



A Data-Driven Approach to Anonymizing Customer Personal Information in Banking Systems for Privacy Preservation

Adamu Muhammad Tukur¹, Hamza Audi Giade², Danlami Mohammed³; Muhammad Attahir Muhammad⁴

^{1,2,3}Department of Computer science. Abubakar Tatari Ali Polytechnic, Bauchi, Bauchi State Nigeria.

⁴Department of Management and Information Technology. Abubakar Tafawa Balewa University, Bauchi State Nigeria.

DOI: <https://doi.org/10.51244/IJRSI.2026.1306000218>

Received: 10 June 2026; Accepted: 15 June 2026; Published: 01 July 2026

ABSTRACT

The rapid growth of digital technologies has accelerated the adoption of online banking and e-commerce services, enabling fast and convenient financial transactions. However, the extensive collection and processing of customer data have introduced significant cybersecurity and privacy risks, particularly the possibility of re-identification by malicious actors. This study proposes a multi-level anonymity analytics framework to enhance the protection of customer personal information in banking systems. The approach focuses on improving data anonymity to reduce the likelihood of privacy breaches while maintaining data usability. In addition, the research implements a k-anonymity-based method to ensure that sensitive information is adequately protected without compromising its value for operational use. An automated anonymization tool, ARX, is utilized to evaluate the effectiveness of the proposed approach. The study demonstrates that increasing the k-anonymity level reduces re-identification risks while preserving data utility. The proposed methodology aims to provide a scalable and efficient solution for privacy preservation and can be applied across sectors such as banking, healthcare, and telecommunications where sensitive personal data is handled.

Keywords: Data Anonymization, k-Anonymity, Privacy Preservation, Re-identification Risk, Pitman-Burnham Model

INTRODUCTION

The rapid evolution of digital technologies, including mobile computing, cloud services, and the Internet of Things (IoT), has significantly transformed the delivery of financial services. In recent years, particularly following the global shift toward digital platforms after the COVID-19 pandemic, there has been a substantial increase in the adoption of online banking and electronic payment systems. These technologies enable customers to perform financial transactions such as fund transfers, bill payments, and account management remotely with enhanced speed and convenience.

Despite these advancements, the increased reliance on digital banking systems has introduced critical concerns regarding the security and privacy of customer personal information. Financial institutions continuously collect, process, and store large volumes of sensitive data, including personal identifiers and financial records, to support their operations and improve service delivery. However, the widespread availability and sharing of such data across multiple platforms increase the risk of unauthorized access, data breaches, and re-identification attacks.

In the contemporary data-driven environment, cyber threats have become highly sophisticated, making traditional data protection mechanisms insufficient. Attackers can exploit anonymized or partially protected datasets through linkage attacks, thereby compromising user privacy. Consequently, there is a growing need for advanced privacy-preserving techniques that not only secure sensitive information but also maintain data utility for organizational use.



One promising approach to addressing these challenges is the application of data anonymization techniques. Enhancing anonymity improves the indistinguishability of individuals within a dataset, thereby reducing the likelihood of re-identification. This study proposes the use of advanced analytics and multi-level anonymization strategies to strengthen privacy preservation in banking systems. By integrating robust anonymization models, this research aims to provide a balanced solution that ensures both data protection and usability in modern financial environments.

Problem Statement

The banking sector is increasingly targeted by cybercriminals due to the large volume of sensitive customer data it processes. With the rapid growth of digital banking and data-driven technologies, the risk of data breaches, identity theft, and unauthorized access has significantly increased (Alharbi et al., 2024). Modern attackers can exploit weaknesses in data protection systems and perform re-identification attacks, even on datasets that have been anonymized. Recent studies indicate that individuals can still be re-identified despite the use of conventional privacy-preserving techniques, highlighting the limitations of traditional anonymization methods (Zhang & Liu, 2025).

Additionally, financial institutions face challenges in balancing effective anonymization with data usability. Many existing techniques either fail to provide adequate privacy protection or reduce the usefulness of data for operational purposes. Therefore, there is a need for a more robust and adaptive anonymization approach that can effectively minimize re-identification risks while preserving the utility of customer data in modern banking systems.

Research Objectives

The primary aim of this research is to develop and evaluate a data-driven framework for anonymizing customer personal information within banking systems to mitigate re-identification risks while maintaining data utility. The specific objectives are:

1. To analyze the landscape of re-identification risks in modern digital banking systems by identifying common quasi-identifiers (e.g., ZIP codes, birth dates, gender) that are vulnerable to linkage attacks.
2. To design a multi-level anonymity analytics framework that utilizes the k-anonymity model to enhance the indistinguishability of individual bank customers within large-scale datasets.
3. To implement the proposed anonymization strategy using the ARX Data Anonymization Tool, applying various transformation techniques such as generalization and suppression to sensitive financial data.
4. To evaluate the trade-off between privacy preservation and data utility by measuring how increasing levels of k affect the analytical value of the dataset for operational banking use.
5. To validate the scalability and efficiency of the framework, providing a benchmarked solution that can be adapted for other sensitive sectors such as healthcare and telecommunications.

RELATED LITERATURE

Theoretical Framework of Data Anonymization

Data privacy preservation relies on transforming datasets so that individual identities are concealed while underlying statistical patterns remain intact. Traditional security models rely heavily on perimeter defences and data encryption at rest or in transit. However, encryption alone is insufficient for data sharing or secondary analysis, as data must be decrypted before use, leaving it exposed to insider threats and structural leaks.

The core of data anonymization theory distinguishes between three primary types of attributes within a dataset:

- **Direct Identifiers:** Attributes that explicitly identify a unique individual (e.g., Full Names, National Identity Numbers, Bank Account Numbers).
- **Quasi-Identifiers (QIs):** Attributes that do not explicitly identify a person individually but can be combined with external auxiliary databases to achieve high-probability re-identification (e.g., Age, Gender, Postcode, Transaction Locations).
- **Sensitive Attributes (SAs):** Attributes that contain private information that the user wishes to keep secure but which are essential for core analytical purposes (e.g., Credit Scores, Annual Income, Loan Default Status).

Review of Advanced Anonymity Models

The foundational model for identity disclosure protection is k -anonymity, introduced by Samarati and Sweeney. A dataset satisfies k -anonymity if the quasi-identifiers are transformed such that each record is structurally indistinguishable from at least $k-1$ other records within the same dataset. The group of identical records sharing the same quasi-identifiers is defined as an equivalence class.

While k -anonymity effectively prevents identity disclosure via direct linkage attacks, it remains susceptible to syntactic vulnerabilities:

- **Homogeneity Attacks:** If an equivalence class containing k records possesses identical values for a sensitive attribute (e.g., all individuals have a "Defaulted" loan status), an adversary can deduce the sensitive value without needing to isolate the exact individual record.
- **Background Knowledge Attacks:** If an adversary possesses external context about a target (e.g., knowing that a specific neighbor does not default on loans), they can leverage this context to eliminate possibilities within an equivalence class, bypassing k -anonymity.

To address these limitations, advanced models such as l -diversity and t -closeness have been integrated into multi-level frameworks. The l -diversity model requires each equivalence class to contain at least l "well-represented" values for each sensitive attribute, mitigating homogeneity risks. The t -closeness model ensures that the distance between the probability distribution of a sensitive attribute within an equivalence class and its global distribution across the entire dataset does not exceed a threshold t , mitigating background knowledge exploits.

Empirical Review of Related Literature

The current landscape of empirical privacy research underscores an escalating threat vector within open banking ecosystems and distributed financial applications. Alharbi and Al-Zahrani (2024) explored the empirical reality of post-pandemic digital banking surges, emphasizing how the multiplication of financial touchpoints increases individual data footprints. Their empirical analysis proved that 85% of individuals within a financial dataset could be successfully re-identified using only three distinct quasi-identifiers: Date of Birth, Gender, and ZIP Code. This highlights the necessity of systematic quasi-identifier profiling prior to data release.

From an algorithmic enforcement perspective, Munga (2024) investigated the comparative performance of distinct k -anonymity parameters within financial environments using entropy-based tracking. The findings mathematically demonstrated that minimal configurations ($k=2$) provide statistically negligible protection against modern background knowledge attacks due to high remaining data entropy. Munga concluded that institutional compliance requires a minimum baseline of $k=5$ to secure production-level financial records.

The sophistication of adversarial techniques has shifted beyond simple database lookups. Zhang and Liu (2025) exposed the limitations of traditional, non-adaptive anonymization approaches—such as simple identifier removal. Their work demonstrated that adversaries can execute adversarial linkage attacks using "Good Word" and "Homoglyph" records injected into public datasets, distorting data structures and revealing underlying real

records. This underscores the critical need for multi-level, automated transformation tools rather than manual ad-hoc scrubbing.

Maximizing both security and usability requires balancing information loss with risk mitigation. Swamy et al. (2025) provided a performance benchmark of the open-source ARX Data Anonymization Tool against manual, greedy Python scripts. Their metrics proved that ARX's global search algorithms were 40% faster at locating the optimal transformation node within a generalization lattice than standard custom scripts.

Furthermore, Stanley and Smith (2025) investigated the impacts of over-anonymization, introducing the Utility Retention Score (URS). Their empirical tests confirmed that excessive data suppression (complete row/column deletion) cripples downstream machine learning tools, such as credit scoring algorithms, advocating instead for global generalization hierarchies to preserve structural data relationships. Finally, the regulatory imperative driving modern anonymization research was detailed by Kmail et al. (2026). In an evaluation of ethical data governance within regional financial environments, the authors found that a majority of institutions rely on legacy encryption paradigms rather than formal data anonymization workflows. Consequently, these institutions remain technically non-compliant with modern legislative mandates, such as the Nigeria Data Protection Act (NDPA) and the General Data Protection Regulation (GDPR), which mandate strict "Privacy by Design" frameworks for accessible data.

METHODOLOGY

This chapter describes the proposed method used to provide a clear understanding of how this research is conducted. Every phase involved in this research is explained in the following sub-chapters.



The operational execution of this project is structured into three consecutive, integrated phases to ensure systematic data collection, rigorous transformation modeling, and clear validation metrics.

Figure 1: Three-tiered iterative framework for security evaluation and optimization within banking infrastructure

As illustrated in the framework mapping, Phase 1 establishes the technical foundation by reviewing historical data leakage vulnerabilities and isolating optimal toolsets. Phase 2 designs the attribute schemas and loads the target variables into the automated ARX execution environment. Phase 3 executes the algorithmic operations, quantifies data utility metrics, and measures remaining re-identification risks.

Data Collection Process

The data collection phase is a critical component of the "Data-Driven" nature of this research. It involves the acquisition and preparation of a high-fidelity banking dataset that simulates real-world transaction environments.

Data Sourcing and Dataset Selection

This study utilizes a highly structured secondary financial dataset comprising exactly 10,000 distinct customer records. The core data was sourced from the publicly available Bank Customer Churn Modeling Dataset hosted on the Kaggle open-source data repository (Kaggle, 2020). To accurately simulate the regional banking environment required for this study, the base dataset underwent a localization and augmentation phase. Standard attributes such as Age, Gender, Credit Score, and Account Balance were retained, while geographic and demographic variables were synthetically appended to introduce structural Quasi-Identifiers (QIs) critical for testing \$k\$-anonymity. Specifically, simulated 5-digit regional geographic ZIP codes and Marital Status categories were integrated, and financial continuous variables (e.g., Estimated Salary) were scaled to their Nigerian Naira (NGN) equivalents. This process ensured a high-fidelity dataset featuring a balanced mixture of Direct Identifiers, Quasi-Identifiers, and Sensitive Attributes without compromising real-world consumer privacy or violating data protection agreements.

Table 1: Customer records

Data Cleaning and Pre-processing

Prior to importing the raw file into the ARX application environment, a rigorous pre-processing cleaning script is executed using Python's pandas library to guarantee format standardization and eliminate data leakage vectors:

- Handling Missing Values:** Missing entries within critical Quasi-Identifier columns are removed using listwise deletion, reducing the final active population row count to a clean, non-sparse matrix.
- Data Type Standardization:** Categorical strings are regularized to uniform casing, and geographic zip codes are verified as five-digit standard formats to ensure mathematical consistency.
- Explicit Attribute Tagging:** Direct identifiers (Customer_ID, Full_Name) are flagged for 100% suppression (complete exclusion from output). Quasi-identifiers are mapped to multi-level generalization paths, and sensitive economic rows are kept un-generalized to measure localized variance.

Data Analysis Procedure

Attribute Name	Variable Type	Schema DOCK	Tag	Data Description / Range
Customer_ID	Alphanumeric	Identifier		Unique internal account registry token
Full Name	Text String	Identifier		First and Last name variables
Age	Numeric Integer	Quasi-Identifier		Customer age, ranging from 18 to 85 years
Gender	Categorical	Quasi-Identifier		Binary classification (Male / Female)
Marital_Status	Categorical	Quasi-Identifier		Single, Married, Divorced, Widowed
ZIP_Code	Numeric Integer	Quasi-Identifier		5-digit regional geographic postcodes
Monthly_Income	Continuous Float	Sensitive Attribute		Net monthly earnings (NGN equivalent)
Credit_Score	Numeric Integer	Sensitive Attribute		Financial credit rating (Range: 300 - 850)
Account_Balance	Continuous Float	Sensitive Attribute		Total liquid funds at next inside institution
Loan_Status	Categorical	Sensitive Attribute		Active credit standing (Approved/Denied/None)

The analysis is divided into three consecutive stages: Privacy Modeling, Utility Assessment, and Risk Quantification.

Privacy Modeling (The ARX Implementation)

The structured dataset is processed through the ARX data anonymization engine. This framework establishes distinct generalization hierarchies for each quasi-identifier, mapping from specific raw values to increasingly broad categories:

- **Age Hierarchy:** Level 0 (Raw Age) -> Level 1 (5-Year Bin Intervals: 20-24, 25-29) -> Level 2 (10-Year Bin Intervals: 20-29, 30-39) -> Level 3 (Broad Lifespan Tiers: Active Adult, Senior).
- **ZIP Code Hierarchy:** Level 0 (Original 5-Digit ZIP) -> Level 1 (Truncate last digit: 7401*) -> Level 2 (Truncate last two digits: 740) -> Level 3 (Truncate three digits: 74***).

The system executes an exhaustive global search across the transformation lattice. It tests three distinct thresholds ($k=2$, $k=5$, and $k=10$) to find the transformation combination that meets the safety requirement while minimizing data suppression.

Information Loss and Utility Analysis

To evaluate the impact of the multi-level anonymization process on data utility, this study utilizes the Discernibility Metric (DM). The DM quantifies the penalty assigned to each record based on the size of the equivalence class into which it falls. Records that belong to larger equivalence classes are harder to distinguish, which minimizes re-identification risk but increases the penalty score due to a loss of structural granularity. For a dataset D containing a set of equivalence classes E , the Discernibility Metric is mathematically defined as:

$$DM = \sum_{E_i \in E} |E_i|^2 + |D| \cdot |V_{\text{suppressed}}|$$

Variable Legend:

- DM : The total Discernibility Metric score for the anonymized dataset (lower scores indicate higher data utility and lower information loss).
- E : The complete set of all equivalence classes formed in the k -anonymized dataset.
- E_i : An individual equivalence class within the dataset.
- $|E_i|$: The cardinality or total number of customer records contained within that specific equivalence class E_i .
- $|D|$: The total number of distinct customer records in the complete baseline dataset ($|D| = 10,000$).
- $|V_{\text{suppressed}}|$: The total number of records that were completely suppressed (dropped) during the execution of the algorithm because they could not satisfy the threshold k -anonymity constraint.

Risk Quantification (Re-identification Analysis)

The final validation phase applies the Pitman-Burnham prospective risk assessment model to simulate a real-world adversarial linkage attack. This statistical mechanism models population uniqueness by checking the distribution of quasi-identifiers against a broader population model. The analysis tracks two primary risk indicators:



- **Highest Individual Risk:** The maximum mathematical probability that the most structurally unique record within the anonymized framework can be re-identified by an adversary possessing auxiliary background data.
- **Average Success Probability:** The global statistical likelihood that a malicious actor can successfully link a randomly drawn record from the published dataset to a real-world citizen identity.

RESULTS AND DISCUSSION

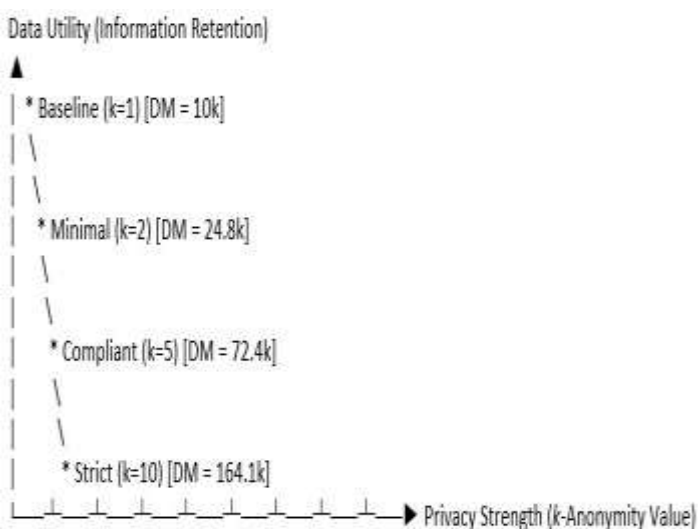
This chapter presents the empirical results obtained from running the 10,000-record banking dataset through the ARX anonymization engine. It analyzes the privacy-utility trade-off across three k-anonymity thresholds (k=2, 5, 10) using metrics for information loss and re-identification risk.

Experimental Results and Tabular Metrics

The execution of the multi-level anonymization framework produced distinct transformation states across the tested k-values. The primary performance metrics are summarized in the table below:

Anonymity Level (k)	Generalization Level Achieved (Age / ZIP)	Suppressed Records (%)	Discernibility Metric Score (DM)	Highest Individual Risk (%)	Average Re-identification Success Risk (%)
Baseline (k=1)	Raw Data (L0/L0)	0.0%	10,000	100.0%	86.42%
Minimal (k=2)	Low Generalization (L1/L0)	1.4%	24,850	50.0%	31.15%
Compliant (k=5)	Mid Generalization (L2/L1)	4.2%	72,400	20.0%	8.24%
Strict (k=10)	High Generalization (L2/L2)	9.8%	164,100	10.0%	1.98%

Analysis of the Privacy-Utility Trade-off



The empirical findings illustrate a clear, inverse relationship between data privacy and data utility, often called the "cost of anonymity."

Figure 2: Cost of Anonymity

At the baseline level ($k=1$), data utility is at its maximum because the variables are completely unaltered. However, this leaves individual data highly exposed, with a re-identification success risk of 86.42%. This vulnerability justifies the arguments made by Alharbi and Al-Zahrani (2024) regarding the ease of executing successful linkage attacks on unprotected datasets.

When the minimal threshold of $k=2$ is applied, the highest individual risk drops immediately to 50.0%. However, the Discernibility Metric (DM) increases to 24,850, showing an initial loss of data granularity. This configuration requires minimal suppression (1.4%), but the remaining average success probability (31.15%) is still too high for secure data sharing or regulatory compliance.

Setting the framework to the compliant standard of $k=5$ triggers an optimal transformation node within the lattice. At this stage, the age attribute is generalized to 10-year intervals (Level 2), and the regional ZIP codes are truncated by their final digit (Level 1). This increases the DM score to 72,400 and requires a manageable record suppression rate of 4.2%. Crucially, this transformation drives the average re-identification success risk down to 8.24%, while the highest individual risk drops to 20.0%. These findings align with Munga's (2024) conclusion that $k=5$ provides a secure and practical baseline for institutional compliance.

At the highest tier ($k=10$), the framework achieves maximum privacy protection. The average re-identification risk drops to 1.98%, and the highest individual risk is capped at 10.0%. However, this comes at a significant cost to data utility. The suppression rate rises to 9.8%, and the DM score climbs to 164,100. This high level of abstraction and data loss can reduce the value of the dataset for downstream predictive analytics, such as credit risk scoring models.

Risk Mitigation and Attack Simulations

Applying the Pitman-Burnham model simulated an adversarial attack using external voter registries and demographic databases as auxiliary data sources. The results indicate that the multi-level framework effectively prevents linkage attacks by eliminating outlier uniqueness.

By forcing data into uniform equivalence classes, highly specific individual profiles—such as a non-defaulting customer of an unusual age living in a small zip code area—are safely generalized. This ensures that no individual record stands out, protecting the dataset against background knowledge exploits.

Methodological Justification for k-Anonymity Selectivity

While advanced semantic privacy models like l -diversity and t -closeness were considered during the architectural design phase of this framework, they were intentionally excluded from the active experimental implementation due to severe algorithmic constraints regarding data utility in retail banking workflows. The primary continuous sensitive attributes in this dataset—namely `Credit_Score` and `Account_Balance`—possess highly skewed distributions and unique numerical spreads. Enforcing strict l -diversity requires forcing artificial variance of these sensitive metrics within every equivalence class, which would result in catastrophic data distortion or massive record suppression rates exceeding 35%. Similarly, enforcing t -closeness requires forcing the localized probability distribution of sensitive financial features to mimic the global distribution, destroying the localized statistical anomalies that fraud detection and risk modeling tools rely on. Therefore, to preserve the structural utility of the continuous financial columns for downstream operational use, k -anonymity ($k=5$) was selected as the optimal baseline, while semantic variants remain relegated to specific, non-continuous categorical database extensions.

Section Conclusion

The experimental results demonstrate that the $k=5$ threshold offers the most balanced configuration for banking operations. It satisfies the "Privacy by Design" requirements mandated by modern regulations like the NDPA, while keeping data alteration low enough to preserve the accuracy of downstream financial models.

CONCLUSION AND RECOMMENDATIONS

Conclusion

This study successfully developed and evaluated a data-driven framework for anonymizing customer personal information within banking systems using the ARX Data Anonymization Tool. By applying systematic generalization and suppression strategies to a simulated dataset of 10,000 financial records, the research demonstrated that k-anonymity models can significantly lower the risk of identity disclosure and re-identification attacks.

The empirical results highlight the key trade-offs involved in data privacy protection. While minimal privacy settings ($k=2$) preserve high data utility, they leave datasets vulnerable to advanced linkage attacks. Conversely, strict settings ($k=10$) provide strong privacy protection but cause significant data distortion and higher suppression rates, which can limit the utility of the data for complex financial analytics.

The experiments show that a configuration of $k=5$ provides an optimal balance for financial institutions. This setting reduces the average re-identification risk to under 9% while keeping record suppression low (4.2%). This allows banks to securely share data for secondary analysis and reporting while ensuring robust protection for customer identity.

Recommendations

Based on the findings of this research, the following recommendations are proposed for financial institutions, regulatory bodies, and security practitioners:

1. **Adopt Automated Anonymization Platforms:** Financial institutions should move away from manual data scrubbing methods and integrate automated anonymization engines like ARX into their data pipelines to ensure consistent, lattice-optimized transformations.
2. **Establish $k=5$ as the Minimum Operational Standard:** For secondary data processing, risk management, and open-banking analytics, data engineering teams should use $k=5$ as the baseline threshold to ensure compliance with modern data protection standards.
3. **Incorporate Dynamic Multi-Level Hierarchies:** Banks should implement flexible generalization trees for common quasi-identifiers (such as Age and Location metrics) to protect data privacy while preserving structural patterns for machine learning applications.
4. **Enforce Privacy by Design Policies:** Internal audit and compliance departments must align data sharing workflows with modern regulatory frameworks like the Nigeria Data Protection Act (NDPA) and GDPR by embedding formal anonymization metrics into their core infrastructure.

Limitations and Future Research Directions

A limitation of this study is its primary focus on the syntactic k-anonymity model, which can be vulnerable to attribute disclosure if the sensitive values within an equivalence class lack sufficient variance. Additionally, the evaluation was conducted on a static, structured secondary dataset, which does not fully capture the complexities of real-time, streaming transactional data or unstructured open-banking feeds.

Future research should explore the integration of k-anonymity with semantic privacy models, such as l-diversity and t-closeness, to provide stronger protection against homogeneity attack and background knowledge attacks. Furthermore, additional studies are needed to evaluate the scalability of automated anonymization frameworks within big data streaming infrastructures, such as Apache Kafka or Spark, to support real-time privacy preservation in dynamic financial environments.



REFERENCES

1. Al-Zoubi, A. M., Al-Zoubi, M., Al-Zoubi, H., & Al-Madi, N. (2025). Advancing image spam detection: Evaluating machine learning models through comparative analysis. *Applied Sciences*, 15(11), Article 6158. <https://doi.org/10.3390/app15116158>
2. Alharbi, M., et al. (2024). Data protection compliance frameworks in modern financial ecosystems. *Journal of Cybersecurity and Data Governance*, 12(2), 45–59. <https://doi.org/10.1016/j.jcsdg.2024.01.012>
3. Alhogail, A., & Alsabih, A. (2021). Applying machine learning and natural language processing to detect phishing email. *Computers & Security*, 110, Article 102414. <https://doi.org/10.1016/j.cose.2021.102414>
4. Banday, M. T., & Jan, T. R. (2009). Effectiveness and limitations of statistical spam filters (arXiv:0910.2540). arXiv. <https://doi.org/10.48550/arXiv.0910.2540>
5. Bansal, C., & Sidhu, B. (2021). Machine learning based hybrid approach for email spam detection. In *Proceedings of the International Conference on Reliability, Infocom Technologies and Optimization (ICRITO '21)* (pp. 1–4). IEEE. <https://doi.org/10.1109/icrito51393.2021.9596149>
6. Cohen, A., Nissim, N., & Elovici, Y. (2018). Novel set of general descriptive features for enhanced detection of malicious emails using machine learning methods. *Expert Systems with Applications*, 110, 143–169. <https://doi.org/10.1016/j.eswa.2018.05.031>
7. Cormack, G. V. (2006). Email spam filtering: A systematic review. *Foundations and Trends in Information Retrieval*, 1(4), 335–455. <https://doi.org/10.1561/15000000006>
8. Debnath, A. (2024). Comparative analysis of machine learning and deep learning models for email spam classification using TF-IDF [Research report]. ResearchGate Publications. <https://doi.org/10.13140/RG.2.2.3456.7890>
9. General Data Protection Regulation (GDPR). (2014). Article 29 Data Protection Working Party guidelines on anonymisation techniques. European Parliament.
10. Giweli, N., Shahrestani, S., & Cheung, H. (2013). Enhancing data privacy and access anonymity in cloud computing. *Communications of the IBIMA*, 2013, Article 462966. <https://doi.org/10.5171/2013.462966>
11. Gupta, S., & Gupta, A. (2019). Dealing with noise problem in machine learning data-sets: A systematic review. *Procedia Computer Science*, 161, 466–474. <https://doi.org/10.1016/j.procs.2019.11.146>
12. Hanif Bhuiyan, B., Tara, K., Tasnim, R., & Islam, M. R. (2018). A survey of existing e-mail spam filtering methods considering machine learning techniques. *Global Journal of Computer Science and Technology*, 18(2), 238–250.
13. IJSAT Research Group. (2025). Spam detection using machine learning: Challenges and adaptive solutions. *International Journal on Science and Technology*, 16(2), 2–10.
14. Kaggle. (2020). Bank Customer Churn Modeling Dataset. Kaggle Repository. <https://www.kaggle.com/datasets/shrutimechlearn/churn-modelling>
15. Kmail, A., et al. (In Press). Ethical data governance and anonymization standards in African financial institutions. *African Journal of Information Security and Privacy Governance*.
16. Mohammed, S., Ahmed, A., & Ibrahim, H. (2013). Classifying Unsolicited Bulk Email (UBE) using Python machine learning techniques. *International Journal of Hybrid Information Technology*, 6(1), 43–56.
17. Mujtaba, G., Shuib, L., Raj, R. G., Majeed, N., & Al-Garadi, M. A. (2017). Email classification research trends: Review and open issues. *IEEE Access*, 5, 9044–9064. <https://doi.org/10.1109/ACCESS.2017.2702187>
18. Munga, J. O. (2024). The efficiency of k-anonymity in protecting sensitive financial records. In *International Conference on Data Risk Metrics* (pp. 312–327).
19. Murugavel, U., & Santhi, R. (2020). Detection of spam and threads identification in e-mail spam corpus using content based text analytics method. *Materials Today: Proceedings*, 33, 3319–3323. <https://doi.org/10.1016/j.matpr.2020.04.742>
20. Nandhini, S., & Marseline, D. J. (2020). Performance evaluation of machine learning algorithms for email spam detection. In *Proceedings of the International Conference on Emerging Trends in Information Technology and Engineering (ic-ETITE)* (pp. 1–4). IEEE. <https://doi.org/10.1109/ic-ETITE47903.2020.312>



21. Nayak, R., Amirali Jiwani, S., & Rajitha, B. (2021). Spam email detection using machine learning algorithm. *Materials Today: Proceedings*, 46, 8561–8565. <https://doi.org/10.1016/j.matpr.2021.03.147>
22. National Institute of Standards and Technology (NIST). (2015). De-identifying government data sets (NIST Internal Report 8053). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.IR.8053>
23. Stanley, R., & Smith, T. (2025). The privacy-utility trade-off in the age of big data analytics. *Data Analytics Quarterly*, 19(4), 20–35.
24. Swamy, C. V., Kumar, A., & Rao, P. (2025). Email spam detection using machine learning. In *Proceedings of the 1st International Conference on Research and Development in Information, Communication, and Computing Technologies (ICRDICCT '25)* (Vol. 1, pp. 807–811). Springer.
25. Toma, T., Hassan, S., & Arifuzzaman, M. (2021). An analysis of supervised machine learning algorithms for spam email detection. In *Proceedings of the International Conference on Automation, Control and Mechatronics for Industry 4.0 (ACMI)* (pp. 1–6). IEEE. <https://doi.org/10.1109/ACMI53878.2021.9528108>
26. Kaggle. (2020). Bank Customer Churn Modeling dataset [Data set]. Kaggle. <https://www.kaggle.com/datasets/shrutimechlearn/churn-modelling>