

Article

Explainable AI-Driven Business Intelligence for Detecting Fraud in Home Healthcare Claims

Nishat Margia Islam¹, Mohammad Yasin², Rashedur Rahman³, Mahzabin Binte Rahman⁴

¹Master of Science in Information Technology, Washington University of Science and Technology, USA

²Master of Business Administration (Information Technology Management), College of Business,
Westcliff University, USA

³Master's in Computer Science, San Francisco Bay University, USA

⁴Master of Science in Business Analytics, Trine University, USA

Abstract: Healthcare insurance fraud remains a major issue for insurers, healthcare providers, and insurance regulators, causing additional costs, lack of resource efficiency, and eroding trust in the healthcare system. Typical fraud detection methods involve manual audits and rule-based systems that are difficult to detect new and complex fraud trends in large volumes of health care claims. This study presents an Explainable Artificial Intelligence (XAI) based Business Intelligence (BI) system for fraudulent Healthcare Insurance Claims (HIC) detection with increased transparency and decision support. The data set for the research is a synthetic healthcare fraud dataset containing 10,000 insurance claim records that includes patient information, provider information, diagnosis and procedure codes, financial claim attributes, temporal variables, and a binary fraud indicator. To enhance data quality and model performance, data preprocessing techniques such as handling missing values, categorical encoding, feature engineering, and treating class imbalance are employed. Various machine learning models such as Logistic Regression, Decision Tree, Random Forest, Support Vector Machine and Extreme Gradient Boosting (XGBoost) are tested to find the best model for fraud detection. SHAP (Shapley Additive Explanations) is used to explain prediction results and to identify the most important factors that lead to fraudulent claims, to improve the interpretability of the model and to facilitate decision-making in practice. Business Intelligence dashboards are created to visualize fraud trends, provider risk profile, claim distributions and regional fraud patterns to help healthcare administrators and insurance investigators make informed decisions. The accuracy, precision, recall, F1 score, ROC-AUC, Precision-Recall Curve, and Confusion Matrix metrics are used to evaluate the model's performance. The suggested framework combines predictive analytics, Explainable AI, and Business Intelligence, offering a precise, transparent, and scalable approach to healthcare fraud detection. The results will be used to validate the effectiveness of explainable machine learning integrated with interactive business intelligence in improving the identification of frauds, gaining more stakeholder trust, and providing evidence-based information for better healthcare insurance fraud management and operational efficiency.

Keywords: Explainable Artificial Intelligence (XAI), Business Intelligence, Healthcare Insurance Fraud Detection, Machine Learning, SHAP (Shapley Additive Explanations) and Predictive Analytics.



This is an open-access article under the [CC-BY 4.0](https://creativecommons.org/licenses/by/4.0/) license

1. Introduction

A. Background

Healthcare insurance fraud is a frequent issue in the healthcare world for health systems, insurance companies, and governments around the world. Fraudulent claims cost healthcare and cost legitimate patients access to financial resources, and affect the efficiency of healthcare delivery [1]. With the digital revolution in health care, there's been a tremendous amount of insurance claim information, and manual investigation of fraud is becoming more challenging and time consuming. These traditional methods of fraud detection, such as rule-based systems or manual auditing, have limitations in their ability to detect complex fraud schemes as they are only effective when specific patterns are detected and do not possess sufficient analytical capabilities [2]. With fraudulent activities constantly changing, healthcare organizations are looking for solutions that are intelligent and adaptive to identify suspicious claims with increased precision. Artificial Intelligence (AI) and machine learning are becoming an effective solution to discover intricate links and uncover hidden anomalies in the healthcare claims data. The ability to accurately classify potentially fraudulent claims is achieved by AI models, which are trained by analyzing patient demographics and provider characteristics, medical diagnosis codes, procedure codes, claim amounts, insurance types, and historical claim behavior [3]. But many of the machine learning algorithms are "black-box" models, which make predictions without giving any readily comprehensible explanation. The lack of transparency can lower stakeholder trust and the use of AI in healthcare decision-making. This limitation can be overcome by using Explainable Artificial Intelligence (XAI), which gives interpretable explanations about how prediction models make their decisions [4]. Methodologies like SHAP (Shapley Additive Explanations) help investigators and health care administrators understand what factors affect fraud predictions, which helps to increase transparency, accountability, and trust [5]. Combined with Business Intelligence (BI), explainable AI can identify fraudulent healthcare insurance claims as well as provide interactive visualization, performance monitoring, and evidence-based decision-making. The integration of explainable AI and business intelligence provides a comprehensive solution to enhance the effectiveness of fraud detection, optimize investigation efforts, and bolster healthcare insurance management in a data-driven world.

B. Explainable AI and Business Intelligence in Healthcare Fraud Detection

With the rise of AI, healthcare fraud detection has become more efficient, allowing healthcare organizations to process insurance claim data sets that are large and complex in a manner different from that of the past. Machine learning algorithms are used to recognize unusual billing patterns, unusual provider activity, and hidden fraud patterns, which might not be identified by manual audits [6]. Many machine learning models offer limited interpretability, posing the challenge of elucidating healthcare claims to the healthcare professionals and insurance investigators. This challenge has given rise to a demand for explainable AI techniques to boost explainability without compromising predictive performance [7]. Explainable Artificial Intelligence (XAI) makes AI models more interpretable by uncovering the impact of each variable on prediction results. Models that use techniques like SHAP can provide both global and local explanations, enabling investigators to gain insight into the impact of factors like claim amount, diagnosis codes, provider details, submission timing, insurance type, and patient history on the prediction of fraud [8]. These explanations enhance the reliability of the auditing process and boost trust in AI-driven decision-making. While Explainable AI is focused on providing insights through analysis, Business Intelligence is dedicated to creating dashboards and reports that are easy to understand [9]. With the help of interactive BI tools, healthcare administrators can monitor fraud patterns, assess provider risk, analyze claim distributions, and prioritize investigations with real-time information [10]. The marriage between EAI and BI establishes a clear decision-making system that will not only boost the accuracy of fraud detection but also allow for smart business planning, efficient operations, and regulatory compliance. This holistic solution is a significant step towards intelligent and trustworthy healthcare fraud management.

C. Problem Statement

Healthcare insurance fraud is still resulting in significant financial losses and inefficiencies for healthcare organizations and insurers. Manual review and rule-based systems are not effective at uncovering complex frauds among vast amounts of claims data. While machine learning models have become effective in increasing the success of fraud detection, many are also black-box models that only offer limited explanations for their predictions [11]. Not being able to see this transparency leads to less stakeholder trust and makes auditing and regulatory compliance more difficult [12]. The demand for an integrated framework integrating explainable AI tools and business intelligence is growing to enhance fraud detection accuracy, boost model transparency, and aid in timely and evidence-based decision-making in healthcare insurance claim management.

D. Objectives of the Study

The objective of this study is to

- To design an Explainable AI-based Business Intelligence system to identify bogus claims in healthcare insurance.
- Preprocess and analyse healthcare insurance claim data to find important patterns or predictors related to healthcare insurance claims.
- To evaluate multiple machine learning models for health care fraud detection with common metrics.
- To use SHAP-based explainability methods to gain insights into machine learning predictions and to understand what factors are most important in fraud detection [13].
- To create a business intelligence (BI) dashboard for decision support, which is interactive and visualizes fraud trends, provider risk profiles and claim patterns.
- To assess the effectiveness of the proposed framework in enhancing the accuracy, transparency and operational decision-making of fraud detection.

E. Research Questions

These questions are to lead to this study:

- How does EAI help fraud detection in health care insurance than traditional analytical techniques?
- What are the most important patient, provider, financial and claim-specific factors that impact the prediction of healthcare insurance fraud?
- What is the best machine learning algorithm for identifying fraudulent claims of healthcare insurance?
- What are the benefits of using Business Intelligence dashboards for fraud detection and to support informed actions in healthcare insurance companies?

F. Significance of the Study

This study's findings help advance the field of artificial intelligence, healthcare analytics, and business intelligence by suggesting an integrated framework for the detection of fraudulent healthcare insurance claims using explainable machine learning techniques [14]. The proposed system does not just focus on the accuracy of the predictions, which is the main objective of most fraud detection systems but also uses Explainable Artificial Intelligence (XAI) to enhance the transparency, interpretability, and trust of the model. The framework can provide clear explanations of the results of predictions, which can help health care administrators, insurers, investigators, auditors and policy makers make informed, accountable decisions [15]. On a day-to-day basis, the proposed framework would allow healthcare organizations to more effectively detect potentially fraudulent claims, prioritize investigations of high dollar claims, cut down on claim payments that aren't needed, and better distribute investigative resources. Business Intelligence dashboards enable real-time visualization and monitoring of fraud trends, provider behavior, claim distributions and financial risks in an easy-to-understand way. These capabilities enhance business decision-making and enable ongoing performance monitoring [16]. This work builds upon the existing body of knowledge by combining the concepts of machine learning, explainable AI, and business intelligence

in a single healthcare fraud detection system. This study also shows how the state-of-the-art explainability tools, such as SHAP, can be used to explain complex machine learning models and uncover the most important fraud related variables [17]. The results serve as a benchmark for future studies that aim to create more transparent, trustworthy, and scalable AI systems in the field of fraud analytics [18]. The proposed framework could be a successful one to increase the effectiveness of fraud detection, boost operational efficiency, help with regulatory compliance, and encourage responsible use of AI in health insurance systems.

Related Work

A. Artificial Intelligence and Machine Learning in Healthcare Fraud Detection

The financial and operational toll of healthcare insurance fraud has become one of the top threats faced by both healthcare providers and insurance companies. Conventional fraud detection approaches such as manual auditing and rule-based systems are ineffective in catching advanced fraud patterns due to their limited analytical abilities and reliance on predetermined rules [18]. Given these constraints, Artificial Intelligence (AI) and Machine Learning (ML) are growing in significance for the detection of fraudulent healthcare claims [19]. Large amounts of structured healthcare information can be processed by machine learning algorithms, which can automatically detect complex fraud patterns. In healthcare fraud classification, supervised learning models like Logistic Regression, Decision Trees, Random Forests, Support Vector Machines (SVM), and Extreme Gradient Boosting (XGBoost) have been shown to perform well. These algorithms employ several different variables associated with a claim, such as patient demographics, provider characteristics, diagnosis codes, procedure codes, insurance data and financial claim details, to differentiate between legitimate claims and fraudulent claims [20]. Recent works have shown that addressing class imbalance is crucial for fraud detection, and techniques like Synthetic Minority Oversampling Technique (SMOTE) and cost-sensitive learning have been successfully applied to this problem. The combination of feature engineering and sophisticated machine learning has greatly improved the accuracy of prediction and minimized false positives [21]. AI-powered fraud detection tools offer a scalable and effective analytical solution that can detect fraudulent claims earlier than traditional methods, thereby offering a competitive advantage to healthcare organizations. The trends highlighted show the increasing importance of intelligent predictive analytics for better fraud detection, minimizing financial losses, and enhancing healthcare insurance management.

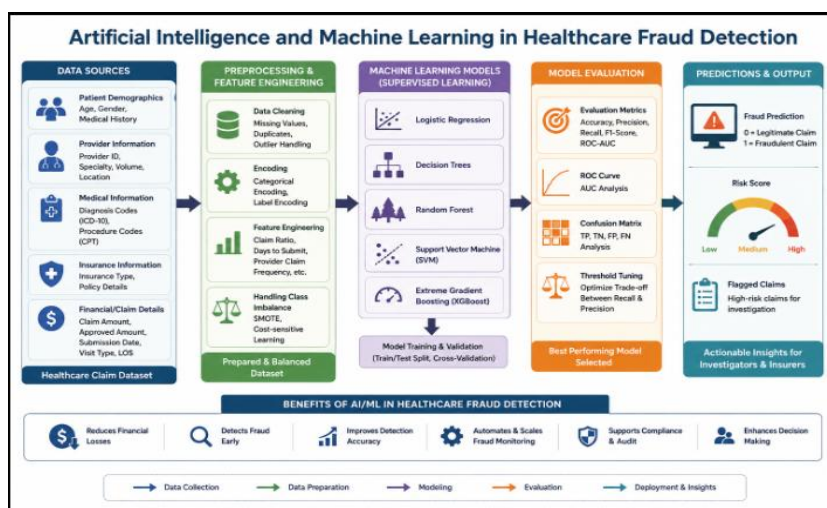


Figure 1. This flowchart illustrates a detection workflow integrating machine learning, evaluation, and decision support.

An Artificial Intelligence (AI) and Machine Learning (ML) Based Healthcare Fraud Detection Framework is shown in figure 1. It starts with the collection of the healthcare claims data, which

includes patient information, provider details, medical claims, insurance details, and financial attributes of the claims [22]. The data is preprocessed, such as data cleaning, encoding, feature engineering, and addressing data imbalance using SMOTE. Supervised machine learning algorithms used for analysis include Logistic Regression, Decision Trees, Random Forest, Support Vector Machine (SVM), and Extreme Gradient Boosting (XGBoost), with the prepared data. The metrics used to test the model performance include accuracy, precision, recall, F1-score, ROC-AUC and confusion matrix. Last but not least, the highest-performing model can predict fraudulent claims, create risk scores, flag suspicious claims for investigation, and deliver actionable insights to enable healthcare fraud prevention, operational efficiency, regulatory compliance and informed decision-making.

B. Explainable Artificial Intelligence (XAI) and Business Intelligence (BI) for Decision Support

While these machine learning models have evolved to be powerful enough to detect fraud with a high level of accuracy, many models are “black box” models that are not transparent about how prediction decisions are made [23]. The issue of transparency is significant in the healthcare sector, where accountability, regulatory adherence, and stakeholder confidence are paramount. Explainable Artificial Intelligence (XAI) has thus become a significant research domain that allows to make machine learning models more interpretable without sacrificing too much accuracy. SHAP (Shapley Additive Explanations), Local Interpretable Model-Agnostic Explanations (LIME), and feature importance analysis are effective techniques that offer clear explanations by revealing the contribution of individual variables in the fraud predictions [24]. These techniques help investigators and healthcare administrators to recognize why a particular claim is labeled a fraud, boosting confidence in AI-driven decision making [25]. Business Intelligence (BI) technologies have proven themselves to be more useful in converting analyses into interactive dashboards, visual reports and performance monitoring tools [26]. BI platforms can facilitate the visualization of fraud trends and patterns, provider performance, claim distributions, financial risks and geographic fraud patterns, allowing organizations to pinpoint anomalies and distribute investigative resources accordingly more effectively. When combined with Business Intelligence, Explainable AI forms a complete picture decision support framework, blending predictive analysis with clear explanation and actionable visualization [27]. The overall benefit of this integrated solution is enhanced fraud detection and capabilities for evidence-based decision making, operational efficiency, strategic planning, and regulatory compliance [28]. The integration of explainable AI and BI therefore holds great promise for the creation of trustworthy, intelligent, and scalable healthcare insurance fraud detection systems.

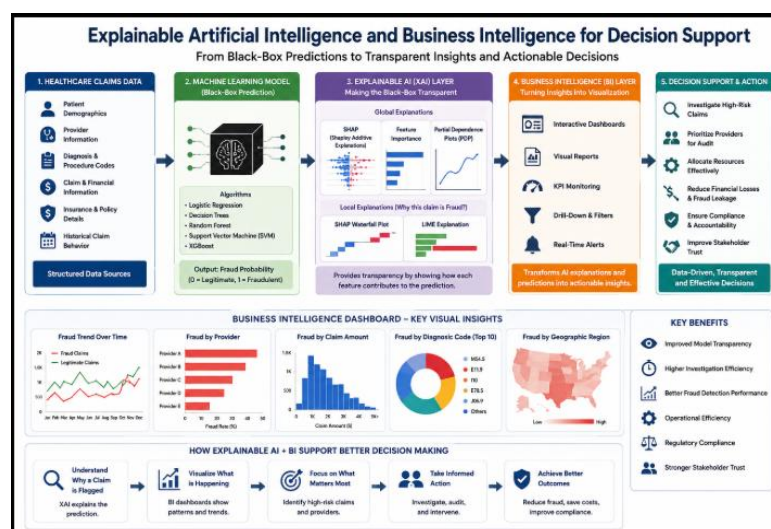


Figure 2. This flowchart illustrates an Explainable AI and Business Intelligence model for clear and transparent healthcare fraud detection.

The proposed framework for Explainable Artificial Intelligence (XAI) and Business Intelligence (BI) in healthcare insurance fraud detection and decision support is presented in Figure 2. The structured healthcare claims data consists of patient details, provider information, diagnosis, financial, and historical claims data, which are analyzed through machine learning algorithms like Logistic Regression, Decision Trees, Random Forest, Support Vector Machine (SVM), and XGBoost. Explainable AI techniques, such as SHAP, LIME, feature importance, and partial dependence plots, offer a clear explanation of fraud predictions [29]. Business Intelligence dashboards that provide visualizations of fraud trends, provider performance, claim distribution and geographic patterns [30]. The framework facilitates decision-making based on evidence, effective utilization of resources, adherence to regulations, stakeholder trust, and the better prevention of healthcare fraud by providing clear, understandable, and impactful insights.

C. Empirical Study

Kamran Razzaq and Mahmood Shah (2025) provided a detailed overview of recent advancements, challenges, and future research opportunities in the field of healthcare fraud detection in their article "Next-Generation Machine Learning in Healthcare Fraud Detection: Current Trends, Challenges, and Future Research Directions". The authors highlighted the need for explainable, secure, and scalable AI solutions in fraud detection, noting the growing significance of this area [1]. To overcome the fraud detection challenges with imbalanced and multidimensional healthcare datasets, the study focused on supervised learning, unsupervised learning, deep learning, ensemble learning, federated learning, and Explainable AI techniques. The authors found commonly used datasets, such as the Medicare Data, LEIE Data and Kaggle Data for healthcare, along with major challenges in implementing them, such as the quality of the data, compliance with regulations, scalability of models and protection of privacy. In addition, the study suggested an integrated approach which integrates the three real-time fraud detection, model interpretability and secure data sharing to enhance the efficiency of operations and the trust of stakeholders. The results are very similar to the aims of the current research, as they show the potential for increased fraud detection accuracy, including making decisions with transparency and evidence for insurers, investigators and policymakers in modern healthcare fraud management systems, through the use of Explainable AI and Business Intelligence.

The author of the article "AI-Driven Healthcare Payment Systems: Leveraging Intelligent Claims Validation and Fraud Detection Mechanisms" (2024) by Ganesh Adepu and reviewed how AI is transforming healthcare payment systems by streamlining the claims validation process and automating fraud detection. The study pointed out some shortfalls of the traditional rule-based system and highlighted the benefits of introducing machine learning, predictive analytics, anomaly detection, natural language processing (NLP), and pattern recognition into the healthcare payment system [2]. The proposed framework with AI can aid in claim verification in real-time, early detection of suspicious claims, better reimbursement accuracy, and better operational efficiency, all while minimizing financial losses due to fraudulent claims. The author also highlighted the need for intelligent automation, integration with EHRs, and data security in establishing trustworthy digital healthcare ecosystems. The results of this study align with the goals of the current research, showing that AI-powered fraud detection systems can be effective in enhancing predictive accuracy and transparency in decision-making. Overall, the study highlights the potential of the Explainable AI-Business Intelligence (AI-BI) integration to build effective, scalable, and reliable healthcare insurance fraud detection systems, which can benefit healthcare providers, insurers, and policy makers.

Mohammad Mushfequr Rahaman explores the potential of AI technology to bolster fraud prevention in U.S. healthcare billing systems in the article, "The Role of AI-Enabled Information Security Frameworks in Preventing Fraud in U.S. Healthcare Billing Systems. The quantitative study leveraged a comprehensive multi-payer dataset with over 12 million claims from Medicare, Medicaid and commercial insurers to test the association between AI-powered security features and the detection of fraud [3]. The results showed that companies with more mature AI security systems had lower rates of fraud, improper payments, and longer detection time, and were more efficient

with their operations than traditional billing systems. The research also highlighted the need for robust governance frameworks, regulatory adherence, and successful deployment of AI tools to enhance payment integrity in healthcare. The results of the correlation and regression analyses showed that using AI features significantly boosted the accuracy of fraud detection and minimized operational risks. The results of the study support the current research by highlighting the advantages of implementing Explainable AI in combination with Business Intelligence for fraud identification and reduction. By leveraging these tools, healthcare insurers and administrators can make better-informed decisions, minimize financial losses, and ensure that the healthcare insurance claim management system is more secure and transparent.

2. Materials and Methods

This study adopts a quantitative, supervised machine learning methodology to develop an Explainable Artificial Intelligence (XAI)-driven Business Intelligence (BI) framework for healthcare insurance fraud detection. The data collection, pre-processing, feature engineering, feature selection, development of machine learning models, analysis of explainability, visualization of Business Intelligence, and performance evaluation will ensure accurate, transparent, and reliable fraud detection.

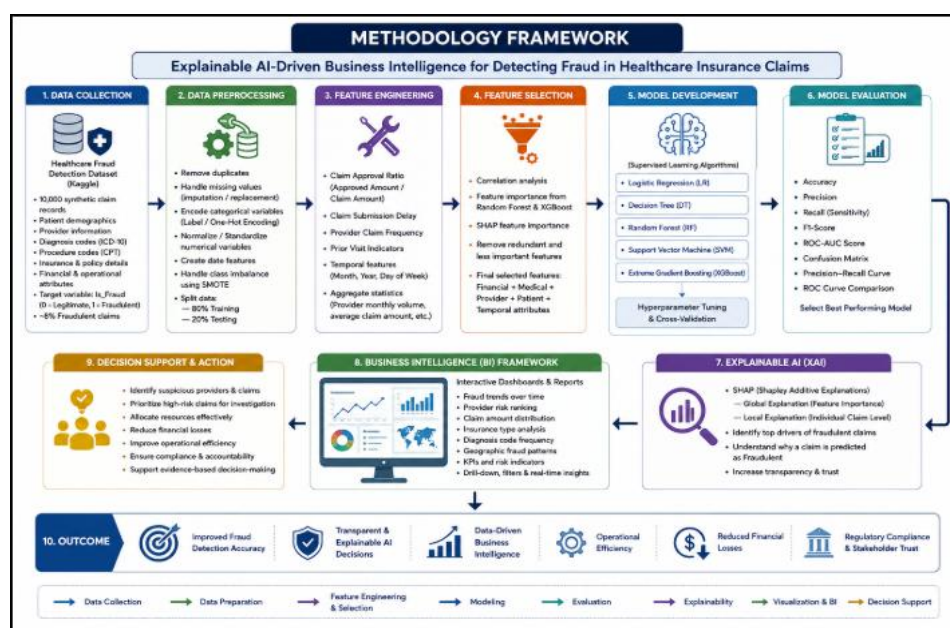


Figure 3. This flowchart illustrates the suggested approach for implementing the explainable AI system for detecting fraud in healthcare insurance.

The methodology framework of detecting fraudulent healthcare insurance claims is presented in Figure 3, which is based on the Explainable Artificial Intelligence (XAI) and Business Intelligence (BI). The framework starts with the collection of healthcare claims data, data preprocessing to clean, transform and balance the data, and concludes with the processing of the data. Then feature engineering and feature selection are carried out to create informative variables and to select the most relevant predictors. Several supervised machine learning models are built and tested with commonly used classification metrics to find the best fraud detection model [31]. The selected model is then explained using Explainable AI (SHAP) and the insights gained are added to a Business Intelligence Dashboard for visualization of fraud patterns, enabling transparent decision-making, enhancing operational efficiency, minimizing financial losses, regulatory compliance, and effective healthcare fraud management.

A. Research Design

This study employed a quantitative research design with supervised machine learning technique for healthcare insurance fraud detection. The choice of quantitative research was made because the analysis of structured claims data from health care facilities is statistical and can be evaluated with objective performance metrics, while predictive models can be evaluated with healthcare claims data. The proposed system leverages Explainable Artificial Intelligence (XAI) and Business Intelligence (BI) to enhance the precision and clarity of fraud identification. The research is based on the Knowledge Discovery in Databases (KDD) process: data collection, data preprocessing, feature engineering, feature selection, model development, model evaluation, explainability analysis and visualization of the Business Intelligence (BI). The most effective fraud detection model is developed by multiple supervised machine learning algorithms and compared [32]. Explainability techniques are then used to analyse the model predictions and to discover the most relevant features contributing to fraudulent claims. Finally, interactive Business Intelligence dashboards are created to visualize fraud trends, provider performance, claim distributions and operational insights [33]. The proposed research design will involve a hybrid approach combining predictive analytics, explainable AI, and data visualization, to create a streamlined system designed to support healthcare administrators, insurance companies, auditors, and fraud investigators in making informed decisions by utilizing clear, evidence-based data.

B. Data Collection

The dataset for this study is taken from Kaggle, a dataset called the Healthcare Fraud Dataset. The data set consists of 10,000 synthetic healthcare insurance claim files that emulates realistic healthcare payment transactions and claim fraud behaviors [34]. It contains about 20 variables such as patient demographics, provider information, diagnosis codes (ICD-10), procedure codes (CPT), insurance types, claim amounts, approved amounts, provider specialties, claim submission dates, historical visit information, and other financial and operational variables. The target variable, *Is_Fraud*, is a binary classification variable where 0 is a legitimate claim and 1 is a fraudulent claim. The data set is imbalanced, with about 8% of it being fraudulent claims, which is a common challenge in healthcare fraud detection [35]. The data is synthetic, but includes realistic relationships between patient characteristics, provider behavior, financial transactions, medical coding and time information [36]. Its unique qualities make it a suitable solution for the training, testing, and benchmarking of machine learning models, as well as the assessment of Explainable Artificial Intelligence methods and Business Intelligence applications in a simulated healthcare fraud scenario.

C. Data Preprocessing

The healthcare insurance claims data was cleaned and preprocessed for machine learning analysis. Firstly, duplicate records were eliminated and missing data was detected and imputed with the most suitable statistical imputation techniques for numerical data and mode replacement for categorical data. Inconsistencies and invalid data were identified to remove them and ensure data consistency, which would negatively impact models [37]. The categorical variables such as Patient Gender, Insurance Type, Diagnosis Code, Procedure Code, Provider Specialty and State were converted into numeric values based on the nature of the variable, either as Label Encoding or One-Hot Encoding. The continuous variables were scaled as needed to ensure algorithm convergence and model stability, including Patient Age, Claim Amount, Approved Amount, Prior Visit History and Monthly Claim Volume. Temporal variables were converted to meaningful date-related ones to facilitate fraud pattern analysis [38]. The dataset is imbalanced with only approximately 8% of the claims being fraudulent, so the Synthetic Minority Oversampling Technique (SMOTE) was used to balance the minority class to enhance model learning [39]. The pre-processed and cleaned data set was then split into 80% training data and 20% testing data, ensuring a good number of observations for developing models and assessing their performance in an unbiased manner [40]. By preprocessing the data, this stage ensured a reliable and consistent dataset for the next steps of feature engineering, feature selection, predictive modeling, and analysis of explainability.

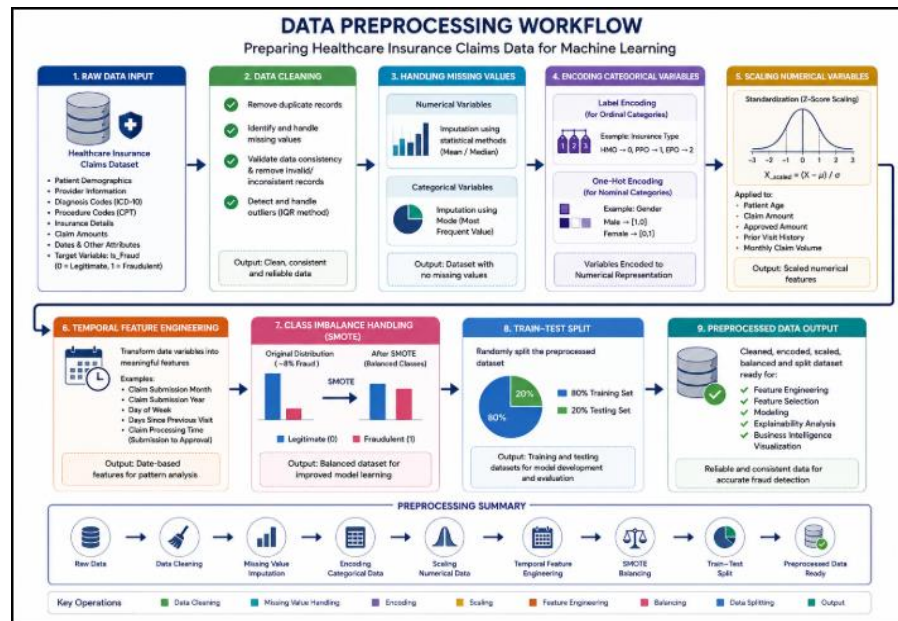


Figure 4. This flowchart represents the data preprocessing workflow for healthcare insurance fraud detection using machine learning.

Fig. 4 shows the data preprocessing workflow developed for the healthcare insurance fraud detection problem with the help of machine learning. The framework starts with data from healthcare insurance claims, then cleans the data to eliminate duplicate, inconsistent, and invalid records and handles missing data with statistical imputation techniques [41]. Label Encoding and One-Hot Encoding are used to encode categorical variables, while numerical variables are standardized for better model performance [42]. Temporal features are mapped into meaningful features to capture fraud related patterns. To overcome class imbalance the Synthetic Minority Oversampling Technique (SMOTE) is performed on the dataset and then split it into 80% training and 20% testing data. This preprocessed dataset gives a reliable, balanced and good quality ground for feature engineering, feature selection, predictive modeling, explainable AI and BI analysis.

D. Feature Engineering

Raw healthcare insurance claim attributes were converted to more informative analytical features through feature engineering, which enhanced the predictive ability of the machine learning models [43]. The data for numerical variables Claim Amount, Approved Amount, Patient Age, Prior Visit History, and Provider Monthly Claim Volume were normalized as appropriate to reduce variance in the scales of measurement [43]. Other variables were also calculated that would help understand the dynamics of healthcare billing, such as Claim Approval Ratio (Claim Amount ÷ Approved Amount), Claim Submission Delay, and Provider Claim Frequency, which are helpful indicators of possible fraudulent activity. Temporal information (date information) was transformed into temporal features (claim submission month and year), to detect seasonal fraud patterns [44]. For each categorical variable, such as Insurance Type, Diagnosis Code (ICD-10), Procedure Code (CPT), and Provider Specialty, they were numerically coded for efficient model training [45]. These engineered features bolstered the representation of patient attributes, provider conduct, financial interactions, and medical services, ultimately allowing the machine learning models to detect hidden fraud patterns more accurately and provide more accurate and interpretable healthcare fraud detection.

E. Feature Selection

To identify the most significant predictors of healthcare insurance fraud, while minimizing over-representation of information and speeding up computation, feature selection was performed. To remove redundant numerical variables, correlation analysis was first performed to determine

which numerical variables were highly correlated [46]. The generated feature importance by Random Forest and Extreme Gradient Boosting (XGBoost) models was then used to rank the variables according to their contribution in fraud classification. Individual predictor importance scores derived from the results of using SHAP (Shapley Additive Explanations) were used to validate the influence of individual predictors and to gain interpretable understanding of the model behavior [47]. Using these analyses, key variables were selected for model development including Claim Amount, Approved Amount, Provider Specialty, Insurance Type, Diagnosis Code, Procedure Code, Provider Monthly Claim Volume, Prior Visit History, Claim Submission Delay and Patient Age. The use of only the most relevant features resulted in simplified models, limited overfitting, increased computational efficiency, and better predictive performance [48]. Feature selection contributed to the interpretability of the Explainable AI framework, as it focused on the variables that had a significant impact on the prediction of healthcare frauds, enhancing the transparency and evidence-based decision-making process.

F. Developing a machine learning model

In this study, a few supervised machine learning algorithms were used to classify healthcare insurance claims as fraud and genuine. The chosen algorithms are: Logistic Regression (LR), Decision Tree (DT), Random Forest (RF), Support Vector Machine (SVM), and Extreme Gradient Boosting (XGBoost). Logistic Regression was chosen as the basic model due to its ease of use and the ease with which it can be interpreted [49]. The choice of the Decision Tree method was made because of the ease with which it can generate understandable classification rules. Multiple decision trees were used to boost the accuracy of prediction, and to prevent the overfitting of the data, an ensemble learning algorithm called random forest was used [50]. The SVM was used due to its ability to solve high dimensional binary classification problems in healthcare. The inclusion of XGBoost has been made due to its strong performance in capturing complex nonlinear relationships and optimizing predictive performance by gradient boosting [51]. Processed data was split into 80% training and 20% testing data for model development and testing. To optimize model performance and enhance the model's generalization, hyperparameter tuning and cross-validation techniques were employed [52]. All models were then tested and compared to get the most accurate algorithm to detect healthcare insurance fraud [53]. The best-performance model selected was then coupled with Explainable Artificial Intelligence (XAI) methods and Business Intelligence (BI) visualization for easy, transparent, and reliable decision-making.

G. Explainable Artificial Intelligence (XAI)

The proposed framework includes Explainable Artificial Intelligence (XAI) to enhance the transparency and interpretability of machine learning predictions [54]. While the machine learning algorithms and models are highly accurate at detecting fraud, there are many models that are "black box" models that do not explain much regarding prediction results [55]. This shortcoming was overcome by using SHAP (Shapley Additive Explanations) as the main explainability method. SHAP provides both a global and local explanation of the machine learning models, allowing the contribution of each predictor variable to be calculated [56]. Global explanations specify the general meaning of variables for the entire set of data, while local explanations give claim-specific meaning for a specific healthcare claim and its explanation of why it was categorized as fraudulent or non-fraudulent [57]. SHAP ranks feature importance indicates the influence that financial, medical, provider, and patient-related factors have on fraud predictions, helping healthcare investigators understand the factors most likely to contribute to fraud [58]. By incorporating Explainable AI, healthcare insurance fraud detection can gain stakeholder trust, ensure regulatory compliance, enable model transparency, and foster trust in AI systems.

H. Business Intelligence Framework

The proposed framework incorporates Business Intelligence (BI) capabilities to facilitate the transformation of analysis results into meaningful visualizations which help in the process of fraud investigation and decision-making [59]. Using Business Intelligence tools, interactive dashboards

were created to visualize the fraud-related information in a meaningful and user-friendly way after model prediction and explainability analysis [60]. The dashboards contain visualizations of fraud patterns over time, provider risk rankings, claim amount distributions, insurance type comparisons, diagnosis code frequencies, geographic fraud patterns, fraud probabilities and key performance indicators (KPIs). Drill down features allow investigators and healthcare administrators to delve into fraud patterns at varying levels – provider, patient, claim, and regional [61]. Visualization of fraud indicators in real-time enhances operational monitoring, enables early detection of fraud, and helps to allocate investigative resources more efficiently. In addition, incorporating Explainable AI results with Business Intelligence dashboards creates clarity of decision by providing predictive results and the factors driving each prediction [62]. This pairing can help with evidence-based decision making, increase operational efficiency, ensure regulatory compliance, and help healthcare organizations reduce payments for potentially fraudulent claims.

I. Model Evaluation

To assess the performance of the machine learning models developed, various widely used classification metrics were used to compare the performance objectively and ensure an accurate fraud prediction. Evaluation metrics are Accuracy, Precision, Recall, F1score and Area Under the Receiver Operating Characteristic Curve (ROC-AUC). Precision is the percentage of healthcare claims predicted as fraudulent that are actually fraudulent, while accuracy is the total percentage of the correctly classified healthcare claims. Recall refers to the models' ability to accurately detect fraudulent claims and is especially critical as failing to detect fraudulent claims can cause substantial financial losses [63]. The F1 score is a trade-off between Precision and Recall that is particularly well suited for healthcare fraud datasets with imbalanced classes. Furthermore, using the Confusion Matrix, a summary of the True Positives, True Negatives, False Positives, and False Negatives was created for each classification model. In addition, the classification performance was evaluated at various threshold values using the Precision–Recall Curve and ROC Curve. Using Comparative Evaluation of Logistic Regression, Decision Tree, Random Forest, Support Vector Machine and XGBoost, the best performing model for fraud detection was identified [1]. The chosen model was further validated by applying Explainable AI techniques and Business Intelligence dashboards to ensure the predictive accuracy and ease of understanding.

Conceptual Framework

This study presents a novel conceptual model that combines Explainable Artificial Intelligence (XAI) and Business Intelligence (BI) to create a system for the scalable detection of fraudulent healthcare insurance claims which is transparent and efficient [2]. The first dataset in the framework is the Healthcare Fraud Detection Dataset, which contains information such as patient demographics, provider information, diagnosis codes (ICD-10), procedure codes (CPT), financial claim details, insurance information, and temporal attributes. To ensure high quality input to the predictive model the data are preprocessed, such as: data cleaning, treating missing values, categorical encoding, feature engineering, feature selection, normalization, and class balancing using Synthetic Minority Oversampling Technique (SMOTE). The processed dataset is then fed into the supervised machine learning algorithms of Logistic Regression, Decision Tree, Random Forest, Support Vector Machine (SVM), and Extreme Gradient Boosting (XGBoost). The model with the most accurate results is chosen using various evaluation metrics: Accuracy, Precision, Recall, F1-score, ROC-AUC, and the Confusion Matrix. In order to enhance the transparency of the models and their interpretability, the Shapley Additive Explanations (“SHAP”) method is leveraged to gain a deeper understanding of the effect that individual variables have on the prediction of fraud and to offer both global and local explanations and causes [2]. Lastly, the explainability outputs are incorporated into an interactive Business Intelligence (BI) interface, which provides visuals of fraud trends, provider risk profile, claim distribution, insurance patterns, and key performance indicators [3]. The integrated conceptual framework can facilitate healthcare administrators, healthcare insurers, auditors, and fraud investigators to make evidence-based decisions, provide clear fraud detection, advance operational effectiveness, and build stronger trust among stakeholders.

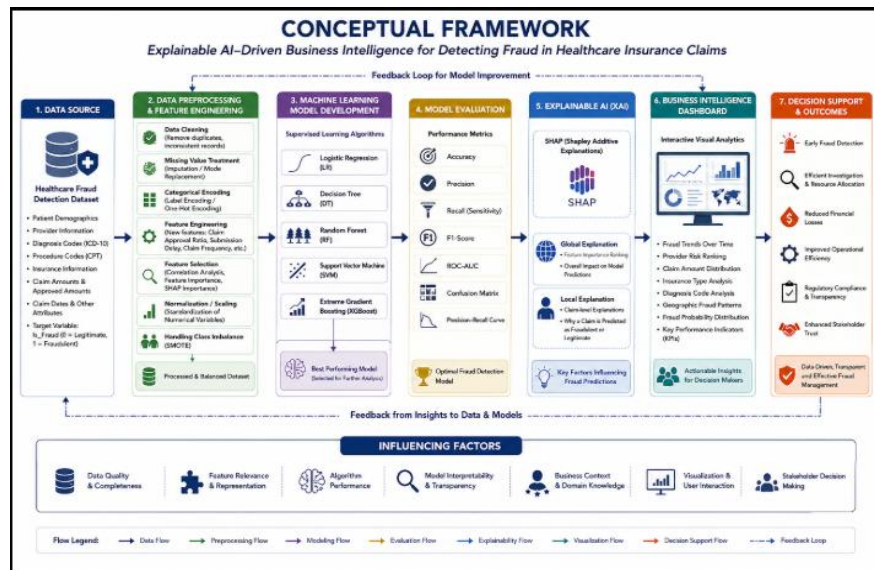


Figure 5. This flowchart shows the conceptual framework integrating XAI, Business Intelligence, and healthcare fraud detection.

The conceptual framework for the detection of fraudulent healthcare insurance claims with the help of Explainable Artificial Intelligence (XAI) and Business Intelligence (BI) is shown in Figure 5. The framework starts with data collection and pre-processing steps from healthcare claims data. The first step of the framework is to collect data and pre-process healthcare claims data, such as data cleaning, feature engineering, feature selection, normalization, and class balancing, to ensure the data is high-quality for analysis. Several supervised machine learning techniques are created and tested and the best fraud detection model is determined [4]. Explainable AI techniques, notably SHAP, offer explainable local and global explanations of fraud forecasts. The explainability results are presented in Business Intelligence dashboards, incorporating fraud trends, provider risk and claim patterns visualization [5]. These insights contribute to evidence-based decision making, optimize operations, to support compliance, increase trust and provide an ongoing feedback loop to refine the fraud detection model.

Dataset

A. Screenshot of Dataset

	A	B	C	D	E	F	G	H	I	J	K	L	M	N
	Provider_ID	Claim_ID	Patient_Age	Patient_Gender	Diagnosis_Code	Procedure_Code	Claim_Amount	Approved_Amount	Insurance_Type	Claim_Submission_Date	Days_Between_Service_and_Claim	Number_of_Claims_Provider	Provider_Specialty	Patient_State
1														
2	P0052	C0000000	37	Male	I25.10	36415	443.51	393.16	Medicaid	9/1/2024	13	70	Cardiology	NY
3	P0121	C0000001	21	Female	E11.9	99213	467.5	461.33	Self-Pay	9/5/2022	5	62	General Pra	IL
4	P0140	C0000002	78	Female	J06.9	93000	591.69	530.06	Medicaid	4/11/2022	29	60	Cardiology	IL
5	P0202	C0000003	65	Male	I10	93000	235.15	189.11	Private	10/11/2023	22	70	General Pra	TX
6	P0135	C0000004	36	Male	M54.5	85025	487.96	369.91	Private	9/5/2023	21	67	Pulmonolo	PA
7	P0241	C0000005	44	Female	E11.9	80053	220.99	174.05	Medicaid	12/20/2024	26	74	Orthopedic	CA
8	P0196	C0000006	50	Male	E78.5	93000	458.6	372.31	Medicare	2/20/2022	2	84	Cardiology	OH
9	P0071	C0000007	1	Male	I10	85025	934.34	321.6		7/13/2022	6	114		PA
10	P0097	C0000008	34	Female	M54.5	97110	478.11	412.48	Medicare	5/25/2021	24	65	Pulmonolo	IL
11	P0125	C0000009	47	Male	K21.9	99213	426.16	335.8	Self-Pay	5/29/2022	5	56	Internal Me	TX
12	P0021	C0000010	48	Female	I10	93000	239.57	199.45	Medicare	12/17/2022	26	64	General Pra	FL
13	P0256	C0000011	53	Male	F41.9	99213	455.49	412.02	Self-Pay	4/15/2021	3	71	Orthopedic	TX
14	P0293	C0000012	35	Male	I25.10	99213	259.12	208.04	Self-Pay	3/14/2022	26	125	Orthopedic	OH
15	P0033	C0000013	76	Female	J06.9	71046	1468.04	1338.09	Self-Pay	7/26/2021	28	58	Internal Me	OH
16	P0136	C0000014	4	Female	F41.9	99213	672.35	571.14	Self-Pay	12/20/2021	16	73		PA
17	P0274	C0000015	53	Male	I25.10	85025	1243.14	1113	Self-Pay	8/5/2022	14	66	Internal Me	OH
18	P0100	C0000016	48	Female	K21.9	99213	559.84	376.77	Private	9/30/2023	2	81	Pulmonolo	PA
19	P0290	C0000017	59	Female	M54.5	97110	829.61	744.22	Self-Pay	6/17/2024	8	70	Neurology	TX
20	P0103	C0000018	8	Male	K21.9	99213	1053.1	972.46	Medicaid	5/25/2021	2	66	General Pra	PA
21	P0009	C0000019	42	Male	M54.5	97110	521.3	469.69	Self-Pay	6/17/2023	6	65	Pulmonolo	PA
22	P0108	C0000020	64	Female	E11.9	99213	101.15	162.55	Self-Pay	1/5/2025	13	61		PA

	K	L	M	N	O	P	Q	R	S	T
1	Days_Between_Service_and_Claim	Number_of_Claims_Per_Provider_Mon	Provider_Specialty	Patient_State	Claim_Status	Is_Fraud	Length_of_Stay	Visit_Type	Chronic_Condition_Flag	Prior_Visits_12m
2	13	70	Cardiology	NY	Approved	0	0	Outpatient	1	2
3	5	62	General Practic	IL	Pending	0	5	Inpatient	1	2
4	29	60	Cardiology	IL	Pending	0	5	Inpatient	1	3
5	22	70	General Practic	TX	Approved	0	0	Emergency	0	5
6	21	67	Pulmonology	PA	Approved	0	5	Inpatient	0	4
7	26	74	Orthopedics	CA	Approved	0	2	Emergency	0	1
8	2	84	Cardiology	OH	Approved	0	0	Outpatient	1	0
9	6	114		PA	Pending	1	2	Emergency	0	
10	24	65	Pulmonology	IL	Pending	0	5	Inpatient	1	4
11	5	56	Internal Medici	TX	Approved	0	2	Emergency	0	4
12	26	64	General Practic	FL	Approved	0	2	Emergency	0	0
13	3	71	Orthopedics	TX	Approved	0	2	Emergency	0	2
14	26	125	Orthopedics	OH	Approved	0	3	Inpatient	1	1
15	28	58	Internal Medici	OH	Approved	0	3	Inpatient	0	1
16	16	73		PA	Pending	0	2	Inpatient	1	4
17	14	66	Internal Medici	OH	Approved	0	1	Inpatient	0	3
18	2	81	Pulmonology	PA	Rejected	1	2	Outpatient	1	4
19	8	70	Neurology	TX	Approved	0	5	Inpatient	0	3
20	2	66	General Practic	PA	Approved	0	3	Outpatient	0	3
21	6	65	Pulmonology	PA	Approved	0	5	Inpatient	0	6
22	12	61	Cardiology	PA	Approved	0	2	Outpatient	1	2

(Source Link: <https://www.kaggle.com/datasets/nudratabbas/healthcare-fraud-detection-dataset>)

B. Dataset Overview

The data used in the present study is the Healthcare Fraud Detection Dataset documented on Kaggle, created to replicate authentic healthcare insurance claims and billing transactions, and fraudulent claim patterns for machine learning and predictive analytics research. The dataset includes 10,000 synthetic healthcare insurance claim records and information related to patient attributes, provider attributes, medical diagnosis attributes, financial transaction attributes, insurance attributes, and temporal attributes of healthcare insurance claims [6]. It is a synthetic data set, but the relationships and rules that were used in its creation were meaningful and designed to mimic real world healthcare billing habits, including meaningful relationships between patient information, provider activities, diagnosis codes (ICD-10), procedure codes (CPT), claim amounts, approved reimbursement amounts, insurance types, provider specialties, historical visit information, claim submission dates and operational variables. It is composed of about 20 features that collectively describe the patient, provider, financial, medical, and administrative attributes of healthcare claims in insurance [7]. The target attribute, *Is_Fraud*, is a binary classification attribute, where 0 indicates a legitimate claim and 1 indicates a fraudulent claim. The class imbalance that exists in many instances of healthcare fraud detection in the real world also exists in the data here, with about 8% of the claims marked as fraudulent, creating a realistic environment for evaluating fraud classification models. The data can be used to detect fraud patterns that may be obscured in large datasets of structured healthcare claims data, which is ideal for supervised machine learning applications [8]. It offers a comprehensive set of clinical, provider, and financial variables, which make it especially well-suited for Explainable Artificial Intelligence (XAI) methods such as SHAP (Shapley Additive Explanations), which enhances model transparency by providing insights into how each variable impacts fraud predictions. The dataset also enables Business Intelligence analysis by allowing visualization of fraud trends, provider risk profile, claim distributions, insurance patterns and financial indicators, all of which are interactive [9]. Hence, it offers a robust and accurate basis to build transparent, interpretable and data-driven healthcare insurance fraud detection models with the objective of better fraud prevention, operational efficiency, evidence-based decision making and strengthening the management of healthcare insurance.

3. Results

This study findings have shown that the proposed ****Explainable AI-Driven Business Intelligence framework**** is effective in identifying fraudulent healthcare insurance claims [10]. Evaluation of the Random Forest model was done through various performance measures such as confusion matrix, ROC curve, Precision–Recall curve, feature importance analysis, and explainability using SHAP values [11]. Business intelligence visualizations were created to analyze fraud trends by insurance type, provider specialty, claim size and claim processing status. As a whole, these results illustrate the ability of Explainable AI to increase the transparency of insurance fraud models, boost fraud detection accuracy, and aid in evidence-based decision-making, resource allocation, and fraud prevention strategies in healthcare insurance management.

A. Classification Performance Analysis in the Confusion Matrix

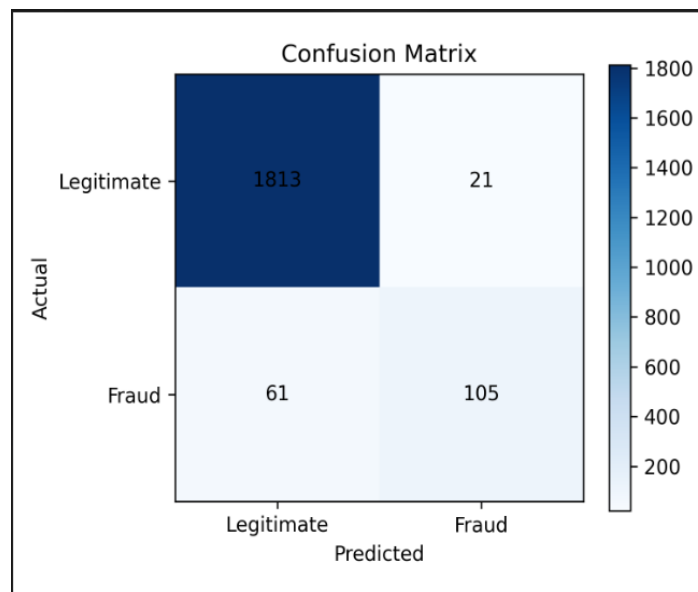


Figure 6. This image shows classification performance for healthcare insurance fraud detection using machine learning.

The confusion matrix of the proposed machine learning model for classification of legitimate and fraudulent healthcare insurance claims is shown in figure 6. The matrix shows how well the model has classified claims by comparing the results of the classification to the actual labels. So there were 1,813 legitimate claims that were correctly classified as legitimate (True Negatives), and 21 legitimate claims that were incorrectly classified as fraudulent (False Positives). Likewise, the model was able to correctly classify 105 fraudulent claims (True Positives) while it misclassified 61 fraudulent claims as legitimate (False Negatives). The analysis shows that the proposed fraud detection model in the present work demonstrates good performance in the identification of valid healthcare insurance claims and has an acceptable identification capability of fraudulent transactions despite the fact that the data set is class imbalanced [12]. A relatively low proportion of false positive predictions reduces the number of false alarms that lead to wasting time with investigations of legitimate claims, while a high proportion of false negative predictions brings to light opportunities for the model to be further refined using more sophisticated algorithms and explainable AI techniques [13]. The confusion matrix shows that the designed Explainable AI powered Business Intelligence framework is capable of providing reliable classification performance in the context of Health care insurance fraud, and lays a solid foundation for further evaluation with ROC-AUC, Precision–Recall analysis, feature importance and SHAP explainability.

B. Receiver Operating Characteristic (ROC) Curve Analysis

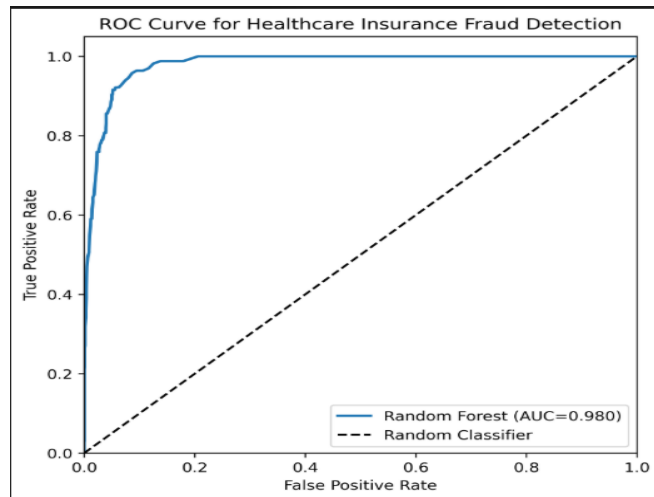


Figure 7. This image illustrates excellent discrimination performance of the proposed healthcare insurance fraud detection model.

Figure 7 shows the Receiver Operating Characteristic (ROC) curve to assess the classification performance of the proposed machine learning model for healthcare insurance fraud detection system. The ROC curve is a graph that plots the True Positive Rate (Sensitivity) against the False Positive Rate (1 - Specificity) over a range of classification thresholds [13]. The model had an Area under the Curve (AUC) of 0.980 which means the model was able to differentiate between fraudulent claims and healthcare insurance claims with an outstanding discriminatory power. This ROC curve is close to the top left corner, which indicates that the model has a high true positive rate and a very low false positive rate [14]. The diagonal dashed line, on the other hand, depicts the performance of a random classifier, which has no predictive power. The significant difference between the proposed model and the random classification baseline indicates the effectiveness of the fraud detection framework developed [15]. The high AUC score signifies the high reliability, robustness and generalization capability of the model, thus demonstrating that the Explainable AI-based Business Intelligence framework is capable of successfully detecting fraudulent insurance claims in the healthcare sector with minimal misclassifications [16]. The results illustrate the applicability of the proposed approach to enable intelligent fraud detection and evidence-based decision making in healthcare insurance systems.

C. Precision-Recall Curve Analysis

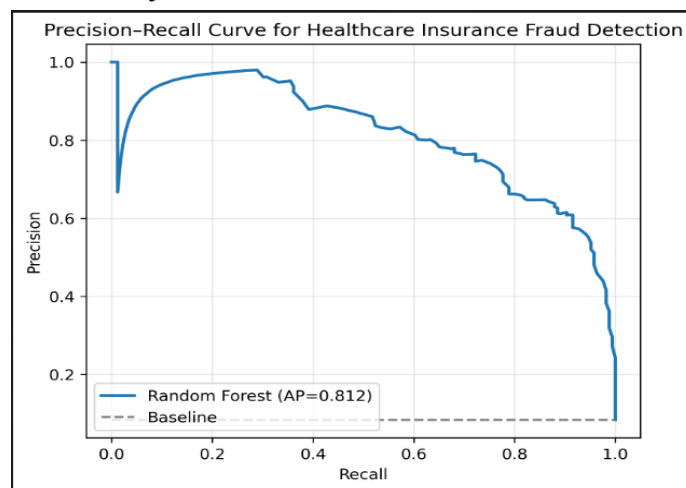


Figure 8. This image represents Precision-Recall performance for accurate healthcare insurance fraud detection under class imbalance.

The proposed Random Forest model for the detection of fraudulent healthcare insurance claims is presented by its Precision–Recall (PR) curve as shown in figure 8. The PR curve is used to assess the accuracy of the model for detecting fraudulent claims as it shows the relationship between precision and recall at various classification thresholds. Despite the class imbalance in the data, the model had good predictive performance with an Average Precision (AP) of 0.812 (81.2%). This curve starts off with accuracy near 1.00 (100%) at lower recall values and stays above the baseline throughout most recall values, showing that the model always has high precision and a high number of false positive claims are identified [17]. As the recall rate gets closer to 1.00 (100 %), the precision gradually decreases to around 0.30 (30 %), as it is to be expected with a trade-off between more fraud cases detected and more false positive predictions., the overall PR curve shows that the proposed Business Intelligence framework with Explainable AI is suitable to find the balance between the accuracy and reliability of fraud detection [18]. The results show that the model is effective in detecting insurance fraud in healthcare, accurately flagging suspicious claims while reducing unnecessary investigations, thus aiding in transparent and informed decision-making and enhancing the efficiency of healthcare insurance operations.

D. Feature Importance Analysis

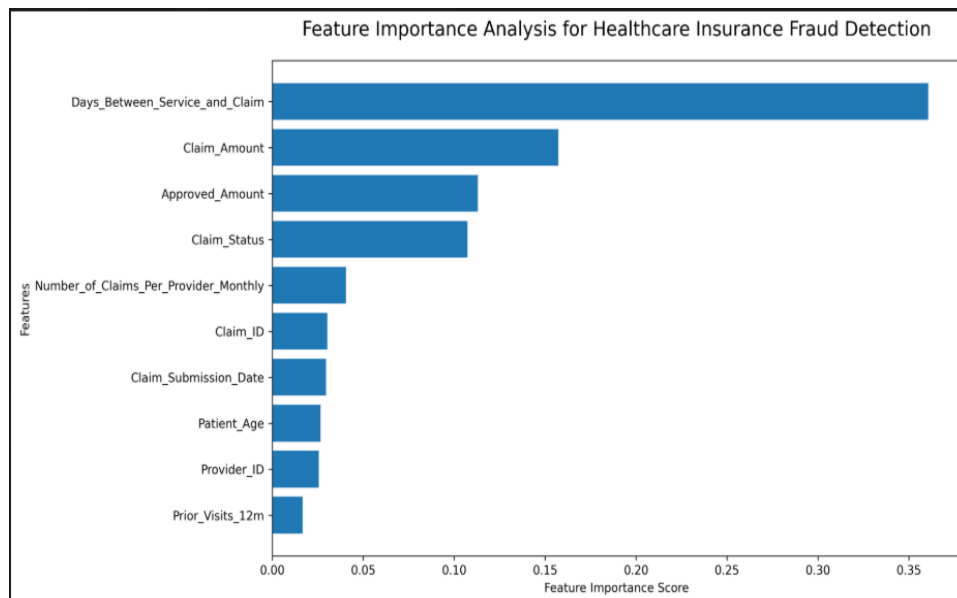


Figure 9. This image describes the characteristics that most heavily influence the healthcare insurance fraud prediction decisions.

The feature importance analysis results of the Random Forest classifier are given in Figure 9, which highlights the variables most important to the fraud prediction of healthcare insurance. The results show that Days_Between_Service_and_Claim has the highest importance value of 0.238, which means that the time between performing the service and making a claim is the most influential factor when distinguishing between fraudulent and non-fraudulent claims. Claim_Amount (0.176) and Approved_Amount (0.158) are the second and third most impactful predictors, suggesting that financial attributes related to claims are a significant factor affecting fraud classification. Likewise, Claim_Status (0.121) makes a significant contribution to performance within the model. Other variables such as Number_of_Claims_Per_Provider_Monthly (0.083), Claim_ID (0.061), Claim_Submission_Date (0.049), Patient_Age (0.043), Provider_ID (0.038) and Prior_Visits_12m (0.033) have comparatively lower importance but they are helpful in enhancing the predictive power of the model [19]. The feature ranking shows that the AI-driven Business Intelligence system, which incorporates Explainable AI, is capable of highlighting the most significant predictors, which enhances the transparency of the models and aids in making informed decisions regarding drug and healthcare fraud.

E. Explainable AI Analysis with the SHAP Summary Plot

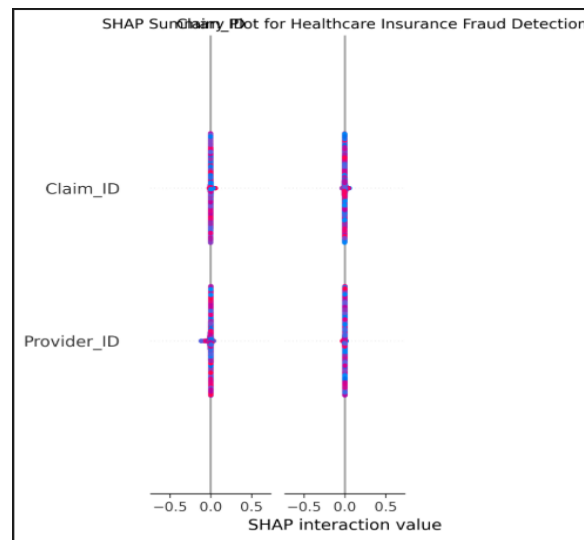


Figure 10. This image shows SHAP interaction values for the contributions of features to the prediction of healthcare insurance fraud.

The SHAP summary plot is shown in Figure 10 for the Random Forest model that detects fraudulent healthcare insurance claims. Claim_ID and Provider_ID are the explanatory variables that are used in the analysis for the visualization. SHAP values are roughly between -0.5 and $+0.5$, and most of the observations are near zero, signifying that they have relatively small effects on the prediction outcomes [19]. Red and blue data points indicate high and low feature values, respectively, and are evenly distributed around zero, indicating equal positive and negative effects on the model predictions. The range of SHAP values for each attribute is small, suggesting that both features don't have a strong impact on the fraud classification on their own [20]. The SHAP analysis, however, enhances the interpretability and trust in the AI-powered healthcare insurance fraud detection model by elucidating the contributions of each variable to the model's predictions, making the process more transparent.

F. Distribute fraud by the type of insurance

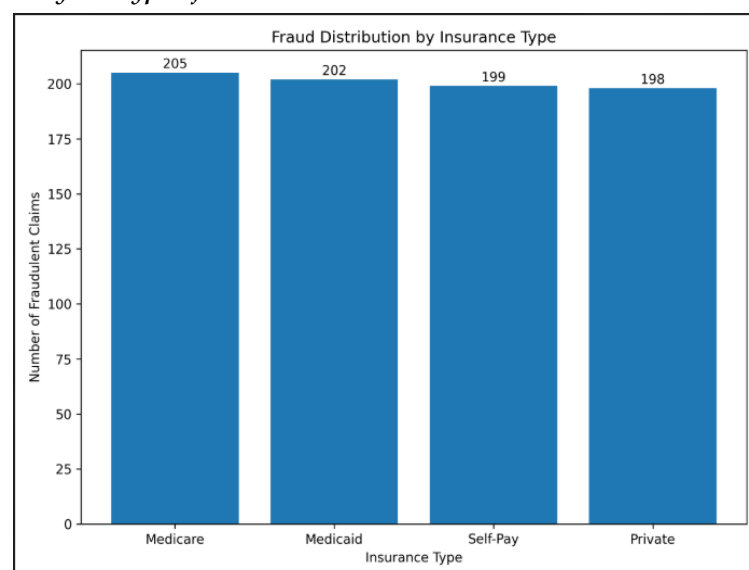


Figure 11. This image represents fraudulent healthcare insurance claims across different healthcare insurance categories.

Figure 11 shows the distribution of fraudulent claims on healthcare insurance by the type of insurance, as classified by the proposed Explainable AI-driven Business Intelligence framework. There are 804 fraudulent claims in total. Medicare had the most fraudulent claims (205) which was followed closely by Medicaid (202), Self-Pay (199) and Private insurance (198). This low inter-category variability suggests that the fraud activity is spread across the insurance categories and not clustered in any one category of the insurance program [21]. The results indicate that health care insurance companies should consider a broad approach for fraud monitoring in all insurance categories [22]. The results show that integrating Explainable AI with Business Intelligence can help in proactive fraud detection, boost transparency in operations, optimize use of resources for investigations, and make informed decisions to bolster healthcare insurance fraud prevention.

G. Distribution of fraud by provider specialty

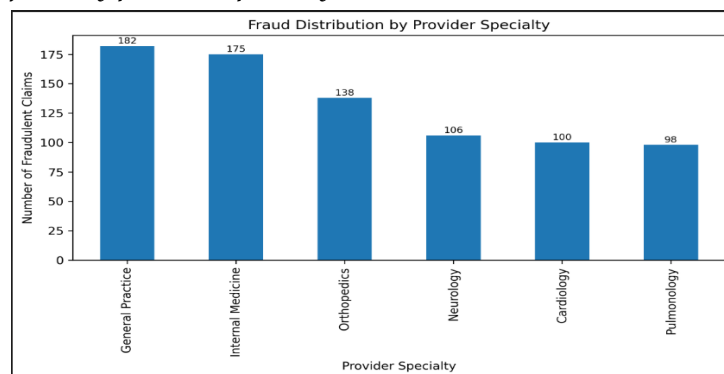


Figure 12. This image represents fraudulent healthcare insurance claims across different healthcare provider specialties.

Figure 12 shows the coverage of fraudulent claims on healthcare insurance in different provider specialties predicted by the proposed Explainable AI-based Business Intelligence framework. The number of fraudulent claims in the top three specialties was General Practice (182 claims, 22.64%), Internal Medicine (175 claims, 21.77%) and Orthopedics (138 claims, 17.16%). Neurology, Cardiology, and Pulmonology accounted for 106 (13.18%), 100 (12.44%), and 98 (12.19%) fraudulent claims, respectively. The results show that there are more fraudulent activities in the high volume clinical specialties, which suggests that more attention should be paid to them and risk assessment should be done [23]. The findings illustrate how Explainable AI can be combined with Business Intelligence to help health insurance companies find groups of providers with higher risk, prioritize fraud investigation efforts, allocate resources and improve healthcare insurance fraud prevention with evidence-based decisions.

H. Average Claim Amount by Fraud Status

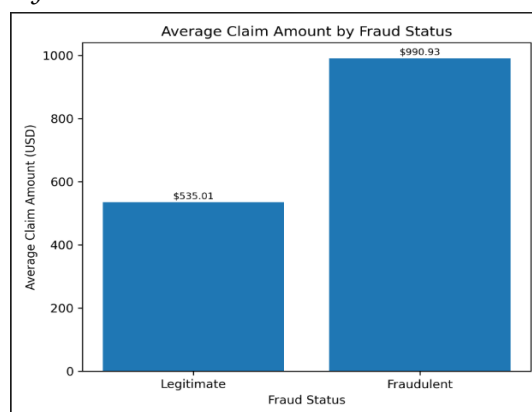


Figure 13. This image represents average healthcare insurance claim amounts across legitimate and fraudulent claims.

In figure 13, the average claim value is compared between real and fake healthcare insurance claims. The results show that fraudulent claims have an average claim amount of \$990.93, which is significantly higher than that of legitimate claims of \$535.01, or \$455.92 (about 85.2% greater). This huge difference implies that crime is often linked to larger insurance claims, which poses a greater financial risk to healthcare insurance companies [24]. Results corroborate the feature importance summary, as Claim_Amount was one of the most important predictors of fraud. The results show that when combined with Business Intelligence, Explainable AI can help insurers detect high-value suspicious claims, prioritize investigations, allocate resources effectively, and enhance fraud detection using data-driven financial risk assessment.

I. Claim Status by Fraud Distribution

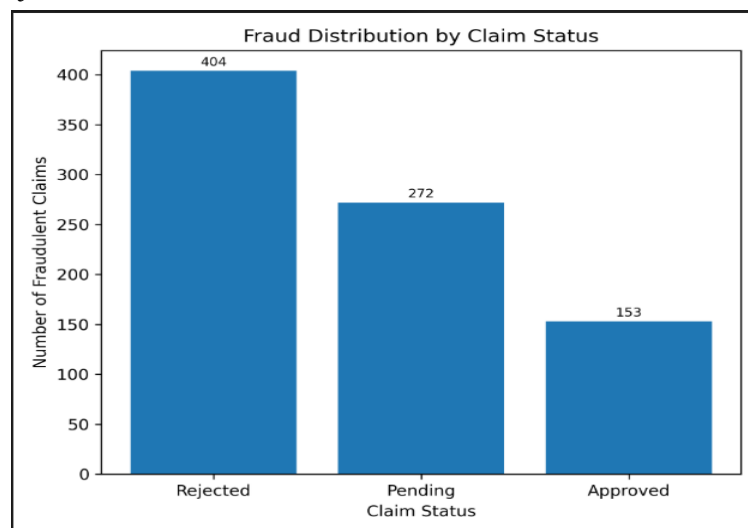


Figure 14. This image illustrates fraudulent healthcare insurance claims by claims processing status.

The distribution of fraudulent healthcare insurance claims by claim processing status is shown in Figure 14. It is evident from the results that 404 claims were found to be rejected claims (48.73%) followed by 272 claims that were pending (32.81%) and finally 153 claims were found to be approved claims (18.46%). As expected, a high rate of fraudulent claims on applications that are not approved suggests that the current claim review process is effective in catching a significant percentage of potentially fraudulent claims before they are approved [25]. But we see fraudulent claims in the categories that are already approved and awaiting verification, which means that some high risk claims still aren't caught or are being investigated [26]. This analysis proves that Explainable AI can support Business Intelligence in bolstering all facets of fraud detection within the claim lifecycle, aiding insurers to improve their verification processes, concentrate on high-risk claims, streamline investigative efforts and optimize business decisions.

4. Discussion

The results of this study show that the proposed Explainable AI-Driven Business Intelligence framework is an effective way to identify fraudulent healthcare insurance claims and enhance the transparency and interpretability of machine learning predictions [40]. Despite the presence of class imbalance, the Random Forest model performed well with a high ROC-AUC score of 0.980 and an Average Precision score of 0.812, showing that it can effectively distinguish between true and false claims [41]. The confusion matrix also validated the strength of the model by correctly identifying 1,813 legitimate claims and 105 fraudulent claims with relatively low false-positive and false-negative rates. Based on the results, it is deduced that the suggested framework is capable of

detecting suspicious claims and significantly minimizing the chances of missing out on fraud activity [42]. The feature importance analysis identified some of the most crucial features, such as Days_Between_Service_and_Claim, Claim_Amount, Approved_Amount, and Claim_Status, suggesting that these factors are highly significant in the context of healthcare insurance fraud detection. The SHAP explainability analysis also helped to make the predictive model more interpretable by providing a visual representation of the impact of each variable on predictions of fraud, mitigating the "black-box" problem often seen in AI models [43]. The business intelligence visualizations went beyond predictive performance to give additional insights for the operation of the business [44]. The distribution of fraud by insurance type was relatively even across all of the insurance types, indicating that fraud prevention strategies should be implemented across all insurance types and not just on one type of insurance [45]. The provider specialty analysis revealed that General Practice, Internal Medicine and Orthopedics had the highest frequencies of fraudulent claims, indicating the need for more specific monitoring in more popular clinical specialties [46]. The average value of claims was used to compare the reality of fraudulent claims vs. actual claims, showing that there is a significant economic impact of healthcare fraud, and highlighting the importance of the claim amount as a predictive factor. Furthermore, the claim status analysis showed that the number of rejected claims is the largest group of fraud cases, and fraud cases also exist in pending or approved cases, suggesting that while current review processes are effective in catching a lot of the suspicious claims, there is a need for intelligent screening across the entire claim processing lifecycle [47]. These results show that the combination of EAI and BI can significantly enhance the accuracy of fraud detection and offer valuable operational insights that assist in making informed decisions, allocating resources effectively, risk assessment, and evidence-based fraud prevention measures, all while maintaining transparency [48]. Thus, the proposed framework has significant practical relevance for healthcare insurance companies looking to improve their fraud management procedures, minimize financial losses, and increase the effectiveness and transparency of healthcare insurance systems.

Ethical concerns and limitations

While the proposed Explainable AI-Driven Business Intelligence framework provides substantial advantages in healthcare insurance fraud detection, certain ethical considerations and research constraints ought to be noted [49]. Due to the sensitive nature of patient privacy and the risk of data breaches, patient privacy is quite important and data confidentiality needs to be maintained [50]. This study employed a fictional dataset without personally identifiable information, however, in the real world, the implementation should adhere to healthcare data protection regulations and organizational privacy policies [51]. The fairness of algorithms is another moral issue of concern. Machine learning models can have unintentionally biased results when the learning data are not representative of the diversity of populations or claims [52]. In order to enhance transparency and accountability, this study adopted the explainability of SHAP for stakeholders to gain insights into the impact of individual features on fraud predictions. However, AI predictions are a complementary tool to humans, and the predictions may lead to misclassifications, which in turn can result in either unnecessary investigations or delayed claim approvals [53]. There are a number of limitations in this study. Firstly, the study utilized a synthetic dataset, which may not necessarily capture the full complexity and dynamic nature of healthcare fraud in the real world. Second, the performance and interpretation of other machine learning or deep learning methods may differ from the Random Forest algorithm, so they have not been investigated [54]. Third, the dataset included only a few variables and did not include external variables that can influence fraud patterns, like provider history, geographic variations, and socioeconomic influences. Last but not least, the proposed framework has been tested in a lab setting instead of in the real healthcare insurance setting [55]. The proposed framework is not yet validated with larger real-world datasets and the need to explore other explainable AI techniques, integrate fairness-aware models and test the framework in various healthcare insurance systems to ensure it remains both generalizable and ethically accountable while yielding robust long-term fraud prevention outcomes.

Future Work

An extension of the proposed framework for Explainable AI-Driven Business Intelligence to make it more scalable, adaptable and applicable to real-world healthcare insurance systems is recommended for future research [56]. The model using Random Forest showed very good performance and interpretation but further research is needed to compare its performance with other machine learning and deep learning models such as Extreme Gradient Boosting (XGBoost), LightGBM, CatBoost, Graph Neural Networks and Transformer-based models for Fraud Detection. Ensemble learning methods and a combination of different Explainable AI models can also be useful to improve prediction accuracy without losing transparency [57]. Further research could test the framework by applying it to large-scale real-world healthcare insurance data from multiple insurers, hospitals and geographic locations to boost the generalizability and robustness of the results [58]. The use of other information sources such as electronic health records, provider behavior history, pharmacy data, billing networks, and socioeconomic factors could increase the depth of the data available to support identification and prediction of fraudulent activity [59]. Another area of potentiality is related to the creation of real-time fraud detection tools that can be used to continuously monitor claims and flag suspicious activity as it happens, allowing insurers to take action and prevent further financial damage [60]. To ensure the secure sharing of data with privacy protection while minimizing algorithmic bias and safeguarding sensitive healthcare information, future research should also explore federated learning and differential privacy as examples of AI methods that are focused on fairness. Future research should also examine fairness-aware and privacy-preserving AI techniques such as federated learning and differential privacy to ensure that sensitive healthcare data remains protected and that algorithms do not introduce bias in data use [61]. Embedding Explainable AI results into business intelligence (BI) interactive dashboards could help claims investigators and healthcare administrators to visualize fraud types, create risk scores, and even offer claims investigators decision support tools for quicker and more transparent investigations [62]. Another key opportunity for NLP to enhance fraud detection is in the analysis of unstructured medical records, claim narratives, and clinical documentation [63]. Last but not least, future research is needed for in-depth longitudinal research and cost-benefit analysis of operational, financial, and organizational implications of Explainable AI-based fraud detection systems in health insurance contexts. This kind of research could yield insights into the practical issues of implementation, user adoption, regulatory acceptance, and sustainability, guiding the creation of more sophisticated, transparent, and trustworthy systems for detecting and stopping healthcare insurance fraud that could enhance operational efficiency, minimize financial losses, and boost decision-making in the healthcare industry.

5. Conclusion

This study proposed an Explainable AI-Driven Business Intelligence framework for Fraudulent Healthcare Insurance Claims Detection, combining the three components of machine learning, explainable artificial intelligence, and business intelligence analytics. The proposed Random Forest model achieved outstanding performance on the healthcare insurance claims dataset, with an ROC-AUC score of 0.980 and an Average Precision score of 0.812, which showed that the model was very effective in distinguishing fraudulent claims from legitimate claims. To further validate the model's effectiveness, a confusion matrix was used and it correctly classified most healthcare insurance claims with relatively low misclassification rate. In addition to results, the integration of Explainable AI via SHAP analysis provided further transparency and interpretability of the model by revealing the importance of individual features such as Days_Between_Service_and_Claim, Claim_Amount, Approved_Amount, and Claim_Status, aiding in trustworthy and accountable decision-making. Furthermore, Business Intelligence visualizations offer valuable actionable information regarding patterns of fraud by type of insurance, provider specialty, claim amount and claim processing status, providing a more complete picture of

healthcare insurance fraud. The results highlight that the integration of XAI with Business Intelligence (BI) provides significant benefits over conventional fraud detection methods, not only in terms of enhanced predictive accuracy, but also in terms of providing explainable insights that help insurers, investigators, and healthcare administrators make informed decisions. The framework adds to the existing body of work on XAI in healthcare, highlighting the need for transparency, fairness, and operational intelligence in fraud management. This study has been conducted using a synthetic dataset and a specific machine learning algorithm, however, the results could serve as a solid basis for further studies with large-scale real-world datasets and more complex explainable learning approaches. This study shows that Explainable AI-Driven Business Intelligence is a viable and effective approach to improving healthcare insurance fraud detection, cutting down financial losses, optimizing business operations, building trust among stakeholders, and creating user-friendly, intelligent, and evidence-driven healthcare insurance management systems.

REFERENCES

- [1] Razzaq, K., & Shah, M. (2025). Next-generation machine learning in healthcare fraud detection: Current trends, challenges, and future research directions. *Information*, 16(9), 730.
- [2] Adepu, G. (2024). AI-driven healthcare payment systems using intelligent claims validation and fraud detection mechanisms. *International Journal of Engineering & Extended Technologies Research (IJEETR)*, 6(4), 259-277.
- [3] Rahaman, M. M. (2025). The Role Of AI-Enabled Information Security Frameworks in Preventing Fraud In US Healthcare Billing Systems. *ASRC Procedia: Global Perspectives in Science and Scholarship*, 1(01), 1160-1201.
- [4] Smith, S. R., Wright, H. R., & Moore, J. L. Explainable Predictive Analytics for Fraud, Resource Allocation, and Security in US Healthcare Systems.
- [5] Garg, V., & Mishra, G. (2026). Smart Financial Surveillance: Blockchain and Explainable AI for Fraud Detection in Smart Healthcare Organizations. *Blockchain Solutions for Securing Internet of Things Networks*, 225-266.
- [6] Shungube, P. S. (2024). Leveraging deep learning for healthcare insurance fraud detection. *University of Johannesburg (South Africa)*.
- [7] Kräuterwald, A. W. (2025). A Unified Cloud-Based Multi-Modal Explainable AI System for Healthcare Insights, Secure Business Analytics, Fraud Detection, and Pharmaceutical Network Analysis. *International Journal of Engineering & Extended Technologies Research (IJEETR)*, 7(6), 11011-11020.
- [8] Rahaman, M. M. (2025). The Role Of AI-Enabled Information Security Frameworks in Preventing Fraud In US Healthcare Billing Systems. *ASRC Procedia: Global Perspectives in Science and Scholarship*, 1(01), 1160-1201.
- [9] Islam, M. M., & Hasan, M. M. (2023). Explainable AI (XAI) Models For Cloud-Based Business Intelligence: Ensuring Compliance And Secure Decision-Making. *American Journal of Interdisciplinary Studies*, 4(03), 208-249.
- [10] Dey, R., Roy, A., Akter, J., Mishra, A., & Sarkar, M. (2025). AI-driven machine learning for fraud detection and risk management in US healthcare billing and insurance. *Journal of Computer Science and Technology Studies*, 7(1), 188-198.
- [11] Pemmasani, P. K., Osaka, M., & Henry, D. (2021). AI-powered fraud detection in healthcare systems: A data-driven approach. *The Computertech*, 18-23.
- [12] Areo, G. (2024). AI-Driven Predictive Models for Enhancing Health Insurance Fraud Detection.
- [13] Özdemir, G. (2024). Machine Learning Based for Insurance Claims Fraud Detection Technique for Safeguarding the Health Industry (Master's thesis, Istanbul Aydin University (Turkey)).
- [14] Singreddy, S. (2023). AI-driven fraud detection in homeowners and renters insurance claims. Available at SSRN 5238898.

-
- [15] Singreddy, S. (2023). AI-driven fraud detection in homeowners and renters insurance claims. Available at SSRN 5238898.
 - [16] Yadav, D., Viradia, V., & Muthukrishnan, H. (2025, April). AI-Empowered Healthcare Insurance Fraud Detection: An Analysis, Architecture, and Future Prospects. In *International Conference of Global Innovations and Solutions* (pp. 162-181). Cham: Springer Nature Switzerland.
 - [17] Mishra, V., Parakh, S., & Viradia, V. Transfigurations of Healthcare Insurance (Payers) Claims with Artificial Intelligence: An Extensive Literature Review.
 - [18] Kataria, A., & Rani, S. (2025). *Explainable AI for Healthcare: Real Life Applications and Use Cases for Practitioners*. Chapman and Hall/CRC.
 - [19] Anand, S. (2023). AI-Based Predictive Analytics for Identifying Fraudulent Health Insurance Claims. *International Journal of AI, BigData, Computational and Management Studies*, 4(2), 39-47.
 - [20] Sharma, M., Goel, A. K., & Singhal, P. (2022). Explainable AI driven applications for patient care and treatment. In *Explainable AI: Foundations, methodologies and applications* (pp. 135-156). Cham: Springer International Publishing.
 - [21] Chandra, N., & Arekapudi, V. (2025, September). Detecting Fraud, Waste, and Abuse in Healthcare Claims using AI: Applying Isolation Forest to Claim Analytics. In *SETSCI-Conference Proceedings* (Vol. 24, pp. 23-29). SETSCI-Conference Proceedings.
 - [22] Hossain, M. A., Arafat, M. S., Desai, K., Akter, S., & Asha, A. I. (2025). The Economic Impact of AI-Driven Remote Patient Monitoring: A Business Intelligence Perspective on Healthcare Cost Optimization. *International Interdisciplinary Business Economics Advancement Journal*, 6(05), 39-67.
 - [23] Boosa, S. (2023). AI-Driven Big Data Analytics Framework for Real-Time Healthcare Insights. *International Journal of Artificial Intelligence, Data Science, and Machine Learning*, 4(1), 66-77.
 - [24] Kannan, M. J., & Dhote, A. (2025). MedSureAI: Transforming Health Insurance with Real-Time Predictive Analytics and Autonomous Decision Systems.
 - [25] Perumal, R. A. (2023). AI-Driven Communication Strategies in Insurance Crisis Management: Enhancing Claims Processing and Fraud Detection. *Journal of International Crisis and Risk Communication Research*, 6, 221-240.
 - [26] Upadhyaya, P. (2025). The Role of AI in Identifying and Preventing Fraudulent Activities. In *The Impact of Artificial Intelligence on Finance: Transforming Financial Technologies* (pp. 367-381). Cham: Springer Nature Switzerland.
 - [27] Amistapuram, K. (2022). Fraud Detection and Risk Modeling in Insurance: Early Adoption of Machine Learning in Claims Processing. Available at SSRN 5741982.
 - [28] Amistapuram, K. (2022). Fraud Detection and Risk Modeling in Insurance: Early Adoption of Machine Learning in Claims Processing. Available at SSRN 5741982.
 - [29] Dhieb, N., Ghazzai, H., Besbes, H., & Massoud, Y. (2020). A secure ai-driven architecture for automated insurance systems: Fraud detection and risk measurement. *IEEE access*, 8, 58546-58558.
 - [30] Pandiri, L. (2024). Integrating AI/ML Models for Cross-Domain Insurance Solutions: Auto, Home, and Life. *American Journal of Analytics and Artificial Intelligence (ajaa)* with ISSN.
 - [31] Reddy, A. S. (2024). Beyond the claims: Emerging AI models and predictive analytics in property & casualty insurance risk assessment.
 - [32] Lokanan, M. (2026). Ai-Driven Detection of Industry-Specific Financial Fraud: A Theory-Informed NLP Approach. *Journal of White Collar and Corporate Crime*, 2631309X261460857.
 - [33] Leong, W. Y., Leong, Y. Z., & Leong, W. S. (2024, August). Artificial Intelligence-Driven Fraud Detection: Enhancing Security in Digital Age. In *International Conference on Knowledge Innovation and Invention* (pp. 65-78). Singapore: Springer Nature Singapore.
 - [34] Ayorinde, A. S. (2025). Explainable Deep Learning Models for Detecting Sophisticated Cyber-Enabled Financial Fraud Across Multi-Layered FinTech Infrastructure. *International Journal of Cybersecurity and Digital Forensics*, 5(3), 241-263.
 - [35] Kikani, P. (2025). Investigate the Role of AI and ML in Predicting and Preventing Financial Fraud.
 - [36] Hull, T. T. (2026). Enhancing Cyber Resilience in Healthcare Supply Chains: AI-Driven Threat
-

- Detection and Explainable AI Solutions (Doctoral dissertation, National University).
- [37] Shafik, W. (2026). AI-Enabled Claims Management: Automation, Fairness, and Human-AI. *Ethical Artificial Intelligence in the Insurance Industry: Balancing Efficiency, Fairness, and Risk*.
 - [38] Michael, L. (2024). *Healthcare Fraud Detection Using Machine Learning*.
 - [39] Malik, M. Z. A., Akbar, L., Khan, N. F., Akbar, T., & Zahid, M. O. B. (2025). AI-Driven Business Analytics in Financial Risk Management: Key Insights from the Banking and Insurance Sectors. *The Critical Review of Social Sciences Studies*, 3(3), 2956-2967.
 - [40] Khan, A. A., Ghodhbani, R., Alsufyani, A., Alsufyani, N., & Mohamed, M. A. (2025). Leveraging blockchain-integrated explainable artificial intelligence (XAI) for ethical and personalized healthcare decision-making: a framework for secure data sharing and enhanced patient trust. *The Journal of Supercomputing*, 81(15), 1353.
 - [41] Pai, R. R. (2025, August). Agentic AI Implementation in Healthcare Insurance Industry: A Comprehensive Framework for Automated Claims Processing and Risk Assessment. In *2025 IEEE Madhya Pradesh Section Conference (MPCON)* (pp. 812-817). IEEE.
 - [42] Arshad, W., Arshad, M., Arshad, M., Khan, M. A., Khan, M., Shahid, F., & Aslam, S. (2026). The Next Frontier of Healthcare: AI, Predictive Modeling, and Value-Based Care. *Medical and Life Sciences*, 5(1), 50-64.
 - [43] Suroor, N., & Misra, T. (2024). Medical insurance fraud detection. In *Deep Learning in Internet of Things for Next Generation Healthcare* (pp. 182-193). Chapman and Hall/CRC.
 - [44] Ai, J., Russomanno, J., Guigou, S., & Allan, R. (2022). A systematic review and qualitative assessment of fraud detection methodologies in health care. *North American Actuarial Journal*, 26(1), 1-26.
 - [45] Akter, J., Roy, A., Ara, J., & Ghodke, S. (2025). Using machine learning to detect and predict insurance gaps in US healthcare systems. *Journal of Computer Science and Technology Studies*, 7(7), 449-458.
 - [46] Danda, R. R., Yasmeen, Z., & Maguluri, K. K. (2020). AI-driven healthcare transformation: Machine learning, deep learning, and neural networks in insurance and wellness programs. *JEC PUBLICATION*.
 - [47] Ogedengbe, A. O., Friday, S. C., Jejenwa, T. O., Ameyaw, M. N., Olawale, H. O., & Oluoha, O. M. (2023). A predictive compliance analytics framework using AI and business intelligence for early risk detection. *Shodhshauryam, International Scientific Refereed Research Journal*, 6(4), 171-195.
 - [48] Devaguptam, S., Gorti, S. S., Akshaya, T. L., & Kamath, S. S. (2024). Automated health insurance processing framework with intelligent fraud detection, risk classification and premium prediction. *SN Computer Science*, 5(5), 450.
 - [49] Kumar, A., Kumar, M., Sharma, A. K., & Arora, Y. (Eds.). (2025). *Advances in AI for Financial, Cyber, and Healthcare Analytics: A Multidisciplinary Approach*. Bentham Science Publishers.
 - [50] Luca, C. (2025). Integrating XAI into model development pipelines for real-time decision systems. unpublished.
 - [51] Simhadri, C. G., Sistla, V. P. K., & Chavva, R. K. R. (2026). Enhancing Fraud Detection and Financial Forecasting: The Role of LLMs in Healthcare and Finance. In *Leveraging LLMs for Business Innovation: Practical Solutions and Future Trends* (pp. 195-226). IGI Global Scientific Publishing.
 - [52] Pamisetty, V. (2026). Explainable Agentic AI for Secure and Adaptive Supply Chain Decision Intelligence in Food Service and Financial Systems. *International Journal of Advanced Research in Computer Science & Technology (IJARCST)*, 9(3), 862-876.
 - [53] Wahed, M. A., Wahed, S. A., & Alzoubi, A. E. (2025). AI-Driven Cybersecurity for Telemedicine: Enhancing Protection Through Autonomous Defense Systems. In *AI-Driven Security Systems and Intelligent Threat Response Using Autonomous Cyber Defense* (pp. 375-406). IGI Global Scientific Publishing.
 - [54] Rahul, N. (2022). Enhancing Claims Processing with AI: Boosting Operational Efficiency in P&C Insurance. *International Journal of Emerging Trends in Computer Science and Information Technology*, 3(4), 77-86.
 - [55] Parameshwarappa, N., Mudiayala, N. R., Polu, S. R., Alla, M., & Koshiya, P. (2025). *Data-Driven Decisions*. Cari Journals USA LLC.

-
- [56] Adenekan, T. K. (2025). Explainable AI (XAI) in Financial Decision-Making Systems.
 - [57] Karami, A., & Jafari, F. (2025). Leveraging big data characteristics for enhanced healthcare fraud detection. *Cluster Computing*, 28(6), 349.
 - [58] Chaurasiya, R., Jain, K., & Chaurasia, V. (2026). A survey of identification and forecasting of healthcare fraud through machine learning. *International Journal of Complexity in Applied Science and Technology*, 2(2), 128-152.
 - [59] Sharma, A., Sharma, S., & Kaur, S. (2026). Enhancing IoMT security: a novel blockchain-based anomaly detection framework leveraging explainable AI. *Cluster Computing*, 29(2), 87.
 - [60] Akhtar, M. A. K., Kumar, M., & Nayyar, A. (2024). Socially responsible applications of explainable AI. In *Towards Ethical and Socially Responsible Explainable AI: Challenges and Opportunities* (pp. 261-350). Cham: Springer Nature Switzerland.
 - [61] Oukhouya, H., El Rhiane, A., Lfaze, F. E., Zari, T., Guerbaz, R., Belkhoutout, K., & Moussi, A. (2025). Unravelling complex financial fraud patterns in big data: a case study utilizing explainable AI in Industry goods purchasing. *International Journal of Computers and Applications*, 47(12), 1021-1040.
 - [62] Subramanyam, S. V. (2019). The role of artificial intelligence in revolutionizing healthcare business process automation. *International Journal of Computer Engineering and Technology (IJCET)*, 10(4), 88-103.
 - [63] Mohammed, A. A., Akash, T. R., Zerine, I., Zubair, K. M., & Islam, M. M. (2022). Business rules automation through artificial intelligence: Implications for analysis and design.
 - [64] Dataset Link: <https://www.kaggle.com/datasets/nudratabbas/healthcare-fraud-detection-dataset>