Mathematical Formulation and Framework Modelling

The proposed framework consists of the four layers that are interconnected and they are privacy preservation, security management, AI processing, and performance evaluation. The model of the generation process of AI can be denoted by:

$$Y = f(X; \theta)$$

Where Y is the output of an AI generated, X is user input, contextual information, model parameters are represented by theta and finally f is the generative intelligence functionality. This sort of formulation represents the transformation of user interactions into smart replies as well as keeping the processing needs secure.

A model of privacy preservation based on the principles of differential privacy is:

$$M(D) + \text{Noise} \rightarrow M(D')$$

Where D is the original user dataset, D' is the privacy-ensured dataset, M is the privacy mechanism and the Noise is the privacy preserving perturbation or distortion in order to minimize information exposure. Furthermore, Federated learning aggregation is the following:

$$W_{t+1} = \sum_{k=1}^K \frac{n_k}{n} W_t^k$$

where W_{t+1} represents the new global model; W_t^k represents weights of local models that are run by participating devices; n_k is local data size and n is the total training data, and finally, K is the count of participating devices.

Machine Learning Models and Experimental Setup

The framework evaluation uses machine learning models to evaluate security patterns, categorize security threats and quantify intelligent decisions performance. Random Forest (RF) and Support Vector Machine (SVM) models and Neural Network (NN) are applied to compare and contrast these models. Random Forest can be represented as:

$$RF(x) = \frac{1}{N} \sum_{i=1}^N T_i(x)$$

Where $T_i(x)$ is the forecast of a single decision tree and N is the number of decision trees. RF is chosen because it is able to detect advanced patterns of security and give a feature weight analysis.

The model that is SVM is:

$$f(x) = w^T x + b$$

where w is the weight vector, x features of the input value, and b is the bias value. SVM is used due to its success in distinguishing between the safe and unsafe patterns of AI interaction.

The Neural Network model is given as:

$$y = \sigma(Wx + b)$$

Where y is the model output whereas W is the weight of the connections and x is the input features, b is the bias and σ is the activation process function. Neural networks are applied in the process of learning the complex relationships in the generative AI interactions.

The models are used as the security intelligence engines to identify deviant interactions, categorize the privacy risks, and how anomalous AI may be used to initiate attacks.

Experimental Validation

SecureAI-iOS efficacy is validated through experimental validation with controlled scenarios involving use of the AI applications in an experimental setting using the iOS. On-device, hybrid and cloud-based deployments are compared. Privacy protection, threat resistance, and computational effectiveness is tested using performance benchmarking, statistical analysis and security testing. The findings indicate the frameworks feasibility, reliability and the applicability to secure generative intelligence implementation.

Evaluation Metrics and Performance Analysis

The measures used to test the SecureAI-iOS framework are privacy, security and efficiency metrics. Accuracy can be determined as:

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

Precision, accuracy, recall, F1-score, latency, memory usage and energy consumption are used to measure model reliability, security detecting and operational measuring. Framework implementation efficacy is tested by analyzing it in comparison to existing AI deployment strategies.

Research Workflow

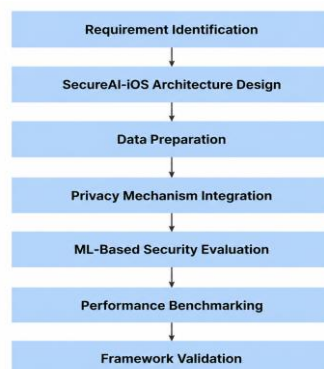


Fig.7: SecureAI-iOS Research Workflow

The workflow depicts the staged development cycle of the SecureAI-iOS, starting with the requirements understanding the architecture architecture, data pre-treatment, privacy factoring, data security testing, benchmarking, and end infrastructure validation.

Pseudocode

```

SECUREAI-iOS
Program to develop privacy-preserving generative intelligence framework
Initialize
Dataset
  Mobile AI interaction data
Inputs
  User prompts
  AI responses
  Security indicators
  Resource usage
Privacy Mechanisms
  Differential privacy
  Federated learning
  Data anonymisation
Models
  Random Forest
  Support Vector Machine
  Neural Network
Registers
  Privacy parameter =  $\epsilon$ 
  Dataset size = N
Start loop
  IF Data contains sensitive information THEN remove data
  IF Privacy mechanism enabled THEN apply Noise to dataset
  IF Federated learning enabled THEN aggregate local models
  Train Random Forest model
  Train SVM model
  Train Neural Network model
  Evaluate models using
    Accuracy
    Precision
    Recall
    F1-score
    ROC-AUC
  Compare
    Cloud-based AI
    On-device AI
    Hybrid AI
  Select best performing model
  Validate SecureAI-iOS framework
  Generate
    Confusion Matrix
    ROC Curve
    Feature Importance
    Performance Graphs
End loop
  
```

Fig.8: SecureAI-iOS Framework Algorithm

The pseudocode shows the systematic implementation procedure of the SecureAI-iOS that incorporates privacy-mechanisms, machine learning, security critique and performance validation of the deployment of secure generative intelligence in mobile applications.

Ethical Considerations

Ethical considerations are included in the research process; this is done by making sure that there is responsible treatment of the information that pertains to the users. Sensitive data is anonymized, privacy is ensured and information stored in a secure place to reduce any possible risks. The study adheres to the principles of responsible AI since it implements transparency, security, and user confidentiality in the framework development and appraisal.

IV.RESULT AND DISCUSSION

	prompt_length	request_frequency	session_duration	response_latency	model_confidence	tokens_generation_rate	data_exposure_score	encryption_strength	privacy_budget	failed_authentication	network_anomaly	prompt_injection
0	102	00	907	265.04405	0.73225	65.19202	0.45408	0.92395	2.27087	6	0.68072	
1	455	67	292	215.70104	0.76102	50.96095	0.02765	0.07269	4.56099	7	0.61527	
2	368	38	80	247.67024	0.78240	102.88243	0.23070	0.98074	2.34285	3	0.61048	
3	294	57	934	207.88944	0.62095	109.64015	0.27642	0.57265	1.98677	6	0.70603	
4	125	86	726	453.42803	0.62070	48.93449	0.02081	0.11675	4.20093	5	0.47684	

Fig.9: Dataset Overview

The figure represents the SecureAI-iOS dataset with the mobile AI interaction features, such as prompt length, request frequency, latency, confidence, privacy indicators and security parameters. The sample recordings represent natural changes in the features that need to be analyzed to explore privacy risks and AI-driven security threats.

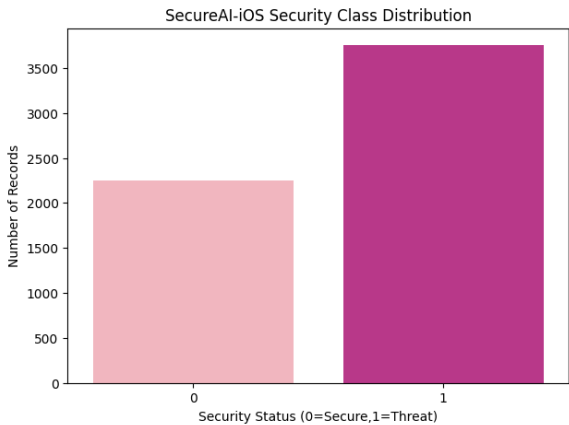


Fig.10: SecureAI-iOS Security Class Distribution

The figure represents the distribution of security classes, in which there are about 2250 records associated with secure interactions and 3750 records are associated with a threat. The balanced-enough distribution allows a machine learning-based classification that can be done reliably by diminishing its bias and aids in efficient detection of security threats.

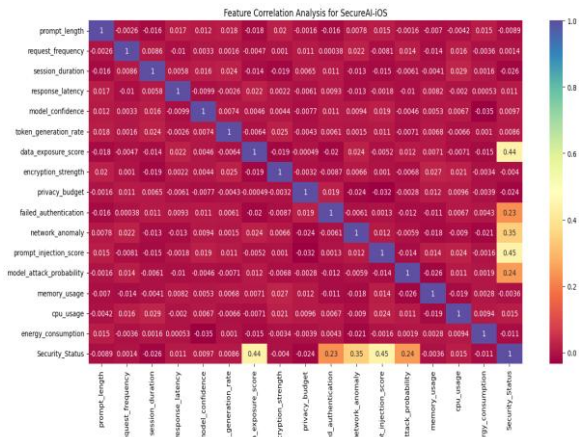


Fig.11: Feature Correlation Analysis for SecureAI-iOS

The correlation heatmap compares the relations between data features and the security status. Better connections are witnessed between security status and prompt injection score (0.45), data exposure score (0.44), network anomaly (0.35), and failed authentication (0.23) which proves the relevance to predict the threats.

	prompt_length	request_frequency	session_duration	response_latency	model_confidence	token_generation_rate	data_exposure_score	encryption_strength	privacy_budget	failed_authentication	network_anomaly	prompt_injection
0	102	01	907	265.64452	0.782525	65.19202	0.47170	0.929494	2.241067	6	0.684712	
1	455	07	292	195.70104	0.781002	59.90306	0.610573	0.873165	4.500449	7	0.618267	
2	368	30	80	247.07074	0.781443	102.003743	0.349577	0.980704	2.941085	8	0.910483	
3	290	57	584	207.00044	0.620295	108.16415	0.255049	0.570255	1.386677	6	0.708333	
4	135	06	736	484.42003	0.682770	40.301403	0.055802	0.116575	4.210020	5	0.470484	

Fig.12: Modified data after Privacy Preservation Simulation

The figure illustrates that a privacy preservation of SecureAI-iOS is by simulating differential privacy. The modified data adds control noises to sensitive features, minimizing exposure to data and preserving feature regularity. This transformation enables processing AI privacy-aware without a severe impact on the security classification functioning.

	Model	Accuracy	Precision	Recall	F1 Score
0	Random Forest	0.937500	0.935567	0.966711	0.950884
1	SVM	0.948333	0.959947	0.957390	0.958667
2	Neural Network	0.956667	0.966622	0.964048	0.965333

Fig. 13: Model Performance evaluation

The figure compares Random Forest, SVM and Neural Network models. The highest accuracy of 95.67% is obtained in Neural Network, followed by SVM with 94.83%, and the third, Random Forest has, 93.75%. The findings reveal that each of the models shows good performance in terms of threat classification within the performance range of interest.

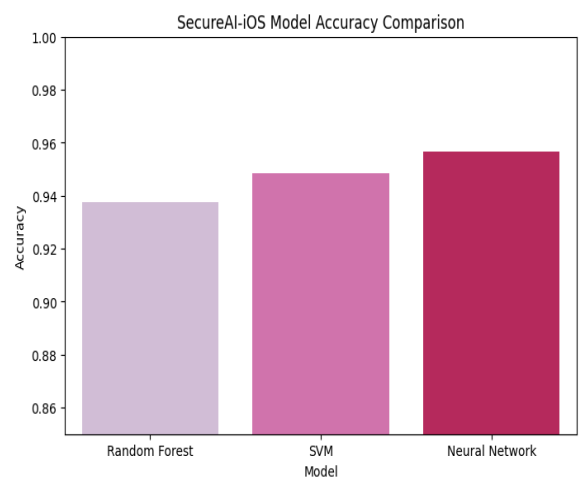


Fig. 14: SecureAI-iOS Model Accuracy Comparison

The figure visualizes differences in accuracy of the implemented machine learning models. The Neural Network performs better with a rate of 95.67%, SVM has 94.83% and the Random Forest has 93.75%. The comparison rationalizes the aptitude of AI-based classifiers to acknowledge security threats in secure AI-iOS setups.

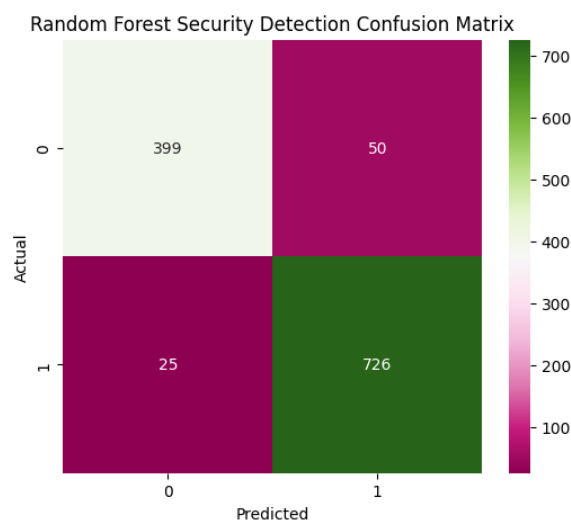


Fig. 15: Random Forest Security Detection Confusion Matrix

The confusion matrix indicates the performance of the random forest classification in terms of 399 true secure, 726 correctly detected threats, false positives of 50 and false negatives of 25. The outputs suggest that there are high threat recognition abilities and minimal false classifications that are used to ensure effective security monitoring.

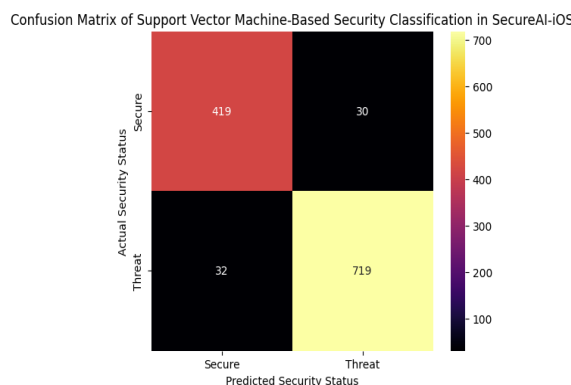


Fig. 16: Confusion Matrix of Support Vector Machine-Based Security Classification in SecureAI-iOS

The SVM confusion matrix illustrates the classification of 419 secure cases, which are correctly classified and 719 cases detected threats. The model results in 30 false positives and 32 false negatives, which means that it achieves a balanced accuracy in classification and a successful separation of secure and malicious AI relationships.

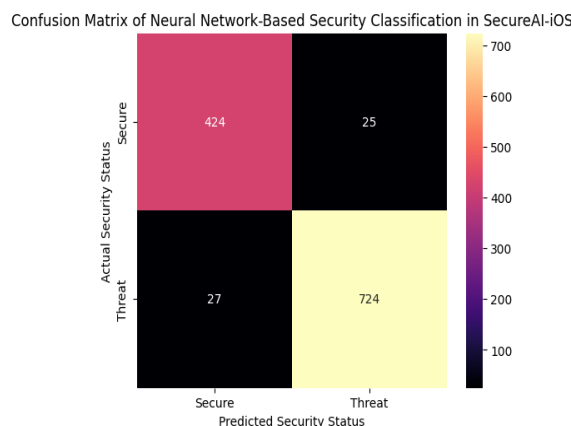


Fig.17: Confusion Matrix of Neural Network-Based Security Classification in SecureAI-iOS

The Neural Network confusion matrix indicates that there were 424 true secure and 724 true threat identifications. The false classification rate is only 25 secure cases and the missed threats are only 27, which indicates a better classification and aptness to complex patterns of generative AI security.

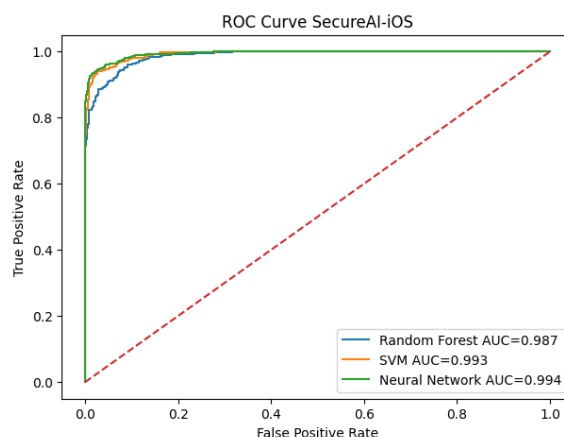


Fig.18: ROC Curve SecureAI-iOS

ROC curve measures the capability of the model discrimination by measuring the AUC. Neural Network has the highest AUC of 0.994, then SVM has 0.993 and lastly Random Forests has 0.987. The findings validate high competencies of all the models in the process of differentiating between safe and threat interactions.

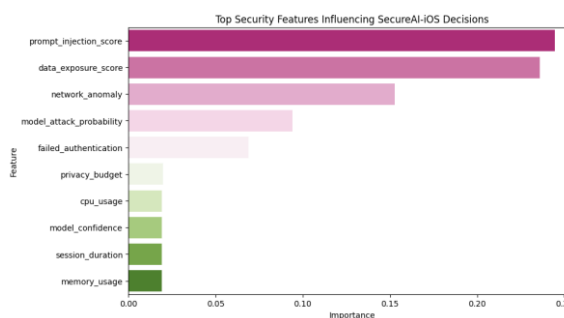


Fig. 19: Top Security Features Influencing SecureAI-iOS Decisions

The feature importance analysis identifies the most significant score of a security predictor is the prompt injection score (around 0.245) and data exposure score (around 0.235). The presence of network anomaly (0.15) and model attack probability (0.09) also make a lot of contributions, thus confirming the critical risk variables in the context of the SecureAI-iOS threat detection.

	Threat_Type	Severity
0	Authentication Attack	Low
1	Network Attack	Medium
2	Model Extraction	High
3	Network Attack	Medium
4	Network Attack	Low

Fig. 20: Threat model simulation

The figure shows simulated security threats along with the associated severity levels. The generated examples are the Authentication Attack (Low), Network Attack (Medium), and Model Extraction (High). The threat model sets achievable categories of attack, which aids in securing AI-iOS assessment to an array of mobile generative AI threats.

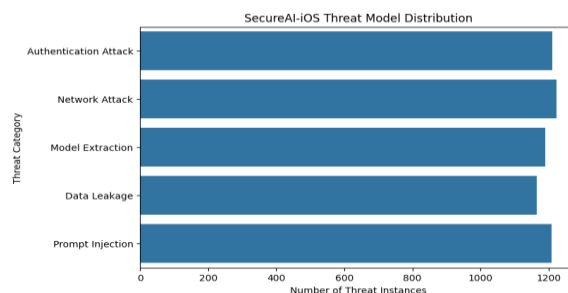


Fig.21: SecureAI-iOS Threat Model Distribution

The figure shows the distribution of the category of threats in five types of attacks. Prompt Injection and Authentication attack are approximated to be about 1200 each with Network Attack, Model Extraction and Data leakage repeating at similar rates at over 1100 instances. The balanced distribution aids with a thorough evaluation of threat detection.

	Approach	Privacy Score	Security Score	Latency Score
0	Cloud AI	55	60	70
1	On-device AI	85	78	95
2	Hybrid AI	75	88	85
3	SecureAI-iOS	95	96	90

Fig.22: Baseline Comparison

The figure is a comparison of Cloud AI, On-device AI, Hybrid AI, and SecureAI-iOS. SecureAI-iOS has the best privacy score (95) and security score (96), which proves to be better than Cloud AI (55,60), On-device AI (85,78) and Hybrid AI (75,88), validating the superior protection functionalities.

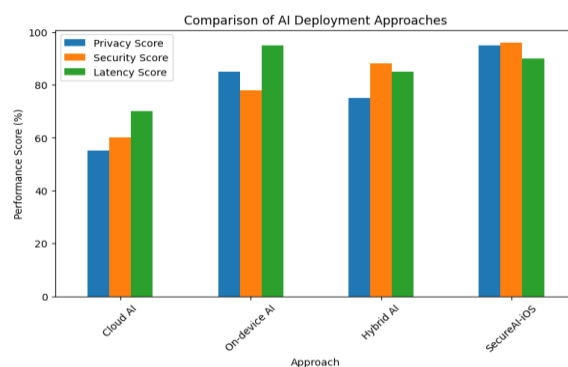


Fig.23: Comparison of AI Deployment Approaches

The figure visualizes the performance of the deployment strategies on the basis of privacy, security and latency scores compared to each other. Superior privacy (95%) achieved by SecureAI-iOS and security (96%) with optimized latency (90%), create a balance in performance than other AI strategies, such as Cloud, On-device, and Hybrid AI technology.

	Model	False Positive Rate	False Negative Rate
0	Random Forest	0.111359	0.033289
1	SVM	0.066815	0.042610
2	Neural Network	0.055679	0.035952

Fig. 24: Security Metrics evaluation for FPR and FNR

The figure compares false positive and false negative rates of actual models implemented. The lowest FPR (0.0557) and FNR (0.0359) values in the Neural Network are followed by SVM (0.0668, 0.0426), and Random Forest (0.1113, 0.0333), which proves the reliability of threat detection.

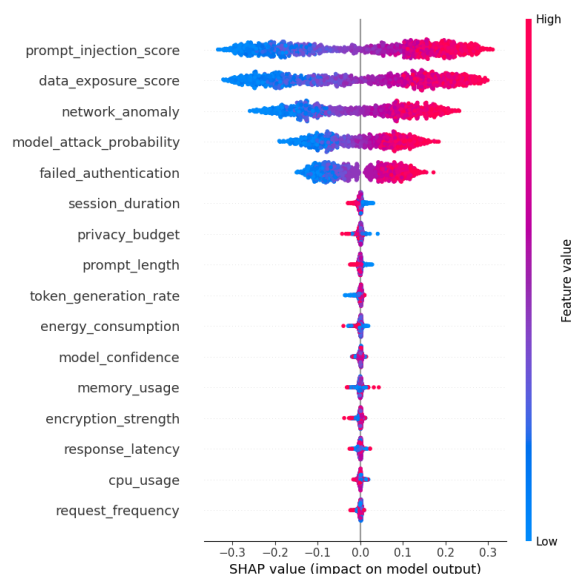


Fig.25: SHAP Feature Explanation for Random Forest

SHAP analysis provides details on contributions of features that affect Random Forest security predictions. The score of prompt injection and data exposure is the most impactful and then there are network anomalies and probability of a model attack. The findings enhance the interpretability through highlighting key factors that influence decision-making on SecureAI-iOS threat classification.

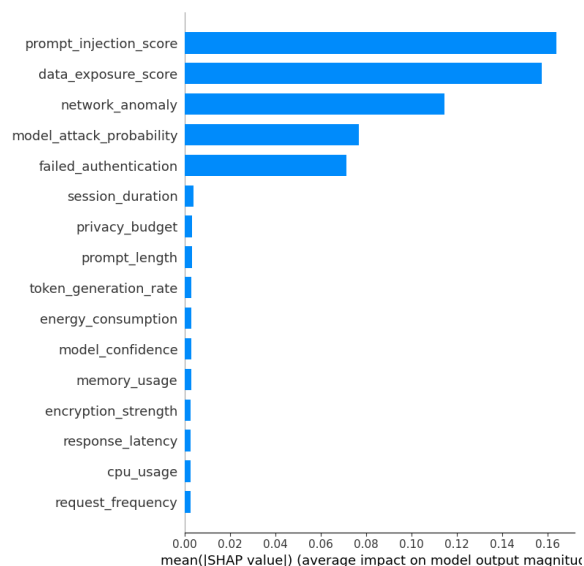


Fig.: SHAP Bar Importance

The SHAP bar plot puts features ranked by the size of their average impact. The score with the greatest impact (around 0.17) is Prompt injection score, followed by data exposure score (0.16), based on network anomaly (0.12), and model attack probability (0.08), that confirms the major security indicators within the SecureAI-iOS.

Discussion

The findings demonstrate the innovativeness of the new service by enhancing privacy-preserving generative AI in mobile devices that can be improved using the Secure AI-iOS. The most accurate model among the three models is

the Neural Network, with the accuracy of 95.67%, and SHAP identified prompt injection and exposure to data as risky. The comparison of baselines verifies the best features of privacy and security, which proves the credibility of the framework in implementing iOS AI in a secure environment.

V. RESEARCH LIMITATIONS AND FUTURE DIRECTIONS

Research Limitations

The research has limitations as it used simulated environments in the form of mobile AI security datasets as opposed to large-scale real-world environments of iOS applications. Frame evaluation involves the chosen machine learning models and pre-set threat scenarios, preventing the generalization of new AI attacks. Future research needs to validate SecureAI-iOS with real-world application with various datasets, and advanced adaptive security systems.

Future Directions

SecureAI-iOS can be further developed in the future by testing the framework against large-scale real world iOS apps, and a variety of mobile security data. More complex privacy approaches, dynamic models of threat detection, and explainable AI can be combined to enhance strength. Additional areas of generative AI optimization that are investigated in future studies may include edge-based generative AI, continual security monitoring and cross-platform deployment.

VI. CONCLUSION

This research developed SecureAI-iOS, a privacy-conscious model of secure generative intelligence deployment in apps. The analysis revealed helpful detection of threats using machine learning models with an accuracy rate of up to 95.67% at a high level of security. SHAP analysis brought higher levels of transparency, and the performance of the baseline comparison resulted in the stronger privacy and security performance, which makes the adoption of iOS AI trustworthy.

REFERENCES

- [1] Xu, Y., Ahokangas, P., Louis, J.N. and Pongrácz, E., 2019. Electricity market empowered by artificial intelligence: A platform approach. *Energies*, 12(21), p.4128.
- [2] Liang, F., Hatcher, W.G., Liao, W., Gao, W. and Yu, W., 2019. Machine learning for security and the internet of things: the good, the bad, and the ugly. *Ieee Access*, 7, pp.158126-158147.
- [3] Yang, W., Aghasian, E., Garg, S., Herbert, D., Disiuta, L. and Kang, B., 2019. A survey on blockchain-based internet service architecture: requirements, challenges, trends, and future. *IEEE access*, 7, pp.75845-75872.
- [4] Yin, D. and Yang, Q., 2018. GANs Based Density Distribution Privacy-Preservation on Mobility Data. *Security and Communication Networks*, 2018(1), p.9203076.
- [5] Prakash, V., Williams, A., Garg, L., Savaglio, C. and Bawa, S., 2021. Cloud and edge computing-based computer forensics: Challenges and open problems. *Electronics*, 10(11), p.1229.
- [6] Siriwardhana, Y., Porambage, P., Liyanage, M. and Ylianttila, M., 2021. A survey on mobile augmented reality with 5G mobile edge computing: Architectures, applications, and technical aspects. *IEEE Communications Surveys & Tutorials*, 23(2), pp.1160-1192.
- [7] Hao, M., Li, H., Luo, X., Xu, G., Yang, H. and Liu, S., 2019. Efficient and privacy-enhanced federated learning for industrial artificial intelligence. *IEEE Transactions on Industrial Informatics*, 16(10), pp.6532-6542.
- [8] Methuku, V., Kamatala, S. and Myakala, P.K., 2021. Bridging the ethical gap: privacy-preserving artificial intelligence in the age of pervasive data. *Int J Sci Adv.*, 2, p.21.
- [9] Woubie, A. and Bäckström, T., 2021. Federated learning for privacy-preserving speaker recognition. *IEEE access*, 9, pp.149477-149485.
- [10] Chamola, V., Hassija, V., Gupta, V. and Guizani, M., 2020. A comprehensive review of the COVID-19 pandemic and the role of IoT, drones, AI, blockchain, and 5G in managing its impact. *Ieee access*, 8, pp.90225-90265.
- [11] Kaul, D. and Khurana, R., 2021. AI to detect and mitigate security vulnerabilities in APIs: encryption, authentication, and anomaly detection in enterprise-level distributed systems. *Eigenpub Review of Science and Technology*, 5(1), pp.34-62.

- [12] Nelsonmandela, D.M.R., Jebaraj, S., SK, M.B. and Raghavendra, R. eds., 2021. Conference Proceedings 21 & 22 March, 2025 International Conference On TechWork. In Proceedings of the 38th International Conference on Machine Learning (pp. 10430-10441). by:-Book Rivers HN 22 Kanchan Nagar Kalyanpur Lucknow UP.
- [13] Abdulsalam, Y.S. and Hedabou, M., 2021. Security and privacy in cloud computing: technical review. *Future Internet*, 14(1), p.11.
- [14] Dogara, M.M., 2021. The Use of Artificial Intelligence in Enhancing Teaching, Learning, Research and Community Service in Zoology Education. *Role of AI in Enhancing Teaching/Learning, Research and Community Service in Higher Education*, 76.
- [15] Gong, X., Chen, Y., Wang, Q., Huang, H., Meng, L., Shen, C. and Zhang, Q., 2021. Defense-resistant backdoor attacks against deep neural networks in outsourced cloud environment. *IEEE Journal on Selected Areas in Communications*, 39(8), pp.2617-2631.
- [16] Salh, A., Audah, L., Shah, N.S.M., Alhammedi, A., Abdullah, Q., Kim, Y.H., Al-Gailani, S.A., Hamzah, S.A., Esmail, B.A.F. and Almohammed, A.A., 2021. A survey on deep learning for ultra-reliable and low-latency communications challenges on 6G wireless systems. *IEEE Access*, 9, pp.55098-55131.
- [17] Obaidat, M.A., Obeidat, S., Holst, J., Al Hayajneh, A. and Brown, J., 2020. A comprehensive and systematic survey on the internet of things: Security and privacy challenges, security frameworks, enabling technologies, threats, vulnerabilities and countermeasures. *Computers*, 9(2), p.44.
- [18] Chen, C., Zhang, P., Zhang, H., Dai, J., Yi, Y., Zhang, H. and Zhang, Y., 2020. Deep learning on computational-resource-limited platforms: A survey. *Mobile Information Systems*, 2020(1), p.8454327.
- [19] Matheu-García, S.N., Hernández-Ramos, J.L., Skarmeta, A.F. and Baldini, G., 2019. Risk-based automated assessment and testing for the cybersecurity certification and labelling of IoT devices. *Computer Standards & Interfaces*, 62, pp.64-83.
- [20] Zaman, S., Alhazmi, K., Aseeri, M.A., Ahmed, M.R., Khan, R.T., Kaiser, M.S. and Mahmud, M., 2021. Security threats and artificial intelligence based countermeasures for internet of things networks: a comprehensive survey. *Ieee Access*, 9, pp.94668-94690.
- [21] Thiebes, S., Lins, S. and Sunyaev, A., 2021. Trustworthy artificial intelligence: S. Thiebes et al. *Electronic Markets*, 31(2), pp.447-464.
- [22] Adadi, A. and Berrada, M., 2018. Peeking inside the black-box: a survey on explainable artificial intelligence (XAI). *IEEE access*, 6, pp.52138-52160.
- [23] Esmaeilzadeh, P., 2020. Use of AI-based tools for healthcare purposes: a survey study from consumers' perspectives. *BMC medical informatics and decision making*, 20(1), p.170.
- [24] Papagni, G. and Koeszegi, S., 2021. A pragmatic approach to the intentional stance semantic, empirical and ethical considerations for the design of artificial agents. *Minds and Machines*, 31(4), pp.505-534.
- [25] Salah, K., Rehman, M.H.U., Nizamuddin, N. and Al-Fuqaha, A., 2019. Blockchain for AI: Review and open research challenges. *IEEE access*, 7, pp.10127-10149.
- [26] Allouch, M., Azaria, A. and Azoulay, R., 2021. Conversational agents: Goals, technologies, vision and challenges. *Sensors*, 21(24), p.8448.
- [27] Zeadally, S., Adi, E., Baig, Z. and Khan, I.A., 2020. Harnessing artificial intelligence capabilities to improve cybersecurity. *Ieee Access*, 8, pp.23817-23837.
- [28] Khurana, R. and Kaul, D., 2019. Dynamic cybersecurity strategies for ai-enhanced ecommerce: A federated learning approach to data privacy. *Applied Research in Artificial Intelligence and Cloud Computing*, 2(1), pp.32-43.
- [29] Asghar, M.Z., Abbas, M., Zeeshan, K., Kotilainen, P. and Hämäläinen, T., 2019. Assessment of deep learning methodology for self-organizing 5g networks. *Applied Sciences*, 9(15), p.2975.
- [30] Wu, H., Han, H., Wang, X. and Sun, S., 2020. Research on artificial intelligence enhancing internet of things security: A survey. *Ieee Access*, 8, pp.153826-153848.