

Cloud Security for Nigerian Critical Infrastructure: a Five-Pillar Implementation Framework and Synthetic-Data Demonstration

¹Olawale Olalekan Onalaja, ¹Sulaiman Adesegun Kukoyi, ²Oluwaseun Ayodeji Odusanya,
¹Abayomi Opeoluwa Kehinde and ^{*3}Olugbenga Gabriel Obadina

¹ Ogun State Polytechnic of Health and Applied Science, Ilese-Ijebu, Ogun State, Nigeria

²D. S. Adegbenro ICT Polytechnic, Itori, Ewekoro, Ogun State, Nigeria

³Olabisi Onabanjo University, Ago-Iwoye, Ogun State, Nigeria.

*Corresponding email: olugbenga.obadina@yahoo.com

ABSTRACT

Cloud-supported critical services require coordinated technical controls, accountable governance, and demonstrable recovery arrangements. This article proposes a Nigeria-focused cloud-security implementation framework and evaluates the behaviour of an accompanying implementation-gap index using synthetic data. The framework organises existing practices into five pillars: technical and operational controls, policy and governance, capacity building, collaboration, and threat intelligence. It links these pillars to control ownership and auditable evidence, with reference to international guidance and Nigerian policy instruments. The computational demonstration uses 100 fictional organisational profiles per replication, nine indicators with fixed marginal distributions, and 1,000 replications at each of three latent dependence settings. No professionals were surveyed, no experts were interviewed, and no incident records were analysed. Although the equal-pillar mean index remains 58.625 under every setting by construction, the simulated mean number of profiles with co-existing unsupported-software and authentication gaps increases from 32.42 to 47.99 as latent dependence increases. The corresponding number with an index of at least 75 increases from 20.67 to 48.02. Median overlap between alternative weighting schemes and the equal-pillar top-20 selection ranges from 13 to 20 profiles. These results demonstrate that fixed headline percentages do not identify the concentration of implementation gaps, and that review priorities depend on joint structure and weighting choices. The contribution is a transparent implementation and evaluation approach, not evidence of Nigerian incident prevalence, terrorist attribution, regulatory non-compliance, or intervention effectiveness. Independent field evaluation is required before operational adoption.

Received: 12 May 2026

Accepted: 25 May 2026

Published: 24 Sept. 2026

Keywords: Cloud Security, Critical Infrastructure, Nigeria, Synthetic Data, Cybersecurity Governance, Sensitivity Analysis, Shared Responsibility

INTRODUCTION

Cloud computing enables on-demand access to shared computing resources through service models that distribute operational responsibilities differently between providers and customers (Mell and Grance, 2011). For an organisation delivering an essential service, moving an application or dataset to the cloud does not eliminate responsibilities for identity management, data governance, supplier oversight, or continuity. Rather, it changes where controls must operate and which parties must demonstrate that they are effective. Cloud-security research has long identified the difficulties introduced by outsourced infrastructure, service dependencies, and the need to maintain security and compliance across organisational boundaries (Hashizume et al., 2013).

These issues have particular institutional relevance to Nigerian critical infrastructure. The Designation and Protection of Critical National Information Infrastructure Order, 2024 identifies covered systems across sectors that include power and energy, information and communications, and banking and finance. It also provides for protection planning and a Trusted Information Sharing Network (Federal Republic of Nigeria, 2024b). A useful cloud-security framework should therefore connect organisational safeguards with existing coordination arrangements rather

than propose a separate national architecture without considering its relationship to established institutions.

The original motivation for this work included concern about terrorist threats. However, perceived intent, technical capability, and attribution of a particular incident are distinct evidential questions. An implementation assessment cannot establish that a terrorist organisation conducted ransomware operations merely because a service is vulnerable or respondents consider it a potential target. The present article consequently adopts a threat-agnostic resilience perspective. Deliberate disruption remains within scope, but no attack is attributed to a named group, and no estimate of cyberterrorism probability is made.

A second challenge concerns the interpretation of incomplete evidence. Aggregate statements such as the proportion of systems with unsupported software or missing authentication controls do not reveal whether those gaps occur in the same organisations. Likewise, a composite score can conceal how much a prioritisation decision depends on the chosen indicator definitions and weights. A locally relevant framework should expose these assumptions instead of turning a convenient summary into an apparently validated estimate of security risk.

Cloud Responsibility and Service Resilience

Cloud security requires explicit service boundaries. Responsibilities for guest operating systems, application configuration, privileged identities, physical infrastructure, and data handling are not identical in infrastructure-, platform-, and software-as-a-service arrangements. The Cloud Security Alliance's Cloud Controls Matrix (CCM) version 4.1 provides a cloud-specific control framework and supporting ownership and applicability guidance (Cloud Security Alliance, 2026). The pNCSF uses this shared-responsibility perspective: each relevant control should identify both its accountable owner and the evidence available from the party implementing it.

The NIST Cybersecurity Framework (CSF) 2.0 provides a complementary outcome-oriented structure organised around Govern, Identify, Protect, Detect, Respond, and Recover. It is intentionally non-prescriptive and supports context-specific current and target profiles (NIST, 2024). ISO/IEC 27001:2022 addresses the management-system requirements surrounding information security (ISO/IEC, 2022). These instruments serve different purposes. A mapping to their themes does not demonstrate conformity, certification, or equivalent coverage.

Incident response and recovery must also extend beyond purchasing preventive tools. NIST SP 800-61 Revision 3 relates incident-response considerations to the wider CSF 2.0 risk-management process (Nelson et al., 2025). Where cloud-supported services interface with operational technology, performance, reliability, and safety constraints must influence control selection and change management (Stouffer et al., 2023). An unsupported component may require a documented treatment plan and compensating safeguards rather than an untested update applied without regard to operational consequences.

Nigerian Institutional and Policy Context

The policy context is not adequately described by the Cybercrimes Act 2015 and the Nigeria Data Protection Regulation 2019 alone. Relevant instruments include the Cybercrimes legislation with its 2024 amendment, the Nigeria Data Protection Act 2023, the 2024 critical-infrastructure designation order, and the Nigeria Data Protection Commission's General Application and Implementation Directive (GAID) 2025. Table 1 identifies their relevance to this framework without treating the table as a comprehensive legal opinion.

Table 1: Selected Policy and Standards Context for the Proposed Framework

Instrument or reference	Design implication	Interpretation limit
Cybercrimes legislation with 2024 amendment (Federal Republic of Nigeria, 2024a)	Record applicable cybersecurity and critical-infrastructure responsibilities.	The legislation's existence is not a measurement of enforcement effectiveness.
CNII designation order, 2024 (Federal Republic of Nigeria, 2024b)	Connect protection planning and information sharing to designated assets and existing coordination arrangements.	Designation does not establish that a particular operator is secure or compliant.
NDPA 2023 and GAID 2025 (Federal Republic of Nigeria, 2023; NDPC, 2025)	Document applicable personal-data duties, controller/processor roles, assessment and breach-management procedures.	An implementation-gap index is not a legal-compliance determination.
National Cybersecurity Policy and Strategy 2021 (ONSA, 2021)	Use the national strategy as a coordination reference.	No policy-performance rating is inferred from synthetic data.
National Digital Cloud Policy announcement, 17 August 2026 (FMCIDE, 2026)	Track government-cloud governance and data-classification requirements relevant to the service.	An announcement is not the complete implementing rulebook; public/private applicability differs.
NIST CSF 2.0; CSA CCM v4.1; ISO/IEC 27001:2022	Relate outcome management, cloud control ownership and information-security governance.	The authors' crosswalk does not certify conformity or equivalent coverage.

Sources: The Named Official Instruments, Ministry Announcement and Standards Bodies, Listed in the References. This Purposive Contextual Summary is not Exhaustive Legal Advice. The 2026 Cloud-Policy Source is Explicitly an announcement. ONSA = Office of the National Security Adviser

The National Cybersecurity Policy and Strategy 2021 remains a named strategic reference in the national policy record (Office of the National Security Adviser, 2021). More recently, the Federal Ministry of Communications, Innovation and Digital Economy announced the National Digital Cloud Policy on 17 August 2026. Its announcement addresses government cloud transformation, investment, and proportionate safeguards for defined government and regulated data; it expressly does not describe a blanket localisation requirement for all commercial data (FMCIDE, 2026). This recent policy direction reinforces the need for an updatable obligations register and clear separation between public-sector requirements and obligations applicable to a particular private operator.

The present documentary review is purposive rather than systematic. It prioritises official instruments and technical guidance relevant to control ownership, personal-data protection, infrastructure coordination, and recovery. The policy context was checked during preparation on 7

September 2026. Where a source is a policy announcement rather than the full implementing instrument, that distinction is retained. Applicability, commencement, subsequent directives, and sector-specific obligations must be checked by the responsible organisation before deployment.

Research Questions and Scope of the Contribution

This article addresses three questions. First, how can established cloud-security practices be organised into a Nigeria-focused control-owner-evidence framework? Second, how much can the concentration of implementation gaps vary when their marginal frequencies are held constant? Third, how sensitive are the profiles selected for review to alternative, explicitly stated weighting schemes?

The contribution is deliberately bounded. The proposed Nigeria Cloud Security Framework, abbreviated pNCSF, is an implementation and coordination proposal, not an official Nigerian standard. Its five pillars do not constitute new security controls. The accompanying contribution is a

reproducible demonstration of how aggregate indicators, joint assumptions, and weights affect implementation-gap review. This is a synthetic-data framework study, not a mixed-methods field investigation or an evaluation of national cybersecurity performance.

MATERIALS AND METHODS

Research Design: Framework Demonstration Versus Empirical Validation

Design-oriented research distinguishes the development of an artefact from evidence that the artefact is useful in practice (Peppers et al., 2007). Here, the artefact consists of the five-pillar organisation, a responsibility-and-evidence record, and a transparent implementation-gap calculation. The computational experiment evaluates internal behaviour under declared assumptions. It does not constitute a completed field-validation stage.

Simulation reporting benefits from specifying aims, data-generating mechanisms, quantities of interest, methods, and performance measures in advance of interpretation (Morris et al., 2019). This principle is applied below. In particular, synthetic profiles are not presented as anonymised respondents or as a reconstruction of unavailable microdata.

The study makes no claim that its chosen distributions approximate the Nigerian population.

The Proposed Five-Pillar Implementation Framework

The pNCSF is designed to make four matters explicit: the cloud service and essential function under review; the parties responsible for each control; the evidence used to assess implementation; and the uncertainty remaining when evidence is incomplete. It retains five complementary pillars but treats incident response, recovery, and supplier coordination as cross-cutting activities. Figure 1 summarises the proposed workflow, and Table 2 gives examples of responsibilities and evidence.

An assessment record should identify the service boundary, deployment and service model, applicable requirement, control owner, implementing party, evidence artefact, assessment date, unresolved dependency, and next review action. For example, an authentication entry would specify which human accounts are in scope and how implementation was verified. Workload identities require their own credential and authorisation controls; an instruction to apply interactive multi-factor authentication to every machine identity would not be technically appropriate.



Figure 1: Proposed Implementation and Evidence Workflow. The Five-Pillar framework Connects Policy and Service Context to Accountable Control Ownership and Auditable Evidence. The Numerical Index is a Limited Review Aid Within this Workflow, not a Compliance Certificate, an Attack-probability Model or an Official Nigerian Framework

Technical and operational controls cover asset and dependency visibility, supported software or documented treatment of legacy exposure, human-account authentication, configuration monitoring, protected backups, recovery testing, and response capability. The accountable cloud-service owner must distinguish operator-managed components from provider-managed services. Evidence should establish the relevant control's scope and test status, not merely the existence of a product licence.

Policy and governance connect obligations to operational responsibilities. A control responsibility matrix, supplier agreement, data-classification decision, personal-data assessment, and exception register make the governance pillar inspectable. The framework does not infer statutory compliance from a self-rated maturity score. It instead requires each legal or contractual obligation to be assessed against the appropriate evidence and accountable function.

Capacity building addresses whether the people performing cloud administration, incident response, procurement, and oversight can fulfil their assigned roles. Training attendance is an input, whereas demonstrated competence and completed

exercises are more informative evidence of capability. Resource constraints should guide feasible delivery arrangements, but geographic labels alone should not determine an organisation's assessed capability.

Collaboration addresses coordinated action among service owners, providers, integrators, regulators, and incident-response bodies. Contact arrangements, escalation responsibilities, information-sharing safeguards, and joint exercises are preferable to an untested claim that collaboration exists. National proposals should take account of the Trusted Information Sharing Network provided for in the 2024 designation order rather than assume that another central platform is necessarily required (Federal Republic of Nigeria, 2024b).

Threat intelligence concerns the evaluation and use of relevant threat information. The proposed evidence record includes provenance, relevance to the service, confidence, the response decision, and whether the item remains current. Unverified attribution is not upgraded to a fact because it appears in an intelligence feed. Information sharing must also respect confidentiality and personal-data obligations.

Table 2: Proposed Pillars, Accountable Functions and Illustrative Evidence

Pillar	Accountable functions	Illustrative evidence	Related CSF 2.0 functions
Technical and operational controls	Cloud-service owner; operator and provider teams within agreed boundaries.	Asset/dependency record; scoped MFA evidence; legacy treatment plan; restore and response-test records.	Identify; Protect; Detect; Respond; Recover
Policy and governance	Leadership/CISO; procurement; service owner; data-protection function.	Control responsibility matrix; obligations register; supplier agreement; documented exceptions.	Govern; Identify
Capacity building	Role owners; security and training leads.	Role-specific competency requirements; practical assessment and exercise records.	Protect
Collaboration	Service owner; provider response lead; relevant coordination contacts.	Escalation and notification matrix; joint exercise; information-sharing safeguards.	Govern; Respond; Recover
Threat intelligence	Security-operations or threat-analysis lead.	Source provenance; relevance and confidence assessment; triage decision; action record.	Identify; Detect; Respond

From Evidence to Review Priorities

The index developed in Section 2.6 is a compact demonstration of review prioritisation, not a substitute for the full evidence record. A high value means that more of the selected implementation gaps are present under specified weights. It does not quantify expected loss, probability of attack, service criticality, or the likelihood that an adversary will succeed.

Unknown evidence must remain distinct from both an established gap and a verified control. Similarly, a genuinely inapplicable control should not automatically be scored as either fully implemented or absent. The numerical experiment assumes all nine indicators are applicable and observed. An operational extension would need a documented applicability rule, a consistent treatment of unknowns, and independent assessment of measurement validity before scores could be compared across organisations.

Synthetic Data: Purpose and Provenance

The experiment evaluates sensitivity to assumptions while preserving the same headline inputs. It uses 100 organisational profiles per replication. This size allows fixed counts to correspond directly to percentages and makes the

supplied example dataset easy to inspect; it is not a statistically powered sample of organisations. No person was recruited, surveyed, or interviewed. No operational network, cloud account, incident log, or confidential organisational record was accessed.

Four numerical targets retain illustrative values from the predecessor manuscript: 65 profiles with an unsupported-software gap, 50 with a human-authentication gap, a policy-translation level averaging 2.8, and an incident-response level averaging 2.5. These targets are scenario assumptions, not independently verified observations. The ordinal constructs and their complete level distributions are specified anew. Additional indicator totals are also assumed explicitly. No ransomware-incidence variable, terrorist-attribution variable, or legal-compliance outcome is generated.

Each example profile has a label and an identifier. Sector labels comprise 34 finance, 33 telecommunications, and 33 energy profiles, assigned independently of the security indicators. These labels provide scenario context only: they do not establish a sampling frame, sectoral prevalence, or a sector-specific risk ordering. Table 3 specifies every marginal distribution used in the experiment.

Table 3: Fully Synthetic Scenario Inputs: Exact Marginal Distributions in Each 100-Profile Dataset

Indicator	Coding	Fixed scenario input	Pillar
Unsupported-software gap	Binary: 1 = gap	65 of 100 profiles	Technical / operational
Human-account MFA gap	Binary: 1 = gap	50 of 100 profiles	Technical / operational
Restore-test gap	Binary: 1 = gap	45 of 100 profiles	Technical / operational
Incident-response level	Ordinal: 1–5	Counts: 20, 35, 25, 15, 5 Mean: 2.50	Technical / operational
Ownership-documentation gap	Binary: 1 = gap	50 of 100 profiles	Policy / governance
Policy-translation level	Ordinal: 1–5	Counts: 15, 25, 30, 25, 5 Mean: 2.80	Policy / governance
Role-training gap	Binary: 1 = gap	60 of 100 profiles	Capacity
Joint-response-exercise gap	Binary: 1 = gap	55 of 100 profiles	Collaboration
Intelligence-triage gap	Binary: 1 = gap	70 of 100 profiles	Threat intelligence

Data Generation: Exact Margins with Alternative Dependence Structures

For profile i and indicator j , latent gap propensities are generated as

$$Z_{ij} = \sqrt{\rho} U_i + \sqrt{1 - \rho} E_{ij}, \quad U_i, E_{ij} \sim N(0,1). \quad (1)$$

The latent variables are otherwise independent. For each binary indicator, a gap is assigned to the exact number of profiles with the largest latent values required by Table 3. For each ordinal indicator, lower capability levels are assigned to higher latent gap propensities according to the fixed level

counts. Thus, every replication has identical marginal counts and ordinal means, while their joint arrangement can change. Three settings are examined: $\rho = 0.0, 0.5$, and 0.9 . The parameter ρ describes the common latent dependence before rank allocation. It is not the observed correlation of the resulting binary or ordinal variables. Moreover, fixing the margins introduces constraints across profiles within a replication; the generated records should not be interpreted as independent Bernoulli observations from a population survey. The master seed is 20260907. The illustrative downloadable dataset uses a separate seed stream with $\rho = 0.5$. The main

experiment uses 1,000 replications for each dependence setting, for 3,000 artificial datasets in total. Streams are identified by the master seed, condition index, and replication index, allowing every result to be regenerated without selecting a favourable random realization.

Implementation-Gap Index

Let the binary gaps denote unsupported software (u_i), human-account MFA (a_i), restore testing (b_i), ownership documentation (o_i), role training (c_i), joint response exercises (l_i), and intelligence triage (t_i). The policy-translation level P_i and incident-response level R_i range from 1 to 5, with higher values representing stronger assumed capability. They are converted to gaps by

$$p_i = \frac{5-P_i}{4}, \quad r_i = \frac{5-R_i}{4}. \quad (2)$$

This conversion assumes equal spacing between ordinal levels. It is a modelling convenience, not a validated property

of a questionnaire or maturity scale. The technical and governance pillar gaps are

$$T_i = \frac{u_i + a_i + b_i + r_i}{4}, \quad G_i = \frac{o_i + p_i}{2}. \quad (3)$$

The remaining pillar gaps are $C_i = c_i$, $L_i = l_i$, and $I_i = t_i$. For nonnegative weights summing to one, the implementation-gap index is

$$S_i = 100(w_T T_i + w_G G_i + w_C C_i + w_L L_i + w_I I_i). \quad (4)$$

The default specification gives each pillar a weight of 0.2. Its indicator weights are therefore unequal: each of the four technical indicators receives 0.05, each governance indicator receives 0.10, and each of the other three indicators receives 0.20. This distinction is important when comparing equal-pillar and equal-indicator weighting. Table 4 gives all alternatives; none was elicited from experts or fitted to incident outcomes.

Table 4: Assumed Weighting Schemes for Sensitivity Analysis

Weighting scheme	Technical	Governance	Capacity	Collaboration	Intelligence
Equal pillars (reference)	0.20	0.20	0.20	0.20	0.20
Equal indicators	4/9	2/9	1/9	1/9	1/9
Technical emphasis	0.40	0.20	0.15	0.15	0.10
Governance emphasis	0.15	0.40	0.15	0.15	0.15

Weights sum to one. Equal-indicator weighting accounts for four indicators in the technical pillar, two in governance, and one in each remaining pillar

Quantities Examined, Computational Checks and Reproducibility

For every replication, the analysis records the mean index, the profile-level 90th percentile, the number of profiles with both unsupported-software and MFA gaps, the number with an index of at least 75, and the percentage of total index mass located in the 20 highest-scoring profiles. Quantiles use linear interpolation between ordered values. The value 75 is an illustrative reporting threshold, not a validated high-risk classification or a regulatory boundary. Likewise, selecting 20 profiles represents an assumed review capacity, not an empirically optimised budget.

Weight sensitivity is measured by the overlap between the equal-pillar top-20 set and each alternative top-20 set. Scores are rounded to 12 decimal places for tie identification, and equal scores are ordered by fictional profile identifier. The tie rule is reproducible but has no substantive security meaning. Jaccard overlap and Spearman rank correlation are supplied in the supplementary results as additional descriptive measures.

For replicate-level metrics, Monte Carlo standard errors are calculated as the sample standard deviation across replications divided by the square root of 1,000. Reported simulation ranges are the 2.5th and 97.5th percentiles across generated datasets, not confidence intervals for Nigerian organisations. No population hypothesis tests, survey sampling errors, or causal effects are estimated.

The supplied code checks the exact marginal distributions, reproducibility of the example stream, the analytical mean, index bounds, and monotonicity when a pillar gap is reduced. These are implementation checks and straightforward consequences of the scoring design; they are not independent demonstrations of predictive validity. The software-version manifest, configuration, complete replication summaries, and example dataset accompany the manuscript.

RESULTS AND DISCUSSION

Fixed Means Do Not Identify Co-Occurrence

The mean pillar gaps are 0.55625 for technical and operational controls, 0.525 for policy and governance, 0.600 for capacity, 0.550 for collaboration, and 0.700 for threat intelligence. Consequently, the equal-pillar mean index is exactly 58.625 under every dependence setting. More generally,

$$\bar{S} = 100 \sum_{k=1}^5 w_k \bar{H}_k, \quad (5)$$

where H_k denotes a pillar gap. Since the marginal means and weights are fixed, rearranging the indicators among profiles cannot change this mean. The invariance is a property of the construction, not evidence that different joint vulnerability structures are equally consequential.

The unsupported-software and MFA margins also illustrate an identification limitation. For gap sets A and B , the possible count with both gaps satisfies

$$\max(0, n_A + n_B - N) \leq n_{A \cap B} \leq \min(n_A, n_B). \quad (6)$$

With 65 and 50 gaps among 100 profiles, the intersection can range from 15 to 50 without altering either marginal percentage. The latent dependence settings illustrate part of this range; they do not exhaust all admissible joint arrangements.

Dependence Changes the Concentration of Gaps

Table 5 and Figure 2 show the computed results. At $\rho = 0$, the mean count with both unsupported-software and MFA gaps is 32.42. It rises to 40.08 at $\rho = 0.5$ and 47.99 at $\rho = 0.9$. The corresponding Monte Carlo standard errors are approximately 0.074, 0.070, and 0.041 profiles. These are properties of the specified generators, not estimates of co-occurrence in Nigerian organisations.

The upper tail changes markedly despite the constant mean. The average profile-level 90th percentile increases from 82.43 to 96.28 and 99.88 across the three settings. The mean number of profiles at or above the illustrative threshold of 75 increases from 20.67 to 37.11 and 48.02. The concentration of total index mass in the highest-scoring 20 profiles similarly increases from 28.46% to 32.81% and 33.83%.

Table 5: Experiment: Mean Replicate-Level Metrics and Monte Carlo Standard Errors

Latent ρ	Both software + MFA gaps (count)	90th percentile of index	Index ≥ 75 (count)	Top-20 index mass (%)
0.0	32.42 (0.074)	82.43 (0.072)	20.67 (0.095)	28.46 (0.019)
0.5	40.08 (0.070)	96.28 (0.040)	37.11 (0.087)	32.81 (0.011)
0.9	47.99 (0.041)	99.88 (0.011)	48.02 (0.065)	33.83 (0.001)

Each Condition Contains 1,000 Datasets of 100 Profiles. Parentheses Contain Monte CARLO Standard Errors, Not Survey Standard Errors. The Mean Index is 58.625 in Every Dataset by Construction. Top-20 Index Mass is Concentration of the Constructed Score, not Preventable Incidents.

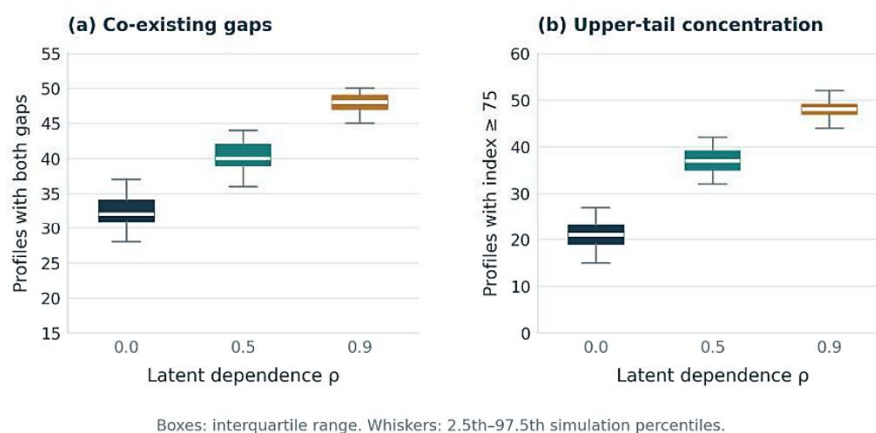


Figure 2: Dependence Experiment with Unchanged Marginal Inputs. Each Box Summarises 1,000 Independently Generated Datasets of 100 Profiles. Panel (a) Counts Co-Existing Unsupported-Software and Human-MFA Gaps; Panel (b) Counts Index Values of at Least 75. The Parameter ρ is Latent, Not the Observed Indicator Correlation. Boxes Show Interquartile Ranges, Central Lines Medians, and Whiskers Central 95% Simulation Ranges. Neither Panel Estimates Nigerian Prevalence

The last measure describes concentration under the chosen index. It is not the proportion of incidents preventable by reviewing 20 organisations. Selecting the highest scores necessarily concentrates that index by construction; this should not be interpreted as evidence that the index outperforms another method against an independent security outcome.

The downloadable example dataset is only one separate realisation at $\rho = 0.5$. It contains 38 profiles with co-existing software and MFA gaps and 39 profiles at or above 75, with a median index of 59.375. Its mean remains 58.625. These

example values are not substituted for the replication-level results or presented as field findings.

Priorities Depend on Weighting Choices

Table 6 and Figure 3 summarise top-20 selection overlap. Under latent independence, the median overlap with equal-pillar weighting is 13 profiles for equal-indicator weighting, 15 for technical emphasis, and 14 for governance emphasis. At $\rho = 0.5$, these medians are 18, 19, and 18. At $\rho = 0.9$, all three alternatives have a median overlap of 20, although some replications still select different profiles.

Table 6: Weight Sensitivity: Top-20 Overlap with Equal-Pillar Weighting

Alternative weighting	$\rho = 0.0$	$\rho = 0.5$	$\rho = 0.9$
Equal indicators	13 [10, 16]	18 [15, 20]	20 [18, 20]
Technical emphasis	15 [12, 18]	19 [16, 20]	20 [18, 20]
Governance emphasis	14 [11, 17]	18 [16, 20]	20 [19, 20]

Entries are Median Overlap Counts [2.5th, 97.5th Simulation Percentiles] Across 1,000 Replications Per Condition. The Maximum possible Overlap is 20. Ranges Describe Conditional Simulation Variation, Not Confidence Intervals for Real Organisations. Ties Use the Deterministic Fictional-ID Rule.

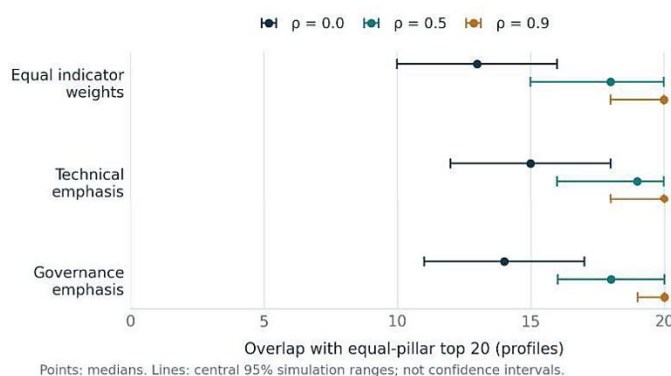


Figure 3: Sensitivity of Review Priorities to Weights. Points Show Median Overlap with the Equal-Pillar top-20 Selection; Horizontal Lines Show the Central 95% Simulation Range. Results Summarise 1,000 Replications for Each Weight/Dependence Combination. Higher Agreement Under Stronger Latent Dependence is a Property of these Artificial Scenarios, Not Independent Evidence of Framework Validity

Stronger common dependence makes profiles more consistently strong or weak across pillars, which reduces disagreements among positive weighting schemes in these scenarios. Agreement under strong dependence is not external validation: it reflects the assumed joint structure. Conversely, disagreement under weaker dependence shows why weights should be declared and decision-makers should inspect which controls drive an assessment.

Changing weights also changes the index definition. The fixed mean is 56.94 under equal-indicator weighting, 57.00 under technical emphasis, and 57.09 under governance emphasis, compared with 58.625 under equal-pillar weighting. None of these lower means represents an improvement in an organisation's security; the underlying indicators are unchanged.

Incomplete Evidence Should Produce Explicit Uncertainty

No values are missing in the experiment. For an eventual assessment with unknown indicator gaps, however, the linear construction provides simple bounds. If v_j is an indicator's effective weight and x_j its normalised gap, then

$$S_{\text{low}} = 100 \sum_{j \in O} v_j x_j, \quad S_{\text{high}} = S_{\text{low}} + 100 \sum_{j \in M} v_j. \quad (7)$$

Here, O and M denote observed and missing indicators, respectively; observed gaps are treated as known, and each missing applicable gap may range from zero to one. If intelligence triage alone is unknown under equal-pillar weights, the interval width is 20 index points. This bound makes the uncertainty visible instead of treating absent evidence as an implemented control. It does not address measurement error in observed indicators or determine how genuinely inapplicable controls should be handled.

Interpretation and Contribution

The experiment supports a methodological conclusion: marginal indicators and a composite mean cannot determine how implementation gaps are distributed across profiles. The same stated totals can produce substantially different co-occurrence, upper-tail concentration, and review selections. This is especially relevant when a proposed framework is motivated by headline percentages but the underlying organisational records are unavailable.

The pNCSF's practical contribution is to connect the selected indicators to service boundaries, accountable owners, evidence, and review actions. Its contribution is contextual

integration and transparent assessment, not a claim that five pillars are intrinsically superior to the six functions of NIST CSF 2.0 or to a cloud-specific control catalogue. The crosswalk should help operators navigate those resources rather than introduce another competing compliance badge. The results also illustrate the distinction between implementation coverage and security outcomes. An index can be reproducible, bounded, and monotonic without being calibrated to loss, disruption, or adversary success. Its apparent precision should not conceal ordinal-scale assumptions, choices of indicator granularity, or unequal effective weights. A threshold or ranking is useful only for a clearly specified decision with appropriately assessed consequences.

Implementation Implications for Nigeria

A practical deployment should begin with a service-scoping and obligations register rather than a score. The organisation would identify the essential function, dependencies, relevant cloud service models, responsible parties, applicable personal-data requirements, and any sector-specific constraints. These requirements should be mapped to an evidence record and reviewed as policy changes. The 2026 cloud-policy announcement makes it particularly important not to assume that the obligations of a Federal Ministry, a private financial institution, and a cloud provider are interchangeable (FMCIDE, 2026).

The next stage would establish a baseline of identity, software-support, backup, monitoring, and response evidence. Customer and provider responsibilities should be recorded separately. In operational-technology environments, verification must respect safety and availability constraints; a numerical legacy-gap flag does not authorise an operational change without appropriate engineering review (Stouffer et al., 2023).

Subsequent review should test cross-organisational response, confirm the availability of relevant evidence from providers, and document how actionable threat information enters decisions. Where connectivity or staffing is constrained, the framework should support proportionate evidence collection, maintained offline response arrangements, and feasible testing schedules. These are design considerations, not findings that all rural installations are less secure than urban ones.

The final stage is independent evaluation. Operators should assess whether the framework improves clarity of responsibility, completeness of evidence, closure of agreed actions, and the quality of response exercises. Operational measures such as recovery-test outcomes and incident-handling performance should be assessed directly. A decrease in the index alone would not establish reduced incident risk, and a change caused solely by reweighting would not constitute an improvement at all.

Limitations

The most important limitation is external validity. Every profile is generated from assumed margins and a simple latent dependence mechanism. No Nigerian organisation's condition, staffing, policy compliance, cloud architecture, or incident history is measured. The findings cannot establish national or sectoral prevalence, urban-rural differences, policy effectiveness, or terrorist capability. Synthetic replication reduces Monte Carlo error conditional on the generator; it does not reduce uncertainty about whether that generator resembles reality.

Second, the assessment instrument is provisional. Binary indicators simplify controls that may differ in scope and strength, while the ordinal transformations assume equal distances between capability levels. The use of a separate restore-test indicator and an incident-response level may also represent overlapping constructs. Expert content review, inter-rater assessment, alternative ordinal transformations, and evaluation of redundancy are needed before operational scoring.

Third, all weighting schemes and the top-20 review capacity are assumed. They do not reflect stakeholder preference elicitation, measured intervention cost, mission impact, or an optimisation criterion. Strong ranking agreement under high dependence is expected for similarly ordered profiles and must not be presented as proof of robustness across all realistic settings.

Fourth, the generators cover independent and positively associated latent propensities with exact margins. They omit negative associations, provider-level clustering, temporal change, conditional service dependencies, and failures that propagate through networks. The analytical overlap bounds are broader than the three mechanisms investigated. Future scenarios could examine those structures, but more elaborate assumptions would still require empirical justification.

Fifth, the documentary assessment is selective. It records the versions and scope of sources used but is not an exhaustive legal analysis or a certification mapping. The 2026 ministry announcement is treated as an announcement, not as evidence that every implementing mechanism is already operational. Implementation must be checked against current instruments and competent regulatory guidance.

Field-Evaluation Agenda

A subsequent empirical study should recruit identifiable units of analysis through a documented sampling strategy while protecting participant and organisational confidentiality. It should distinguish the number of respondents from the number of organisations, record service models and sector composition, and obtain the ethics and organisational permissions appropriate to the planned data collection. None of those activities is claimed in the present article.

The evaluation should first examine whether different assessors reach similar judgments from the same evidence and whether the indicators represent the intended constructs. It should then test decision utility against outcomes not defined by the index itself, such as the identification of independently

verified control deficiencies or the completeness of provider-operator response responsibilities. Where feasible, prospective observation or a carefully designed comparative study could assess operational outcomes, with attention to confounding and disclosure restrictions.

The package can support that preparation by testing data-entry rules, reporting templates, sensitivity procedures, and analysis code before confidential records are introduced. It cannot replace the field evidence needed to validate the framework or justify institutional adoption.

CONCLUSION

This article proposes a five-pillar cloud-security implementation framework for Nigerian critical infrastructure and demonstrates an accompanying implementation-gap index using entirely synthetic data. The framework emphasises explicit responsibility, auditable evidence, coordination with existing institutions, and transparency about uncertain or inapplicable controls. The computational results show that identical marginal indicators and an invariant mean index can conceal markedly different concentrations of gaps. They also show that review priorities can change with weights, particularly when indicators are less strongly associated. These are conditional findings about a specified scoring system and generator, not observations of Nigerian security conditions or estimates of attack likelihood. The pNCSF should therefore be treated as a research proposal and a reproducible assessment prototype. Its operational value, measurement validity, and relationship to service outcomes require independent evaluation. Transparent simulation is useful for developing that evaluation, provided that observations are never represented as survey, interview, incident, or compliance evidence.

REFERENCES

- Cloud Security Alliance. (2026). Cloud Controls Matrix and CAIQ v4.1. <https://cloudsecurityalliance.org/artifacts/cloud-controls-matrix-v4-1>
- Federal Ministry of Communications, Innovation and Digital Economy (FMCIDE). (2026, August 17). Federal Government unveils National Digital Cloud Policy to drive investment, digital sovereignty and government transformation [Policy announcement]. Official ministry policy announcement
- Federal Republic of Nigeria. (2023). Nigeria Data Protection Act, 2023. Federal Republic of Nigeria Official Gazette. https://cert.gov.ng/ngcert/resources/Nigeria_Data_Protection_Act_2023.pdf
- Federal Republic of Nigeria. (2024a). Cybercrimes (Prohibition, Prevention, etc.) Act, 2015 with Amendment Act 2024. Consolidated publication hosted by the Office of the National Security Adviser/ngCERT. https://cert.gov.ng/ngcert/resources/CyberCrime_Prohibition_Prevention_etc_Act_2024.pdf
- Federal Republic of Nigeria. (2024b). Designation and Protection of Critical National Information Infrastructure Order, 2024. Federal Republic of Nigeria Official Gazette, 111(107). <https://cert.gov.ng/ngcert/resources/cnii-gazette.pdf>
- Hashizume, K., Rosado, D. G., Fernández-Medina, E., & Fernandez, E. B. (2013). An analysis of security issues for cloud computing. *Journal of Internet Services and*

Applications, 4, Article 5. <https://doi.org/10.1186/1869-0238-4-5>

International Organization for Standardization/International Electrotechnical Commission (ISO/IEC). (2022). ISO/IEC 27001:2022: Information security, cybersecurity and privacy protection—Information security management systems—Requirements. <https://www.iso.org/standard/27001>

Mell, P., & Grance, T. (2011). The NIST definition of cloud computing (NIST SP 800-145). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.800-145>

Morris, T. P., White, I. R., & Crowther, M. J. (2019). Using simulation studies to evaluate statistical methods. *Statistics in Medicine*, 38(11), 2074–2102. <https://doi.org/10.1002/sim.8086>

National Institute of Standards and Technology (NIST). (2024). The NIST Cybersecurity Framework (CSF) 2.0 (NIST CSWP 29). <https://doi.org/10.6028/NIST.CSWP.29>

Nelson, A., Rekhi, S., Souppaya, M., & Scarfone, K. (2025). Incident response recommendations and considerations for cybersecurity risk management: A CSF 2.0 Community Profile (NIST SP 800-61 Rev. 3). National Institute of

Standards and Technology. <https://doi.org/10.6028/NIST.SP.800-61r3>

Nigeria Data Protection Commission (NDPC). (2025). Nigeria Data Protection Act (NDP Act) 2023: General Application and Implementation Directive (GAID) 2025 (NDPC/NDP ACT-GAID/01/2025). <https://ndpc.gov.ng/wp-content/uploads/2025/07/NDP-ACT-GAID-2025-MARCH-20TH.pdf>

Office of the National Security Adviser. (2021). National cybersecurity policy and strategy. https://cert.gov.ng/ngcert/resources/NATIONAL_CYBERSSECURITY_POLICY_AND_STRATEGY_2021.pdf

Peffers, 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. <https://doi.org/10.2753/MIS0742-1222240302>

Stouffer, K., Pease, M., Tang, C., Zimmerman, T., Pillitteri, V., Lightman, S., Hahn, A., Saravia, S., Sherule, A., & Thompson, M. (2023). Guide to operational technology (OT) security (NIST SP 800-82 Rev. 3). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.800-82r3>



©2026 This is an Open Access article distributed under the terms of the Creative Commons Attribution 4.0 International license viewed via <https://creativecommons.org/licenses/by/4.0/> which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is cited appropriately.