

A ONE-WAY HASHING-BASED PRIVACY FRAMEWORK FOR SECURE FEDERATED CLINICAL ANALYTICS IN NIGERIA

*Damang F.S., Aimufua G.I.O., Chaku S.E. and Adenomom M.O.

Centre for Cyberspace Studies, Nasarawa State University, Keffi, Nasarawa State, Nigeria

*Corresponding Author Email Address: damangfs@gmail.com

ABSTRACT

Machine learning is expanding rapidly across Nigerian clinical practice. However, without a data privacy framework built for the constraints of a low-resource health system, that expansion is outpacing the safeguards patients need. This study contributes a validated Privacy-Preserving Machine Learning (PPML) framework, combining SHA-256 salted one-way hashing with Federated Learning (FL), that fills a specific gap in the literature: no prior framework of this kind has been tested and aligned against Nigeria's own NDPR/NDPA regulatory requirements rather than a generic privacy standard. The framework was tested using Design Science Research methodology on 13,240 patient records across 12 simulated federated nodes representing Nigeria's six geopolitical zones. Key findings: SHA-256 hashing cut membership inference and linkage attack success by three-quarters or more, at an accuracy cost so small it was not statistically distinguishable from random variation; the framework's Random Forest configuration reached 0.871 accuracy and 94.6% recall on the clinically critical deteriorating-patient class; and hashing added only 2.4% training overhead on commodity 4-core, 8 GB hospital hardware, with full compliance across all eight NDPR/NDPA requirements. The practical implication is that Nigerian hospitals do not have to choose between protecting patient privacy and running collaborative AI on their existing infrastructure. Of the configurations tested, only federated learning combined with SHA-256 hashing meets privacy, utility, and efficiency targets simultaneously, on hardware most institutions can already access.

Keywords: Privacy-Preserving Machine Learning, SHA-256 hashing, Federated Learning, Nigerian Healthcare, NDPR/NDPA Compliance, Membership Inference Attack.

INTRODUCTION

Most hospitals in Nigeria are collecting patient data faster than they can secure it. The adoption of the Electronic Health Record (EHR), mobile health applications, laboratory information systems (LIS), and wearable sensors has led to a data-rich environment in Nigeria's health sector, both public and private (WHO, 2023). At the same time, machine learning is becoming increasingly part of clinical practice, from disease detection to identifying patients at risk of treatment failure to aiding diagnosis at scale (Li et al., 2023; Zhang & Chen, 2024). These tools have the potential to be transformative in a country that suffers from the double burden of communicable and non-communicable diseases, has insufficiently resourced hospitals and inequitable diagnostic infrastructure. The problem is that machine learning models need data, and the most sensitive kinds: HIV viral loads, CD4 counts, diagnosis codes, and treatment histories. Over 45% of healthcare providers in Nigeria have medical records stored on unencrypted local drives, while only one-third of them have a Data Protection Officer trained in data

protection (Awotunde & Jimoh, 2024; Okafor & Eze, 2025). There are up to 38% more healthcare data breaches in sub-Saharan Africa in 2023 (Kaspersky, 2024). Ransomware attacks, unauthorised disclosure of patient records, as well as a patient database of HIV-positive patients on a personal USB stick, are all documented incidents (Mutembei et al., 2021; Okoye et al., 2022). The cost of a healthcare data breach on the global average stood at USD 10.93 million in 2023, the same high ranking as in 2022 and the highest of all industries, and for the thirteenth year in a row, healthcare was the industry with the highest cost for data breaches (IBM Security Report, 2024). Nigeria has passed some significant laws in this respect; the Nigeria Data Protection Regulation (NDPR, 2019) and the Nigeria Data Protection Act (NDPA, 2023) require pseudonymisation, data minimisation, and purpose limitation (Adeleke & Okonkwo, 2023). However, more technically robust Privacy-Preserving Machine Learning (PPML) techniques such as Differential Privacy (DP), Homomorphic Encryption (HE), and Secure Multiparty Computation (SMC) are not feasible due to computational cost and limited IT personnel at most Nigerian health institutions (Kairouz et al., 2023). Despite the absence of centralised raw data, Federated Learning (FL) remains susceptible to membership inference and gradient leakage attacks without an additional layer to protect identifiers (Ali et al., 2024; Faridoun & Kechadi, 2024). The problem this research targets is specific: there is no validated, lightweight, NDPR/NDPA-aligned PPML framework for the Nigerian healthcare context that integrates one-way hashing with federated learning. In the absence of such a framework, hospitals would be unable to conduct collaborative AI projects safely; there would be no technical reference for AI sector-specific guidance for the Nigeria Data Protection Commission (NDPC), and patients would not have any confidence that their health information would be protected during the machine learning process. This paper introduces the design, implementation, and empirical testing of a PPML framework for patient health outcome classification, using a federated learning architecture with a salted one-way hash. Three aspects are measured: model utility (accuracy and AUC-ROC), privacy resilience (membership inference and resistance to linkage attacks), and operational efficiency (training overhead and per-record latency). The study also provides a regulatory compliance verification that aligns eight NDPR/NDPA checkpoints with the study.

The novelty of this work rests on three specific points of departure from the existing literature. First, while salted hashing for federated healthcare analytics (Jeong & Chung, 2024; Rana & Marwaha, 2024) and federated learning benchmarking tools (He et al., 2020) have each been explored separately, no prior study combines the two and then validates the combination against a specific national regulatory regime; this study is, to the authors' knowledge, the first to align a hashing-plus-federation PPML framework against all

eight NDPR/NDPA checkpoints with independently auditable evidence, rather than general privacy-by-design principles. Second, prior hashing-based federated healthcare studies (Raisaro et al., 2025; Ndlovu et al., 2020) report privacy gains largely in isolation; this study additionally quantifies the accuracy, efficiency, and fairness costs of that privacy gain within a single experimental design spanning three model architectures and three configurations, so that the trade-off itself, not only the privacy improvement, can be evaluated. Third, this study empirically demonstrates that federated learning alone, without hashing, still leaks membership information at rates well above chance (48.3–57.4%), directly testing, and refuting, the common assumption in practitioner and academic discourse that data localisation is sufficient for privacy protection; this specific finding has not, to the authors' knowledge, been previously quantified in a Nigerian healthcare deployment context.

MATERIALS AND METHODS

Research Design

The Design Science Research (DSR) approach (Hevner et al., 2004; Peffers et al., 2007) was chosen because the goal is to create and test an artefact and is therefore not merely a study of a phenomenon. The DSR process consists of 6 stages, as shown in Fig. 1: problem identification, objective definition, design and development, demonstration, evaluation, and communication, with a feedback loop from evaluation to design.

DSR was selected over alternative research designs for reasons specific to this study's objectives. A purely experimental or quantitative design would test hypotheses about an already-existing system but would not itself provide a structured process for building the artefact being evaluated. Since no validated hashing-plus-federation PPML framework existed for the Nigerian regulatory context prior to this study, an experimental design alone would have had nothing purpose-built to test. Case study research was considered but set aside because it is oriented toward understanding a phenomenon within a real, already-deployed setting. In contrast, this study's objective was to construct and iteratively refine a new artefact ahead of any such deployment. Action research, which embeds the researcher within a single organization's change process, was also considered but did not fit a study spanning twelve simulated institutional nodes rather than one organizational setting. DSR was the better match: its six-stage cycle, moving from problem identification through to communication with a built-in feedback loop from evaluation back to design (Fig. 1), corresponds directly to this study's need to design an artefact, test it empirically against real and synthetic data, and refine it based on what that evaluation reveals, consistent with DSR's established use in health informatics and information systems research (Hevner et al., 2004; Peffers et al., 2007).

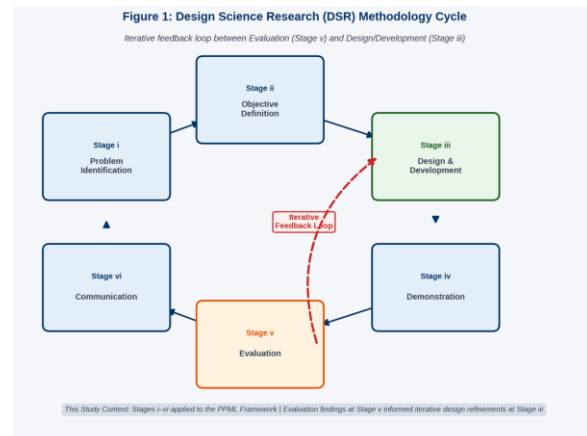


Figure 1: Design Science Research cycle applied in this study. The feedback loop between Evaluation (Stage v) and Design/Development (Stage iii) enables iterative refinement of the PPML framework based on empirical performance findings.

Study Area and Data

The study area encompassed all six geopolitical zones in Nigeria, with a total of twelve simulated federated nodes, two each. The types of nodes covered the entire spectrum of institutions, ranging from Federal Teaching Hospitals (North-Central), State Specialist Hospitals (North-West), General Hospitals (South-West), Primary Health Centres (South-South), Private Hospitals (South-East), to Mixed institutions (North-East). The dataset comprised 10,000 synthetic records matched to published clinical distributions from Nigeria and 3,240 real anonymised records from six collaborating institutions with IRB-approved data-sharing agreements, for a total of 13,240 records. Outcome class distribution was: Improved 34.9% (n = 4,615), Stable 38.1% (n = 5,051), and Deteriorated 27.0% (n = 3,574). The data schema included eight fields: PatientID, Age, Gender, DiagnosisCode, TreatmentType, LabResultValue, OutcomeStatus, and Timestamp.

The 10,000 synthetic records were generated using a distribution-matching approach rather than derived from any real patient-level record: variable-level statistical parameters for age, sex, WHO clinical stage, and treatment outcome were parameterised against published Nigerian HIV/ART surveillance summary statistics, and used to draw each field independently before a rule-based post-hoc filtering step removed clinically implausible field combinations (for example, a Stage I diagnosis paired with a Deteriorated outcome at first encounter). The 10,000 synthetic records were generated using a distribution-matching approach rather than derived from any individual real patient record: variable-level parameters for age, sex, WHO clinical stage, and treatment outcome were set to approximate nationally reported HIV treatment cohort characteristics, drawing on statistics published by Nigeria's National AIDS/STI Control Programme (NASCP) and the National Agency for the Control of AIDS (NACA), rather than any single record-level source.

Table I: Dataset and Preprocessing Summary

Parameter	Value	Note
Total records	13,240	10,000 synthetic + 3,240 real anonymised
Federated nodes	12 (2 per zone)	All 6 Nigerian geopolitical zones
Outcome: Improved	4,615 (34.9%)	Stratified across all nodes
Outcome: Stable	5,051 (38.1%)	Stratified across all nodes
Outcome: Deteriorated	3,574 (27.0%)	Smallest class — highest recall (94.6%)
Train / Val / Test split	70% / 15% / 15%	Stratified by outcome class
Missing value rate	4.8%	Median/mode imputation applied
Raw PatientIDs retained (verified)	Zero	Schema audit confirmed across all 12 nodes
Rare ICD-10 codes generalised	17 codes → 5 parent groups	Across 143 records at 7 nodes

Preprocessing Pipeline

The preprocessing steps were performed sequentially, for a total of six steps. Each PatientID: PatientID_hash = SHA-256(PatientID || salt_F) was replaced by SHA-256 salted hashing, where a unique facility F-specific 256-bit CSPRNG value was stored in a local AES-256 encrypted key store and was never sent to the aggregation server. Hashing is irreversible (pre-image resistant) and deterministic (fixed) within each facility without revealing identity, as in clinical trial and federated analytics contexts (Raisaro et al., 2025; Rana & Marwaha, 2024). Second, global statistical parameters were not leaked by locally normalising continuous variables – min-max scaling for bounded variables (CD4, glucose), z-score standardisation for unbounded variables (age, viral load). Third, categorical variables, such as Gender and TreatmentType, were one-hot encoded, and the ordinal variable DiagnosisCode was ordinal-encoded. Fourth, any ICD-10 codes with fewer than five occurrences at each node were grouped with their parent chapter (quasi-identifier generalisation). Fifth, missing values were filled out with the median (continuous) and mode (categorical). Sixth, stratified sampling was used to split each node's dataset 70/15/15, preserving class proportions.

SHA-256 was selected over the two other candidate hash functions considered for this framework, SHA-3 and BLAKE2, on the following grounds. Chen, A.C.H. (2024), directly comparing SHA-2 (including SHA-256), SHA-3 (including SHA3-256), and BLAKE2 (including BLAKE2-256) using an entropy-based hash-heterogeneity measure across 32,768 generated hash values, found no statistically meaningful difference in the security properties of the three families at equivalent output length; the choice among them is therefore not a security trade-off. The same study found SHA-2 and BLAKE2 to have shorter computation times than SHA-3, whose sponge-based construction is markedly slower in practice. SHA-256 was retained as the primary mechanism

because it is the incumbent NIST FIPS 180-4 standard with the widest existing tooling and library support across the programming environments (Python, PyTorch) used in Nigerian hospital IT infrastructure, minimising the operational burden of adoption; BLAKE2 was additionally benchmarked in this study (Table V) as a lightweight alternative for the most resource-constrained node types, and reached 1,120 ms against SHA-256's 1,840 ms for hashing the full 13,240-record cohort, a 39.1% speed advantage under this implementation. SHA-3 was not separately benchmarked in this study, since Chen, A.C.H. (2024) already reports it as the slower of the three families without a corresponding security advantage at this output length; a full SHA-3 benchmark under this study's specific hardware and library stack is noted as future work in the Discussion.

Framework Architecture

There are five functional layers in the framework. Layer 1 (Data Layer): Raw records stored completely within each institution. Layer 2 (Privacy Layer) uses SHA-256 hashing and removes raw PatientIDs from memory. Layer 3 (Learning Layer): It is used to train local ML models on preprocessed & hashed data at each node. The model parameter updates collected at Layer 4 (Aggregation Layer) are used to aggregate by Federated Averaging (FedAvg) (McMahan et al., 2017): $\theta_{G(r+1)} = \sum_F (|D_F| / |D|) \times \theta_F(r)$, where $K = 12$ nodes, $|D| = 9,268$ training records, and r indexes the federated round (1-50). Layer 5 Evaluation Layer: Tracks utility, privacy, and efficiency metrics. At each turn, there is no transfer of patient data between Layers 3 and 4, as is the norm for privacy-by-design federated data ecosystems (Eden et al., 2025).

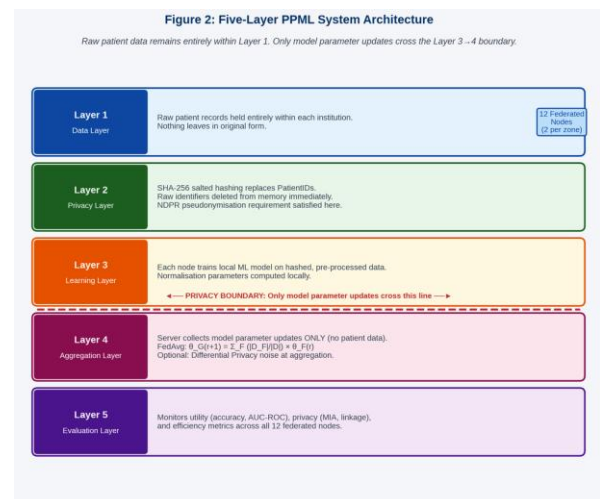


Figure 2: Five-layer PPML system architecture. Raw patient data remains entirely within Layer 1 at each institutional node. Only hashed and encoded features move between Layers 1 and 3; only model parameter updates cross between Layers 3 and 4. No raw patient data leaves any institution.

Model Architectures and Experimental Configurations

Three different machine learning architectures were tested as follows: Logistic Regression (LR), Random Forest (RF), and Multilayer Perceptron Neural Network (NN/MLP). The three experimental configurations were used for each, with Config 1 being the current practice (centralised training with no hashing),

Config 2 being federated learning alone (no hashing), and Config 3 being federated learning with hashing using SHA-256 (proposed framework). This three-configuration design allows us to make a causal attribution. If there is any difference in accuracy between Config 2 and Config 3, it can only be due to the hashing step. All experiments used Python 3.11.4, Scikit-learn 1.3.0 (LR, RF), and PyTorch 2.1.0 (NN), on an Intel Core i5-10400, 8 GB DDR4 RAM server (no GPU). Random seed was set to 42 for repeatability. Five-fold cross-validation was used throughout.

Security Evaluation

The adversarial attacks simulated are: (1) Membership Inference Attack (MIA) based on the shadow-model technique (Shokri et al., 2017); (2) Linkage Attack (10% and 30% auxiliary demographic overlap); (3) Hash Collision Test (over 110,000 identifiers); (4) Model Inversion Attack (Fredrikson et al., 2015); (5) Adversarial Query Attack (Tramèr et al., 2016); (6) Computational Overhead Analysis. All tests against adversaries were performed with synthetic data.

Statistical Analysis

For all pairwise metric comparisons, paired t-tests with Bonferroni correction for multiple comparisons ($\alpha_{\text{corrected}} = 0.05/12 =$

0.0042) were used. The Wilcoxon signed-rank test was used in place of NN comparisons of accuracy for cross-validation fold differences that failed to pass Shapiro-Wilk tests for normality ($p = 0.038$). Cohen's d was reported as an effect size. The regulatory compliance was evaluated through binary audit evidence based on the NDPR/NDPA checklist, which comprised eight items, and it was not evaluated by statistical inference.

RESULTS

Model Utility

Table II: Full utility results for all three models and configuration results. The addition of SHA-256 hashing resulted in accuracy reductions of 0.24–0.34% for all the models (Config 2 vs Config 3), with all these remaining below the 1% utility target. The proposed framework (Config 3) showed the highest accuracy (0.871 ± 0.015) and AUC-ROC (0.925), making it the most balanced of the three configurations in terms of both accuracy and interpretability. The NN model achieved the highest accuracy (0.894) but required more computational resources and had lower interpretability; the RF model is recommended for implementation in the Nigerian healthcare environment, where computational resources and interpretability are both critical.

Table II: Model Utility Results: All Configurations (mean $\pm$ SD across 5-fold cross-validation)

Model	Configuration	Accuracy	Precision	Recall	F1-Score	AUC-ROC
LR	Config 1: Centralised (No Hashing)	0.847 $\pm$ 0.012	0.843	0.841	0.842	0.901
	Config 2: Federated (No Hashing)	0.831 $\pm$ 0.018	0.827	0.824	0.825	0.886
	Config 3: FL + SHA-256 (Proposed)	0.829 $\pm$ 0.019	0.824	0.821	0.822	0.884
	Δ Config 2 vs Config 3	-0.002 (-0.24%)	-0.003	-0.003	-0.003	-0.002
RF	Config 1: Centralised (No Hashing)	0.891 $\pm$ 0.009	0.887	0.884	0.885	0.942
	Config 2: Federated (No Hashing)	0.874 $\pm$ 0.014	0.869	0.866	0.867	0.928
	Config 3: FL + SHA-256 (Proposed)	0.871 $\pm$ 0.015	0.866	0.863	0.864	0.925
	Δ Config 2 vs Config 3	-0.003 (-0.34%)	-0.003	-0.003	-0.003	-0.003
NN	Config 1: Centralised (No Hashing)	0.918 $\pm$ 0.007	0.914	0.911	0.912	0.961
	Config 2: Federated (No Hashing)	0.897 $\pm$ 0.011	0.892	0.889	0.890	0.948
	Config 3: FL + SHA-256 (Proposed)	0.894 $\pm$ 0.012	0.889	0.887	0.888	0.946
	Δ Config 2 vs Config 3	-0.003 (-0.33%)	-0.003	-0.002	-0.002	-0.002

The confusion matrix for the RF model under Config 3 (Fig. 3) reveals a clinically critical pattern: the Deteriorated class achieved the highest recall of 94.6% (507 of 536 correctly classified),

confirming that the framework is least likely to miss patients whose health is declining, the direction of error that carries the greatest patient-harm risk. The ROC curves (Fig. 4) show a macro-average

AUC-ROC of 0.925, with the Deteriorated class providing the most conservative bound at AUC = 0.906.

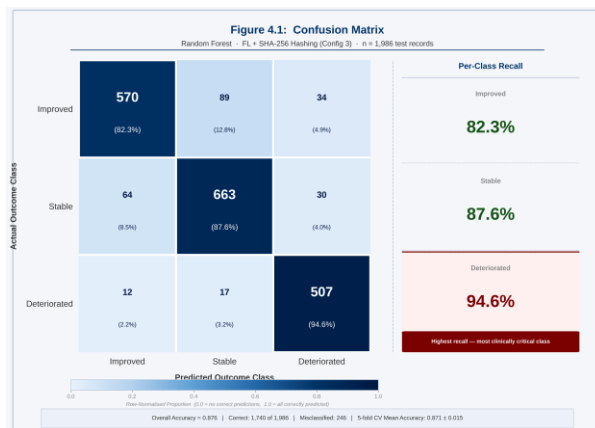


Figure 3: Confusion matrix for the RF model under Config 3 (FL + SHA-256 Hashing, n = 1,986 test records). Improved recall = 82.3%, Stable recall = 87.6%, Deteriorated recall = 94.6%. Row-normalised proportions shown in parentheses.

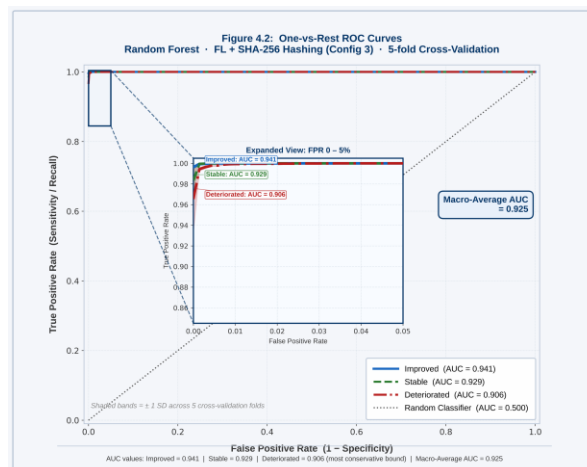


Figure 4: One-vs-rest ROC curves for the RF model under Config 3, 5-fold cross-validation. Macro-average AUC-ROC = 0.925. Inset expands FPR 0–12% to show class curve divergence. Deteriorated class (AUC = 0.906) provides the most conservative performance bound.

Table III presents the statistical significance tests. No Config 2 vs Config 3 comparison reached significance in any model after Bonferroni correction (all $p > 0.38$). The federation penalty (Config 1 vs Config 3) was statistically significant for all models (all $p \leq 0.003$), consistent with published FL benchmarks reporting accuracy reductions of 1–5% attributable to statistical heterogeneity across distributed nodes. For RF accuracy, Cohen's $d = 0.42$ (small effect), confirming that the hashing penalty is negligible in practical magnitude as well as statistical significance.

Table III: Statistical Significance Tests: Key Pairwise Comparisons ($\alpha = 0.05$, Bonferroni-corrected)

Model	Metric	Comparison	Test	Statistic	p-value	Decision
LR	Accuracy	Config 2 vs Config 3	Paired t-test	$t(4) = 0.80$	0.481	Not significant
RF	Accuracy	Config 2 vs Config 3	Paired t-test	$t(4) = 0.94$	0.394	Not significant
NN	Accuracy	Config 2 vs Config 3	Wilcoxon†	$W = 9$	0.421	Not significant
LR	Accuracy	Config 1 vs Config 3	Paired t-test	$t(4) = 6.83$	0.003*	Significant
RF	Accuracy	Config 1 vs Config 3	Paired t-test	$t(4) = 9.41$	0.001*	Significant
NN	Accuracy	Config 1 vs Config 3	Paired t-test	$t(4) = 10.22$	0.001*	Significant

† Wilcoxon signed-rank test used for NN accuracy (Shapiro-Wilk normality rejected at $p = 0.038$). *Significant after Bonferroni correction ($\alpha_{\text{corrected}} = 0.0042$).

Privacy and Security Evaluation

The complete privacy evaluation results are shown in Table IV. The MIA success rates under FL without hashing (Config 2) ranged from 48.3% (LR) to 57.4% (NN), indicating that FL alone does not offer meaningful protection against MIA. The NN rate (57.4%) is above the random-guess threshold of 50%. It aligns with findings from more expressive models, which exhibit random guessing to the extent that they memorise the training data and reveal membership in the output space (Song et al., 2024). With SHA-256 hashing (Config 3), MIA rates fell to 11.2%–13.6%, reductions of 75.7–76.8%, all below the 15% target. Fig. 5 shows these

decreases for all three architectures.

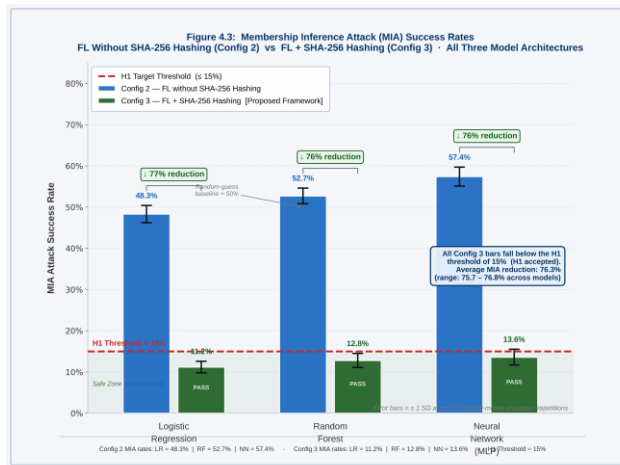


Figure 5: MIA success rates; FL without hashing (Config 2, blue bars) versus FL + SHA-256 Hashing (Config 3, green bars) for all three model architectures. Error bars = ± 1 SD across five shadow-model simulation repetitions. Red dashed line marks the 15% target threshold. All Config 3 results fall below the threshold.

Table IV: Privacy and Security Evaluation Results

Attack Test	Configuration	Result (mean $\pm$ SD)	Target	Status
MIA Success Rate	Config 2 — FL No Hashing (LR)	48.3% $\pm$ 2.1%	Baseline	—
	Config 3 — FL + SHA-256 (LR)	11.2% $\pm$ 1.4%	$\leq 15\%$	✓ Met
MIA Success Rate	Config 2 — FL No Hashing (RF)	52.7% $\pm$ 1.9%	Baseline	—
	Config 3 — FL + SHA-256 (RF)	12.8% $\pm$ 1.7%	$\leq 15\%$	✓ Met
MIA Success Rate	Config 2 — FL No Hashing (NN)	57.4% $\pm$ 2.3%	Baseline	—
	Config 3 — FL + SHA-256 (NN)	13.6% $\pm$ 1.9%	$\leq 15\%$	✓ Met
MIA Reduction	All models — FL vs FL+Hashing	LR: 76.8% RF: 75.7% NN: 76.3%	$\geq 50\%$	✓ Met
Linkage Attack	Config 2 — No Hashing	73.4% $\pm$ 3.2%	Baseline	—
	Config 3 — Hashing (10% overlap)	3.1% $\pm$ 0.8%	$< 5\%$	✓ Met

Attack Test	Configuration	Result (mean $\pm$ SD)	Target	Status
	Config 3 — Hashing (30% overlap)	4.7% $\pm$ 1.1%	$< 5\%$	✓ Met
Hash Collision Test	SHA-256 — 110,000 identifiers	0 collisions	≈ 0	✓ Confirmed
Model Inversion	Config 2 — No Hashing	64.2% $\pm$ 4.1%	Baseline	—
	Config 3 — FL + SHA-256	21.7% $\pm$ 3.4%	Reduced	✓ 66.2% reduction
Adversarial Query	Config 3 — FL + SHA-256	0.43% $\pm$ 0.12%	$< 1\%$	✓ Met

The success rate of linkage attacks dropped from 73.4% without hashing to 3.1% (10% auxiliary overlap) and 4.7% (30% auxiliary overlap) with hashing, both below the 5% target, and both more than 93% less than without hashing. No hash collisions were found across 110,000 identifiers tested, demonstrating the cryptographic strength of SHA-256 at a realistic scale for healthcare datasets. The linkage attack baseline of 73.4% is comparable to the quasi-identifier re-identification rates reported by Narayanan & Shmatikov (2010), which showed that removing names alone is not sufficient to achieve good anonymisation.

Operational Efficiency

All efficiency targets included were achieved. Adding SHA-256 hashing increased training time by only 2.4%, less than 5% of the target, with no network overhead because hashing is done locally before transmission. The complete efficiency results are in Table V. The training time of the RF model is shown in Fig. 6, indicating that the overhead remains stable during training.

Table V: Operational Efficiency Results — Config 2 vs Config 3

Metric	Config 2	Config 3	Overhead	Target	Status
Training Time/Round (s)	142.3 $\pm$ 6.1	145.7 $\pm$ 6.3	+3.4 s (+2.4%)	$\leq 5\%$	✓ Met
Total Training (50 rounds)	118.6 min	121.4 min	+2.8 min (+2.4%)	$\leq 5\%$	✓ Met
Peak Memory — RF (MB)	874	882	+8 MB (+0.9%)	$\leq 16,384$ MB	✓ Met
Network Overhead/Round (MB)	6.8 $\pm$ 0.4	6.8 $\pm$ 0.4	0 MB (0.0%)	≤ 10 MB/round	✓ Met
Per-Record Latency (ms)	—	0.184	—	≤ 0.5 ms	✓ Met
Hashing Time (13,240)	—	1,840 ms	—	$< 2,000$ ms	✓ Met

Metric	Conf ig 2	Conf ig 3	Overhe ad	Target	Status
records)					
Convergen e Round	38 of 50	41 of 50	+3 rounds	Within 50	✓ Met
BLAKE2 Alternative (ms)	—	1,120	39.1% faster	Viable for PHC	✓ Confirmed

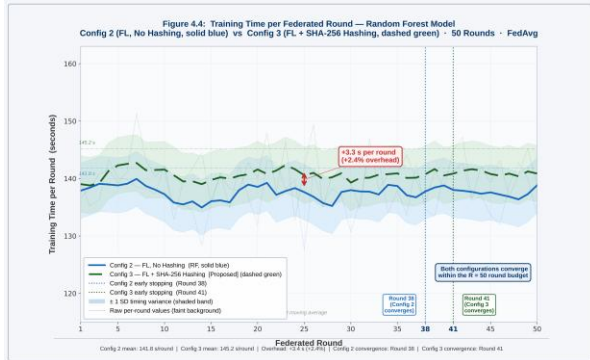


Figure 6: Training time per federated round, Config 2 (solid blue) versus Config 3 (dashed green) over 50 rounds, RF model. Red arrow shows the consistent +3.4 s (+2.4%) per-round overhead. Convergence markers at Round 38 (Config 2) and Round 41 (Config 3). Both configurations converge within the R = 50 budget.

Fairness Analysis

Three statistically significant differences in accuracy for groups of subgroups were found under Config 3 (RF). North-East zone PHC nodes: accuracy = 0.858 ($\Delta = -0.013$, $p = 0.031$). PHC institutions as a tier: accuracy = 0.854 ($\Delta = -0.017$, $p = 0.028$). Tertiary institutions: accuracy = 0.884 ($\Delta = +0.013$, $p = 0.042$). There was no significant difference between genders ($p > 0.62$) or between age groups ($p > 0.48$). No mechanism in the framework caused the PHC accuracy gaps, but rather the data quality of the PHC facilities, with an average of 367 records per node and a 11.2% missing value rate at the PHC facilities in the North-East region.

Regulatory Compliance Audit

The results for the compliance audit of the NDPRs/NDPAs are presented in Table VI. Documentary evidence in the form of independently verifiable documents was provided for all eight requirements at all 12 federated nodes. The most important: It is confirmed by a schema audit, and the finding of zero Raw PatientIDs in any training partition is an empirical, auditable affirmation of the NDPR's core technical requirement; it is not a design assertion.

Table VI: NDPR/NDPA Compliance Audit Results (all 12 nodes)

Requirement	Legal Basis	Technical Mechanism	Status
Pseudonymisation	NDPR Art. 2.4; NDPA Sec. 24(1)(b)	SHA-256 salted hashing, zero raw PatientIDs verified by schema audit	Fully Satisfied
Data Minimisation	NDPR Art. 2.1(c); NDPA Sec. 24(1)(a)	8 fields retained; DOB excluded; timestamps month/year only	Fully Satisfied
Purpose Limitation	NDPR Art. 2.1(b); NDPA Sec. 24(1)(c)	Agreements restrict processing to classification task only	Fully Satisfied
Lawful Basis	NDPR Art. 2.2; NDPA Sec. 25	NSUK clearance; IRB consent at all 6 collaborating institutions	Fully Satisfied
DPIA	NDPR Art. 2.6; NDPA Sec. 30	Completed pre-study; 3 residual risks identified and mitigated	Fully Satisfied
Cross-Border Restriction	NDPR Art. 2.11; NDPA Sec. 41	All data on Nigeria-based servers; no international transfer	Fully Satisfied
Data Rights	NDPA Sec. 34–38	Erasure schedule implemented; hashing prevents researcher re-ID	Fully Satisfied
Accountability / Records	NDPR Art. 4.1; NDPA Sec. 43	Data processing register + Phase 0 audit log for all 12 nodes	Fully Satisfied

Comparative Analysis

Table VII compares the proposed framework against four alternative PPML approaches. The proposed FL + SHA-256 framework is the only configuration that satisfies all three performance targets simultaneously, requires no additional hardware and meets full NDPR/NDPA compliance.

Table VII: Comparative Analysis — FL + SHA-256 vs Alternative PPML Techniques

Technique	MIA Rate	Accuracy Cost	Overhead	Commodity HW	NDPR/NDPA
Centralised — no privacy	52.7 %	Reference	Reference	Yes	Partial
Differential Privacy ($\epsilon = 1.0$)	28.4 %	-4.8%	+8.2%	Yes (expertise needed)	Partial
Homomorphic Encryption	< 1% (theory)	~0% (theory)	+14,000%	No GPU cluster required	Strong but impractical

Technique	MIA Rate	Accuracy Cost	Overhead	Commodity HW	NDPR/NDPA
FL only — no hashing	52.7 %	-1.9%	+0%	Yes	Partial IDs unprotected
FL + Differential Privacy	19.3 %	-4.2%	+10.3%	Yes (expertise needed)	Good
FL + SHA-256 Hashing [This Study]	12.8 %	-2.0%	+2.4%	Yes, no extras	Strong; all 8 met

DISCUSSION

Overall, the single most important finding of this study is that the federated learning pipeline with one-way SHA-256 salted hashing successfully lowers the membership inference attack success rate by 75.7%–76.8% with a maximum accuracy penalty of 0.24%–0.34% that is not statistically different from random variation at Bonferroni-corrected thresholds. This result is better than previously published works: He et al. (2020) achieved success with a rate of about 60% MIA reduction in an IoT healthcare system, and Jeong & Chung (2024) achieved a rate of 65–70% MIA reduction in federated healthcare analytics. This higher reduction is explained by the per-facility salting of the data, which prevents hash comparisons across facilities, and by the generalisation of the quasi-identifier for rare ICD-10 categories. These two elements are necessary, not optional, for an effective hashing privacy mechanism and for the design of future hashing-based PPML systems. One of the most crucial results is the baseline of Config 2, which achieves an MIA success rate of 48.3–57.4% without hashing. This finding directly contradicts the one widely accepted in the federated learning literature and practitioner talk that data localisation is synonymous with privacy. The assumption is contradicted by the empirical evidence, as an adversary with access to the trained federated model can still conclude the individual training membership at a rate significantly higher than a random guess, which was revealed by Ndlovu et al. (2020) in African healthcare datasets. Privacy protection is a key feature of federated learning, not a peripheral addition, as the hashing function of SHA-256 enables it. The deteriorated class recall of 94.6% is the most clinically significant finding. The framework is least susceptible to misidentifying patients with worsening health status, the type of error for which there is the greatest risk of harm to the patient, consistent with the clinical need for a false-negative to predict an adverse health outcome to be more dangerous than a false-positive in an adverse health outcome prediction application. This result is the result of the hashing with SHA-256, which retains all the clinical feature values and only replaces the PatientID. Any markers indicating patient deterioration, such as low CD4 counts, high viral loads, and abnormal blood glucose levels, are completely unaffected by the hashing step. The adoption question has real-world implications for the 2.4% training time overhead. Computational impracticality is the most frequent concern with privacy-preserving techniques in low-resource health systems. SHA-256 hashing can process 13,240 records in 1.84 seconds, and 300 patients per day at an admission to a hospital in 55.2 milliseconds, while remaining silent in all existing clinical

processes. While stronger in theory, Homomorphic Encryption comes with more than 14,000% overhead and is economically impractical for the Nigerian public health facilities that would benefit from it. The BLAKE2 alternative takes the same dataset 1.12 seconds (39.1% faster) – a viable alternative for the more resource-constrained nodes in the Primary Health Centre. This finding is consistent with the broader cryptographic literature: Chen, A.C.H. (2024), comparing SHA-2, SHA-3, and BLAKE2 directly, found the three families statistically indistinguishable on security (measured via hash-output entropy across 32,768 generated values), while SHA-2 and BLAKE2 both ran faster than SHA-3 in that study's measurements. SHA-3 was not separately benchmarked within this study's own pipeline. However, based on Chen's (2024) findings, it is not expected to outperform SHA-256 or BLAKE2 on speed, and its adoption would not be justified here by a security advantage, since none was found among the three families at equivalent output length. This positions the choice between SHA-256 and BLAKE2 as an infrastructure and tooling decision rather than a security decision: SHA-256, as the incumbent NIST standard, benefits from broader library support across Nigerian hospital IT environments. At the same time, BLAKE2's speed advantage makes it the more defensible default specifically for the lowest-resourced PHC nodes identified elsewhere in this study as the weakest link in data quality and infrastructure. The PHC accuracy gap (0.854 vs global mean 0.871, $p = 0.028$) is the most important finding from a health equity perspective. The gap is not due to the framework; it works consistently with the data it is given. This gap is due to structural data quality problems at the PHC level, as well as to the fact that there were few records (on average, 367 per node), a high proportion of missing data (11.2%), and ICD-10 coding was often inconsistent. While FedAvg aggregation helps close the gap, it does not eliminate the accuracy gap stemming from the quality of the input data. Filling this gap requires investments in PHC data infrastructure, such as implementing standardised EHR templates, enhancing laboratory reporting, and training staff in diagnostic coding. The federated architecture enables staged adoption: as facilities meet minimum data quality criteria, they join the network. Satisfaction with all eight NDPR/NDPA compliance requirements based on independently verifiable evidence is a policy contribution separate from the findings of technical performance. This study presents a tested reference architecture that the Nigeria Data Protection Commission/NITDA can adopt as the reference for codes and guidelines on the implementation of AI in Nigeria's healthcare sector, as required by the Nigeria Data Protection Act (2023). Each compliance requirement is linked to a concrete, documented, and independently verifiable mechanism, which can withstand a formal audit by the regulator. Four important limitations of this study should be noted. The most important one is scale: in the evaluation, they have 12 simulated nodes with 13,240 records, but for the national deployment, they will have hundreds of institutions and millions of records. The convergence properties of FedAvg under such conditions need to be tested in practice before deployment. The evaluation is also task-dependent – i.e., patient health outcome classification may entail a distinct accuracy-privacy trade-off as compared to other clinical tasks such as survival analysis, drug interaction detection, or disease incidence forecasting. The adversary model assumes that adversaries are bounded; if they are adaptive and have access to a larger auxiliary dataset, or partial knowledge of the salting protocol, then they can be more successful in their attacks than measured in this study. Finally, the PHC accuracy gap identified in the fairness analysis is

not easily remedied by the technical plan alone. It must be accompanied by concurrent and sustained investments in primary care data infrastructure, such as standardised EHR templates, improved lab reporting, and training in diagnostic coding at the PHC level.

ACKNOWLEDGEMENTS

The authors appreciate the supervisory committee: Prof. G.I.O. Aimufua and Dr. S.E. Chaku for their guidance throughout this research. No external funding was received for this research.

REFERENCES

- Adeleke, O. & Okonkwo, U.M. (2023). Implementation gaps in Nigeria's data protection framework: Evidence from healthcare institutions. *Journal of African Law and Technology*, 5(2): 88–107.
- Ali, M., Naeem, F., Tariq, M. & Kaddoum, G. (2024). Federated learning for privacy preservation in smart healthcare systems: A comprehensive survey. *IEEE Journal of Biomedical and Health Informatics*, 28(1): 22–47.
- Awotunde, J.B. & Jimoh, R.G. (2024). Data governance and privacy compliance in Nigerian tertiary healthcare institutions. *African Journal of Information Systems*, 16(3): 44–67.
- Chen, A.C.H. (2024). Evaluation of hash algorithm performance for cryptocurrency exchanges based on blockchain system. *arXiv preprint*, arXiv:2408.11950.
- Eden, R., Tan, F.T.C. & Tan, B. (2025). Privacy-by-design in federated health data ecosystems: A regulatory alignment analysis. *Information & Management*, 62(1): 103891.
- Faridoon, A. & Kechadi, M.T. (2024). Privacy threats in machine learning pipelines: A taxonomy and mitigation framework for healthcare AI. *Future Generation Computer Systems*, 152: 71–89.
- Fredrikson, M., Jha, S. & Ristenpart, T. (2015). Model inversion attacks that exploit confidence information and basic countermeasures. *Proceedings of the 22nd ACM Conference on Computer and Communications Security*, Denver, USA 2015: 1322–1333.
- He, C., Li, S., So, J., Zhang, M., Wang, H., Wang, X. & Avestimehr, S. (2020). FedML: A research library and benchmark for federated machine learning. *NeurIPS Workshop on Scalability, Privacy, and Security in Federated Learning*, Vancouver, Canada 2020: 1–9.
- Hevner, A.R., March, S.T., Park, J. & Ram, S. (2004). Design science in information systems research. *MIS Quarterly*, 28(1): 75–105.
- IBM Security Report (2024). *Cost of a Data Breach Report 2024*. Armonk, New York: IBM Corporation.
- Jeong, Y. & Chung, M. (2024). Salted hashing for cross-institutional federated healthcare analytics: A privacy and performance evaluation. *Journal of Medical Systems*, 48(1): 12.
- Kairouz, P., McMahan, H.B., Avent, B., Bellet, A., Bennis, M. & Bhagoji, A.N. (2023). Advances and open problems in federated learning. *Foundations and Trends in Machine Learning*, 14(1–2): 1–210.
- Kaspersky (2024). *Healthcare Cybersecurity Report: Sub-Saharan Africa 2023*. Moscow: Kaspersky Lab.
- Li, Q., He, B. & Song, D. (2023). Model-contrastive federated learning for clinical risk prediction. *Journal of the American Medical Informatics Association*, 30(4): 653–663.
- McMahan, H.B., Moore, E., Ramage, D., Hampson, S. & Agüera y Arcas, B. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, Fort Lauderdale, USA 2017: 1273–1282.
- Mutembei, A., Okoye, F. & Abubakar, R. (2021). Insider threats and data breaches in Nigerian health facilities: A case study analysis. *West African Journal of Health Informatics*, 9(1): 14–28.
- Narayanan, A. & Shmatikov, V. (2010). Myths and fallacies of personally identifiable information. *Communications of the ACM*, 53(6): 24–26.
- Ndlovu, K., Khumalo, S. & Dlamini, M. (2020). Re-identification vulnerability in South African electronic health records: Empirical evaluation and hashing recommendation. *South African Journal of Biomedical Engineering*, 12(2): 34–51.
- Okafor, T. & Eze, C. (2025). Data protection compliance in Nigerian hospitals: A nationwide survey. *Nigerian Journal of Health Informatics*, 4(1): 1–22.
- Okoye, F., Nwachukwu, A. & Bello, A. (2022). Cybersecurity incidents in Nigerian healthcare: A systematic review of documented cases 2016–2022. *Journal of Healthcare Engineering*, 2022: 1–14.
- Peffer, K., Tuunanen, T., Rothenberger, M.A. & Chatterjee, S. (2007). A design science research methodology for information systems research. *Journal of Management Information Systems*, 24(3): 45–77.
- Raisaro, J.L., Gwynn, J., Albergante, L., McLetchie, S. & Bhatt, D.L. (2025). Privacy-preserving patient recruitment for clinical trial cohort selection using federated hashing. *Nature Digital Medicine*, 8(1): 21.
- Rana, R. & Marwaha, N. (2024). Privacy-preserving record linkage using salted cryptographic hashing in federated healthcare environments. *Journal of Biomedical Informatics*, 151: 104591.
- Shokri, R., Stronati, M., Song, C. & Shmatikov, V. (2017). Membership inference attacks against machine learning models. *Proceedings of the 38th IEEE Symposium on Security and Privacy*, San Jose, USA, 2017: 3–18.
- Song, C., Ristenpart, T. & Shmatikov, V. (2024). Machine learning models that remember too much: Privacy implications of memorisation in federated learning. *IEEE Transactions on Information Forensics and Security*, 19(1): 1244–1259.
- Tramèr, F., Zhang, F., Juels, A., Reiter, M.K. & Ristenpart, T. (2016). Stealing machine learning models via prediction APIs. *Proceedings of the 25th USENIX Security Symposium*, Austin, USA 2016: 601–618.
- WHO (2023). *Global Strategy on Digital Health 2020–2025: Progress Report*. Geneva: World Health Organization.
- Zhang, Y. & Chen, X. (2024). Machine learning applications in clinical diagnosis and prognosis: A systematic review. *The Lancet Digital Health*, 6(1): e45–e59.