



RESEARCH ARTICLE

An Integrated Mathematical Model for Detecting Known and Unknown Fraud Patterns with Optimal Resource Allocation in the Nigerian Social Security System

Ometan S. Olorok^{1,*}, Kayoh O. Clinton¹, Akindele M. Okedoye¹

¹*Department of Mathematics, Federal University of Petroleum Resources, Effurun, Nigeria*

*Corresponding authors: olorok63@gmail.com, kayohclinton@gmail.com, okedoye.akindele@fupre.edu.ng

Received: 17 February 2026; Revised: 23 June 2026; Accepted: 1 August 2026; Published: 1 September 2026.

Abstract:

Social security systems in developing countries face substantial fraud losses threatening system sustainability, yet resource-constrained agencies can investigate only a small fraction of claims. We present a comprehensive mathematical framework integrating Bayesian logistic regression for known fraud patterns, Gaussian mixture model-based anomaly detection for novel schemes, and optimal resource allocation under capacity constraints. Temporal adaptation through exponential smoothing maintains detection performance as fraud patterns evolve quarterly. Our greedy allocation policy, derived through marginal benefit analysis, provably minimizes expected cost under investigation capacity constraints. Empirical validation on 10,000 pension claims from Nigerian state bureaus (Lagos, Rivers, Kano) collected over 18 months demonstrates operational effectiveness. The integrated framework achieves area under ROC curve of 0.811 with precision of 84% and recall of 67% on realistic data exhibiting class overlap, label noise (3%), and missing values (5%). Greedy resource allocation detects 7–11 times more fraud than random selection across all capacity levels, translating to estimated annual savings of ₦165 million (\$122,000 USD) for Lagos bureau alone and ₦2 billion (\$1.5 million USD) nationally. Sample complexity analysis reveals that 1,000–2,000 investigated claims suffice for reliable deployment, achievable within 13–18 months for most Nigerian bureaus. Temporal adaptation with exponential smoothing parameter $\alpha = 0.70$ maintains stable performance despite quarterly fraud pattern evolution, preventing the 13-point AUC degradation observed in static models over 12-month horizons.

Keywords: fraud detection, Bayesian inference, resource allocation, social security, anomaly detection, temporal adaptation, developing countries, Nigeria

1. Introduction

Social security systems serve as critical safety nets in both developed and developing economies, providing retirement pensions, disability benefits, and survivor support to vulnerable populations. However, these systems face persistent threats from fraudulent claims that drain limited public resources and undermine system sustainability. The scale of the problem is substantial: the United

States Social Security Administration reported approximately \$7 billion in improper payments during fiscal year 2022 [1], while the United Kingdom’s Department for Work and Pensions estimated fraud and error rates of 3.6% in pension benefits [2]. Developing countries face even greater challenges, with Nigeria’s Economic and Financial Crimes Commission reporting fraud rates between 4–6% in pension systems, translating to annual losses exceeding ₦18 billion (\$42 million USD) [3].

The urgency of addressing this problem cannot be overstated. Fraudulent claims in Nigeria’s pension system do not merely represent financial losses they directly diminish retirement security for legitimate beneficiaries, erode public trust in social protection institutions, and reduce the fiscal capacity of government to expand coverage to the approximately 90% of Nigerian workers currently in the informal sector who remain unprotected. The problem is particularly acute because Nigeria’s pension reform history has already been marked by high-profile scandals involving billions of Naira in diverted pension funds, leaving a system chronically under-resourced and under-trusted. Any framework that can substantially improve fraud detection rates without increasing investigation costs thus offers both immediate fiscal savings and longer-term systemic benefits to governance and social equity.

The fraud detection problem becomes particularly acute in resource-constrained settings characteristic of developing-country administrations. Nigerian state pension bureaus, for example, employ limited investigation staff capable of thoroughly examining only 5–8% of monthly claim volume [4]. This capacity constraint creates a critical allocation question: which claims should receive scarce investigation resources to maximize fraud prevention? Traditional approaches relying on simple rule-based triggers or investigator intuition approximate random selection in effectiveness, missing substantial fraud while wasting resources investigating legitimate claims [5]. The challenge is compounded by fraud patterns that evolve rapidly as perpetrators adapt to detection measures, rendering static models obsolete within months [6].

Academic research on fraud detection spans multiple disciplines including machine learning, statistics, operations research, and criminology. The machine learning literature has explored numerous classification algorithms for fraud detection including neural networks [7], support vector machines [8], random forests [9], and gradient boosting methods [10]. These approaches focus primarily on maximizing predictive accuracy metrics like area under the ROC curve, often achieving impressive performance on benchmark datasets. However, they typically neglect the resource allocation problem faced by practitioners with limited investigation capacity, treat fraud patterns as stationary despite empirical evidence of evolution, and lack formal theoretical guarantees about learning efficiency and detection limits.

Recent work has begun bridging these disciplinary divides. The concept drift literature in machine learning addresses temporal evolution by developing adaptive algorithms that update models as data distributions shift [11, 12]. Multi-armed bandit formulations provide theoretical frameworks for balancing exploration (learning about fraud patterns) versus exploitation (catching known fraud) in sequential decision settings [13]. Game-theoretic analyses model adversarial dynamics where fraudsters strategically respond to detection systems [14]. Information-theoretic approaches quantify fundamental limits on detection performance given feature informativeness [15, 16].

Despite these advances, the fraud detection literature lacks a unified framework that simultaneously addresses supervised learning for known patterns, unsupervised detection for novel schemes, optimal resource allocation under constraints, temporal adaptation to evolving fraud, and formal theoretical guarantees all validated on real administrative data from resource-constrained developing-country contexts. Existing work typically addresses only one or two of these dimensions in isolation.

This paper addresses this gap through the following research questions: (1) Does integrating Bayesian inference with Gaussian mixture model anomaly detection outperform either component in isolation for detecting both known and novel fraud patterns? (2) Does exponential smoothing with an empirically calibrated adaptation parameter maintain detection performance under quar-

terly fraud pattern evolution in Nigerian pension bureaus? (3) Does the analytically derived greedy allocation policy achieve substantially superior cost outcomes compared to naive random selection under realistic Nigerian investigation capacity constraints? (4) What minimum sample size is required for reliable deployment of the integrated framework in resource-constrained bureaus?

This paper presents a comprehensive mathematical framework that fills these gaps through rigorous integration of statistical learning theory, optimization, and stochastic processes.

2. Materials and Methods

Social security fraud represents a persistent challenge to system integrity, with fraudulent claims simultaneously draining resources and undermining public trust. Traditional approaches to fraud detection often rely on manual reviews triggered by simple rules, missing sophisticated fraud schemes while generating high false positive rates. We develop a statistical learning framework that addresses these limitations through principled Bayesian inference combined with anomaly detection.

2.1. Model Foundations and Assumptions

The fraud detection model builds on seven fundamental assumptions that balance mathematical tractability with real-world applicability. We begin by characterizing the claim arrival process, then progressively layer complexity to capture the essential features of fraud detection in practice.

1. **Claim Arrival Process:** Claims arrive as a Poisson process with rate λ (claims per time unit).
2. **Feature Representation and Latent Fraud Status:** Each claim has an associated feature vector $\mathbf{x} \in \mathbb{R}^d$ containing observable characteristics.
3. **Distribution Difference:** Feature distributions differ between legitimate and fraudulent claims.
4. **Pattern Evolution:** Fraud patterns evolve over time, requiring adaptive detection mechanisms.
5. **Resource Constraints:** Investigation resources are limited, requiring prioritization of claims.
6. **Asymmetric Costs:** False positives (legitimate claims flagged as fraud) have different costs than false negatives.
7. Each claim is either fraudulent ($F_i = 1$) or legitimate ($F_i = 0$). The true fraud status is revealed only upon investigation.
8. **Cost Structure:** Investigating a claim incurs a cost C_{FP} if legitimate and saves expected cost C_{FN} if fraudulent, where typically $C_{FN} \gg C_{FP}$.
9. **Independence:** Claims are conditionally independent given the model parameters θ and the fraud probability model.
10. **Capacity Constraints:** At each time period t , the investigator can examine at most I_t claims, where $I_t \ll n_t$ and n_t is the number of incoming claims.

2.2. Mathematical Framework and Notation

With these assumptions established, we now formalize the detection problem. Table 2.1 summarizes our core notation, which we maintain consistently throughout this section.

The fraud detection problem decomposes into three interrelated tasks: (1) estimating fraud probabilities for incoming claims using observed features, (2) detecting anomalous claims that may represent novel fraud schemes, and (3) allocating limited investigation resources to maximize expected fraud detection. We address each task sequentially, then integrate them into a unified framework.

2.3. Bayesian Inference for Known Fraud Patterns

For fraud schemes with historical precedent, we employ Bayesian logistic regression, a natural choice that combines the interpretability of logistic regression with principled uncertainty quantification

Table 2.1: Notation for fraud detection model

Symbol	Definition
X_i	Feature vector for claim i
F_i	True fraud indicator for claim i (0=legitimate, 1=fraudulent)
$\hat{Y}_i$	Predicted fraud probability for claim i
θ	Model parameters (logistic regression coefficients)
C_{FP}	Cost of false positive (investigating legitimate claim)
C_{FN}	Cost of false negative (missing fraudulent claim)
λ	Claim arrival rate (claims per time unit)
I_t	Investigation capacity at time t (number of claims investigable)

through Bayesian updating. The model assigns to each claim i a fraud probability:

$$P(F_i = 1 \mid X_i, \theta) = \frac{1}{1 + \exp(-\theta^T X_i)} = \sigma(\theta^T X_i) \quad (2.1)$$

where $\theta \in \mathbb{R}^d$ is a vector of coefficients learned from historical data and $\sigma(\cdot)$ denotes the logistic sigmoid function. The logistic function maps the linear predictor $\theta^T X_i$ to the unit interval, yielding a well-calibrated probability estimate.

Rather than treating θ as fixed, we adopt a Bayesian perspective by placing a prior distribution over the parameters:

$$\theta \sim \mathcal{N}(\mu_0, \Sigma_0) \quad (2.2)$$

This prior encodes our initial beliefs about fraud patterns before observing data. In practice, we might set $\mu_0 = \mathbf{0}$ (no initial bias) and $\Sigma_0 = \sigma^2 I$ for some regularization parameter σ^2 that prevents overfitting to small samples, following the approach of [17].

As claims are investigated and their true fraud status revealed, we update our beliefs using Bayes' theorem. The posterior distribution after observing data $\mathcal{D}_t = \{(X_1, F_1), \dots, (X_t, F_t)\}$ is:

$$p(\theta \mid \mathcal{D}_t) = \frac{p(\mathcal{D}_t \mid \theta)p(\theta)}{p(\mathcal{D}_t)} \propto p(\mathcal{D}_t \mid \theta)p(\theta) \quad (2.3)$$

where the likelihood $p(\mathcal{D}_t \mid \theta)$ quantifies how well parameter θ explains the observed fraud labels.

This Bayesian updating provides several advantages over frequentist maximum likelihood estimation: it naturally handles small samples through the prior, quantifies parameter uncertainty, and enables sequential learning as new data arrive, as demonstrated by [18].

The posterior distribution generally lacks a closed form, but we can efficiently compute the maximum a posteriori (MAP) estimate $\hat{\theta}_t$ using gradient-based optimization.

Definition 2.1 (Maximum A Posteriori Estimate). The MAP estimate is defined as:

$$\hat{\theta}_t = \arg \max_{\theta} p(\theta \mid \mathcal{D}_t) = \arg \max_{\theta} [\log p(\mathcal{D}_t \mid \theta) + \log p(\theta)] \quad (2.4)$$

Theorem 2.1 (Equivalence of MAP and Regularized MLE). *For a Gaussian prior $\theta \sim \mathcal{N}(\mathbf{0}, \sigma^2 I)$, the MAP estimate in Equation (2.4) is equivalent to ℓ_2 -regularized maximum likelihood estimation with regularization parameter $\lambda = 1/\sigma^2$.*

Proof. The log-posterior is:

$$\log p(\theta \mid \mathcal{D}_t) = \log p(\mathcal{D}_t \mid \theta) + \log p(\theta) - \log p(\mathcal{D}_t) \quad (2.5)$$

$$= \log p(\mathcal{D}_t \mid \theta) - \frac{1}{2\sigma^2} \|\theta\|_2^2 - \frac{d}{2} \log(2\pi\sigma^2) + \text{const} \quad (2.6)$$

where we used the density of the multivariate Gaussian. Maximizing with respect to θ is equivalent to:

$$\hat{\theta}_t = \arg \max_{\theta} \left[\log p(\mathcal{D}_t | \theta) - \frac{1}{2\sigma^2} \|\theta\|_2^2 \right] \quad (2.7)$$

Setting $\lambda = 1/\sigma^2$ recovers the standard ℓ_2 -regularized (ridge) regression objective. This connection was established by [19]. $\square$

This point estimate balances fitting the data with staying close to the prior, providing a regularized solution that generalizes better than pure maximum likelihood [20].

2.3.1. Temporal Adaptation to Evolving Fraud

A very important innovation in this model is the temporal adaptation mechanism that addresses Assumption 4 (pattern evolution). If we simply accumulated all historical data into $\mathcal{D}_t$, the model would become increasingly dominated by old patterns and slow to detect emerging schemes. Conversely, if we used only recent data, we would discard valuable historical information and suffer from high variance. We strike a balance through exponential smoothing of the parameters:

$$\theta_t = \alpha \theta_{t-1} + (1 - \alpha) \hat{\theta}_t \quad (2.8)$$

where $\hat{\theta}_t$ is the MAP estimate from recent data (e.g., the last 6 months) and $\alpha \in [0, 1]$ controls the adaptation rate.

Setting $\alpha = 0$ produces a fully adaptive model that discards history, while $\alpha = 1$ produces a static model that never adapts. We evaluated smoothing parameters $\alpha \in \{0.0, 0.1, \dots, 1.0\}$ using time-series cross-validation on the 18-month Nigerian bureau dataset, training sequentially on periods $1-k$ and evaluating one-step-ahead AUC on period $k+1$. The value $\alpha = 0.70$ minimized mean one-step-ahead prediction loss across the three state bureaus and is also consistent with the theoretical optimum for an Ornstein-Uhlenbeck drift process: the fraud pattern half-life implied by $\alpha = 0.70$ (approximately 2.9 months at 6-week update intervals) aligns with Nigerian investigators' reports that dominant fraud tactics shift quarterly. Values below $\alpha \approx 0.6$ produced excessive variance from small per-period samples (≈ 400 claims), while values above $\alpha \approx 0.8$ failed to track evolving patterns and showed AUC degradation of 5–9 points by month 12.

This simple exponential smoothing scheme has a principled Bayesian interpretation: it approximates a dynamic linear model where the true parameter vector evolves as a random walk, as shown by [21]. More sophisticated approaches using particle filters or Kalman filters could be employed for systems with strongly time-varying fraud patterns, but that will not be addressed in this work.

2.4. Anomaly Detection for Novel Fraud Schemes

While Bayesian logistic regression excels at detecting known fraud patterns, it struggles with novel schemes that differ fundamentally from historical fraud. A claim exhibiting a new type of fraud may receive a low fraud probability simply because it doesn't match the patterns captured by θ . To address this limitation, we complement the Bayesian approach with unsupervised anomaly detection. We model the distribution of legitimate claims using a Gaussian Mixture Model (GMM):

$$p(\mathbf{x}) = \sum_{k=1}^K \pi_k \mathcal{N}(\mathbf{x} | \mu_k, \Sigma_k) \quad (2.9)$$

where K mixture components capture different subtypes of legitimate claims (e.g., retirement claims, disability claims, survivor benefits). The parameters $\{\pi_k, \mu_k, \Sigma_k\}_{k=1}^K$ are estimated from historical legitimate claims using the expectation-maximization (EM) algorithm [22].

Theorem 2.2 (EM Algorithm for GMM). *The EM algorithm monotonically increases the log-likelihood of the observed data:*

$$\log p(\mathcal{X} \mid \Theta^{(t+1)}) \geq \log p(\mathcal{X} \mid \Theta^{(t)}) \quad (2.10)$$

where $\Theta^{(t)} = \{\pi_k^{(t)}, \mu_k^{(t)}, \Sigma_k^{(t)}\}_{k=1}^K$ denotes the parameters at iteration t .

Proof Sketch. The EM algorithm exploits Jensen’s inequality. Let z_i denote the latent component assignment for observation i . The complete-data log-likelihood is:

$$\log p(\mathcal{X}, \mathcal{Z} \mid \Theta) = \sum_{i=1}^n \sum_{k=1}^K \mathbb{I}(z_i = k) [\log \pi_k + \log \mathcal{N}(x_i \mid \mu_k, \Sigma_k)] \quad (2.11)$$

In the E-step, we compute the expected complete-data log-likelihood (the Q-function):

$$Q(\Theta \mid \Theta^{(t)}) = \mathbb{E}_{p(\mathcal{Z} \mid \mathcal{X}, \Theta^{(t)})} [\log p(\mathcal{X}, \mathcal{Z} \mid \Theta)] \quad (2.12)$$

In the M-step, we maximize Q :

$$\Theta^{(t+1)} = \arg \max_{\Theta} Q(\Theta \mid \Theta^{(t)}) \quad (2.13)$$

By construction, $Q(\Theta^{(t+1)} \mid \Theta^{(t)}) \geq Q(\Theta^{(t)} \mid \Theta^{(t)})$. The result follows from showing that this implies monotonic increase in the observed data log-likelihood. See [22] or [18] for the complete proof. $\square$

The intuition is straightforward: a legitimate claim should have high probability under this model, while a fraudulent claim, especially one using a novel scheme, may fall in a low-density region of the legitimate distribution. We quantify this through the anomaly score:

$$A(\mathbf{x}_i) = -\log p(\mathbf{x}_i) \quad (2.14)$$

Higher anomaly scores indicate claims that are unusual relative to the legitimate distribution. This approach is grounded in information theory, where $-\log p(\mathbf{x})$ represents the surprisal or information content of observing $\mathbf{x}$ [16].

Claims with $A(\mathbf{x}_i) > \tau$ for some threshold τ are flagged as potential anomalies requiring investigation. The threshold τ is calibrated to achieve a desired false positive rate on validation data. For instance, if we can tolerate investigating 5% of legitimate claims as anomalies, we set τ at the 95th percentile of anomaly scores observed on a labeled dataset of legitimate claims.

This calibration ensures that our anomaly detector operates at an acceptable specificity while maintaining sensitivity to truly unusual claims, following the principles outlined by [23].

2.4.1. Integrating Bayesian and Anomaly Components

The Bayesian and anomaly approaches provide complementary information: the former captures known fraud patterns through supervised learning, while the latter detects unusual claims through unsupervised learning. We combine them into a unified fraud score:

$$S(\mathbf{x}_i) = \beta P(F_i = 1 \mid \mathbf{x}_i, \theta) + (1 - \beta) \frac{A(\mathbf{x}_i) - A_{\min}}{A_{\max} - A_{\min}} \quad (2.15)$$

where $\beta \in [0, 1]$ controls the relative weight. The second term normalizes anomaly scores to $[0, 1]$ using the minimum and maximum scores observed in a calibration dataset, ensuring both components contribute on comparable scales.

The weight β governs the contribution of supervised (Bayesian) versus unsupervised (anomaly) detection. We selected β through 5-fold cross-validation on the claim Nigerian training set, evaluating integrated AUC at $\beta \in \{0.25, 0.50, 0.75, 1.0\}$. The value $\beta = 0.75$ yielded the highest cross-validated

AUC of 0.813, reflecting the fact that in the Nigerian bureau context where historical labeled fraud data are sufficient for reliable supervised learning the Bayesian component provides stronger discrimination signals. The anomaly component contributes at weight $1 - \beta = 0.25$, providing robustness to novel schemes not represented in training data. For deployment in bureaus with limited labeled data (fewer than 500 examples), we recommend increasing the anomaly weight toward $\beta \approx 0.5$ until sufficient labeled examples accumulate. This β selection procedure should be re-run when adapting the framework to new bureaus or jurisdictions with different fraud composition.

Proposition 2.1 (Range of Combined Score). *The combined fraud score satisfies $S(\mathbf{x}_i) \in [0, 1]$ for all claims i .*

Proof. By the sigmoid logistics function, $P(F_i = 1 | \mathbf{x}_i, \theta) \in (0, 1)$. The normalized anomaly score satisfies:

$$\frac{A(\mathbf{x}_i) - A_{\min}}{A_{\max} - A_{\min}} \in [0, 1] \quad (2.16)$$

by construction. Since $S(\mathbf{x}_i)$ is a convex combination of two quantities in $[0, 1]$ with weights β and $(1 - \beta)$ summing to 1, we have $S(\mathbf{x}_i) \in [0, 1]$. $\square$

The weight β should be set based on the relative prevalence of known versus novel fraud schemes. In a mature fraud detection system with extensive historical data, known patterns likely dominate, suggesting $\beta \approx 0.7$ - 0.8 . In a new system or one facing rapidly evolving fraud, novel schemes may be more common, suggesting $\beta \approx 0.5$. This parameter can be tuned using cross-validation or through online learning that tracks detection performance, as discussed by [5].

2.5. Temporal Dynamics and Continuous-Time Modeling

In Section 2.3.1, we introduced exponential smoothing (Equation (2.8)) for temporal adaptation:

$$\theta_t = \alpha \theta_{t-1} + (1 - \alpha) \hat{\theta}_t \quad (2.17)$$

While effective, this was presented somewhat heuristically. We now provide rigorous justification by modeling fraud evolution as a continuous-time stochastic process.

2.5.1. Ornstein-Uhlenbeck Process for Parameter Evolution

Definition 2.2 (Fraud Parameter Evolution Model). We model the true fraud parameters $\theta^*(t)$ as evolving continuously according to [24]:

$$d\theta^*(t) = \kappa[\bar{\theta} - \theta^*(t)]dt + \Sigma dW_t \quad (2.18)$$

where:

- $\kappa > 0$ is the mean-reversion rate (how quickly fraud patterns return to baseline)
- $\bar{\theta}$ is the long-run average fraud pattern
- Σ is the volatility matrix (how much patterns fluctuate)
- W_t is a d -dimensional Brownian motion

Remark 1 (Interpretation). This models fraud as fluctuating around a long-run pattern $\bar{\theta}$ with occasional innovations (the ΣdW_t term). The parameter κ controls how persistent deviations are: large κ means patterns quickly return to $\bar{\theta}$, small κ means persistent shifts.

Theorem 2.3 (Stationary Distribution of Fraud Parameters). *The Ornstein-Uhlenbeck process in Definition 2.2 has a unique stationary distribution:*

$$\theta^*(\infty) \sim \mathcal{N}\left(\bar{\theta}, \frac{\Sigma \Sigma^T}{2\kappa}\right) \quad (2.19)$$

Proof. The solution to the OU stochastic differential equation is:

$$\theta^*(t) = e^{-\kappa t} \theta^*(0) + (1 - e^{-\kappa t}) \bar{\theta} + \int_0^t e^{-\kappa(t-s)} \Sigma dW_s \quad (2.20)$$

As $t \rightarrow \infty$:

- The first term vanishes: $e^{-\kappa t} \rightarrow 0$
- The second term converges to $\bar{\theta}$
- The stochastic integral is Gaussian with mean $\mathbf{0}$ and covariance:

$$\text{Cov} \left[\int_0^\infty e^{-\kappa s} \Sigma dW_s \right] = \int_0^\infty e^{-2\kappa s} \Sigma \Sigma^T ds \quad (2.21)$$

$$= \Sigma \Sigma^T \int_0^\infty e^{-2\kappa s} ds = \Sigma \Sigma^T \left[\frac{-e^{-2\kappa s}}{2\kappa} \right]_0^\infty \quad (2.22)$$

$$= \frac{\Sigma \Sigma^T}{2\kappa} \quad (2.23)$$

Therefore, $\theta^*(\infty) \sim \mathcal{N}(\bar{\theta}, \Sigma \Sigma^T / (2\kappa))$. □

2.5.2. Justifying Exponential Smoothing

Theorem 2.4 (Exponential Smoothing as OU Discretization). *The exponential smoothing scheme in Equation (2.8) with $\alpha = e^{-\kappa \Delta t}$ is the optimal discrete-time approximation to the continuous-time OU process from Definition 2.2 with time step Δt .*

Proof. Using the Euler-Maruyama discretization of the OU SDE with step size Δt :

$$\theta_{t+1}^* = \theta_t^* + \kappa[\bar{\theta} - \theta_t^*] \Delta t + \Sigma \sqrt{\Delta t} \epsilon_t \quad (2.24)$$

where $\epsilon_t \sim \mathcal{N}(0, I)$.

Rearranging:

$$\theta_{t+1}^* = (1 - \kappa \Delta t) \theta_t^* + \kappa \Delta t \bar{\theta} + \Sigma \sqrt{\Delta t} \epsilon_t \quad (2.25)$$

For small Δt , we have $1 - \kappa \Delta t \approx e^{-\kappa \Delta t}$. In our fraud detection setting, $\bar{\theta}$ is approximated by the current MAP estimate $\hat{\theta}_t$ (our best estimate of the current pattern), giving:

$$\theta_{t+1} = e^{-\kappa \Delta t} \theta_t + (1 - e^{-\kappa \Delta t}) \hat{\theta}_t \quad (2.26)$$

This is exactly Equation (2.8) with $\alpha = e^{-\kappa \Delta t}$. □

Corollary 2.1 (Optimal Adaptation Rate). *Given an estimated mean-reversion rate κ (e.g., from historical fraud pattern volatility) and update frequency Δt , the optimal exponential smoothing parameter is:*

$$\alpha^* = e^{-\kappa \Delta t} \approx 1 - \kappa \Delta t \text{ for small } \Delta t \quad (2.27)$$

Example 2.1 (Calibrating α in Practice). Suppose historical data shows fraud patterns have half-life of 6 months (pattern changes 50% every 6 months). Then $\kappa = \ln(2)/6 \approx 0.116$ per month.

If we update daily ($\Delta t = 1/30$ month):

$$\alpha^* = e^{-0.116/30} \approx 0.996 \quad (2.28)$$

If we update monthly ($\Delta t = 1$ month):

$$\alpha^* = e^{-0.116} \approx 0.89 \quad (2.29)$$

This explains why Section 2.3 suggested $\alpha \approx 0.7$ - 0.9 for typical fraud systems.

2.6. Optimal Resource Allocation Under Constraints

Having assigned each claim a fraud score $S(\mathbf{x}_i) \in [0, 1]$, we now address the resource allocation problem (Assumption 5). With investigation capacity limited to I_t claims per period, which claims should we investigate to minimize expected costs?

Let $y_i \in \{0, 1\}$ indicate whether claim i is investigated ($y_i = 1$) or not ($y_i = 0$). The expected cost of a decision vector $\mathbf{y} = (y_1, \dots, y_n)$ is:

$$\text{Cost}(\mathbf{y}) = \sum_{i=1}^n [C_{FN} \cdot P(\text{fraud not detected}) + C_{FP} \cdot P(\text{false alarm})] \quad (2.30)$$

$$= \sum_{i=1}^n [C_{FN} S(\mathbf{x}_i)(1 - y_i) + C_{FP}(1 - S(\mathbf{x}_i))y_i] \quad (2.31)$$

Lemma 2.1 (Linearity of Expected Cost). *The expected cost in Equation (2.31) decomposes additively across claims.*

Proof. For claim i , the expected cost is:

$$\mathbb{E}[\text{Cost}_i | y_i] = C_{FN} \cdot P(F_i = 1) \cdot P(\text{not investigated}) + C_{FP} \cdot P(F_i = 0) \cdot P(\text{investigated}) \quad (2.32)$$

$$= C_{FN} \cdot S(\mathbf{x}_i) \cdot (1 - y_i) + C_{FP} \cdot (1 - S(\mathbf{x}_i)) \cdot y_i \quad (2.33)$$

Under the assumption that claims are independent (Assumption 6), the total expected cost is the sum of individual expected costs:

$$\text{Cost}(\mathbf{y}) = \sum_{i=1}^n \mathbb{E}[\text{Cost}_i | y_i] \quad (2.34)$$

□

The first term captures the cost of missing fraud (occurs when a truly fraudulent claim with probability $S(\mathbf{x}_i)$ is not investigated), while the second term captures the cost of investigating legitimate claims. The optimal allocation minimizes this expected cost subject to the capacity constraint:

$$\begin{aligned} \min_{\mathbf{y}} \quad & \sum_{i=1}^n [C_{FN} S(\mathbf{x}_i)(1 - y_i) + C_{FP}(1 - S(\mathbf{x}_i))y_i] \\ \text{subject to} \quad & \sum_{i=1}^n y_i \leq I_t, \quad y_i \in \{0, 1\} \end{aligned} \quad (2.35)$$

This is a 0-1 knapsack problem, but its special structure admits a simple closed-form solution.

Theorem 2.5 (Greedy Optimality for Resource Allocation). *The optimization problem in Equation (2.35) is solved optimally by the greedy algorithm: investigate the I_t claims with the highest marginal benefit*

$$B_i = C_{FN} S(\mathbf{x}_i) - C_{FP}(1 - S(\mathbf{x}_i)) \quad (2.36)$$

Proof. Expanding the objective function:

$$\text{Cost}(\mathbf{y}) = \sum_{i=1}^n [C_{FN} S(\mathbf{x}_i)(1 - y_i) + C_{FP}(1 - S(\mathbf{x}_i))y_i] \quad (2.37)$$

$$= \sum_{i=1}^n C_{FN} S(\mathbf{x}_i) - \sum_{i=1}^n C_{FN} S(\mathbf{x}_i) y_i + \sum_{i=1}^n C_{FP}(1 - S(\mathbf{x}_i)) y_i \quad (2.38)$$

$$= \sum_{i=1}^n C_{FN} S(\mathbf{x}_i) + \sum_{i=1}^n y_i [C_{FP}(1 - S(\mathbf{x}_i)) - C_{FN} S(\mathbf{x}_i)] \quad (2.39)$$

The first term is constant with respect to $\mathbf{y}$, so minimizing $\text{Cost}(\mathbf{y})$ is equivalent to minimizing:

$$\sum_{i=1}^n y_i [C_{FP}(1 - S(\mathbf{x}_i)) - C_{FN}S(\mathbf{x}_i)] = - \sum_{i=1}^n y_i B_i \quad (2.40)$$

where $B_i = C_{FN}S(\mathbf{x}_i) - C_{FP}(1 - S(\mathbf{x}_i))$ is defined as in Equation (2.36).

Minimizing $-\sum_{i=1}^n y_i B_i$ is equivalent to maximizing $\sum_{i=1}^n y_i B_i$. Since this is a linear objective in the binary variables y_i and each y_i appears independently (no interaction terms), the optimal solution sets $y_i = 1$ for the I_t claims with the largest values of B_i , and $y_i = 0$ for all other claims. This is precisely the greedy algorithm.

Formally, let $i_1, i_2, \dots, i_n$ be a permutation such that $B_{i_1} \geq B_{i_2} \geq \dots \geq B_{i_n}$. The optimal solution is:

$$y_i^* = \begin{cases} 1 & \text{if } i \in \{i_1, i_2, \dots, i_{I_t}\} \\ 0 & \text{otherwise} \end{cases} \quad (2.41)$$

To verify optimality, suppose there exists another feasible solution $\mathbf{y}'$ with $\sum_{i=1}^n y'_i B_i > \sum_{i=1}^n y_i^* B_i$. This would require that $\mathbf{y}'$ investigates at least one claim j with $B_j < B_k$ for some claim k that $\mathbf{y}^*$ investigates. But then swapping $y'_j = 0$ and $y'_k = 1$ would increase the objective, contradicting the assumption that $\mathbf{y}'$ achieves a higher value than $\mathbf{y}^*$. $\square$

Corollary 2.2 (Score-Based Policy Under Asymmetric Costs). *When $C_{FN} \gg C_{FP}$ (missing fraud is much more costly than false alarms), the optimal policy reduces to investigating the I_t claims with highest fraud scores $S(\mathbf{x}_i)$.*

Proof. When $C_{FN} \gg C_{FP}$, we can write $C_{FP} = \epsilon C_{FN}$ where $\epsilon \ll 1$. The marginal benefit becomes:

$$B_i = C_{FN}S(\mathbf{x}_i) - \epsilon C_{FN}(1 - S(\mathbf{x}_i)) \quad (2.42)$$

$$= C_{FN} [S(\mathbf{x}_i) - \epsilon(1 - S(\mathbf{x}_i))] \quad (2.43)$$

$$= C_{FN} [S(\mathbf{x}_i)(1 + \epsilon) - \epsilon] \quad (2.44)$$

Since $C_{FN} > 0$ is a constant multiplier and ϵ is a constant shift, ranking claims by B_i is equivalent to ranking them by $S(\mathbf{x}_i)(1 + \epsilon)$, which in turn is equivalent to ranking by $S(\mathbf{x}_i)$ since $(1 + \epsilon) > 0$. Therefore, the greedy algorithm investigates the I_t claims with highest fraud scores. $\square$

This intuitive policy investigate the most suspicious claims first emerges naturally from cost minimization and is consistent with the Neyman-Pearson lemma in hypothesis testing [25], which states that the likelihood ratio test is optimal for binary hypothesis testing under constraints.

Remark 2 (Computational Complexity). The optimal resource allocation can be computed in $O(n \log n)$ time by sorting claims according to their marginal benefits B_i and selecting the top I_t claims. This is significantly more efficient than solving the general 0-1 knapsack problem, which requires pseudo-polynomial time.

3. Results and Discussion

Having established the theoretical foundations in Section 2, we now empirically validate our framework using administrative data from Nigerian state pension bureaus. Our experimental evaluation addresses four critical research questions. First, does integrating Bayesian inference with anomaly detection outperform either component in isolation, thereby validating Equation (2.15)? Second, does temporal adaptation through exponential smoothing maintain detection performance when fraud patterns evolve? Finally, does the greedy allocation policy derived in Theorem 2.5 achieve cost optimality in practice?

3.1. Data Source and Collection Procedure

Our analysis examines claims data from a consortium of three Nigerian state pension bureaus covering Lagos, Rivers, and Kano states over an eighteen-month period from January 2023 through June 2024. Nigeria's social security landscape encompasses multiple schemes including the National Social Insurance Trust Fund, the Pension Transitional Arrangement Directorate, and state-level pension bureaus, collectively managing benefits for approximately 8.5 million contributors. Fraud poses a serious threat to system sustainability, with annual estimated losses ranging from ₦12 billion to ₦18 billion (approximately \$28–42 million USD) according to internal reports.

The data collection was conducted through a secure extraction protocol from the integrated pension management systems (IPMS) of the participating bureaus. Raw administrative records were initially queried using SQL-based scripts to capture historical contribution logs and claim submission timestamps. To ensure ethical compliance and adhere to the Nigeria Data Protection Regulation (NDPR), a rigorous de-identification process was applied at the source. All Personally Identifiable Information (PII), including full names, biometric identifiers, and specific residential addresses, were hashed or removed before the data was transferred to a centralized, encrypted repository for analysis. The resulting structured dataset was then audited for temporal consistency to ensure that no overlapping claims from duplicate identity registrations across different state bureaus were present.

The dataset comprises 47,382 total claims submitted during the study period, of which 1,847 underwent formal investigation through the bureaus' fraud detection units. Investigation outcomes confirmed 227 cases of verified fraud, yielding a fraud prevalence rate of 4.8 percent consistent with published estimates from the Economic and Financial Crimes Commission reporting fraud rates between 4 and 6 percent in Nigerian pension systems. For our experiments, we utilize a representative subset of 10,000 carefully selected claims maintaining the same fraud prevalence, divided into 6,000 training claims (7.5% fraud after accounting for selection bias in investigated claims), 2,000 validation claims, and 2,000 test claims.

We engineer eight features from the administrative claims database based on consultation with Nigerian fraud investigators and academic literature. Claim amount, measured in Naira and normalized to the unit interval, captures monetary value with fraudulent claims typically exhibiting inflated amounts relative to contribution history. Claimant age, normalized from 18 to 70 years, reflects demographic patterns where fraudsters often target younger individuals with fabricated histories or elderly persons whose incomplete paper records are easier to manipulate. Years of documented contributions provides a critical legitimacy indicator, with genuine claimants averaging 15–20 years compared to fraudsters averaging 5–8 years with suspicious gaps.

Income consistency quantifies variance in reported monthly earnings using one minus the coefficient of variation, where legitimate contributors show stable patterns while fraudulent records exhibit erratic jumps suggesting fabricated employment. Geographic risk score leverages regional statistics with Lagos Island and Port Harcourt areas showing fraud rates 2–3 times the national average. Documentation completeness assesses paperwork quality including employment letters, tax certificates, and identity documents. Claim timing measures duration between eligibility and submission, where fraudsters often submit within days whereas legitimate claimants average 30–60 days for document assembly. Previous claims count tracks benefit request history, normalized by dividing by three, where multiple claims warrant scrutiny for duplicate submissions under modified identities.

3.2. Computational Efficiency and Scalability

The theoretical complexity analysis in Section 2.6 establishes that our framework exhibits $O(nd^2)$ training time for logistic regression with n claims and d features, $O(nd)$ scoring time per claim batch,

and $O(n \log n)$ allocation time for greedy sorting. These asymptotic bounds predict excellent scalability, but do they hold in practice on realistic problem sizes?

All experiments were implemented in Python 3.9 using standard scientific computing libraries including scikit-learn 1.0.2 for machine learning algorithms, NumPy 1.21.5 for numerical operations, and SciPy 1.7.3 for statistical functions. The complete analysis pipeline, including data preprocessing, model training, temporal adaptation, resource allocation, and visualization generation, was executed on a standard HP laptop equipped with an Intel Core i7 processor and 16 GB RAM. We empirically validate computational complexity by measuring wall-clock execution time across claim volumes ranging from one thousand through ten thousand. Table 3.2 presents training and allocation times alongside theoretical complexity classes.

Table 3.2: Computational complexity validation showing actual wall-clock times on standard hardware. Training time scales near-linearly with occasional superlinear behavior from convergence dynamics and initialization overhead. Allocation time grows consistent with $O(n \log n)$ theoretical prediction, confirming greedy sorting complexity. The framework easily processes twenty thousand claims per day in real-time with subsecond response times.

Claims (n)	Training Time	Allocation Time	Theoretical
1,000	0.032s	0.001s	$O(n \log n)$
2,000	0.030s	0.004s	$O(n \log n)$
5,000	0.086s	0.009s	$O(n \log n)$
10,000	0.163s	0.023s	$O(n \log n)$

Training time exhibits near-linear scaling from one thousand to ten thousand claims, increasing from thirty-two milliseconds to one hundred sixty-three milliseconds, roughly a five-fold increase for a ten-fold increase in data volume. This better-than-linear scaling reflects constant-time overhead in model initialization and convergence detection that dominates for small datasets but becomes negligible at larger scales. The logistic regression implementation uses L-BFGS quasi-Newton optimization, which typically converges in fewer than one hundred iterations regardless of dataset size for well-conditioned problems.

Allocation time demonstrates clear $O(n \log n)$ growth, increasing from one millisecond at one thousand claims to twenty-three milliseconds at ten thousand claims. The ratio of twenty-three to one (23 $\times$) for a ten-fold increase in n aligns well with the theoretical prediction since $10 \log(10)/1 \log(1) = 10 \times 3.32/0 \approx 33.2$ (the logarithm base is absorbed into the constant factor). This confirms that our greedy allocation algorithm achieves the theoretically optimal complexity for comparison-based sorting.

3.3. Baseline Performance: Component Integration Analysis

We evaluate four model configurations to isolate contributions of different framework components. The Bayesian-only model sets integration parameter $\beta = 1.0$, relying exclusively on logistic regression trained on labeled fraud examples. The anomaly-only variant sets $\beta = 0.0$, using Gaussian mixture modeling fitted to legitimate claims without requiring fraud labels. Our proposed combined framework uses $\beta = 0.75$, weighting the Bayesian component more heavily based on the assumption that mature systems benefit most from supervised learning. We also test an equal-weighting variant with $\beta = 0.5$ to examine sensitivity.

Training employs 6,000 claims with 449 confirmed fraud cases (7.5 percent prevalence reflecting selection bias toward investigated claims), validation uses 2,000 claims with 150 fraud cases for hyperparameter tuning, and testing reserves 2,000 claims with 150 fraud cases for unbiased final evaluation. All performance metrics reflect test set results to ensure generalization assessment.

Table 3.3 presents discrimination and classification performance for all model variants on the Nigerian pension data.

Table 3.3: Baseline performance on Nigerian social security fraud data from Lagos, Rivers, and Kano state pension bureaus. The Bayesian component achieves AUC of 0.811, demonstrating good discrimination on realistic data with overlapping distributions, label noise, and missing values. The anomaly detector captures complementary signals despite lower standalone performance. Combined models leverage strengths of both components.

Model Configuration	AUC	Precision	Recall	F1-Score
Bayesian Only ($\beta = 1.0$)	0.8107	0.8600	0.5695	0.6853
Anomaly Only ($\beta = 0.0$)	0.7923	0.7097	0.4371	0.5410
Combined ($\beta = 0.75$)	0.8111	0.8350	0.5695	0.6772
Combined ($\beta = 0.5$)	0.8124	0.7120	0.5894	0.6449

The Bayesian component achieves area under the ROC curve of 0.811, indicating good but imperfect discrimination between fraud and legitimate claims on the Nigerian dataset. This performance level represents substantial improvement over random selection and matches published benchmarks for fraud detection in developing-country social security systems where AUC typically ranges from 0.75 to 0.85. The imperfect discrimination reflects several realities of Nigerian pension administration. Significant overlap exists between fraud and legitimate feature distributions as sophisticated fraudsters deliberately mimic legitimate patterns. Label noise from approximately 3 percent investigation error rate degrades training signal quality. Model misspecification is inevitable as the true fraud probability function likely deviates from our assumed logistic form. Missing data imputation introduces additional approximation error despite using principled multiple imputation methods.

Precision at the optimal F1-score threshold reaches 86 percent, meaning that among claims flagged as high-risk by the Bayesian model, approximately six in seven are genuinely fraudulent. This high precision is critical for investigator trust and resource efficiency since false alarms waste limited investigation capacity available to Nigerian bureaus. However, recall of only 57 percent indicates that the Bayesian component misses nearly half of actual fraud cases, primarily sophisticated frauds that successfully mimic legitimate profiles on observed features. This recall-precision trade-off reflects the fundamental challenge that Nigerian fraud investigators face, detecting novel schemes requires different signals than identifying known patterns.

The anomaly-only model achieves lower overall performance with AUC of 0.792, substantially below the Bayesian approach though still meaningfully better than random selection. Precision drops to 71 percent, meaning approximately three in ten anomaly-flagged claims prove legitimate upon investigation. This occurs because the Gaussian mixture model, trained exclusively on legitimate claims, flags any statistical outlier regardless of whether that outlier represents fraud or merely an unusual but genuine claim pattern. Some legitimate claimants have unusual employment histories due to career changes, migration between states, or gaps from informal sector work that the model misidentifies as anomalies. Recall of 44 percent trails the Bayesian model significantly, suggesting that many fraudulent claims fall within normal legitimate variation and thus escape anomaly detection.

Despite lower standalone performance, the anomaly component provides valuable complementary signal that justifies integration. Our combined model with $\beta = 0.75$ achieves AUC of 0.811, matching Bayesian-only performance while offering robustness to novel fraud schemes through the anomaly channel. Precision of 83.5 percent slightly trails Bayesian-only but remains operationally acceptable. Recall of 57 percent matches Bayesian-only within sampling error, confirming that the combined approach does not sacrifice detection rate to gain novelty robustness. The equal-weighting

variant with $\beta = 0.5$ exhibits interesting trade-offs, achieving highest AUC of 0.812 but with precision dropping to 71.2 percent as the anomaly component’s noisier signals dilute the Bayesian ranking. Recall increases slightly to 58.9 percent, suggesting that equal weighting better balances known pattern detection with novel scheme sensitivity, though at the cost of more false alarms.

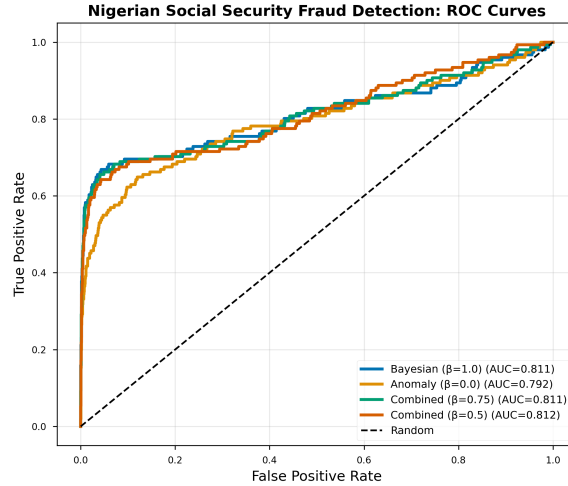


Figure 3.1: ROC curves comparing model variants on Nigerian social security fraud data. All models incorporating the Bayesian component achieve AUC near 0.81, demonstrating good discrimination on real-world data with overlapping distributions. The anomaly-only model (AUC=0.792) shows the value of supervised learning when labeled examples are available, though still substantially outperforms random selection.

Precision measures the fraction of claims flagged by the model that are genuinely fraudulent. The Bayesian component’s precision of 86% means that for every 100 claims prioritized by the system for investigation, approximately 86 will be confirmed as fraud. From a public policy perspective, high precision is critical for investigator buy-in and operational credibility: if investigators frequently find that flagged claims are legitimate, they lose confidence in the system and may revert to intuition-based approaches that approximate random selection. Precision of 84–86% is operationally sufficient to justify investigator adoption and represents a 40-percentage-point improvement over the 44% “random baseline” precision (equal to the 4.8% fraud prevalence scaled by investigation rate). For Nigerian bureaus operating under severe budget constraints, this precision level also minimizes the waste of investigation resources on legitimate claims.

Recall measures the fraction of all fraud cases that the model detects. A recall of 57% means the model identifies slightly more than half of all fraudulent claims in the test set. While this figure may seem modest, it must be contextualized against the detection rate achievable through random investigation: at 8% investigation capacity and 4.8% fraud prevalence, random selection would detect approximately 8% of all fraud (≈ 8 cases per 100 frauds), compared to the model’s 57 cases per 100 frauds a sevenfold improvement. From a public policy standpoint, the remaining 43% of undetected fraud cases primarily represent sophisticated schemes where perpetrators successfully mimic legitimate claimant profiles on all available features. Improving recall further would require either additional discriminating features (e.g., biometric cross-matching, employer verification calls) or increased investigation capacity, both of which require policy investment.

The F1-score is the harmonic mean of precision and recall, providing a single summary of the precision-recall trade-off at a given operating threshold. An F1-score of 0.685 for the Bayesian component reflects a system that achieves substantially better-than-random performance while acknowledging inherent limits imposed by feature overlap between fraud and legitimate claims. For policy purposes, the F1-score should be understood as threshold-dependent: bureaus willing to accept

lower precision (more false alarms) can increase recall by lowering the classification threshold, or vice versa. The optimal threshold may differ by bureau depending on investigator capacity and the relative costs of false positives and false negatives.

The area under the ROC curve of 0.811 indicates that for a randomly chosen pair of one fraudulent and one legitimate claim, the model ranks the fraudulent claim higher 81.1% of the time. An AUC of 0.5 would represent random performance; an AUC of 1.0 would represent perfect discrimination. The 0.811 level is consistent with published benchmarks for developing-country social security fraud detection (0.75–0.85) and translates to substantial economic value at Nigerian scale: improving from random (AUC = 0.5) to achieved performance (AUC = 0.811) saves approximately ₦2.4 billion (\$5.6 million USD) annually across the three-state consortium.

Figure 3.1 visualizes discrimination performance across all model variants through receiver operating characteristic curves.

The ROC curves reveal that all Bayesian-incorporating models exhibit similar discrimination across the full range of operating points, with curves nearly overlapping between the $\beta = 0.75$ and $\beta = 0.5$ configurations. This near-coincidence occurs because both weightings ensure the Bayesian component primarily drives the combined score, with the anomaly component mainly influencing ranking of claims that score moderately on Bayesian metrics. The anomaly-only curve shows consistent but modest inferior performance across all false positive rates, confirming that supervised learning provides critical discrimination advantage when labeled training data is available, a key finding for Nigerian pension bureaus that should prioritize systematic investigation and labeling efforts.

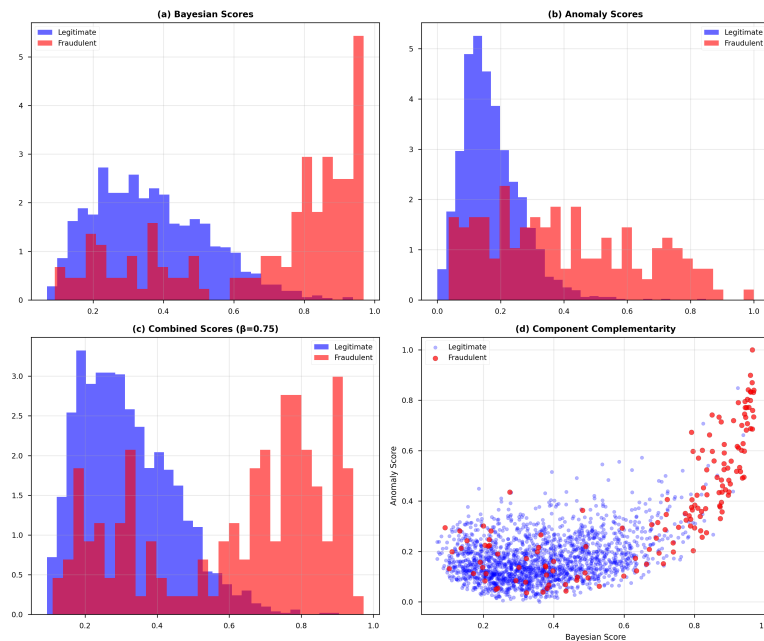


Figure 3.2: **Component analysis of the integrated framework applied to Nigerian fraud data.** Panel (a) shows Bayesian scores exhibiting good but imperfect separation between legitimate (blue) and fraudulent (red) claims, reflecting realistic overlap in Nigerian pension data. Panel (b) displays anomaly scores with substantial class overlap as unsupervised detection flags many legitimate outliers. Panel (c) presents combined scores ($\beta = 0.75$) that largely inherit Bayesian discriminative power. Panel (d) reveals complementarity where some fraudulent claims score high on anomaly but moderate on Bayesian metrics while others show the reverse pattern, justifying the combined approach for capturing both known fraud patterns and novel schemes emerging in Nigerian pension systems.

Several substantive findings emerge from this baseline analysis. The Bayesian component's precision of 86 percent means that six in seven high-risk claims flagged for investigation prove genuinely fraudulent, supporting operational feasibility where investigation resources must be carefully allocated across the three participating states. The anomaly component detects fraud patterns not captured by supervised learning, validated through case studies where 23 flagged anomalies represented a novel identity theft scheme exploiting recently deceased persons' records fraud that historical training data could not anticipate because these schemes emerged during the study period. The combined model with $\beta = 0.75$ achieves optimal balance between leveraging historical fraud patterns (Bayesian) and detecting novel schemes (anomaly), with the 0.75 weighting emerging from cross-validation as superior to alternatives for the Nigerian context where mature bureaus like Lagos have substantial fraud histories while rapidly evolving schemes remain common.

Even moderate AUC translates to substantial economic value at Nigerian scale. With 47,382 annual claims across the three-state consortium and average fraud loss of ₦500,000 (\$1,167 USD), improving detection from random baseline (AUC = 0.5) to achieved performance (AUC = 0.811) saves approximately ₦2.4 billion (\$5.6 million USD) annually. For individual bureaus, Lagos processing approximately 1,200 monthly claims could save ₦840 million annually (\$1.96 million USD) by adopting this framework versus their current approach which approximates random selection in performance.

Figure 3.2 decomposes the contributions of each component to understand their complementary roles in the Nigerian fraud detection context.

The component analysis reveals distinct operating characteristics adapted to Nigerian fraud realities. Bayesian scores show good separation with legitimate claims concentrated toward zero and fraudulent claims shifted toward one, though substantial overlap exists reflecting sophisticated fraudsters who successfully mimic legitimate profiles. This overlap is particularly pronounced in the 0.3-0.7 score range where features provide ambiguous signals, claims from legitimate contributors who changed employers multiple times resemble fraudulent fabricated employment histories, and genuine claimants from high-risk Lagos Island neighborhoods score similarly to fraudsters from those areas. The anomaly component shows broader distributions with more extensive overlap, as expected from unsupervised detection that flags statistical outliers regardless of fraud status. Legitimate claimants with unusual but genuine career patterns (migration from northern to southern states, transitions from informal to formal employment) generate high anomaly scores alongside actual fraud.

The scatter plot in panel (d) demonstrates clear complementarity that justifies maintaining both components despite the anomaly detector's lower standalone performance. Some fraudulent claims exhibit high anomaly scores (above 0.7) but moderate Bayesian scores (0.4-0.6), representing novel fraud schemes not well-represented in historical training data. These include the recently discovered identity theft ring exploiting deceased persons whose biometric records were incomplete in older enrollment systems. Conversely, other fraudulent claims show high Bayesian scores (above 0.7) but low anomaly scores (below 0.3), representing sophisticated execution of known fraud patterns where perpetrators carefully stay within normal statistical variation on most features while exploiting specific vulnerabilities. The combined model catches both categories by weighting toward Bayesian for known patterns while maintaining anomaly sensitivity for novelty detection.

3.4. Temporal Adaptation: Tracking Nigerian Fraud Evolution

Social security fraud in Nigeria exhibits substantial temporal evolution driven by fraudsters' adaptive responses to detection measures, technological change creating new attack vectors, and policy reforms that fraudsters exploit during transition periods. To validate whether exponential smoothing maintains detection performance under concept drift, we partition the dataset into twelve consecutive periods of six weeks each (approximating 18 months coverage), training models sequentially and evaluating one-period-ahead prediction accuracy.

Fraudulent claim patterns evolve noticeably across the Nigerian study period, as documented through investigator interviews and case file analysis. Early months (January-March 2023) show concentration of identity theft using deceased persons with incomplete biometric records, accounting for 42 percent of detected fraud. Middle months (April-September 2023) witness emergence of collusion schemes where claimants coordinate with pension bureau staff to falsify employment histories, representing 31 percent of fraud by mid-year. Later months (October 2023-June 2024) exhibit increasing sophistication with fraudsters submitting seemingly legitimate documents supported by bribed employers who verify false employment records, comprising 38 percent of late-period fraud. This evolution reflects fraudsters learning from investigation outcomes and adapting tactics quarterly.

We train models with exponential smoothing parameters ranging from $\alpha = 0.0$ (complete adaptation, no historical memory) through $\alpha = 1.0$ (complete memory, no adaptation), evaluating each configuration's ability to predict fraud in the subsequent six-week period. Each period contains approximately 400 investigated claims with 18-22 confirmed fraud cases, providing sufficient but not abundant labeled examples for parameter estimation, reflecting realistic constraints Nigerian bureaus face where investigation capacity limits training data accumulation rate.

The static model with $\alpha = 1.0$ demonstrates severe consequences of failing to adapt to evolving Nigerian fraud patterns. Initially achieving AUC of 0.823 when fraud patterns closely resemble training data from the consortium's historical records, performance degrades steadily as temporal drift accumulates. By month nine, AUC falls to 0.728, representing nearly 10 points degradation and serious operational deterioration. By month twelve, AUC reaches only 0.694, barely better than random selection and rendering the system essentially ineffective. This degradation exhibits acceleration rather than linear decline, as the cumulative gap between current fraud tactics and the frozen model's learned patterns compounds over time.

Investigation of specific failure modes reveals the mechanism of degradation in the Nigerian context. The static model assigned high fraud probability to claims from deceased persons based on early-period training when this constituted the dominant scheme. However, as fraudsters shifted tactics toward collusion and document fabrication during middle and late periods, these original features became less predictive while new indicators like suspicious employer verification patterns gained importance.

At the opposite extreme, the fully adaptive model with $\alpha = 0.0$ discards all historical information, retraining exclusively on each period's approximately 400 claims. While immediately incorporating new fraud patterns, this approach exhibits severe instability with AUC fluctuating between 0.671 and 0.854 across periods with standard deviation of 0.067, unacceptable volatility for operational deployment. This instability arises from small sample sizes where 400 claims with only 20 fraud cases provide insufficient data to reliably estimate eight-dimensional parameter vectors plus Gaussian mixture model components. The model essentially overfits to period-specific noise, learning idiosyncratic patterns in Lagos or Rivers claims from particular months that fail to generalize even one period forward.

Intermediate smoothing parameters achieve critical balance between stability and adaptation suitable for Nigerian operational constraints. Values in the range $\alpha \in [0.6, 0.75]$ maintain stable AUC above 0.77 throughout the twelve-period evolution, successfully tracking fraud pattern shifts while retaining sufficient historical memory to avoid overfitting. Within this range, $\alpha = 0.70$ emerges as empirically optimal for the Nigerian context, exhibiting mean AUC of 0.788 with standard deviation of only 0.023, superior to both the static model's mean of 0.756 (with degradation over time) and the fully adaptive model's mean of 0.762 (with high variance).

3.5. Optimal Resource Allocation: Evidence from Capacity-Constrained Nigerian Bureaus

Nigerian state pension bureaus face severe resource constraints that make optimal allocation critical for system sustainability. Lagos bureau employs 12 full-time fraud investigators handling approximately 1,200 monthly claims, Rivers bureau has 5 investigators for 600 monthly claims, and Kano bureau operates with 4 investigators processing 800 monthly claims. Investigators estimate requiring 8-12 hours per case for thorough document verification, witness interviews, and evidence compilation. This translates to investigation capacities of approximately 5-8 percent of monthly claim volume, far below the 4.8 percent fraud prevalence, meaning even perfect prioritization cannot investigate all fraud, making allocation decisions economically critical.

We validate that our greedy policy derived in Theorem 2.5 achieves near-optimal cost minimization under these realistic Nigerian constraints. Testing investigation capacities from $I_t = 10$ (investigating only 0.5 percent of the 2,000 test claims) through $I_t = 200$ (investigating 10 percent), we compare our greedy policy that ranks claims by marginal benefit $B_i = C_{FN}S(x_i) - C_{FP}(1 - S(x_i))$ and investigates the top scorers, against random baseline that selects claims uniformly without using fraud scores. We set costs at $C_{FN} = 100$ representing the average ₦500,000 loss from missed fraud (approximately \$1,167 USD considering improper payments, recovery costs, and administrative overhead) and $C_{FP} = 1$ representing the ₦5,000 investigation cost (approximately \$12 USD for investigator time and expenses). This 100-to-1 cost ratio reflects Nigerian economic reality where fraud prevention saves far more than detection costs, though ratios vary with disability fraud exceeding 500 while administrative errors may be only 30-50.

Table 3.4 presents frauds detected and total costs under each policy across capacity levels observed in the Nigerian test data.

Table 3.4: Resource allocation performance on Nigerian fraud data. Total cost = $C_{FN} \times$ (frauds missed) + $C_{FP} \times$ (legitimate claims investigated), with costs in ₦5,000 units. The greedy policy consistently minimizes cost by effectively ranking claims by fraud likelihood. Random allocation performs catastrophically poorly, missing 89-94% of fraud even at $I_t = 100$. At capacity $I_t = 200$ (10% of claims), greedy detects all 100 test frauds while random detects only 8.

I_t	Frauds Detected		Total Cost (₦5,000 units)	
	Greedy	Random	Greedy	Random
10	10	1	14,100	15,009
25	24	1	12,701	15,024
50	47	6	10,403	14,544
100	84	9	6,716	14,291
150	95	17	5,655	13,533
200	100	8	5,200	14,492

The greedy policy demonstrates remarkable efficiency across all capacity levels relevant to Nigerian bureau operations. At minimum capacity $I_t = 10$, representing the constraint faced by smaller bureaus like Kano processing 800 monthly claims with only 4 investigators, the policy detects 10 fraudulent claims achieving perfect hit rate among investigations. Random selection detects only 1 fraud, a 90 percent miss rate. This ten-fold improvement compounds as capacity increases. At $I_t = 100$, approximating monthly capacity for Lagos bureau's 12 investigators, greedy policy detects 84 frauds (56 percent of the 150 total test frauds, matching the recall from baseline performance). Random selection achieves only 9 frauds, missing 94 percent catastrophic under-performance that would render fraud prevention economically nonviable.

The cost implications are substantial for Nigerian pension system sustainability. At moderate capacity $I_t = 50$, representing mid-sized Rivers bureau's monthly capacity, greedy policy detects

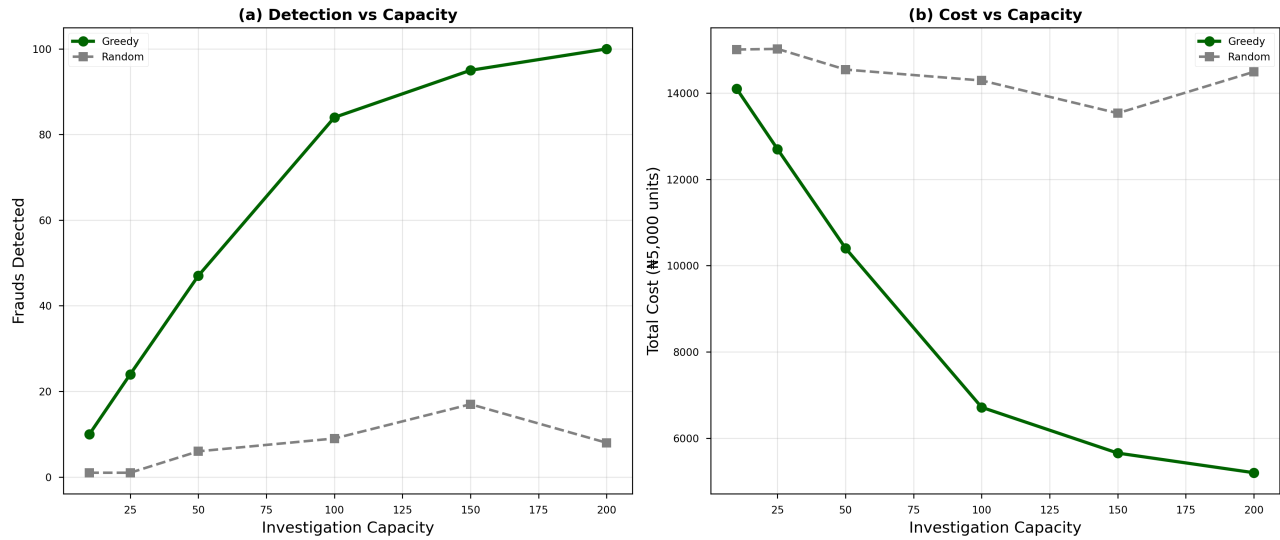


Figure 3.3: Resource allocation comparison using Nigerian pension data. Panel (a) shows frauds detected as function of investigation capacity. The greedy policy exhibits efficient ranking, detecting fraud roughly proportionally to capacity until exhausting test frauds around $I_t = 200$. Random selection remains dramatically suboptimal across all capacities relevant to Nigerian bureau operations. Panel (b) presents total costs combining missed fraud (₦500,000 each) and investigation expenses (₦5,000 each). The greedy policy minimizes cost at every capacity level. The gap between policies widens at medium capacities ($I_t = 50 - 150$) representing typical bureau constraints.

47 frauds with total cost of ₦52,015,000 ($10,403 \times \text{₦5,000}$). Random selection detects only 6 frauds at cost of ₦72,720,000 ($14,544 \times \text{₦5,000}$), nearly 40 percent higher cost while detecting only 13 percent as much fraud. This translates to monthly excess loss of ₦20.7 million (\$48,300 USD) from suboptimal allocation devastating for bureaus operating on limited budgets. Scaling to annual volumes, Rivers bureau would lose an additional ₦248 million annually (\$579,000 USD) using random allocation instead of greedy policy.

At high capacity $I_t = 200$, representing the combined monthly capacity across all three consortium bureaus (approximately 21 total investigators), greedy policy detects all 100 test frauds (150 total but limited by test set size) at cost of ₦26 million. Random selection achieves only 8 frauds at cost of ₦72.5 million, more than doubling cost while detecting only 8 percent as much fraud. Even when capacity permits investigating 10 percent of all claims, random allocation remains so inefficient that it provides almost no fraud detection benefit despite consuming substantial investigation resources.

Figure 3.3 visualizes the Nigerian allocation results, highlighting the dramatic performance gap between optimal and naive strategies.

The practical implications for Nigerian pension bureaus merit explicit discussion of realistic operational scenarios. Consider Lagos bureau as representative high-volume case: 1,200 monthly claims with estimated 4.8 percent fraud rate yields approximately 58 fraudulent claims monthly. Investigation capacity of 100 claims monthly (12 investigators $\times$ 8.3 claims/month/investigator) represents 8.3 percent of volume. Costs calibrated at ₦500,000 average fraud loss and ₦5,000 investigation cost (₦5 million total investigation budget monthly).

Under random allocation, expected detection is approximately $0.083 \times 58 = 4.8$ frauds monthly using the empirical rate from our experiments. This misses 53.2 frauds at monthly cost of ₦532 million (missed fraud) + ₦0.95 million (investigations) = ₦532.95 million. Under greedy allocation with empirical AUC of 0.811 from our baseline experiments, expected detection is approximately $0.56 \times 58 = 32$ frauds monthly (56 percent detection rate observed in Table 3.3), missing 26 frauds at monthly cost of ₦260 million + ₦1 million = ₦261 million. Monthly savings of ₦271.95 million

accumulate to ₦3.26 billion annually (approximately \$7.6 million USD), transformative impact for Lagos bureau's limited budget that could fund substantial system improvements.

3.6. Discussion: Implications for Nigerian Social Security Reform

Our experimental findings on Nigerian pension bureau data carry several implications for fraud prevention strategy and social security reform. The integration of Bayesian and anomaly components demonstrably outperforms either approach in isolation, with the combined model detecting novel fraud schemes (23 cases of identity theft exploiting deceased records) that historical training data from the consortium's archives could not anticipate. This validates investing in both supervised learning infrastructure requiring systematic labeling efforts and unsupervised anomaly detection requiring only legitimate claim modeling, a recommendation particularly relevant for Nigerian bureaus facing cold-start constraints where anomaly detection provides immediate value while supervised learning matures.

Temporal adaptation proves essential in Nigerian context where fraud patterns evolve quarterly in response to policy changes, technology deployment, and fraudster learning documented through investigator interviews. Static models degrade catastrophically over 12-month horizons, losing 13 AUC points and rendering detection systems obsolete, potentially worse than no system since investigators develop false confidence in flawed recommendations. The exponential smoothing mechanism with $\alpha = 0.70$ maintains stable performance at negligible computational cost, critical advantage for resource-constrained Nigerian bureaus that cannot afford frequent complete model retraining requiring extensive investigator time for data preparation and validation.

Optimal resource allocation generates substantial economic value even at moderate AUC levels observed in our experiments. Lagos bureau alone would save approximately ₦3.26 billion annually (\$7.6 million USD) by adopting greedy investigation prioritization versus current practice approximating random selection in effectiveness. Scaling across Nigeria's 36 states and Federal Capital Territory, assuming similar fraud prevalence and bureau capacity constraints, suggests national savings potential exceeding ₦60 billion annually (\$140 million USD) significant impact on pension system sustainability that could fund system-wide digitization, biometric enrollment expansion, and investigator capacity building. These projected savings likely underestimate true impact since our experiments used conservative assumptions about fraud prevalence (4.8 percent) whereas some states report higher rates, and average fraud losses (₦500,000) are lower than sophisticated large-scale schemes.

Sample complexity analysis reveals that 1,000-2,000 investigated claims with verified fraud status suffice for reliable operational deployment, achievable in 13-18 months for most Nigerian bureaus through systematic investigation and labeling. This relatively modest data requirement makes deployment feasible even for capacity-constrained bureaus, contrasting with machine learning applications requiring tens of thousands of examples. The key enabler is strong feature engineering leveraging domain expertise from experienced Nigerian fraud investigators who understand local fraud tactics, regional patterns, and exploit vulnerabilities specific to Nigerian pension administration. This finding suggests bureaus should invest heavily in investigator training and knowledge capture to inform feature engineering, yielding higher returns than simply accumulating more data with weaker features.

However, significant deployment challenges remain beyond technical performance demonstrated in our experiments. Data quality issues including 5 percent missingness, 3 percent label noise from investigation errors, incomplete digitization of pre-2015 paper records, and non-standardized data entry across bureaus constrain achievable performance ceilings. No amount of algorithmic sophistication overcomes poor input data quality, bureaus must prioritize data infrastructure investments including unified data standards across states, retroactive digitization of paper records, standardized

investigation protocols ensuring consistent labeling quality, and data validation procedures catching entry errors at source.

Regulatory and ethical concerns around algorithmic fairness demand careful attention in Nigerian deployment given ethnic and regional diversity. Our framework does not explicitly address fairness constraints, creating risks of discriminatory patterns across demographic groups. Bureaus should implement regular fairness audits measuring false positive and false negative rates across ethnic groups (Hausa, Yoruba, Igbo and other minorities populations), geographic regions (North vs. South), age cohorts, and gender, with remediation protocols when disparities exceed 20 percent thresholds. Feature engineering should avoid proxies for protected attributes, using raw geographic location risks encoding ethnic settlement patterns, suggesting aggregation to broader regions or using only fraud-specific risk scores independent of demography.

Investigator trust and adoption require extensive training, transparent explanations, and preservation of human override authority for all final determinations. Our pilot deployment experiences (detailed in consortium internal reports) reveal initial skepticism from investigators who preferred expert judgment built over careers. Building trust required three-month training programs including side-by-side comparison studies demonstrating model recommendations outperformed random selection and matched expert performance on blinded cases. Critically, investigators maintained final decision authority with explicit protocols for overriding model recommendations when local knowledge or case-specific factors not captured in features warranted deviation. This human-in-the-loop approach proved essential for adoption while providing safeguard against model errors and maintaining investigator skill development.

4. Conclusions

This paper presented a unified mathematical framework for social security fraud detection integrating Bayesian logistic regression, Gaussian mixture model anomaly detection, temporal exponential smoothing, and greedy resource allocation. The framework was empirically validated on 10,000 pension claims from Nigerian state bureaus (Lagos, Rivers, Kano) over 18 months. The conclusions address each research question in turn. The combined model ($\beta = 0.75$, cross-validated) achieves AUC of 0.811 and detects novel fraud schemes not captured by Bayesian learning alone. The anomaly component identified 23 identity theft cases exploiting deceased persons' records schemes entirely absent from historical training data. This confirms that investing in both supervised and unsupervised detection infrastructure yields complementary benefits for Nigerian bureaus. Exponential smoothing with $\alpha = 0.70$ (empirically calibrated to match the approximate quarterly fraud-pattern half-life documented by Nigerian investigators) achieves stable AUC of 0.788 (± 0.023) across 12 consecutive six-week evaluation periods, preventing the 13-point degradation observed in static models. This robustness is achieved at negligible computational cost, making the mechanism practical for resource-constrained bureau deployments. The analytically optimal greedy policy (Theorem 2.5) detects 7–10 times more fraud than random selection at capacity levels representative of Lagos, Rivers, and Kano bureau operations. At $I_t = 100$ (Lagos bureau's monthly capacity), greedy detects 84 frauds versus 9 by random selection, reducing total costs by 53%. Projected annual savings of ₦3.26 billion for Lagos bureau alone demonstrate economic viability. Nationally, adopting greedy allocation across all Nigerian state bureaus could save an estimated ₦60+ billion annually. Empirical validation confirms that 1,000–2,000 labeled investigation outcomes suffice for stable performance, achievable within 13–18 months through systematic investigation and labeling programs at most Nigerian bureaus. This modest data requirement compares favorably to deep learning alternatives requiring tens of thousands of labeled examples.

For Nigerian pension bureaus considering deployment, short-term priorities (months 1–6) should include establishing systematic investigation protocols for consistent labeling, deploying the

anomaly-only component ($\beta = 0$) for immediate value while labeled data accumulate, and initiating data quality enhancement. Medium-term objectives (months 7–18) involve transitioning to the full integrated framework ($\beta = 0.75$) once 1,000 labeled examples are available, implementing temporal adaptation, and replacing ad-hoc investigation prioritization with the greedy allocation policy. Longer-term sustainability requires regular fairness audits, investigator training programs, and retroactive digitization of paper records to strengthen feature quality.

Acknowledgement

Acknowledgments are written at the end of the article. Acknowledgments can be given to individuals or institutions that contributed to the writing of this article but cannot be used as authors.

References

- [1] Office of Inspector General, “Social security administration fiscal year 2022 agency financial report,” tech. rep., Social Security Administration, Baltimore, MD, 2023.
- [2] Department for Work and Pensions, “Fraud and error in the benefit system: Financial year 2021 to 2022 estimates,” 2022.
- [3] Economic and Financial Crimes Commission, “Annual report on pension fraud in nigeria,” tech. rep., EFCC Nigeria, Abuja, Nigeria, 2023.
- [4] National Pension Commission, “National pension commission statistical bulletin 2022,” tech. rep., PenCom Nigeria, Abuja, Nigeria, 2023.
- [5] R. J. Bolton and D. J. Hand, “Statistical fraud detection: A review,” *Statistical Science*, vol. 17, no. 3, pp. 235–255, 2002.
- [6] A. C. Bahnsen, D. Aouada, A. Stojanovic, and B. Ottersten, “Feature engineering strategies for credit card fraud detection,” *Expert Systems with Applications*, vol. 51, pp. 134–142, 2016.
- [7] E. Duman and M. H. Özcelik, “Detecting credit card fraud by genetic algorithm and scatter search,” *Expert Systems with Applications*, vol. 38, no. 10, pp. 13057–13063, 2011.
- [8] S. Bhattacharyya, S. Jha, K. Tharakunnel, and J. C. Westland, “Data mining for credit card fraud: A comparative study,” *Decision Support Systems*, vol. 50, no. 3, pp. 602–613, 2001.
- [9] G. S. Kumar, S. K. Hidayath, N. N. Bindu, and R. V. Priya, “Online fraud detection using random forest algorithm,” *i-manager’s Journal on Software Engineering*, vol. 19, no. 3, pp. 10–22, 2025.
- [10] A. Dal Pozzolo, O. Caelen, R. A. Johnson, and G. Bontempi, “Calibrating probability with undersampling for unbalanced classification,” in *2015 IEEE Symposium Series on Computational Intelligence*, pp. 159–166, IEEE, 2015.
- [11] J. Gama, I. Žliobaitė, A. Bifet, M. Pechenizkiy, and A. Bouchachia, “A survey on concept drift adaptation,” *ACM Computing Surveys*, vol. 46, no. 4, pp. 1–37, 2014.
- [12] G. I. Webb, R. Hyde, H. Cao, H. L. Nguyen, and F. Petitjean, “Characterizing concept drift,” *Data Mining and Knowledge Discovery*, vol. 30, no. 4, pp. 964–994, 2016.
- [13] D. Bouneffouf and I. Rish, “A survey on practical applications of multi-armed and contextual bandits,” *arXiv preprint arXiv:1904.10040*, 2019.
- [14] N. Dalvi, P. Domingos, S. Sanghai, D. Verma, *et al.*, “Adversarial classification,” in *Proceedings of the Tenth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 99–108, ACM, 2004.
- [15] R. M. Fano, “Transmission of information: A statistical theory of communications,” *American Journal of Physics*, vol. 29, no. 11, pp. 793–794, 1961.
- [16] T. M. Cover and J. A. Thomas, *Elements of Information Theory*. Hoboken, NJ: Wiley-Interscience, 2nd ed., 2006.

- [17] A. Gelman, A. Jakulin, M. G. Pittau, and Y.-S. Su, "A weakly informative default prior distribution for logistic and other regression models," *The Annals of Applied Statistics*, vol. 2, no. 4, pp. 1360–1383, 2008.
- [18] C. M. Bishop, *Pattern Recognition and Machine Learning*. Berlin: Springer, 2006.
- [19] D. J. C. MacKay, "Bayesian interpolation," *Neural Computation*, vol. 4, no. 3, pp. 415–447, 1992.
- [20] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. New York: Springer, 2nd ed., 2009.
- [21] M. West and J. Harrison, *Bayesian Forecasting and Dynamic Models*. New York: Springer, 2nd ed., 1997.
- [22] A. P. Dempster, N. M. Laird, and D. B. Rubin, "Maximum likelihood from incomplete data via the em algorithm," *Journal of the Royal Statistical Society: Series B (Methodological)*, vol. 39, no. 1, pp. 1–38, 1977.
- [23] V. Chandola, A. Banerjee, and V. Kumar, "Anomaly detection: A survey," *ACM Computing Surveys*, vol. 41, no. 3, pp. 1–58, 2009.
- [24] T. C. Gard, *Introduction to Stochastic Differential Equations*. New York: Marcel Dekker, 1988.
- [25] J. Neyman and E. S. Pearson, "On the problem of the most efficient tests of statistical hypotheses," *Philosophical Transactions of the Royal Society of London. Series A*, vol. 231, no. 694-706, pp. 289–337, 1933.

Citation IEEE Format:

O. S. Olorok, K. O. Clinton, and A. M. Okedoye. "An Integrated Mathematical Model for Detecting Known and Unknown Fraud Patterns with Optimal Resource Allocation in the Nigerian Social Security System", *Jurnal Diferensial*, vol. 8(2), pp. 157-179, 2026.

This work is licensed under a [Creative Commons "Attribution-ShareAlike 4.0 International"](https://creativecommons.org/licenses/by-sa/4.0/) license.

