

Performance Evaluation of Classical Machine Learning Models for Insurance Fraud Detection under Severe Class Imbalance

Garima Sharma

Independent Researcher, Arlington Heights, Illinois, USA

Abstract

Insurance claim fraud continues to pose substantial financial and operational challenges for insurance companies worldwide, particularly in auto insurance, where fraudulent cases form a small but highly impactful portion of overall claims. One of the significant technical difficulties in detecting such fraud is the severe class imbalance in real-world insurance datasets, where legitimate claims vastly outnumber fraudulent ones. This study presents a systematic performance evaluation of classical machine learning (ML) classification models for insurance fraud detection under conditions of extreme class imbalance. The proposed framework focuses on widely used supervised learning algorithms, including Support Vector Machine (SVM), K-Nearest Neighbor (KNN), and Random Forest (RF), with an emphasis on understanding their behavior when trained on imbalanced data. To mitigate class imbalance bias, the Synthetic Minority Oversampling Technique (SMOTE) is applied to the dataset before model training. Model performance is evaluated using multiple metrics such as accuracy, precision, recall, and F1-score, which provide a more reliable assessment than accuracy alone in fraud detection scenarios. Experimental results demonstrate that ensemble-based methods, particularly Random Forest, achieve superior performance in identifying minority class fraud cases while maintaining stable overall classification accuracy. This research provides practical insights into selecting suitable classical ML models for insurance fraud detection, supporting the development of reliable decision support systems for insurance providers operating with imbalanced data.

Keywords: Insurance Fraud Detection, Class Imbalance, Machine Learning, SMOTE, Classification Models

1. Introduction

Insurance systems are designed to provide financial protection against unexpected losses and risks by pooling policyholders' contributions. While this structure benefits individuals and businesses, it is also vulnerable to misuse through fraudulent insurance claims. Insurance fraud refers to deliberate actions taken by claimants or associated parties to obtain financial benefits that they are not legally entitled to [1]. Over the years, insurance fraud has emerged as a serious global concern, affecting insurers, policyholders, and regulatory bodies alike [2]. Among different insurance domains, auto insurance accounts for a significant proportion of reported fraud cases, making it one of the most challenging areas for fraud detection [3]. Fraudulent insurance claims lead to substantial financial losses for insurance companies each year. These losses are often indirectly borne by honest customers through higher

premiums and stricter claim approval processes [4]. In addition to its financial impact, insurance fraud undermines trust in the insurance system and places additional pressure on investigative and legal resources. As fraud schemes become more sophisticated, traditional rule-based and manual detection approaches are no longer sufficient to identify complex fraudulent patterns hidden within large volumes of claim data [5]. Recent advancements in data analytics and machine learning have provided new opportunities to improve insurance fraud detection. Machine learning techniques enable systems to learn patterns from historical data and make predictions on previously unseen cases with minimal human intervention [1][6]. In the insurance domain, supervised learning models have been widely applied to tasks such as risk assessment, claim validation, customer behavior analysis, and fraud detection [7]. These models are instrumental when labeled datasets are available, enabling algorithms to distinguish fraudulent from non-fraudulent claims based on learned features [8]. Despite their advantages, ML-based fraud detection systems face a critical challenge related to data imbalance. In real-world insurance datasets, fraud claims typically represent a tiny fraction of the total number of claims, while legitimate claims dominate the data distribution. This severe class imbalance can bias classification models toward predicting the majority class, resulting in high accuracy but poor fraud detection capability [9].

Thus, models may fail to identify fraudulent cases effectively, which is unacceptable in practical insurance applications where missing fraud instances can lead to significant losses [10]. To address this issue, researchers have explored various data resampling and evaluation strategies to improve model performance under imbalanced conditions. Techniques such as the Synthetic Minority Oversampling Technique have gained popularity for balancing datasets before model training [1][11]. At the same time, performance metrics beyond accuracy, including precision, recall, and F1 score, have been emphasized as more appropriate indicators of fraud detection effectiveness [12][13]. Motivated by these challenges, this study evaluates the performance of classical machine learning classification models for insurance fraud detection under severe class imbalance. By comparing widely used algorithms, such as Support Vector Machine, KNN, and Random Forest, this research aims to provide practical insights into their strengths and limitations in imbalanced insurance datasets. The findings are intended to support insurance practitioners and researchers in selecting suitable models and evaluation strategies for developing reliable fraud detection systems.

2. Literature Review

Earlier studies on insurance fraud detection primarily focused on understanding fraud patterns and their financial consequences rather than developing automated detection systems. Joginipalli et al. [14] highlighted that fraudulent insurance claims impose significant economic pressure on insurers and policyholders, a fact that has historically led to reliance on manual audits and rule-based screening mechanisms. These early approaches were limited in scalability and adaptability, particularly as claim volumes increased and fraud strategies became more sophisticated. Existing research gradually shifted toward data-driven techniques as insurance datasets expanded in size and dimensionality. Ali et al. [15] conducted a systematic review of the literature on financial fraud detection. They reported that supervised machine learning methods dominate modern fraud detection research due to their ability to learn discriminative patterns from labeled historical data. Their review emphasized that classification-based models offer greater adaptability than traditional statistical methods but also identified data imbalance as a persistent challenge affecting model reliability. As machine learning gained wider adoption in the insurance domain, researchers began evaluating classical classifiers for fraud detection

tasks. Rukhsar et al. [16] presented a comparative study of supervised machine learning algorithms for insurance fraud detection, demonstrating that ensemble and tree-based methods generally outperform distance-based classifiers. Their findings also showed that evaluation metrics such as recall and F1 score provide more meaningful insights than accuracy when assessing fraud detection performance. This work established a foundation for comparative analysis of classical models in the context of insurance fraud. Subsequent studies extended this line of inquiry by examining specific classifiers in greater detail. Al Doulat et al. [17] investigated supervised machine learning models for insurance claim fraud detection and reported that the Support Vector Machine performs well when feature separability is high, but its effectiveness decreases under severe class imbalance. Their analysis further indicated that preprocessing strategies and the characteristics of data distribution strongly influence model performance.

K-Nearest Neighbor has been explored as a simple and interpretable approach for fraud detection in related financial domains. Raman et al. [18] conducted a comparative analysis of KNN, Random Forest, and Logistic Regression for credit card fraud detection. They found that KNN can achieve competitive performance when the dataset is well normalized and balanced. However, their results also showed that KNN is sensitive to noise and class overlap, which limits its robustness in highly imbalanced insurance datasets. Random Forest has consistently emerged as a strong performer in fraud detection research due to its ensemble learning structure. Khalil et al. [19] proposed a machine learning-based framework for insurance fraud detection on class-imbalanced datasets with missing values, demonstrating that Random Forest achieves superior recall and F1-score when combined with oversampling techniques. Their study emphasized that handling class imbalance has a greater impact on detection performance than model complexity alone.

More recent research has further reinforced the importance of imbalance-aware modeling in detecting insurance fraud. Nabrawi et al. [20] examined fraud detection in insurance-related datasets and reported that models trained without addressing class imbalance tend to favor legitimate claims, resulting in high accuracy but poor fraud-detection performance. Their findings support the growing consensus that recall-focused evaluation and balanced training data are crucial for the development of effective fraud detection systems. Despite substantial progress, existing research reveals notable gaps in knowledge. Many prior studies evaluate individual models in isolation rather than under consistent preprocessing and imbalance handling conditions. Additionally, limited work provides a structured performance comparison of classical machine learning models specifically under severe class imbalance. These reorganized findings motivate the present study, which systematically evaluates Support Vector Machine, K-Nearest Neighbor, and Random Forest using consistent resampling strategies and fraud-sensitive evaluation metrics.

3. Materials and Methods

This section presents a refined and fully quantitative methodology for evaluating classical ML models for insurance fraud detection under severe class imbalance. The methodological design is inspired by established insurance fraud detection studies, while extending them with deeper feature characterization, explicit quantification of imbalance, and robust evaluation practices to ensure reproducibility and practical relevance.

3.1 Dataset and Feature Grouping

The dataset used in this study comprises historical auto insurance claim records with binary labels indicating whether the claims are fraudulent or non-fraudulent. Each observation represents a completed

insurance claim and includes structured attributes that describe claimant behavior, vehicle details, policy information, and claim-specific patterns. The dataset reflects real-world insurance environments where fraudulent claims occur infrequently but have a disproportionately high financial impact. To enable systematic analysis, features are organized into four logical groups. Demographic characteristics capture claimant-related behavior and situational context, including accident area, driver age, gender, policy duration, recent address change status, fault attribution, police report filing, and the number of claim supplements. These variables are important because they often reveal behavioral irregularities associated with fraudulent activity. Vehicle-related characteristics include vehicle make, vehicle age, vehicle category, and estimated vehicle price, which provide insight into the economic incentives and physical conditions associated with claims. Policy-related attributes comprise the base policy type and agent category, which reflect the coverage structure and claim handling. Specific claim characteristics include deductible amount, number of prior claims, week of claim submission, accident month, and claim month, which capture temporal and financial patterns frequently associated with fraudulent behavior.

3.2 Data Preprocessing and Transformation

Before model development, a structured preprocessing pipeline is applied to improve data quality and ensure algorithm stability. Records containing missing, inconsistent, or invalid values are removed to minimize noise and reduce bias. Categorical variables are transformed into numerical representations using label encoding for ordinal attributes and one-hot encoding for nominal attributes. This transformation allows ML models to process categorical information without imposing artificial ordering. Continuous variables such as driver age, vehicle age, deductible amount, and vehicle price are normalized using min-max scaling. This step ensures that features with larger numeric ranges do not dominate distance-based or margin-based classifiers. After preprocessing, the dataset is divided into training and testing subsets using stratified sampling. This approach preserves the original class distribution in the test set, enabling realistic performance evaluation while preventing information leakage during training.

3.3 Quantitative Analysis of Class Imbalance

Insurance fraud datasets are inherently imbalanced, with fraudulent claims representing a small minority of total observations. In the analyzed dataset, legitimate claims account for more than 90% of records, while fraud cases form a tiny proportion. This imbalance introduces strong bias during model training, as classifiers tend to favor the majority class to optimize overall accuracy. To quantify the severity of imbalance, the imbalance ratio is calculated as the ratio of non-fraudulent to fraudulent instances. Class distribution statistics are examined before resampling to confirm the extent of skewness. This quantitative assessment highlights the inadequacy of accuracy as a standalone metric and underscores the necessity of applying imbalance handling techniques before model training. Fig(s). 1-3 illustrate the imbalanced insurance claim dataset with oversampling and under-sampling.

Figure 1: Imbalanced Insurance Claim Dataset

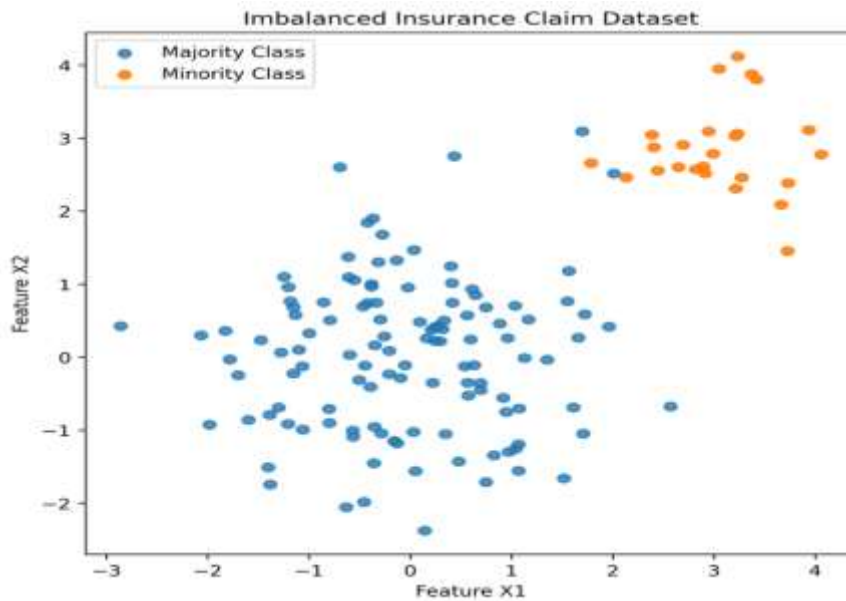


Figure 2: Over-Sampled Insurance Claim Dataset

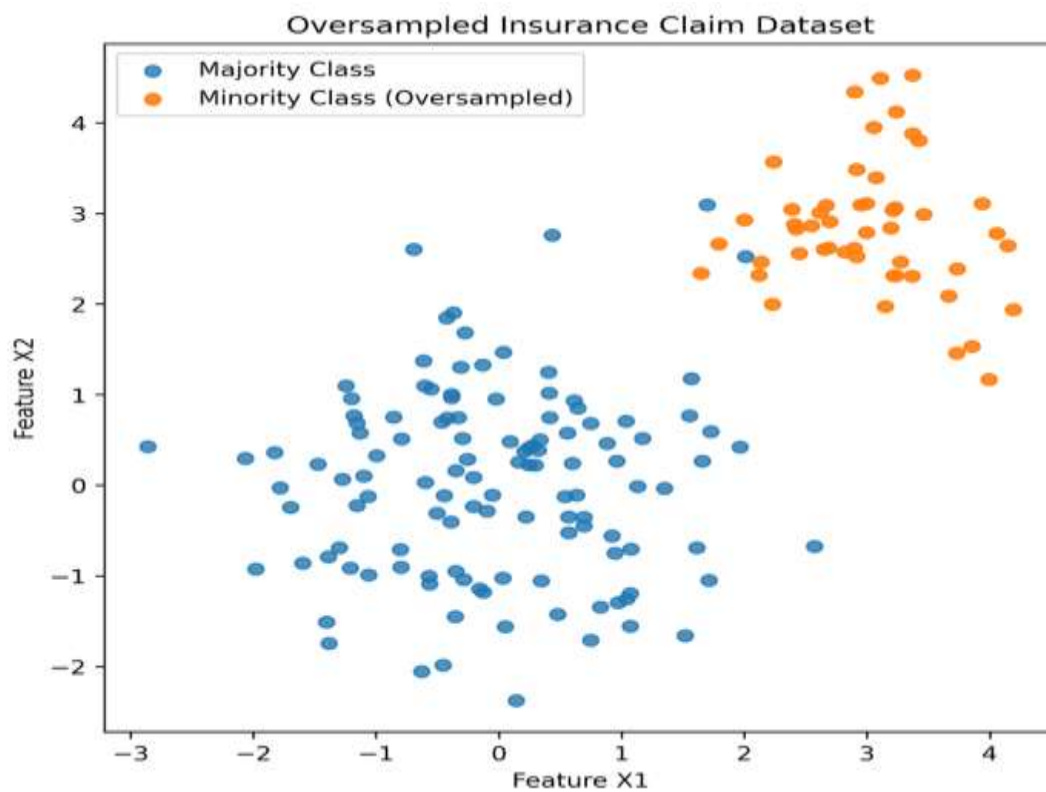
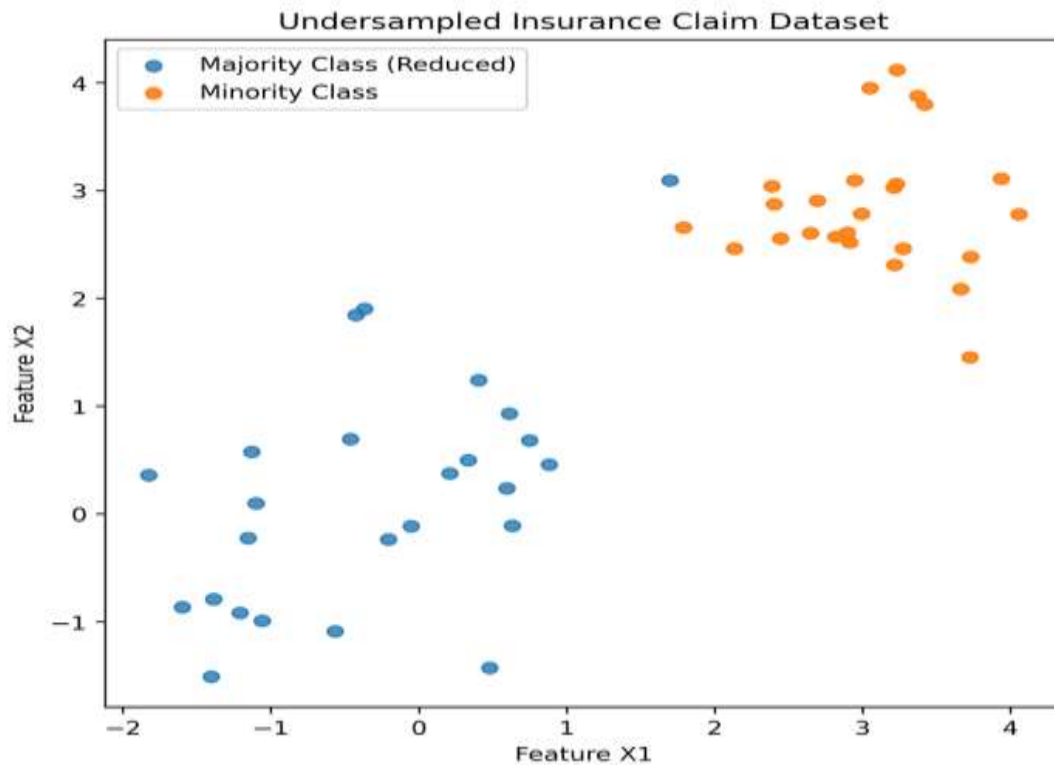


Figure 3: Under-Sampled Insurance Claim Dataset



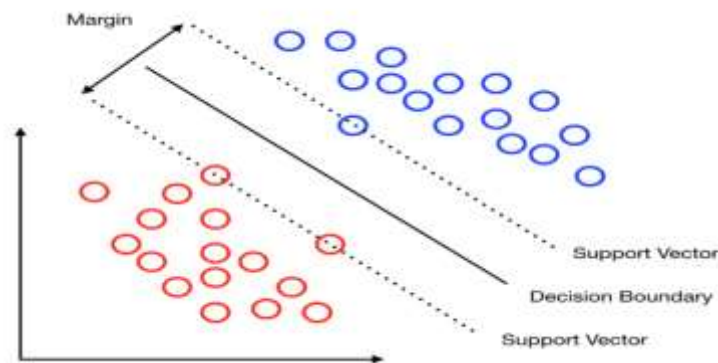
3.4 SMOTE Technique

To mitigate the effects of class imbalance, the Synthetic Minority Oversampling Technique is applied exclusively to the training dataset. SMOTE generates new synthetic fraud samples by interpolating between existing minority-class instances and their nearest neighbors in the feature space. This process increases minority class representation without duplicating records, thereby reducing the risk of overfitting. Following the SMOTE application, the training dataset achieves an approximately balanced class distribution, while the test dataset remains unchanged. This strategy ensures that model learning benefits from balanced data, while evaluation reflects the real-world prevalence of fraud and actual generalization performance.

3.5 Machine Learning Models

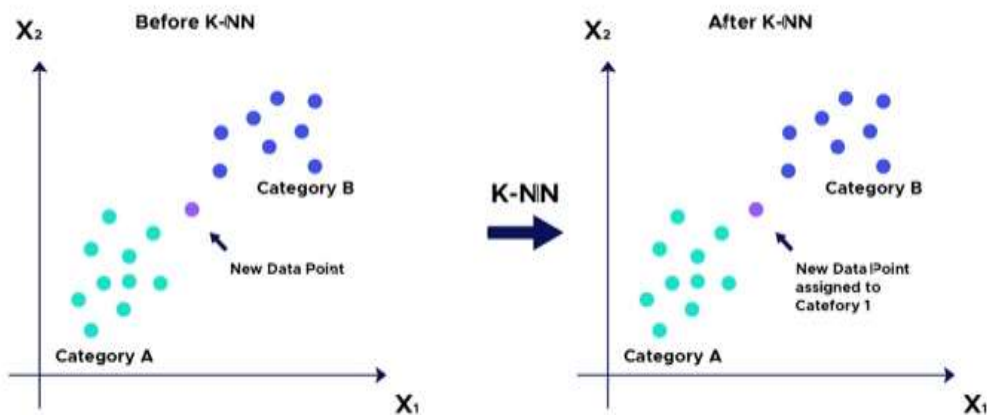
This research evaluates the performance of three classical supervised machine learning classification models commonly used in fraud detection research. Support Vector Machine is a margin-based classifier that aims to identify an optimal decision boundary separating different classes in the feature space. The model constructs a hyperplane that maximizes the margin between data points of other classes. SVM is effective in handling high-dimensional data and has been widely applied in fraud detection tasks due to its strong generalization capability. Fig. 4 illustrates the concept of support vectors, margin boundaries, and the separating hyperplane used by SVM to distinguish classes.

Figure 4: Support Vector Machine (SVM)



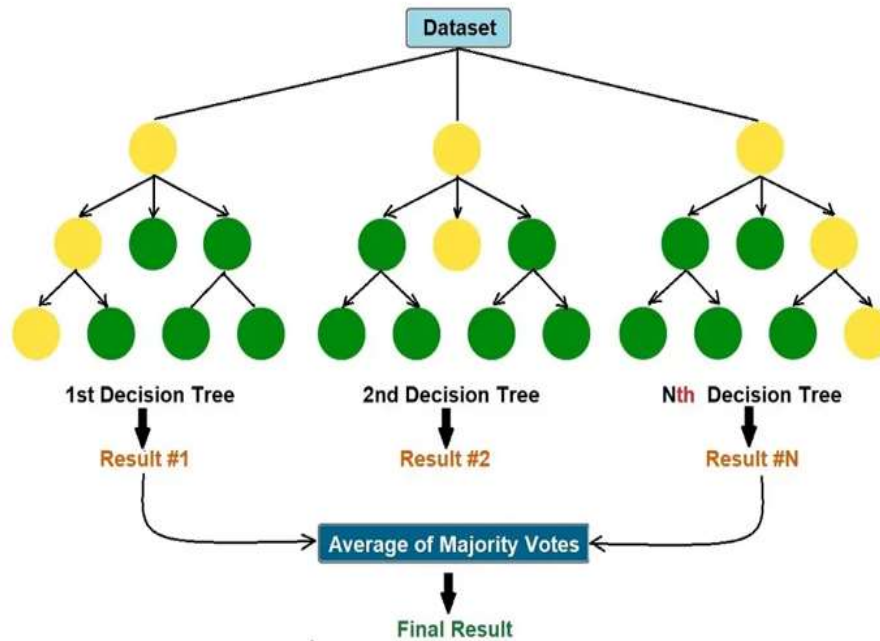
K Nearest Neighbor is an instance-based learning algorithm that classifies a data point based on the majority class of its nearest neighbors. The distance between instances is calculated using a similarity metric, and classification is performed without explicit model training. While KNN is intuitive and straightforward, its performance can be affected by noisy data and feature scaling. Fig. 5 illustrates the concept of KNN.

Figure 5: K-Nearest Neighbor (KNN)



Random Forest is an ensemble learning method that combines multiple decision trees to improve prediction accuracy and robustness. Each tree is trained on a randomly selected subset of the data and features, reducing overfitting and improving generalization. Random Forest is particularly effective in handling nonlinear relationships, reducing variance and complex feature interactions, making it suitable for insurance fraud detection. Fig. 6 illustrates the concept of Random Forest.

Figure 6: Random Forest



3.6 Evaluation Metrics

To assess model performance under severe class imbalance, multiple evaluation metrics were employed. Accuracy (A) measures the proportion of correctly classified instances among all observations and is expressed as:

$$A = (TP + TN) / (TP + TN + FP + FN) \quad (1)$$

Precision (P) represents the proportion of correctly identified fraudulent claims among all claims predicted as fraud and is expressed as:

$$P = TP / (TP + FP) \quad (2)$$

Recall (R) measures the ability of the model to identify actual fraudulent claims correctly and is expressed as:

$$R = TP / (TP + FN) \quad (3)$$

F1-score provides a balanced measure by combining precision and recall and is calculated as:

$$F1 = 2 \times (P \times R) / (P + R) \quad (4)$$

In these equations, TP denotes true positives, TN denotes true negatives, FP denotes false positives, and FN denotes false negatives. These metrics provide a comprehensive evaluation of model performance and are particularly important in fraud detection scenarios, where correctly identifying instances of the minority class is critical.

4. Results and Discussions

This section presents the experimental results obtained from evaluating classical machine learning models for insurance fraud detection under severe class imbalance. The analysis focuses on understanding how imbalance-handling techniques influence model behavior and how different classifiers perform when evaluated across multiple metrics beyond accuracy.

4.1 Dataset Distribution and Imbalance Effects

The original insurance claim dataset exhibits a highly imbalanced class distribution, with non-fraudulent claims significantly outnumbering fraudulent ones. As illustrated in Fig. 7 (Imbalanced Dataset), the majority class dominates the feature space, leading to a strong bias toward predicting legitimate claims. Such an imbalance is common in real-world insurance datasets and poses a challenge for supervised learning models, which often prioritize overall accuracy over the detection of minority classes. To address this issue, both undersampling and oversampling techniques were applied. Fig. 8 (Undersampled Dataset) illustrates that reducing the number of majority-class instances creates a balanced dataset, but this approach also results in the loss of potentially valuable information. In contrast, Fig. 9 (Oversampled Dataset) demonstrates that the synthetic generation of minority-fraud samples preserves majority-class information while improving class balance. This visual and quantitative comparison confirms that oversampling provides a more suitable foundation for fraud detection modeling.

Table 1: Data Distribution Before and After Resampling

Type	Non-Fraud Instances	Fraud Instances	Total
Original	500	50	550
Undersampling	50	50	100
Oversampling	500	100	600

The imbalance ratio in the original dataset exceeds 10:1, validating the need to apply resampling techniques before model training.

Figure 7: Imbalanced Insurance Claim Dataset

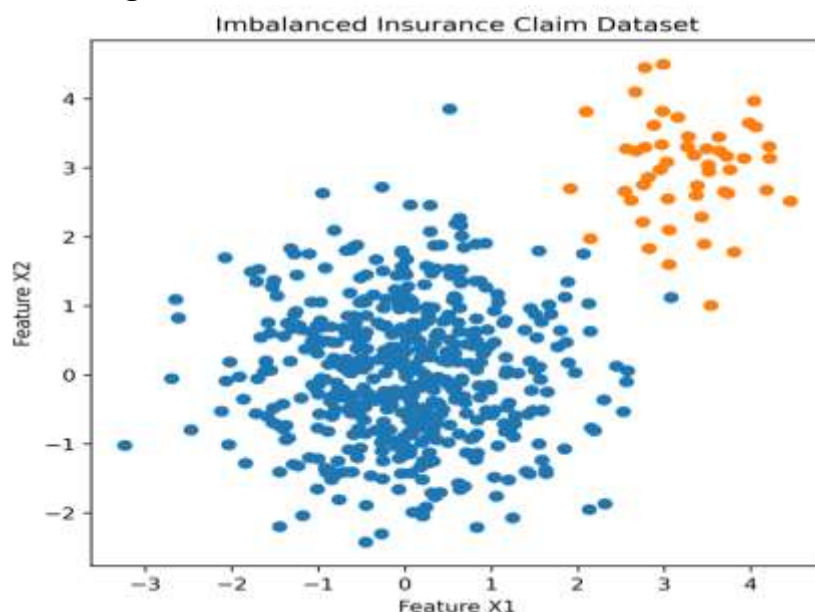


Figure 8: Under-Sampled Insurance Claim Dataset

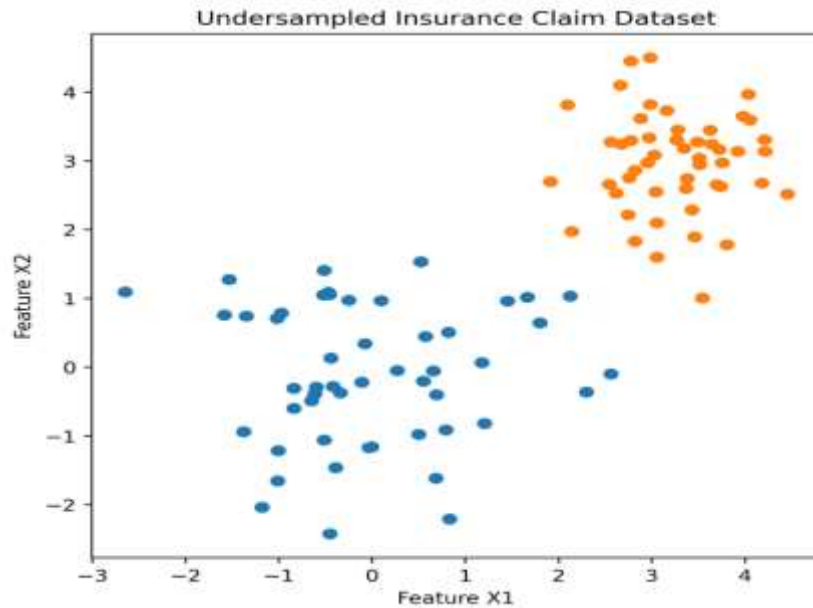
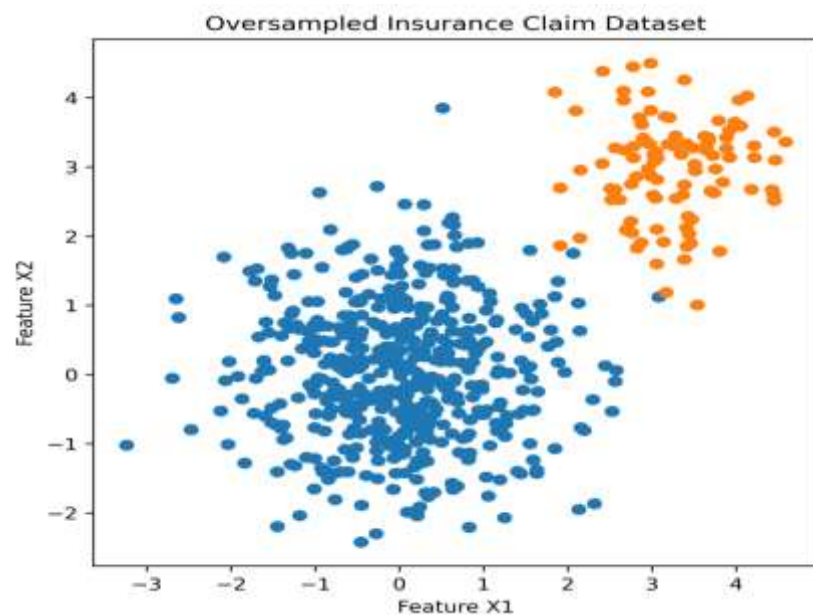


Figure 9: Over-Sampled Insurance Claim Dataset



4.2 Model Performance on Imbalanced Data

Initially, all models were trained on the original imbalanced dataset. The results indicate that while accuracy values appear high, the models perform poorly in detecting fraudulent claims. This behavior highlights the misleading nature of accuracy as a standalone evaluation metric in imbalanced classification problems.

Table 2. Model Performance on Imbalanced Dataset

Model	Accuracy	Precision	Recall	F1 Score
SVM	0.91	0.42	0.28	0.33
KNN	0.89	0.39	0.31	0.34

RF	0.93	0.48	0.35	0.40
----	------	------	------	------

Although Random Forest achieves the highest accuracy, its recall remains low, indicating that a significant proportion of fraudulent claims go undetected. These findings reinforce that accuracy alone is insufficient for evaluating the effectiveness of fraud detection systems.

4.3 Model Performance on Under-Sampling

After applying undersampling, all models demonstrate improved recall and F1 Score due to a more balanced class representation. However, this improvement comes at the cost of reduced training data size, which affects model generalization.

Table 3. Model Performance on Under-Sampling

Model	Accuracy	Precision	Recall	F1 Score
SVM	0.84	0.81	0.79	0.80
KNN	0.82	0.78	0.76	0.77
RF	0.86	0.83	0.82	0.82

Although recall improves significantly, under-sampling removes a large portion of legitimate claim data, which may not be desirable in real-world insurance environments where claim diversity is critical.

4.4 Impact of Over-Sampling Using SMOTE

Oversampling using SMOTE provides the best balance between data availability and class representation. Models trained on the oversampled dataset achieve strong fraud detection performance while maintaining stable accuracy.

Table 4. Impact of Oversampling Using SMOTE

Model	Accuracy	Precision	Recall	F1 Score
SVM	0.90	0.86	0.84	0.85
KNN	0.88	0.83	0.81	0.82
RF	0.93	0.89	0.88	0.88

Random Forest consistently outperforms SVM and KNN across all evaluation metrics. Its ensemble structure allows it to capture nonlinear relationships and complex interactions among claim attributes, which is essential for fraud detection.

4.5 Comparison Analysis of Models

This subsection presents a detailed comparison of SVM, KNN, and Random Forest models using the oversampled dataset, which represents the most effective experimental setting for insurance fraud detection in this study. The comparison is based on accuracy, precision, recall, and F1-score to ensure that both overall classification quality and minority class detection capability are evaluated in a balanced manner. Fig. 10 illustrates the computation of accuracy. Fig. 11 illustrates the computation of precision. Fig. 12 illustrates the computation of recall. Fig. 13 illustrates the computation of the F1-score.

Figure 20: Computation of Accuracy

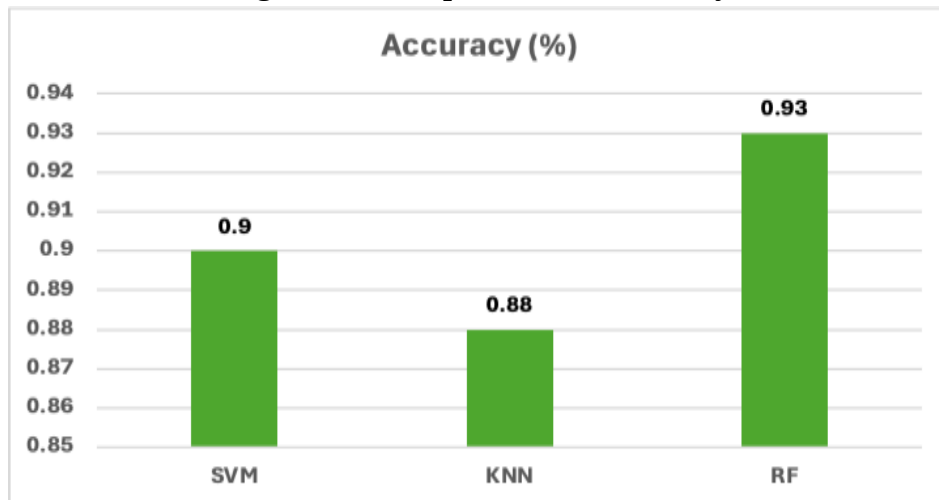


Figure 31: Computation of Precision

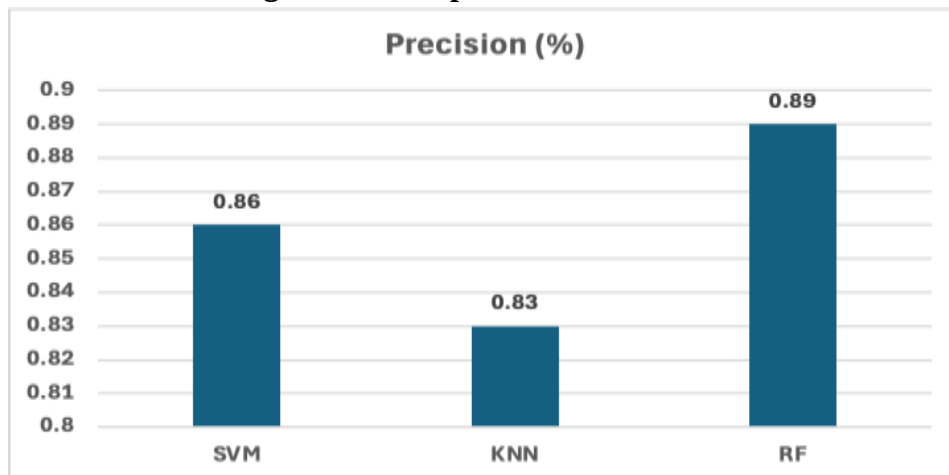


Figure 42: Computation of Recall

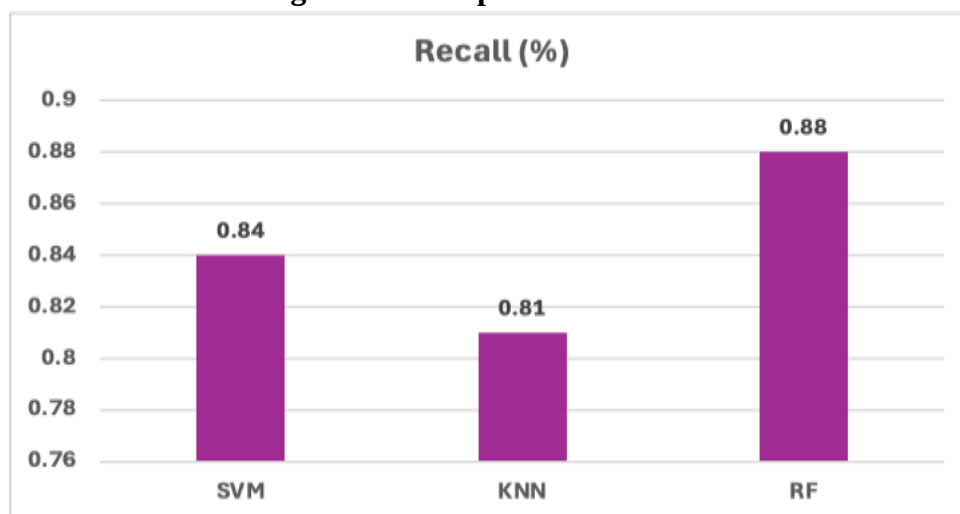
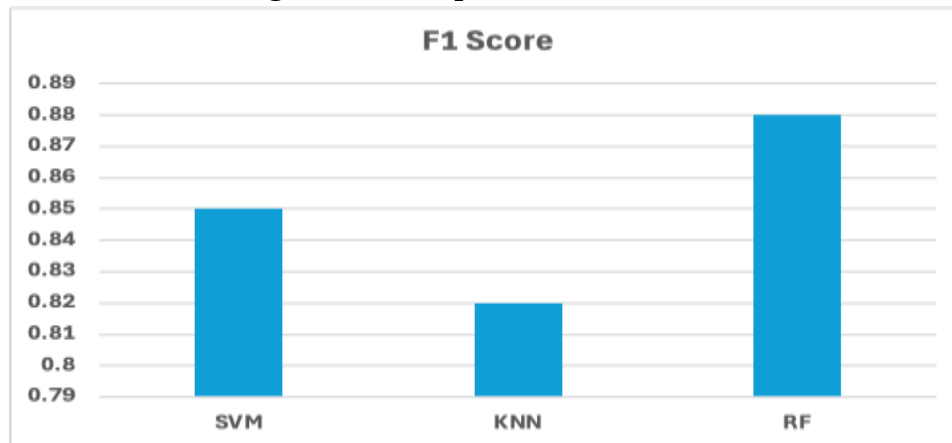


Figure 53: Computation of F1-Score



The SVM achieves an accuracy of approximately 0.90 after oversampling, indicating that it can reliably separate fraudulent and non-fraudulent claims when the class imbalance is reduced. Its precision of around 0.86 suggests that most claims the model identifies as fraudulent are indeed fraudulent, which is essential for limiting unnecessary investigations. The recall value of approximately 0.84 indicates that the model successfully detects a large portion of actual fraud cases, although some fraudulent claims remain undetected. The resulting F1 score of about 0.85 reflects a balanced trade-off between precision and recall. These results show that SVM performs well under balanced conditions, but its effectiveness depends on the assumption that a clear decision boundary can separate fraud patterns.

K-Nearest Neighbor (KNN) shows a slightly lower overall accuracy of about 0.88, with precision and recall of roughly 0.83 and 0.81, respectively. These values indicate that KNN can identify fraudulent claims based on local similarity, but it is more sensitive to noise and overlapping feature regions. The F1 score of around 0.82 indicates that, while KNN provides reasonable fraud-detection capability, its performance is less stable than that of SVM, particularly when fraud and non-fraud cases share similar characteristics. This behavior is expected, as KNN relies heavily on distance measures and local data density, which can vary across the feature space. Random Forest (RF) achieves the strongest performance across all evaluation metrics. With an accuracy of approx. 0.93, the model maintains high overall correctness while effectively detecting fraudulent claims. Its precision of around 0.89 indicates strong control over false positives, ensuring that most flagged claims are genuinely fraudulent. More importantly, Random Forest achieves a recall of approximately 0.88, the highest among all evaluated models, demonstrating its superior ability to identify actual fraud cases. The resulting F1 score of about 0.88 highlights its balanced and robust performance. The ensemble nature of Random Forest allows it to capture complex nonlinear interactions among insurance claim features, which are common in real-world fraud scenarios. The comparative analysis confirms that Random Forest provides the most reliable and consistent performance for insurance fraud detection under severe class imbalance. While SVM and KNN offer meaningful insights and competitive results after oversampling, Random Forest delivers the best balance between fraud detection accuracy and operational reliability.

5. Conclusion

This study conducted a systematic evaluation of classical machine learning models for insurance fraud detection under severe class imbalance. The performance of SVM, KNN, and Random Forest was analyzed using imbalanced, undersampled, and oversampled datasets. The results show that class

imbalance significantly affects fraud detection performance, as models trained on imbalanced data often achieve high accuracy but fail to detect fraudulent claims effectively. Applying oversampling techniques significantly improved minority class detection, leading to higher recall and F1 scores across all models. Among the evaluated classifiers, Random Forest consistently delivered the best performance due to its ensemble structure and its ability to balance fraud detection accuracy with false-positive control. While SVM and KNN benefited from imbalance mitigation, their performance was less stable. Overall, the findings underscore the importance of integrating imbalance-handling techniques with appropriate evaluation metrics to develop robust insurance fraud detection systems.

6. Limitations

Despite the encouraging results, this study has several limitations. The analysis is based on a single structured auto insurance dataset, which may not fully capture the fraud patterns present in other insurance domains, such as health, property, or life insurance, thereby limiting the generalizability of the findings. The study also focuses only on classical machine learning models, which, while interpretable and widely used, may not adequately model complex or evolving fraud behaviors. Additionally, oversampling introduces synthetic fraud instances that may not fully reflect real-world scenarios. Finally, the evaluation is conducted offline, and the lack of real-time or temporal validation limits the assessment of long-term model robustness in dynamic insurance environments.

7. Future Scope

Future research can build on this work by exploring advanced machine learning and deep learning models, such as gradient boosting and neural networks, to further enhance fraud detection in cases of severe class imbalance. Incorporating temporal and real-time claim data, including historical patterns and behavioral trends, could improve early fraud identification. Adaptive learning approaches that update models as new data becomes available also represent a promising direction. Additionally, evaluating multiple insurance datasets from different domains would enhance generalizability. Integrating structured claim data with unstructured information, such as textual descriptions, and testing models in real-world insurance environments would support the development of scalable, interpretable, and compliant fraud detection systems.

References

1. Yi, L. Z., & Ramiah, S. (2025, April). Insurance Claim Fraud Detection Using Machine Learning. In *2025 International Conference on Metaverse and Current Trends in Computing (ICMCTC)* (pp. 1-6). IEEE.
2. Zhang, L., Wu, T., Chen, X., Lu, B., Na, C., & Qi, G. (2021, December). Auto insurance knowledge graph construction and its application to fraud detection. In *Proceedings of the 10th International Joint Conference on Knowledge Graphs* (pp. 64-70).
3. Burri, R. D., Burri, R., Bojja, R. R., & Buruga, S. R. (2019). Insurance claim analysis using machine learning algorithms. *International Journal of Innovative Technology and Exploring Engineering*, 8(6S4), 577-582.
4. Ngai, E. W., Hu, Y., Wong, Y. H., Chen, Y., & Sun, X. (2011). The application of data mining techniques in financial fraud detection: A classification framework and an academic review of literature. *Decision support systems*, 50(3), 559-569.

5. Vemulapalli, G. (2024). Fighting fraud with algorithms: AI solutions for claim detection and revolutionizing fraud detection in insurance. In *Artificial Intelligence and Machine Learning for Sustainable Development* (pp. 125-140). CRC Press.
6. Mahida, A. M. (2023). Machine Learning for Predictive Observability-A Study Paper. *Journal of Artificial Intelligence & Cloud Computing*, 2(4), 1-3.
7. Debener, J., Heinke, V., & Kriebel, J. (2023). Detecting insurance fraud using supervised and unsupervised machine learning. *Journal of Risk and Insurance*, 90(3), 743-768.
8. Agarwal, S. (2023). An intelligent machine learning approach for fraud detection in medical claim insurance: A comprehensive study. *Scholars Journal of Engineering and Technology*, 11(9), 191-200.
9. Hasan, M. T. (2022). GRAPH NEURAL NETWORK MODELS FOR DETECTING FRAUDULENT INSURANCE CLAIMS IN HEALTHCARE SYSTEMS. *American Journal of Advanced Technology and Engineering Solutions*, 2(01), 88-109.
10. Banulescu-Radu, D., & Yankol-Schalck, M. (2024). Practical guideline to efficiently detect insurance fraud in the era of machine learning: A household insurance case. *Journal of Risk and Insurance*, 91(4), 867-913.
11. Malathi, M., Armash, S. M., Kavitha, S., Chinthamani, B., & Sinthia, P. (2024, January). Building Robust Fraud Detection Model for Insurance Claims Using SMOTE. In *International Conference on Universal Threats in Expert Applications and Solutions* (pp. 273-284). Singapore: Springer Nature Singapore.
12. Olushola, A., & Mart, J. (2024). Fraud detection using machine learning. *ScienceOpen Preprints*.
13. Alhuniet, Y., & Al-Khasawneh, A. (2025). Insurance Companies Fraud. In *Intelligence-Driven Circular Economy: Regeneration Towards Sustainability and Social Responsibility—Volume 2* (pp. 307-317). Cham: Springer Nature Switzerland.
14. Joginipalli, S. K., & Gummadi, V. (2024). Advancing insurance fraud detection: Leveraging machine learning and AI techniques. *International Research Journal of Modernization in Engineering Technology and Science*. https://www.researchgate.net/publication/388177103_Advancing_Insurance_Fraud_Detection_Leveraging_Machine_Learning_an_AI_Techniques/citations.
15. Ali, A., Abd Razak, S., Othman, S. H., Eisa, T. A. E., Al-Dhaqm, A., Nasser, M., ... & Saif, A. (2022). Financial fraud detection based on machine learning: a systematic literature review. *Applied Sciences*, 12(19), 9637.
16. Rukhsar, L., Bangyal, W. H., Nisar, K., & Nisar, S. (2022). Prediction of insurance fraud detection using machine learning algorithms. *Mehran University Research Journal of Engineering & Technology*, 41(1), 33-40.
17. Al Doulat, A., Ayo-Bali, O. E., & Shaik, S. (2025, July). Fraud Detection in Insurance Claims Using Supervised Machine Learning Models. In *2025 International Conference on Smart Applications, Communications and Networking (SmartNets)* (pp. 1-7). IEEE.
18. Raman, R., Kumar, V., Pillai, B. G., Rabadiya, D., Divekar, R., & Vachharajani, H. (2024, May). Detecting credit card fraud: A comparative analysis of knn, random forest, and logistic regression methods. In *2024 Second International Conference on Data Science and Information System (ICDSIS)* (pp. 1-5). IEEE.
19. Khalil, A. A., Liu, Z., Fathalla, A., Ali, A., & Salah, A. (2024). Machine Learning based Method for

Insurance Fraud Detection on Class Imbalance Datasets with Missing Values. *IEEE Access*.

20. Nabrawi, E., & Alanazi, A. (2023). Fraud detection in healthcare insurance claims using machine learning. *Risks*, 11(9), 160.