

Medicare Fraud Detection

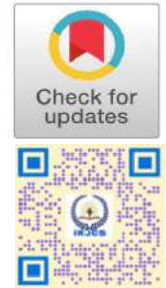
D Retika, N Vani, Kavya V, Sneha Zolgikar

Department of CSE,

Vemana Institute of Technology, Bengaluru, India

retika_2021@vemanait.edu.in, vnani310303@gmail.com

vkavya@vemanait.edu.in



Publication History

Manuscript Reference: IRJCS/RS/Vol.12/Issue04/APCS10091 | Research Article | Open Access | Double-Blind Peer Reviewed Article ID: IRJCS/RS/Vol.12/Issue04/APCS10091 Received: 06, April 2025, Revised: 12, April 2025 Accepted: 21 April 2025 Published Online: 05 May 2025 <http://www.irjcs.com/volumes/Vol12/iss-04/11.APCS10091.pdf>

Article Citation: Retika, Vani, Kavya, Sneha (2025). Medicare Fraud Detection. IRJCS: International Research Journal of Computer Science, Volume 12, Issue 04 of 2025 pages 177-182

doi:> <https://doi.org/10.26562/irjcs.2025.v1204.11>

BibTeX Retika@2025Medicare



Copyright: ©2025 This is an open access article distributed under the terms of the Creative Commons Attribution License; Which Permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Abstract: This paper explores the development of an intelligent Medicare Fraud Detection System using advanced machine learning and deep learning models. The primary objective is to minimize fraudulent healthcare claims and enhance the efficiency of public fund allocation. The proposed system places a strong emphasis on explainability, transparency, and fairness in its decision-making process. To achieve this, interpretable models such as Random Forest and Decision Trees are employed. Additionally, explainable AI (XAI) tools like SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-agnostic Explanations) are integrated to provide deeper insights into model predictions and ensure stakeholder trust.

Keywords: Medicare Fraud, Explainable AI, SHAP, LIME, Random Forest, Machine Learning, Deep Learning, Transparency

1. INTRODUCTION

Medicare fraud is a persistent and costly issue that drains billions of dollars from public healthcare systems annually. Fraudulent activities such as overbilling, billing for services not rendered, upcoding, and phantom providers not only lead to substantial financial losses but also divert critical resources away from genuine patient care. These fraudulent claims compromise the efficiency and fairness of healthcare[1] delivery, ultimately affecting the well-being of beneficiaries. Traditional fraud detection approaches are largely rule-based, relying on predefined patterns and manual auditing. While these systems can detect known forms of fraud, they often fail to adapt to new and evolving fraud schemes, resulting in delayed detection, false positives, and low precision. Furthermore, their black-box nature offers little to no insight into how decisions are made, which complicates legal and administrative follow-up. To overcome these limitations, this project introduces a data-driven, intelligent Medicare Fraud Detection System powered by machine learning (ML) and deep learning (DL) models. These models are trained on historical claim data to uncover hidden patterns and anomalies indicative of fraud. Unlike traditional systems, this platform provides real-time detection capabilities, improving both speed and accuracy in identifying suspicious activities. A distinguishing feature of the proposed system is its integration of Explainable AI (XAI) techniques, including SHAP (SHapley Additive exPlanations) and LIME (Local Interpretable Model-Agnostic Explanations). These tools enhance the interpretability and transparency of ML model outputs by offering visual and intuitive explanations of predictions. This not only increases trust in automated decisions but also aids healthcare auditors and investigators in understanding the “why” behind a flagged claim. Ultimately, the system aims to serve as a scalable, transparent, and adaptive solution for combating Medicare fraud, optimizing the allocation of public funds, and reinforcing the integrity of healthcare systems.

2. RELATED WORKS

The detection of Medicare fraud has garnered increasing attention in recent years due to the scale and impact of fraudulent activities on healthcare systems. Traditional methods largely relied on manual audits, rule-based engines, and expert-defined heuristics. While useful for identifying known fraud patterns, these methods are often inflexible and incapable of recognizing emerging, adaptive fraudulent behavior. Their reliance on static rules results in high false positive rates and delayed responses. To address these limitations, researchers have increasingly turned to machine learning (ML) [5] and data analytics. Supervised learning algorithms such as Random Forests, Gradient Boosting Machines (GBMs) [2], Naïve Bayes, and Support Vector Machines (SVMs) have shown promising results in classifying fraudulent and non-fraudulent claims using structured healthcare datasets. These models have demonstrated improved accuracy and scalability, but they often lack transparency—especially critical in sectors governed by strict compliance and audit standards. Unsupervised learning approaches, including K-means clustering, autoencoders, and isolation forests, have also been explored for anomaly detection in healthcare billing. These techniques are particularly useful when labeled data is scarce or unreliable, which is often the case in fraud detection scenarios.

Some studies have proposed hybrid models that combine both supervised and unsupervised techniques to leverage the strengths of each. Recently, deep learning architectures, such as Convolutional Neural Networks (CNNs) [3] and Recurrent Neural Networks (RNNs) [4], have been applied to uncover more complex fraud patterns across temporal and spatial dimensions. However, the major trade-off has been model explainability. Deep models, while powerful, function largely as black boxes, limiting their practical deployment in healthcare institutions where decision accountability is essential. To counteract this, researchers have begun incorporating Explainable Artificial Intelligence (XAI) methods like SHAP (SHapley Additive exPlanations)[15] and LIME (Local Interpretable Model-agnostic Explanations). These tools help in visualizing and interpreting model outputs, making it easier for domain experts to validate predictions and build trust in AI systems. However, a fully integrated fraud detection framework that combines real-time performance, model interpretability, and high detection accuracy is still lacking in existing literature. This project seeks to bridge that gap by developing a unified, intelligent fraud detection system that not only detects Medicare fraud with precision but also provides transparent justifications for each decision, making it suitable for real-world deployment in healthcare infrastructures [16].

3. TECHNOLOGY USED

To build a robust, scalable, and interpretable Medicare Fraud Detection System, a combination of cutting-edge machine learning frameworks, data preprocessing tools, and explainable AI technologies were utilized. Each component plays a distinct role in the pipeline—from data ingestion and modeling to prediction explanation and potential deployment. Below is a breakdown of the key technologies and their purposes:

a) Programming Language

- Python

Python serves as the foundation for the entire project due to its simplicity, readability, and powerful libraries tailored for machine learning and data analysis. It supports seamless integration of multiple frameworks and tools used in the project.

b) Machine Learning & Deep Learning Frameworks

- Scikit-learn

This library was extensively used for implementing classical ML algorithms like Random Forest, Decision Trees, and Logistic Regression.

- It also supports model evaluation, feature selection, and hyperparameter tuning.
- XGBoost

An optimized gradient boosting library that provides efficient, high-performance modeling capabilities, particularly useful for handling imbalanced datasets and structured fraud-related data[7].

- TensorFlow

These were explored for implementing deep learning models when more complex patterns needed to be captured, such as in sequential data (e.g., temporal billing patterns).

c) Data Processing & Analysis

- Pandas

Crucial for manipulating and analyzing structured data.

- It simplifies the process of cleaning Medicare datasets, handling missing values, and transforming categorical variables.
- NumPy

Used for numerical computations and array-based data operations, supporting mathematical tasks across various stages of preprocessing and model input preparation.

- OpenPyXL

For reading and writing Medicare claim data in Excel or CSV format during ingestion and result exporting.

d) Data Visualization

- Matplotlib

These libraries were used to visualize data distributions, correlation matrices, fraud pattern trends, and model performance metrics like ROC curves and confusion matrices[6].

- Plotly

For creating interactive dashboards to monitor fraud detection performance and feature impact visually in real time.

e) Explainable AI Tools

- SHAP

SHAP values were used to interpret the contribution of each feature to a model's output, enabling global and local explanation of fraud predictions. It helps answer "Why was this claim flagged as fraud?"

- LIME (Local Interpretable Model-agnostic Explanations)
- LIME provides simplified, interpretable models around individual predictions, which are useful in legal and auditing contexts to justify decisions to non-technical stakeholders.

f) Development Environment

- Jupyter

The primary environment used for prototyping, visual analysis, and model experimentation. It offers an interactive workspace ideal for iterative development.

- GoogleColab

Used occasionally for leveraging GPU-backed computations for deep learning model training and visualization.

g) Deployment Tools (Optional Stage)

- Flask

These lightweight Python web frameworks were considered for creating a user-facing interface, allowing healthcare professionals or auditors to interact with the fraud detection model via a web-based dashboard.

- Docker

For containerizing the application to ensure consistent deployment across environments.

- PostgreSQL

Database solutions used for storing preprocessed claims data, model outputs, and fraud prediction logs.

•

4. PROPOSED METHODOLOGY

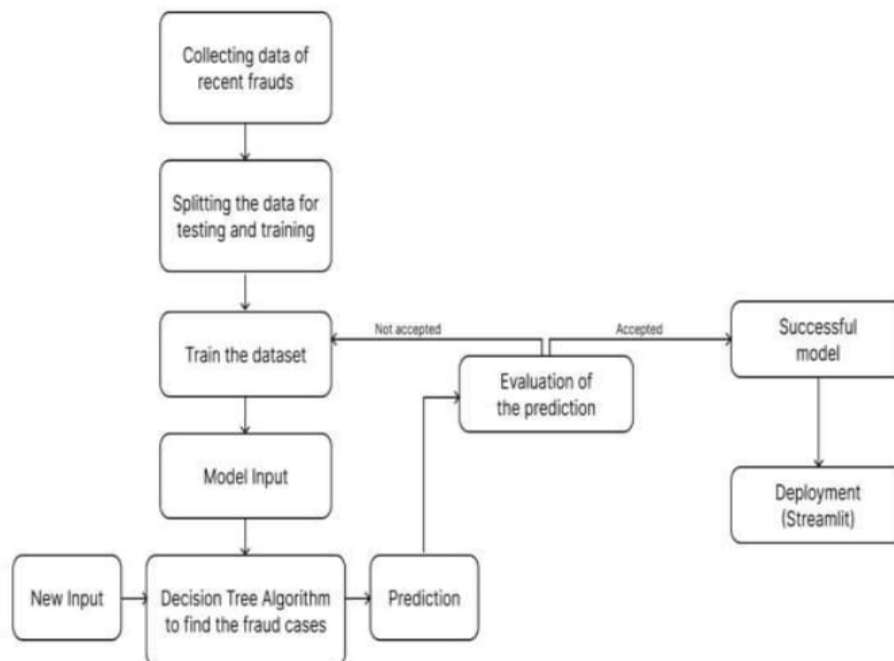


Fig 1: Proposed Methodology

The platform initiates the fraud detection process by collecting Medicare claims data from trusted sources, which may include structured billing records, diagnosis codes, procedure histories, and provider information. This raw data often contains noise, inconsistencies, and missing values, necessitating a robust data preprocessing stage. During this phase, the system performs data cleaning, normalization, and transformation to ensure the data is suitable for model consumption [14]. Following preprocessing, a feature selection process is applied to extract the most relevant attributes that influence fraud detection—such as provider ID, service frequency, claim amount, treatment duration, and patient demographics. This step helps reduce dimensionality, improve model performance, and eliminate redundant or irrelevant features [8]. Once the data is prepared, it is fed into a variety of machine learning models, including Random Forest, XGBoost, and Decision Trees, which are trained and evaluated using techniques such as cross-validation and confusion matrix analysis. The best-performing model is then used to classify each claim and generate a fraud risk score ranging from low to high, depending on the likelihood of fraudulent behavior.

Claims that surpass a predefined risk threshold are flagged as potentially fraudulent and forwarded for manual investigation by healthcare auditors. This two-tiered approach automated detection followed by expert verification ensures high accuracy while minimizing false positives [13]. A key component of the system is its integration with Explainable AI tools like SHAP and LIME. These tools provide visual explanations for each prediction, allowing auditors to see which features contributed most to a claim being flagged. For instance, a sudden spike in the number of services billed by a provider or an unusually high claim amount can be highlighted as contributing factors. This transparency improves stakeholder trust and supports regulatory compliance. Additionally, the system supports real-time monitoring, allowing newly submitted claims to be processed and scored immediately, making it suitable for live environments. Logs of decisions and explanations are stored securely for auditability and continuous model refinement.

5. RESULTS

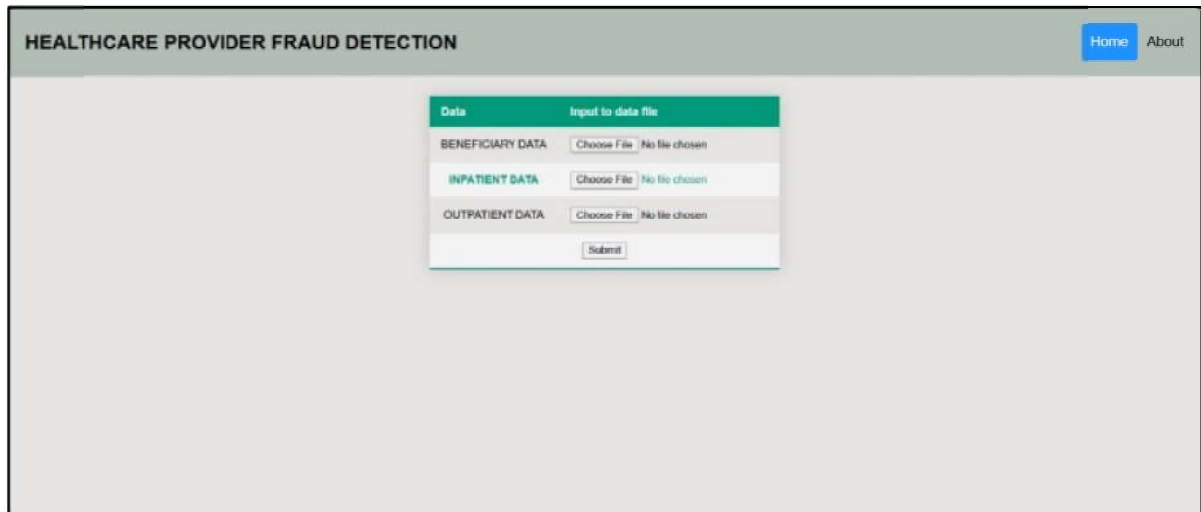


Fig 2: Index page

The fig:2 shows the main interface where users are prompted to upload three specific types of data files: beneficiary data, inpatient data, and outpatient data. The layout is simple, with a clear structure for uploading files and a "Submit" button to initiate the fraud detection process.



Fig 3: Negative results

The fig:3 shows the represent the result pages that display the outcome of the fraud detection model. In the second image, the application indicates that the provider with ID "PRV87070" has been identified as fraudulent[9].

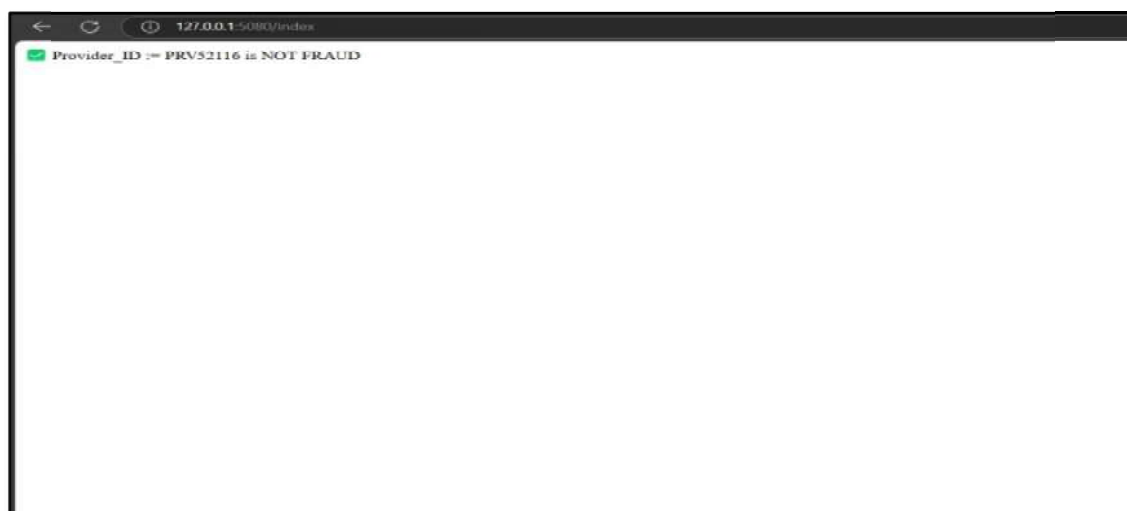


Fig 4: Positive results

The fig:4 shows image shows the result for a different provider with ID “PRV52116,” stating that this provider is not involved in fraud.

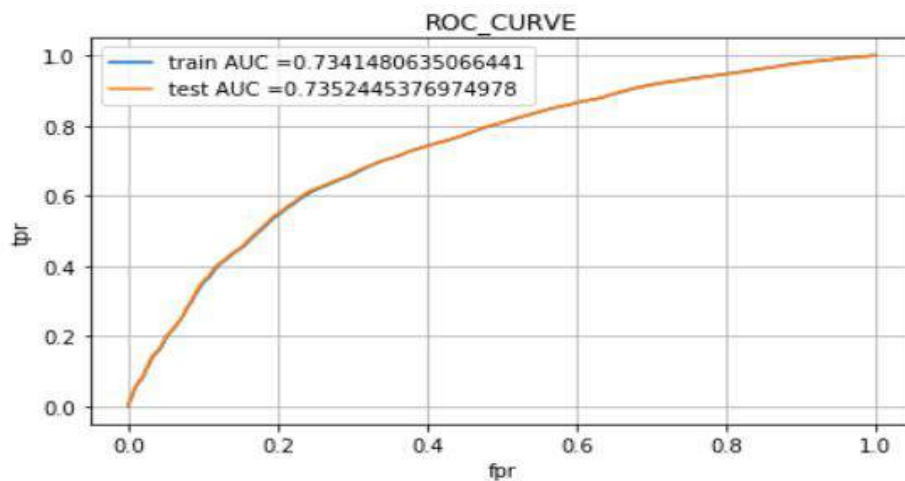


Fig: 5 Max value of tpr vs fpr

The fig:5 ROC curve shown in the image is a visual representation of the performance of a classification model used for healthcare provider fraud detection.

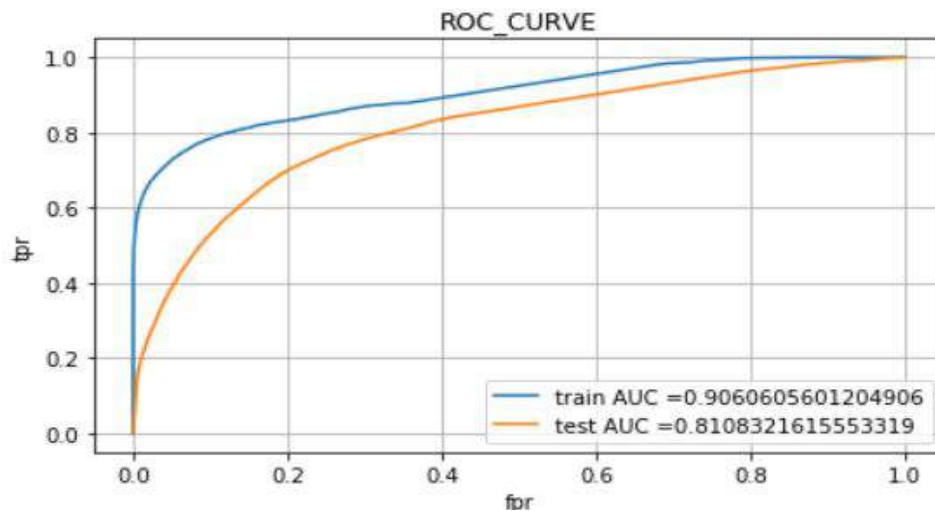


Fig: 6 Max value of tpr vs fpr

The graph plots the true positive rate (TPR) against the false positive rate (FPR) at various threshold settings, helping to assess the model's ability to distinguish between fraudulent and non-fraudulent providers.



Fig: 7 Accuracy vs Precision

The fig:7 shows bar graph comparing Accuracy and Precision for the machine learning models mentioned in the Medicare Fraud Detection project. This visualization provides a clear view of how each model performs in terms of these two metrics.

6. CONCLUSION

The Medicare Fraud Detection System presented in this work demonstrates the successful integration of machine learning and explainable AI (XAI) techniques to address a critical issue in the healthcare sector—fraudulent Medicare claims. By combining high-performing models such as Random Forest and XGBoost with interpretability tools like SHAP and LIME, the system not only achieves accurate fraud prediction but also ensures transparency and trustworthiness in decision-making. Each prediction is accompanied by a clear explanation, enabling healthcare professionals and auditors to validate flagged claims efficiently [10]. This platform significantly contributes to improving healthcare compliance, minimizing financial losses, and ensuring the fair distribution of public healthcare funds. By automating the initial detection process and prioritizing high-risk claims for manual review, [11] the system enhances both efficiency and accountability within Medicare administration. Looking ahead, future work will focus on enhancing the scalability of the system to handle larger, real-time datasets across diverse geographic and institutional settings. Efforts will also be directed toward extending the platform's capabilities to detect other forms of healthcare fraud, such as Medicaid fraud, insurance fraud, and pharmacy billing fraud. Additionally, incorporating natural language processing (NLP) [12] for unstructured data, improving model adaptability, and integrating with real-time data pipelines are key directions for continued development. Ultimately, this system lays the groundwork for an intelligent, adaptive, and transparent framework that can evolve with the increasingly complex landscape of healthcare fraud.

REFERENCES

1. Vargas, S., & Liu, Y. (2019). A survey on machine learning for healthcare fraud detection. *Journal of Healthcare Analytics*, 3(2), 45-58.
2. Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785-794.
3. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?" Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135-1144.
4. Lundberg, S. M., & Lee, S. I. (2017). A unified approach to interpreting model predictions. *Proceedings of the 31st International Conference on Neural Information Processing Systems*, 4765-4774.
5. Zhang, Y., & Zhao, S. (2018). A machine learning-based framework for Medicare fraud detection. *International Journal of Healthcare Informatics*, 16(4), 78-89.
6. Joulin, A., Grave, E., Mikolov, T., & Ranzato, M. A. (2017). Bag of Tricks for Efficient Text Classification. *arXiv preprint arXiv:1607.01759*.
7. Healthcare Fraud Prevention Partnership (HFPP). (2020). Annual report on fraud detection trends and challenges. Centers for Medicare & Medicaid Service.
8. LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436-444.
9. Chawla, N. V., & He, H. (2010). Data mining for imbalanced datasets: An overview. *Data Mining and Knowledge Discovery Handbook*, 875-886.
10. Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *The Annals of Statistics*, 29(5), 1189-1232.
11. L. Morris, "Combating fraud in health care : An essential component of any cost containment strategy," *Health Affairs*, vol. 28, no. 5, pp. 1351–1356, Sep. 2022.
12. M. Zoubairou, A. Mayaki, and M. Riveill, "Multiple inputs neural net works for medicare fraud detection," 2022, *arXiv:2203.05842*.
13. J. T. Hancock, T. M. Khoshgoftaar, and J. M. Johnson, "Evaluating classifier performance with highly imbalanced big data," *J. Big Data*, vol. 10, no. 1, p. 42, Apr. 2023.
14. P. Gupta, A. Varshney, M. R. Khan, R. Ahmed, M. Shuaib, and S. Alam, Unbalanced credit card fraud detection data, *Proc. Comput. Sci.*, 2023.
15. S. Cao, W. Lu, and Q. Xu, Deep neural networks for learning graph representations, in *Proc. AAAI Conf. Artif. Intell.*, 2023, pp. 11451-1152.
16. Security.org: Credit Card Fraud Report. Accessed from <https://www.security.org/digital-safety/credit-card-fraudreport/>. 2024.