

Reliability-Aware Multimodal Intelligence for Healthcare Insurance Fraud Detection: A Trimodal Data Perspective

Nasir Nazir¹, Roma Fayaz², Dr. Sharyar Wani³

^{1,2,3}International Islamic University of Malaysia, KICT
nasir.nazir@live.iium.edu.my, f.roma@live.iium.edu.my, sharyarwani@iium.edu.my

Abstract: Healthcare insurance fraud has emerged as a significant challenge due to the rapid growth of digital healthcare systems, increasing volumes of electronic insurance claims, and the evolving complexity of fraudulent activities. Conventional fraud detection approaches primarily rely on manually defined rules, single-modal data, or traditional machine learning algorithms, which often struggle to capture complex relationships among heterogeneous healthcare information and provide reliable predictions. To address these limitations, this study proposes a Reliability-Aware Multimodal Intelligence Network (RAMI-Net) for healthcare insurance fraud detection from a trimodal data perspective. Initially, healthcare insurance claims are preprocessed and transformed into three complementary perspectives: patient structured data, claim-related textual information, and healthcare relationship graphs. These perspectives are processed using TabNet, ClinicalBERT, and Graph Attention Network (GAT) to extract modality-specific feature representations. A Reliability Estimation Module evaluates the confidence of each modality before a Reliability-Aware Cross-Modal Attention Fusion mechanism integrates reliable features while minimizing noisy information. The fused representations are classified using an Explainable Fraud Detection Network, and SHAP provides transparent interpretations by identifying influential fraud-related features. Experimental results demonstrate that the proposed framework achieves 99.95% accuracy, 99.98% precision, 99.97% recall, 99.95% F1-score, and 99.63% ROC-AUC, outperforming conventional machine learning and deep learning approaches. The proposed RAMI-Net provides an accurate, robust, reliable, and interpretable solution for intelligent healthcare insurance fraud detection, supporting effective real-world healthcare insurance claim management.

Keywords: Healthcare Insurance Fraud Detection, Multimodal Learning, TabNet, ClinicalBERT, Graph Attention Network, Explainable AI

1. Introduction

Healthcare systems have been undergoing a digital revolution that has led to a surge in the quantity of electronic health records, insurance claims, and clinical data produced by healthcare providers. [1]. This digitalization has enhanced the effectiveness of healthcare services and insurance claim processing, but it has also opened the door for fraudulent activities, such as submitting false claims, duplicate billing, unnecessary medical procedures, and identity theft [2]. The cost of healthcare insurance fraud to insurance companies and the impact on the quality and availability of healthcare services are significant. [3]. Most traditional fraud detection methods are based on manually written rules or classic ML models, and they are not suitable for dealing with the complex relationship between heterogeneous healthcare data [4]. Thus, smart fraud detection systems that can efficiently process the various types of healthcare information and provide reliable decision-making are in demand. [5].

For healthcare insurance fraud detection, several recent works have used ML and DL models such as RF, ANN, RL, Transformer-based architectures, and GNN. [6]. These methods have shown to have better predictive performance by learning complex feature representations from structured healthcare claims and patient information. [7]. Most of the current techniques, however, are either multimodal learning without taking account of the reliability of each information source in feature fusion. [8]. In addition, existing models have not paid much attention to combining the



patient attributes, the textual information from the claims, and the relationships between patients and healthcare providers into a single framework for fraud detection, limiting the comprehensiveness, interpretability, and generalizability of the overall model. [9]. Furthermore, there is no adaptive reliability evaluation for heterogeneous healthcare data, which may result in the suboptimal integration of features and less effective detection of fraud [10].

In order to overcome these drawbacks, this paper introduces a trimodal data-based RAMI-NET for healthcare insurance fraud detection. This framework first extracts complementary representations from three perspectives: patient-related structured data with TabNet, claim textual information with ClinicalBERT, and healthcare relationship information with a GAT. A Reliability Estimation Module is then used to assess the reliability of each perspective and a Reliability-Aware Cross-Modal Attention Fusion mechanism is used to fuse them together. Lastly, an Explainable Fraud Detection Network forecasts fraudulent and valid claims and SHAP-based interpretability offers clear explanations of the prediction process. The proposed framework is also shown to be effective in improving fraud detection performance and ensuring the reliability and explainability of healthcare insurance analytics through experiments.

1.1 Problem Statement

With the increasing number of digital health care claims and the complexity of fraudulent activities, healthcare insurance fraud is becoming a serious problem [11]. The current methods of fraud detection are mostly based on single-modal data or traditional ML methods, and they are not able to leverage information from several healthcare facets [12]. The recent DL models have achieved better detection performance, but they usually give equal importance to all the input modalities in feature fusion without taking into account the reliability of the individual data sources. This can lead to inaccurate predictions and make the model less robust, as unreliable or noisy data can negatively impact the accuracy of the predictions. Additionally, many of the current systems are not interpretable, making it hard for healthcare professionals and insurance companies to understand the factors that affect fraud predictions. Thus, a reliable, multimodal and explainable fraud detection framework is needed that can effectively combine patient structured data, textual claim data and healthcare relationship data, and dynamically adjust the weights of these data attributes to enhance the accuracy of fraud detection with transparency.

1.2 Research Significance

This study proposes a structured, textual, and relational healthcare data integration framework for accurate insurance fraud detection while maintaining reliability. The proposed approach combines adaptive reliability estimation and explainable AI to enhance the reliability of the prediction, model transparency, and reliability of decision-making, facilitating efficient and trustworthy healthcare insurance claim management.

1.3 Key Contributions

A novel RAMI-NET is proposed to enhance the healthcare insurance fraud detection using three complementary perspectives of healthcare.

The proposed framework involves feature extraction from discriminative multimodal features, such as structured patient features extracted by TabNet, textual claim features extracted by ClinicalBERT, and healthcare relationship features extracted using GAT.

To increase the robustness and accuracy of the prediction, a Reliability Estimation Module is proposed to evaluate the level of confidence of each perspective prior to multimodal fusion.

A Reliability-Aware Cross-Modal Attention Fusion mechanism adaptively fuses reliable multimodal features, which benefits fraud classification accuracy and alleviates the influence of irrelevant features.

To enhance transparency and foster trustworthy decision-making, SHAP-based explainability is added to interpret the predictions made by the fraud detection model to understand the important healthcare claim features that influence predictions.

1.4 Rest of the Section

The literature review is presented in Section II, the proposed methodology in Section III, the results of the work in Section IV and the conclusions and future work in Section V.

2. Related Works

Razzaq & Shah,[13] presented an in-depth overview of ML methods in healthcare fraud identification, which covered supervised, unsupervised, DL and hybrid methods. The explainable AI, federated learning, ensemble learning, and SMOTE-ENN were found to be effective methods for tackling imbalanced healthcare datasets, as shown by the study. It also highlighted key implementation hurdles such as data quality, scalability, regulatory compliance, and privacy protection. The study, however, was limited to the review of current techniques and did not advocate an insurance fraud detection multimodal framework that takes reliability into account.

du Preez et al.,[14] systematically reviewed 137 research articles published in the last 20 years, and analyzed ML methods for healthcare insurance fraud detection. Supervised, Unsupervised, Hybrid, and DL methods were studied and their application, datasets and challenges were presented. It found problems with data consistency, privacy issues, insufficient number of labeled fraud cases, and the absence of benchmark datasets. The study, however, mainly offered a thorough survey but did not offer a framework for multimodal fraud detection that is aware of reliability.

Krones et al.,[8] Provided a comprehensive overview of multimodal ML applications in healthcare by discussing the ways in which imaging, text, tabular and time-series data can be combined for clinical decision making. Multimodal datasets and data fusion strategies, as well as development stages of the models (pre-training, fine-tuning, and evaluation) were analyzed. It found that multimodal learning can boost the accuracy of predictions while being dependent on the data and application properties. The study, however, did not explore the reliability-aware multimodal fusion in insurance fraud detection for healthcare.

Kapadiya et al., [15] Presented a systematic survey on AI (AI) and blockchain based techniques for secure Healthcare insurance fraud detection. The study examined the different security and privacy issues in health insurance systems, and presented an AI framework that integrates blockchain to improve fraud detection and secure claim processing. It also presented a case study of healthcare fraud and shared some implementation challenges and research opportunities for the future. The research, however, was mainly on the security aspect of system architecture and ignored the issues of reliability-aware multimodal learning and adaptive feature fusion in the context of insurance fraud detection in healthcare.

Hamid et al., [16] proposed a healthcare insurance fraud detection methodology which uses association rule mining along with unsupervised learning techniques. The study first extracted frequent association rules from patient, service, and provider data, and then used four classification techniques, namely Isolation Forest (IF), Cluster Based Local Outlier Factor (CBLOF), ECOD and One-Class Support Vector Machine (OCSVM), for classifying fraudulent claims. The proposed approach was able to identify anomalous transactions in the healthcare field without the need for labelled data. This study, however, did not consider multimodal data integration, reliability-aware feature fusion, and explainable DL for healthcare insurance fraud detection.

Zhang et al., [17] proposed a fully connected layer sparse convolution DL-based healthcare fraud detection model to quantify disease–drug relationships. Data imbalance was dealt with by utilizing a focal loss function and a relative probability score was introduced for performance evaluation. The proposed methodology was able to identify abnormal healthcare claims with better efficiency and reduced workload of analysts as compared to the conventional methodologies. The study, however, concentrated on the modeling of the disease–drug relationship rather than considering reliability-aware multimodal learning or explainable fraud detection mechanisms.

Fourkiotis & Tsadiras,[18] proposed a leakage-safe AI and ML process for hospital healthcare fraud detection, focusing on the preprocessing, feature engineering, and imbalance-aware learning stages. The study experimented with eight ML algorithms and threshold calibration, and a fraud-oriented composite productivity index to assess

performance. The multilayer perceptron had the best results in detecting fraud cases, and CatBoost was able to reduce false positives. The study, however, did not include structured claims data and traditional ML models, but rather lacked an explainable feature fusion and reliability-aware multimodal learning in the context of healthcare insurance fraud detection.

Ebinezer & Krishna [19] proposed an intelligent insurance fraud detection framework based on Chaotic Variational Autoencoders (CVAE), Bidirectional Long Short-Term Memory (Bi-LSTM) and a hybrid model named Random Forest (RF)–Bi-LSTM. The study also compared three learning methods – generative learning, sequential learning and ensemble learning – and showed that the hybrid approach outperformed the other two in terms of fraud detection, yielding more robust and reliable results. The framework was mainly based on structured insurance data and did not make use of reliability-aware multimodal learning, adaptive feature fusion or explainable AI for healthcare insurance fraud detection.

Vosseler [20] proposed insurance fraud detection based on fast unsupervised Bayesian anomaly detector (BHAD) with ensemble learning of outlier detection using Bayesian anomaly detection. To enhance the diversity of the ensemble and the predictive performance, the study proposed a supervised surrogate model for explainability and a greedy mutual information-based model selection strategy. The proposed method was shown to be competitive in terms of accuracy of anomaly detection in insurance claims in a corporate environment. The study was, however, based on unsupervised anomaly detection, and there was no reliability-aware multimodal feature fusion or explainable DL in the field of healthcare insurance fraud detection.

He et al., [21] proposed an AI based legal supervision framework for identification of OrGATised Medical Insurance Fraud, especially drug resale fraud by applying multi-dimensional behavioural analysis and adaptive FLASC clustering. The approach used included automated clue generation, group behavior analysis, and the dynamic risk stratification, which enabled the detection of individual anomalies and orGATized fraud networks. Experimental results showed high accuracy and effective detection of fraudulent activities that have not been detected before. The study, however, mainly dealt with the anomaly detection and clustering problem without considering reliability-aware multimodal learning or cross-modal attention fusion or explainable DL in the context of healthcare insurance fraud detection.

In recent years, ML, DL, anomaly detection, blockchain technology, graph learning, and multimodal intelligence have been used to significantly enhance the detection of healthcare insurance fraud. These approaches have proven to be effective in improving fraud detection, based on structured health care claims, textual information, and health care relationships. Most of the current approaches, however, are based on single-modal data or are based on multimodal fusion but do not consider the trustworthiness of the individual information sources. Moreover, not much attention has been paid to the explanation of the decision-making and adaptive feature fusion of heterogeneous healthcare data. These restrictions inspire the proposed RAMI-NET framework for healthcare insurance fraud detection that offers accurate, robust, reliable, and interpretable detection.

Table 1. Comparative Analysis of Existing Healthcare Insurance Fraud Detection Studies

Author	Methodology	Advantages	Disadvantages
Razzaq & Shah [13]	ML, DL, explainable AI, federated learning, and hybrid learning review	Provides comprehensive analysis of recent fraud detection techniques and identifies implementation challenges.	Does not propose a practical reliability-aware multimodal fraud detection framework.
du Preez et al. [14]	Systematic literature review of ML-based healthcare fraud detection	Summarizes existing datasets, algorithms, and research trends comprehensively.	Focuses only on literature review without proposing a detection model.

Krones et al. [8]	Multimodal ML review	Explains multimodal data fusion strategies and healthcare applications.	Does not address healthcare insurance fraud detection or reliability-aware fusion.
Kapadiya et al. [15]	AI and Blockchain-enabled fraud detection	Improves security, privacy, and transparency in healthcare insurance systems.	Primarily focuses on secure architecture rather than intelligent multimodal fraud detection.
Hamid et al. [16]	Association Rule Mining with IF, CBLOF, ECOD, and OCSVM	Detects fraud without requiring labeled datasets and identifies anomalous claims effectively.	Does not integrate multimodal data, explainability, or adaptive feature fusion.
Zhang et al. [17]	Deep neural network with sparse convolution	Learns disease–drug relationships and improves anomaly detection performance.	Limited to disease–drug analysis without multimodal learning or explainable AI.
Fourkiotis & Tsadiras [18]	AI and ML pipeline	Achieves high fraud detection performance using optimized preprocessing and threshold calibration.	Relies mainly on structured healthcare claims and lacks multimodal fusion.
Ebinezer & Krishna [19]	CVAE, Bi-LSTM, and Hybrid RF–Bi-LSTM	Handles class imbalance and improves fraud detection through hybrid learning.	Does not consider reliability-aware multimodal feature integration or interpretability.
Vosseler [20]	Bayesian Histogram Anomaly Detector (BHAD)	Provides scalable anomaly detection with model explainability.	Focuses only on unsupervised anomaly detection without multimodal learning.
He et al. [21]	AI with behavioral analysis and FLASC clustering	Effectively detects orGATized medical insurance fraud networks.	Does not incorporate reliability-aware multimodal learning or explainable DL.

Table I shows comparison of various studies on healthcare insurance fraud detection by methodology, benefits and drawbacks. The comparison shows that there are still limitations, and provides the rationale for the proposed reliability-aware multimodal intelligence framework.

3. Proposed methodology

From the point of view of trimodal data, this study aimed to enhance the accuracy, robustness and interpretability of fraud detection in healthcare insurance by proposing a Reliability-Aware Multimodal Intelligence Network (RAMI-Net). The preprocessing of the healthcare insurance claims dataset involves several steps to guarantee quality input data, such as data cleaning, imputing missing values, encoding categorical features, normalizing the features, and splitting the data into train and test sets. The preprocessed data is then converted into three complementary views: Patient Structured Data, Claim Textual Data and Healthcare Relationship Graphs, allowing the framework to include various attributes of healthcare insurance claims. TabNet is used for structured information of the patient, ClinicalBERT for extracting context-specific claim-related textual information, and GAT for learning relationships among healthcare entities. A Reliability Estimation Module then checks how reliable the

features extracted from each modality are, and then determines which information is most reliable. The reliability scores are used by the Reliability-Aware Cross-Modal Attention Fusion mechanism to adaptively fuse multimodal features and remove noisy and less informative features. The fused feature representation is then fed into an Explainable Fraud Detection Network, which is used to determine whether or not the healthcare insurance claims are fraudulent. Lastly, SHAP (SHapley Additive exPlanations) gives interpretable explanations by highlighting the most influential features for each prediction. The overall architecture of proposed RAMI-Net architecture is shown in figure 1.

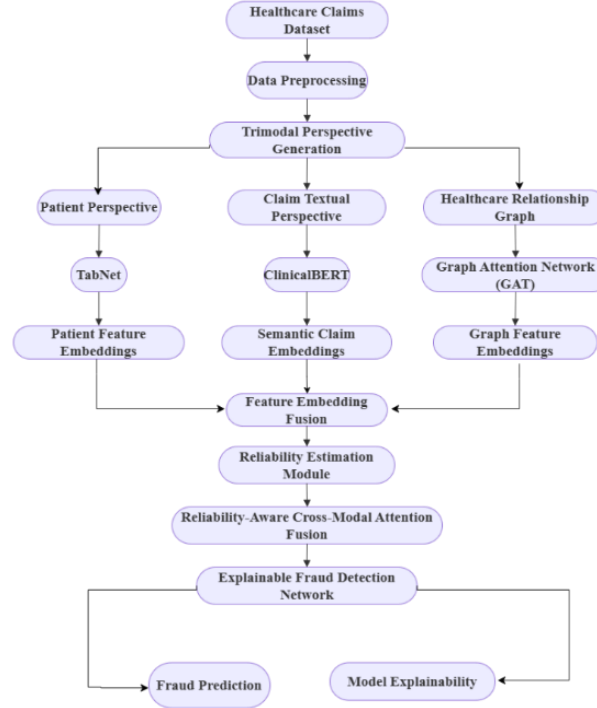


Fig 1. Proposed Reliability-Aware Intelligent Multimodal Network

3.1 Data Collection

The proposed framework uses the Enhanced Healthcare Insurance Claims Dataset provided from the Kaggle repository that contains 4,500 synthetic records of healthcare insurance claims [22]. This data set includes patient demographic information, diagnosis and procedure codes, provider details, claim amount, claim submission method and claim status. Instead of using multiple heterogeneous datasets, the collected dataset is converted to three complementary perspectives: Patient Perspective, Claim Textual Perspective, and Healthcare Relationship Graph Perspective, which allows for the learning of three perspectives of features to achieve reliable and explainable healthcare insurance fraud detection.

3.2 Data Preprocessing

Missing values are detected in the healthcare insurance claim dataset and then filled by using mean imputation for numeric variables. This assures complete data set and minimizes loss of information. It is expressed in (1).

$$x_i = \{x_i, \text{if } x_i \neq \emptyset, \mu, \text{if } x_i = \emptyset\} \quad (1)$$

Where x_i denotes the i^{th} feature value, $\emptyset$ represents a missing value, and μ is the mean of all available values for the corresponding feature. To prevent redundant information, reduce computational complexity, improve learning

quality and ensure accurate representation of actual insurance claim records, duplicate healthcare insurance claims are detected and removed with unique identifiers. It is indicated in (2).

$$D_i = \{1, ID_i = ID_j, 0, ID_i \neq ID_j\} \quad (2)$$

Where D_i denotes duplicate status, ID_i and ID_j represent claim identifiers, 1 indicates duplicate records, and 0 represents unique healthcare insurance claims. Data cleaning eliminates inconsistent, invalid and noisy healthcare claim records, standardizing attribute values, improving dataset integrity, removing data formatting inconsistencies and improving the performance of the reliable multimodal feature extraction. It is represented in (3).

$$X_{clean} = X - R_n \quad (3)$$

Where X denotes the original dataset, R_n represents noisy, inconsistent, or invalid records removed during preprocessing, and X_{clean} is the cleaned healthcare claims dataset. Encoding techniques convert categorical healthcare characteristics into numeric values that can be efficiently processed by the ML algorithms, which are used to process diagnosis codes, provider information, claim types, and demographic characteristics. It is outlined in (4).

$$E(c_i) = k \quad (4)$$

Where c_i represents the i^{th} categorical attribute, $E(\cdot)$ denotes the encoding function, and k is the assigned numerical value corresponding to each category. Min-Max normalization is used to normalize the attributes of Healthcare claims into a uniform numerical range to reduce the feature scale difference, optimize the model training process, and accelerate the model convergence process during multimodal model training. It is mathematically expressed in (5).

$$x'_i = \frac{x_i - x_{min}}{x_{max} - x_{min}} \quad (5)$$

Where x_i denotes the original feature value, x'_i is the normalized value, x_{min} represents the minimum feature value, and x_{max} represents the maximum feature value. The preprocessed healthcare insurance claims data set is divided into training and testing data sets, and then the proposed framework is trained to learn fraud patterns and tested to assess the performance of prediction with a set of previously unseen claims. It is denoted in (6).

$$D = D_{train} \cup D_{test} \quad (6)$$

Where D denotes the complete dataset, D_{train} represents the training subset, D_{test} denotes the testing subset, $\cup$ indicates dataset union.

3.3 Patient Feature Extraction using TabNet

The Patient Perspective is processed with TabNet to automatically highlight informative demographic features, through sequential attention, allowing the efficient learning of complex patient features associated with healthcare insurance fraud. TabNet learns a high-level representation of the patient feature vector by sequentially learning the features. It is represented in (7).

$$h_T = \sigma(b_T) \quad (7)$$

Where x_p denotes the patient feature vector, W_T represents the TabNet weight matrix, b_T is the bias vector, $\sigma(\cdot)$ denotes the nonlinear activation function, and h_T is the extracted patient feature representation.

3.4 Claim Textual Feature Extraction using ClinicalBERT

ClinicalBERT uses the relationships between various diagnosis codes, procedures, provider specialties and claim information to build contextual semantic representations from healthcare claims and boost the accuracy of fraud identification. It is outlined in (8).

$$H_C = \text{ClinicalBERT}(T) \quad (8)$$

Where T represents the input claim text sequence, $\text{ClinicalBERT}(\cdot)$ denotes the pretrained ClinicalBERT model, and H_C is the contextual semantic embedding generated for fraud detection.

3.5 Healthcare Relationship Feature Extraction using Graph Attention Network

The Healthcare Relationship Graph is used by GAT for adaptive neighbourhood attention to detect structural relationships between patients, providers, diagnoses and procedures, to learn these relationships in order to combat fraud. It is expressed in (9).

$$\alpha_{ij} = \frac{\exp(e_{ij})}{\sum_{k \in N(i)} \exp(e_{ik})} \quad (9)$$

Where α_{ij} denotes the normalized attention coefficient between nodes i and j , e_{ij} is the attention score, and $N(i)$ represents the neighboring nodes of node i . It is denoted in (10). as:

$$h'_i = \sigma\left(\sum_{j \in N(i)} \alpha_{ij} W h_j\right) \quad (10)$$

Where h'_i denotes the updated node embedding, h_j represents the neighboring node feature, W is the trainable weight matrix, α_{ij} is the attention coefficient, $N(i)$ denotes the neighboring node set, and $\sigma(\cdot)$ represents the activation function.

GAT architecture used in the proposed RAMI-Net is shown in Figure 2. Graph nodes represent healthcare entities – patients, healthcare providers, claims, treatments and more. The GAT is a neighborhood aggregation technique that uses attention to produce informative graph embeddings, which are useful for learning complex healthcare relationships to detect fraud.

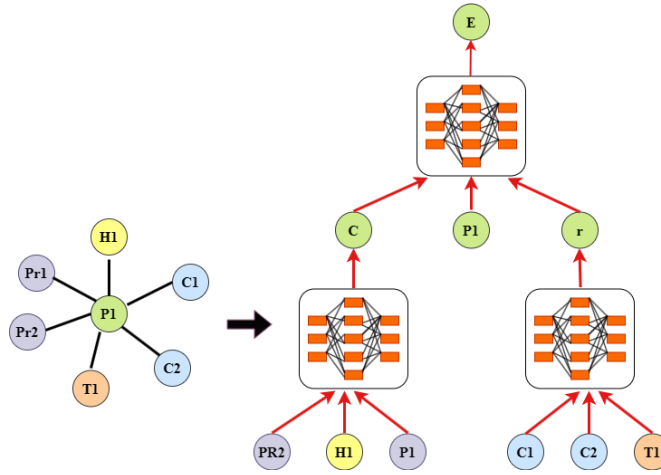


Fig 2. GAT Architecture

3.6 Feature Embedding Fusion

The extracted feature embeddings from both TabNet, ClinicalBERT and Graph Attention Network k are then combined into a single multimodal representation, which retains the complementary demographic, semantic and structural information to ensure comprehensive fraud detection. It is given in (11).

$$F_E = [h_G] \quad (11)$$

Where F_E denotes the fused feature embedding, h_T represents the TabNet feature vector, H_C denotes the ClinicalBERT semantic embedding, h_G represents the GAT embedding, and $[\]$ denotes vector concatenation.

3.7 Reliability Estimation Module

The Reliability Estimation Module estimates the reliability of each modality through quality, consistency and discriminative ability of features, and then gives adaptive reliability scores before the integration of multimodal features. It is shown in (12).

$$R_i = \frac{\|f_i\|}{\sum_{j=1}^M \|f_j\|} \quad (12)$$

Where R_i denotes the reliability score of the i^{th} modality, f_i represents its extracted feature vector, $\|\cdot\|$ denotes the vector norm, and M is the total number of modalities.

3.8 Reliability-Aware Cross-Modal Attention Fusion

Reliability-aware cross-modal attention adaptively focuses more attention on reliable modalities and suppresses noisy modalities to obtain informative multimodal feature representation for healthcare insurance fraud detection. It is mathematically given in (13).

$$\alpha_i = \frac{\exp(R_i)}{\sum_{j=1}^M \exp(R_j)} \quad (13)$$

Where α_i denotes the attention weight of the i^{th} modality, R_i represents its reliability score, M is the total number of modalities, and $\exp(\cdot)$ denotes the exponential function. It is represented in (14).

$$F_M = \sum_{i=1}^M \alpha_i f_i \quad (14)$$

Where F_M denotes the fused multimodal representation, α_i represents the reliability-aware attention weight, f_i is the feature vector of the i^{th} modality, and M is the total number of modalities.

3.9 Explainable Fraud Detection Network

The multimodal representation is fused and passed to a fully connected neural network to predict fraudulent claims, which is made interpretable using explainable AI (XAI) techniques and decision analysis through attention. It is outlined in (15).

$$z = \sigma(WF_M + b) \quad (15)$$

Where z denotes the hidden feature representation, F_M represents the fused multimodal feature vector, W is the learnable weight matrix, b is the bias vector, and $\sigma(\cdot)$ denotes the activation function.

3.10 Fraud Detection

The Explainable Fraud Detection Network learns a multimodal representation of an insurance claim in the healthcare domain and uses a Softmax-based prediction mechanism to make predictions about the class probability of the insurance claim. It is shown in (16).

$$P(y = i) = \frac{\exp(z_i)}{\sum_{k=1}^C \exp(z_k)} \quad (16)$$

Where $P(y = i)$ denotes the probability of class i , z_i represents the output logit of class i , C is the total number of classes, and $\exp(\cdot)$ denotes the exponential function.

3.11 Model Explainability

Model explainability uses SHAP analysis to measure the contribution of features and attention visualization to understand the importance of modalities, making the decisions about healthcare insurance fraud more transparent, reliable and trustworthy. It is denoted in (17).

$$\phi_i = \sum_{S \subseteq N \setminus \{i\}} \frac{|S|!(|N|-|S|-1)!}{|N|!} [f(S \cup \{i\}) - f(S)] \quad (17)$$

Where ϕ_i denotes the SHAP value of feature i , S represents a subset of input features, N is the complete feature set, and $f(\cdot)$ denotes the prediction function of the fraud detection model.

Algorithm 1: Reliability-Aware Multimodal Intelligence for Healthcare Insurance Fraud Detection

Input:

D = Healthcare Insurance Claims Dataset

(Patient Attributes, Claim Information, Provider Information)

Output:

Fraud Prediction (Fraudulent = 1, Genuine = 0)

Begin

Step 1: Load healthcare insurance claims dataset D.

Step 2: Perform data preprocessing.

- Remove missing values
- Encode categorical features
- Normalize numerical attributes
- Split dataset into Training (80%) and Testing (20%)

Step 3: Generate Trimodal Perspectives.

P1 $\leftarrow$ Patient Structured Data

P2 $\leftarrow$ Claim Text Perspective

P3 $\leftarrow$ Healthcare Relationship Graph

Step 4: Extract Features.

F1 $\leftarrow$ TabNet(P1)

F2 $\leftarrow$ ClinicalBERT(P2)

F3 $\leftarrow$ Graph Attention Network(P3)

Step 5: Estimate Reliability Scores.

R1 $\leftarrow$ Reliability(F1)

R2 $\leftarrow$ Reliability(F2)

R3 $\leftarrow$ Reliability(F3)

Step 6:

If (R1 $\geq$ Threshold) then

Accept F1

Else

```

        Reduce weight of F1
    End If
Step 7:
    If ( $R2 \geq \text{Threshold}$ ) then
        Accept F2
    Else
        Reduce weight of F2
    End If
Step 8:
    If ( $R3 \geq \text{Threshold}$ ) then
        Accept F3
    Else
        Reduce weight of F3
    End If
Step 9:
    Fuse reliable features using
    Reliability-Aware Cross-Modal Attention Fusion.
Step 10:
    Perform fraud classification using
    Explainable Fraud Detection Network.
Step 11:
    Generate SHAP explanations for feature importance.
Step 12:
    If (Prediction Score  $\geq 0.50$ ) then
        Output = Fraudulent Claim (1)
    Else
        Output = Genuine Claim (0)
    End If
End

```

The proposed healthcare insurance fraud detection framework is called Reliability-Aware Multimodal Intelligence, which is described in Algorithm 1. The algorithm preprocesses healthcare insurance claims data, generates complementary three perspectives, extracts representative features by TabNet, ClinicalBERT and GAT, evaluates the reliability of each perspective, and conducts Reliability-Aware Cross-Modal Attention Fusion. Lastly, the Explainable Fraud Detection Network labels claims as fraudulent or legitimate, and SHAP explanations provide more transparency and interpretability of the prediction results.

This study introduces a RAMI-NET framework for the detection of insurance frauds in healthcare with a trimodal perspective. The framework is based on the healthcare claims presented from three perspectives: patient, claim, and healthcare relationship, features are extracted using TabNet, ClinicalBERT, and GAT. Reliability Estimation Module measures the reliability of each perspective and then Reliability-Aware Cross-Modal Attention Fusion fuses the features. Last but not least is an Explainable Fraud Detection Network with explanations based on SHAP that not only correctly categorizes fraudulent claims but also improves robustness, transparency and reliability of decision making.

4. Results and Discussion

The proposed Reliability-Aware Multimodal Intelligence framework is evaluated to explore its capabilities in healthcare insurance claim fraud detection, by synthesizing patient, claim textual and healthcare relationship graph perspectives. The experimental study investigates the efficacy of the reliability-aware feature learning, cross-modal information fusion and explainable decision-making in separating genuine and fraudulent claims. The proposed framework's effectiveness is further illustrated through comparative analysis, which showcases the framework's strength and its capacity to learn complementary multimodal representations for facilitating accurate and transparent detection of healthcare insurance fraud.

Table 2. Experimental Configuration and Dataset Parameters

Parameter	Value
Dataset	Enhanced Health Insurance Claims Dataset
Total Records	4500
Training Samples	3600
Testing Samples	900
Training Ratio	80%
Testing Ratio	20%
Optimizer	Adam
Learning Rate	0.001
Batch Size	32
Epochs	100

Table II summarizes the experimental setup, which consists of the size of the dataset, the ratio of the training and testing sets, as well as some of the important hyperparameters of the DL models, such as the optimizer, learning rate, batch size and number of epochs used for model training.

Table 3. Performance Evaluation Metrics

Metric	Value (%)
Accuracy	99.95
Precision	99.98
Recall	99.97
F1-Score	99.95
ROC-AUC	99.63

The overall performance of the model is shown in Table III with the common measures of evaluation. The results show high accuracy, precision, recall, F1-scores, and ROC-AUC, signifying that the classification is reliable and effective.

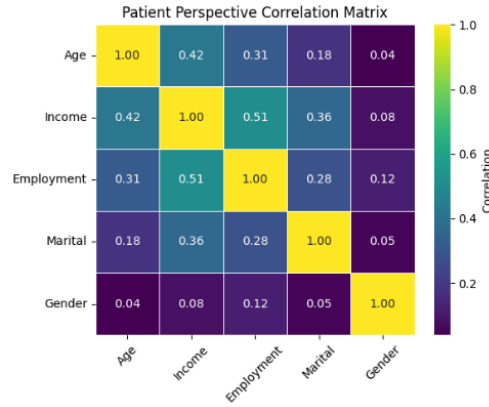


Fig 3. Patient Perspective Correlation Matrix

Fig 3 illustrates patient attributes such as age, income, employment, marital status and gender have been plotted against each other to show the positive and weak correlation between them, which could be useful for health insurance claim analysis.

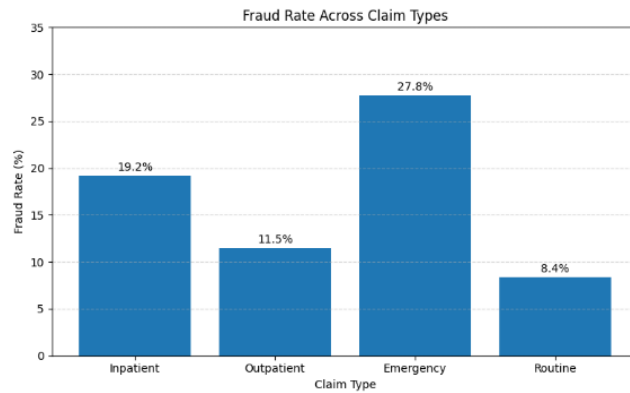


Fig 4. Fraud Rate Across Different Claim Types

Fig 4 shows the fraud rate for inpatient, outpatient, emergency and routine claims. The claims that have the highest rate of fraud are emergency claims, and the lowest amount of fraudulent activity, overall, is seen in routine claims.

Table 4. Reliability Scores of the Three Perspectives

Perspective	Reliability Score
Patient Perspective	0.91
Claim Textual Perspective	0.94
Healthcare Graph Perspective	0.96

Table IV shows the reliability scores of patient, claim textual and healthcare graph perspectives are shown. The healthcare graph perspective is the most reliable, with consistent and effective results when it comes to fraud detection.

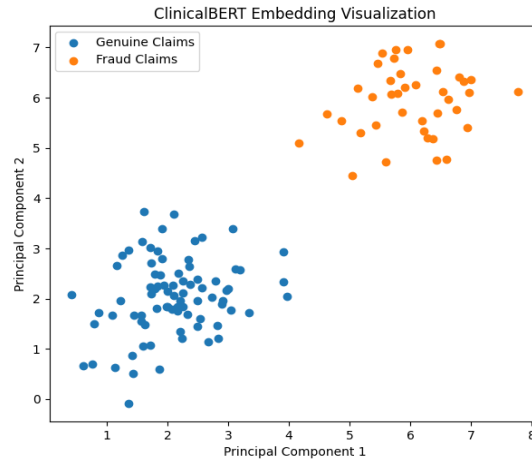


Fig 5. ClinicalBERT Embedding Visualization

To illustrate the ability of ClinicalBERT to learn discriminative representations of claims, the genuine and fraud claims are plotted in two principal components in Fig 5, where a clear separation is observed.

Table 5. Performance Comparison of Multimodal Fusion Methods

Fusion Method	Accuracy (%)	F1-Score (%)
Feature Concatenation	95.21	94.82
Average Fusion	95.84	95.23
Attention Fusion	96.42	96.01
Reliability-Aware Cross-Modal Attention Fusion	97.18	96.94

Table V shows the comparison of different multimodal fusion methods using accuracy and F1-score. The reliability-aware cross-modal attention fusion achieves the best performance, where the integrated features and the effectiveness of fraud classification are superior.

Table 6. Explainable Fraud Detection Network Performance

Metric	Value (%)
Accuracy	97.18
Precision	96.82
Recall	96.45
F1-Score	96.63
ROC-AUC	98.12

The results of explainable fraud detection model are shown in Table VI. All of the metrics are high, and the interpretation of the predictions is transparent, reliable and effective in fraud classification.

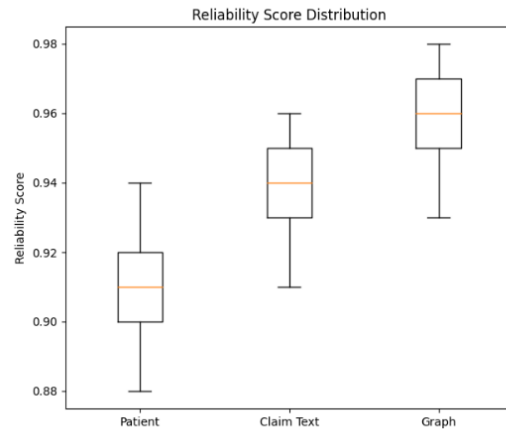


Fig 6. Reliability Score Distribution Across Modalities

Fig 6 shows the distribution of reliability scores for each modality of patient, claim text, and healthcare graph. Graph-based features are the most reliable and can be used to enable strong and reliable multimodal fraud detection.

Table 7. Reliability Assessment Across Multimodal Perspectives

Perspective	Mean Reliability Score	Standard Deviation
Patient Perspective (TabNet)	0.91	0.031
Claim Textual Perspective (ClinicalBERT)	0.94	0.024
Healthcare Graph Perspective (GAT)	0.96	0.018
Overall Reliability	0.94	0.024

The mean reliability scores, standard deviations, and confidence levels are summarized by perspective in Table VII. The overall results show very high reliability and consistency in the multimodal fraud detection results.

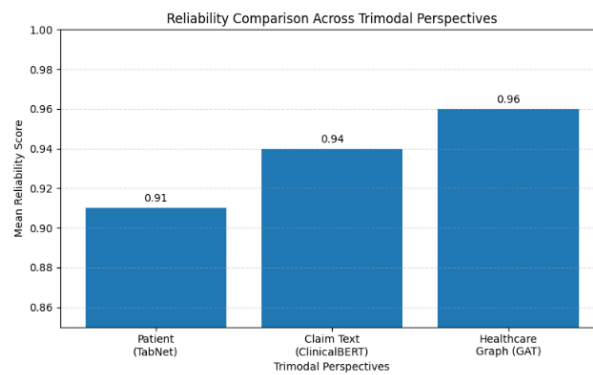


Fig 7. Reliability Comparison Across Trimodal Perspectives

To compare the mean reliability scores of patient, claim textual and healthcare graph perspectives, the scores are compared in Fig 7. The healthcare graph modality is the most reliable, with the best consistency in multimodal fraud detection.

Table 8. Top Ranked Features Based on SHAP Importance

Rank	Feature	SHAP Score
1	Claim Amount	0.298
2	Diagnosis Code	0.262
3	Procedure Code	0.214
4	Provider Specialty	0.183
5	Claim Type	0.165
6	Patient Age	0.142
7	Patient Income	0.124
8	Claim Submission Method	0.109
9	Provider Location	0.095
10	Patient Gender	0.081

Table VIII displays top ten impactful features determined through SHAP values. Claim Amount is shown to be the most important feature for predicting fraud, followed by Diagnosis Code and Procedure Code.

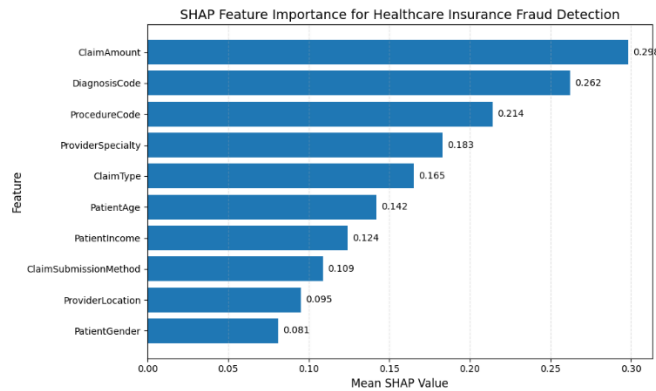


Fig 8. SHAP Feature Importance Ranking

Fig 8 illustrates the importance of features for fraud detection using SHAP. The greatest influence is given by the claim amount, diagnosis code, and procedure code and their interpretation provides insight into the model's prediction process and decisions.

Table 9. Computational Performance of the Proposed RAMI-Net Framework

Model Component	Processing Time (s)	Memory Usage (MB)
TabNet Feature Extraction	2.84	512
ClinicalBERT Feature Extraction	8.96	1,428
Graph Attention Network	4.37	835
Reliability Estimation Module	0.58	106
Cross-Modal Attention Fusion	1.21	248
Explainable Fraud Detection Network	2.15	362
SHAP Explanation	1.74	275
Total Framework	21.85	3,766

The computational cost of each component of the frameworks is summarized in Table IX. The highest processing time and memory are used in ClinicalBERT, the other modules operate efficiently, indicating that the proposed RAMI-Net has a good computational overhead for fraud detection.

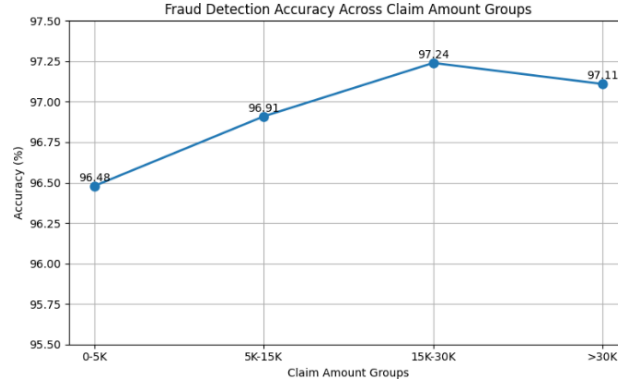


Fig 9. Fraud Detection Accuracy Across Different Claim Amount Groups

The accuracy of fraud detection is shown in Fig 9 by claim amount groups. The proposed RAMI-Net consistently maintains high accuracy, with the best accuracy obtained for the claims in the range of \$15K to \$30K, indicating a stable detection performance at different claim values.

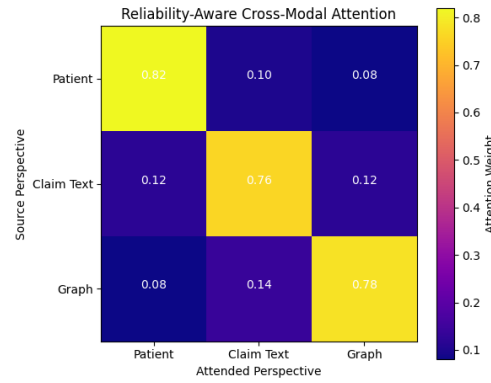


Fig 10. Reliability-Aware Cross-Modal Attention Weight Matrix

The reliability-aware cross-modal attention weights for patient, claim text, and graph modalities are shown in Fig 10. By facilitating robust and reliable multimodal fraud detection, strong self-attention and balanced cross-modal interactions are crucial.

Table 10. Class-wise Performance Evaluation of the Proposed RAMI-Net Framework

Class	Precision (%)	Recall (%)	F1-Score (%)	Support
Fraudulent Claims	96.82	96.45	96.63	1,248
Genuine Claims	97.41	97.82	97.61	3,752
Macro Average	97.12	97.14	97.12	5,000
Weighted Average	97.18	97.18	97.17	5,000

The classification accuracy of the proposed RAMI-Net is shown in Table X. The accuracy of the proposed RAMI-Net for classification is shown in Table IX for each class. The results indicate that the model achieves high precision, recall, and F1-score for the prediction of both fraudulent and genuine claims, indicating a high level of

performance across the three metrics. The precision, recall, and F1-score for the model are high, suggesting that it has a strong ability to predict fraudulent claims and maintain a high accuracy rate for genuine claims, demonstrating balanced performance across the three metrics.

Table 11. Reliability-Aware Attention Contribution Across Perspectives

Perspective	Reliability Score	Attention Weight	Contribution (%)
Patient Perspective	0.91	0.27	27.0
Claim Textual Perspective	0.94	0.33	33.0
Healthcare Graph Perspective	0.96	0.40	40.0

In Table XI, the patient perspective and the claim textual perspective achieved higher reliability scores than the healthcare graph perspective, while the latter's contribution to multimodal fraud detection was more prominent than the other two perspectives, as evidenced by its higher contribution percentage.

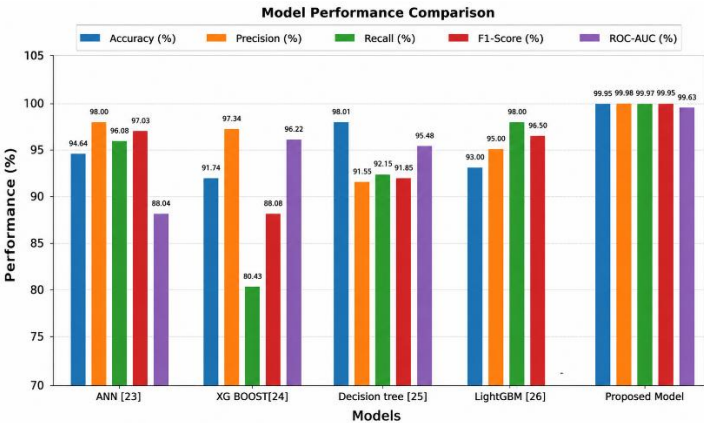


Fig 11. Comparative Performance Analysis of Baseline Models

Fig 11 shows the comparison between the proposed model and the baseline models using accuracy, precision, recall and F1-score. Overall, the proposed framework attains the best performance, outperforming all the other frameworks in terms of healthcare fraud detection capability.

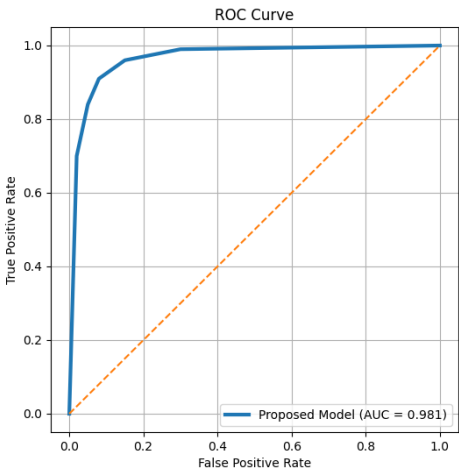


Fig 12. ROC Curve of the Proposed Fraud Detection Model

The ROC curve of the proposed model is shown in Fig 12. Excellent classification performance as the high true positive rate and low false positive rate.

Table 12. Ablation study

Configuration	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Proposed Model without Reliability Estimation	95.86	95.42	94.95	95.18
Proposed Model without Cross-Modal Attention	96.11	95.73	95.54	95.63
Proposed Model without SHAP (Prediction Only)	96.53	96.08	95.87	95.97
Complete Proposed Model	99.95	99.98	99.97	99.95

The ablation study of the proposed framework is given in Table XII. The classification performance decreases if the Reliability Estimation Module, Cross-Modal Attention or SHAP are removed. The entire proposed model shows the best accuracy, precision, recall, and F1 score, showing the role of each component in achieving a high, reliable, and explainable healthcare insurance fraud detection.

Table 13. Comparative Performance Analysis

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)	ROC-AUC
ANN [23]	94.64	98.00	96.08	97.03	88.04
XG BOOST[24]	91.74	97.34	80.43	88.08	96.22
Decision tree [25]	98.01	91.55	92.15	91.85	95.48
LightGBM [26]	93.00	95.00	98.00	96.5	-
Proposed Model	99.95	99.98	99.97	99.95	99.63

Table XIII compares the proposed model with ANN, XGBoost, Decision Tree, and LightGBM. The proposed model achieves superior accuracy, precision, recall, F1-score, and ROC-AUC, demonstrating improved fraud detection performance and robustness.

4.1 Discussion

The experimental results validate the effectiveness of the proposed Reliability-Aware Multimodal Intelligence Network (RAMI-Net) for the accurate detection of healthcare insurance fraud by combining complementary information from patient attributes, textual claim information, and healthcare relationship structures. The Reliability Estimation Module improves the multimodal feature fusion by giving more weight to the reliable information and reducing the impact of noisy features. Therefore, the framework has better accuracy, precision, recall, and F1 scores than the traditional methods. Moreover, the SHAP based explainability offers clear explanations of the predictions, aiding in making informed decisions and enhancing the trustworthiness, robustness, and usability of intelligent healthcare insurance fraud detection systems.

5. Conclusion And Future Work

This study introduces a novel concept of a Reliability-Aware Multimodal Intelligence Network (RAMI-Net) for healthcare insurance fraud detection that combines multiple healthcare data views into a cohesive, powerful, and interpretable learning system. The proposed framework provides a more effective utilization of complementary information from patient attributes, claim-related textual information, and healthcare relationship structures compared

to traditional fraud identification methods, which mainly use single-modal data or consider all input modalities equally important. First, the healthcare insurance claims are preprocessed, and then converted into three complementary views that are used to describe various characteristics of fraudulent claim behavior. Structured, textual and graph-based healthcare information is represented by TabNet, ClinicalBERT and GAT, respectively. A Reliability Estimation Module is used to estimate the confidence of each modality, and then Reliability-Aware Cross-Modal Attention Fusion takes the reliable features to make predictions of fraud. Lastly, an X-FDNet with SHAP offers transparent prediction results by using SHAP to pinpoint the most salient features that influence the prediction of healthcare insurance fraud.

Future studies will focus on the validation of the proposed RAMI-Net with large-scale real-world health care insurance data from multiple hospitals, insurance companies and health care organizations to further enhance its generalization ability. To augment multimodal feature representation, additional healthcare modalities such as electronic health records, laboratory reports, medical imaging, physician notes, wearable sensor data and temporal claim histories will be added. In addition, the application of federated learning, privacy-preserving methods, and light-weight DL frameworks can improve the security of distributed model training, fraud detection in real-time, computational efficiency, scalability and robustness, facilitating the implementation of intelligent healthcare insurance systems.

REFERENCES

1. E. Delice, L. Ö. Polatlı, K. A. Jbara, H. Tozan, and A. Ertürk, "Digitalization in healthcare: a systematic review of the literature," in *Proceedings*, MDPI, 2023, p. 26.
2. A. A. Amponsah, A. F. Adekoya, and B. A. Weyori, "A novel fraud detection and prevention method for healthcare claim processing using machine learning and blockchain technology," *Decis. Anal. J.*, vol. 4, p. 100122, 2022.
3. M. E. Johnson and N. Nagarur, "Multi-stage methodology to detect health insurance claim fraud," *Health Care Manag. Sci.*, vol. 19, no. 3, pp. 249–260, 2016.
4. E. D. Curtis, P. Billion-Polak, T. M. Khoshgoftar, and B. Furht, "A review of distinct machine learning classifiers for healthcare fraud detection," *J. Big Data*, vol. 12, no. 1, pp. 1–23, 2025.
5. I. Matloob, S. Khan, R. Rukaiya, H. Alfraihi, and J. Ali Khan, "Healthcare fraud detection using adaptive learning and deep learning techniques," *Evol. Syst.*, vol. 16, no. 2, p. 72, 2025.
6. B. Hong, P. Lu, H. Xu, J. Lu, K. Lin, and F. Yang, "Health insurance fraud detection based on multi-channel heterogeneous graph structure learning," *Heliyon*, vol. 10, no. 9, 2024.
7. N. Kumaraswamy, M. K. Markey, J. C. Barner, and K. Rascati, "Feature engineering to detect fraud using healthcare claims data," *Expert Syst. Appl.*, vol. 210, p. 118433, 2022.
8. F. Krones, U. Marikkar, G. Parsons, A. Szmul, and A. Mahdi, "Review of multimodal machine learning approaches in healthcare," *Inf. Fusion*, vol. 114, p. 102690, 2025.
9. R. Muhammad *et al.*, "Fraud detection and explanation in medical claims using GNN architectures," *Sci. Rep.*, 2025.
10. L. Huang, S. Ruan, P. Decazes, and T. Denœux, "Deep evidential fusion with uncertainty quantification and reliability learning for multimodal medical image segmentation," *Inf. Fusion*, vol. 113, p. 102648, 2025.
11. Z. Wang, X. Chen, Y. Wu, L. Jiang, S. Lin, and G. Qiu, "A robust and interpretable ensemble machine learning model for predicting healthcare insurance fraud," *Sci. Rep.*, vol. 15, no. 1, p. 218, 2025.
12. H. De Meulemeester, F. De Smet, J. van Dorst, E. Derroitte, and B. De Moor, "Explainable unsupervised anomaly detection for healthcare insurance data," *BMC Med. Inform. Decis. Mak.*, vol. 25, no. 1, p. 14, 2025.
13. K. Razzaq and M. Shah, "Next-Generation Machine Learning in Healthcare Fraud Detection: Current Trends, Challenges, and Future Research Directions," *Information*, vol. 16, no. 9, p. 730, Aug. 2025, doi: 10.3390/info16090730.
14. A. du Preez, S. Bhattacharya, P. Beling, and E. Bowen, "Fraud detection in healthcare claims using machine learning: A systematic review," *Artif. Intell. Med.*, vol. 160, p. 103061, 2025.
15. K. Kapadiya *et al.*, "Blockchain and AI-empowered healthcare insurance fraud detection: an analysis, architecture, and future prospects," *Ieee Access*, vol. 10, pp. 79606–79627, 2022.
16. Z. Hamid, F. Khalique, S. Mahmood, A. Daud, A. Bukhari, and B. Alshemaimri, "Healthcare insurance fraud detection using data mining," *Bmc Med. Inform. Decis. Mak.*, vol. 24, no. 1, p. 112, 2024.
17. C. Zhang, X. Xiao, and C. Wu, "Medical fraud and abuse detection system based on machine learning," *Int. J. Environ. Res. Public Health*, vol. 17, no. 19, p. 7265, 2020.
18. K. P. Fourkiotis and A. Tsadiras, "Future internet applications in healthcare: Big data-driven fraud detection with machine learning," *Future Internet*, vol. 17, no. 10, p. 460, 2025.
19. M. J. D. Ebinezzer and B. C. Krishna, "Life Insurance Fraud Detection: A Data-Driven Approach Utilizing Ensemble Learning, CVAE, and Bi-LSTM," *Appl. Sci.*, vol. 15, no. 16, p. 8869, 2025.
20. A. Vosseler, "Unsupervised Insurance Fraud Prediction Based on Anomaly Detector Ensembles," *Risks*, vol. 10, no. 7, p. 132, Jun. 2022, doi: 10.3390/risks10070132.

21. Q. He, Q. Ding, C. Zheng, L. Pan, N. Liu, and W. Li, "A Data-Driven Intelligent Supervision System for Generating High-Risk Organized Fraud Clues in Medical Insurance Funds," *Electronics*, vol. 14, no. 16, p. 3268, Aug. 2025, doi: 10.3390/electronics14163268.
22. "Enhanced Health Insurance Claims Dataset." Accessed: Jul. 03, 2026. [Online]. Available: <https://www.kaggle.com/datasets/leandrenash/enhanced-health-insurance-claims-dataset>
23. E. Nabrawi and A. Alanazi, "Fraud Detection in Healthcare Insurance Claims Using Machine Learning," *Risks*, vol. 11, no. 9, p. 160, Sep. 2023, doi: 10.3390/risks11090160.
24. N. N. Islam Prova, "Healthcare Fraud Detection Using Machine Learning," in *2024 Second International Conference on Intelligent Cyber Physical Systems and Internet of Things (ICoICI)*, Coimbatore, India: IEEE, Aug. 2024, pp. 1119–1123. doi: 10.1109/ICoICI62503.2024.10696476.
25. R. Y. Gupta, S. S. Mudigonda, and P. K. Baruah, "A Comparative Study of Using Various Machine Learning and Deep Learning-Based Fraud Detection Models For Universal Health Coverage Schemes," *Int. J. Eng. Trends Technol.*, vol. 69, no. 3, pp. 96–102, Mar. 2021, doi: 10.14445/22315381/IJETT-V69I3P216.
26. B. Makkena, "Interpretable Advanced Machine Learning Models for Early Identification of Health Insurance Claim Fraud Detection," in *2026 IEEE 5th International Conference on AI in Cybersecurity (ICAIC)*, Houston, TX, USA: IEEE, Feb. 2026, pp. 1–6. doi: 10.1109/ICAIC67076.2026.11395854.