

Citation: Esmaeili, M., Malekinia, M., & Pourebrahimi, A. (2026). Fraud Detection Analysis in Supplementary Health Insurance Using the LSTM Model. *Digital Transformation and Administration Innovation*, 4(1), 1-9.

Received date: 2025-06-11

Revised date: 2025-09-14

Accepted date: 2025-09-23

Initial published date: 2025-10-20

Final published date: 2026-01-01



Fraud Detection Analysis in Supplementary Health Insurance Using the LSTM Model

Masoumeh Esmaeili ¹, Mohammad Malekinia^{2*}, Alireza Pourebrahimi ³

1. Department of Management, K.i.c., Islamic Azad University, Kish, Iran

2. Department of Management, ST.C., Islamic Azad University, Tehran, Iran

3. Department of Industrial Management, Ka.C., Islamic Azad University, Karaj, Iran

*Correspondence: M_malekinia@azad.ac.ir

Abstract

This study aimed to design and evaluate an advanced Long Short-Term Memory (LSTM) deep learning model to accurately detect fraudulent claims in supplementary health insurance by leveraging sequential data patterns and domain-specific engineered features. An applied research design was used to build a robust fraud detection framework. A dataset of 20,000 health insurance claims was obtained from a supplementary insurance provider, containing both legitimate and fraudulent cases. Raw relational tables — including policy, insured individuals, claims, disease registry, and branch information — were merged into a single structured dataset using SQL. Rigorous data preprocessing was performed: irrelevant variables were removed, highly correlated features were eliminated through correlation analysis, dimensionality reduction was applied via Principal Component Analysis (PCA), and extreme outliers were excluded using the interquartile range (IQR) method. All numerical features were standardized, and class imbalance was addressed by weighting fraudulent cases during training. The processed data were reshaped into sequences suitable for LSTM input and divided into training and testing sets. An LSTM model with 32 hidden units and a Sigmoid output layer was trained using the Adam optimizer and binary cross-entropy loss, with performance validated through k-fold cross-validation. The LSTM model achieved outstanding predictive performance, with an overall accuracy of 100% on the test dataset. Both fraudulent and non-fraudulent claims reached perfect precision, recall, and F1-scores. Macro and weighted averages also recorded 1.00 across all metrics, indicating the model's ability to detect rare fraudulent events without sacrificing specificity. By combining advanced deep learning with systematic data preparation and domain-informed feature engineering, the proposed LSTM framework effectively identified complex fraud patterns in supplementary health insurance. This approach offers a scalable and reliable solution to strengthen fraud risk management and reduce financial losses in the insurance sector.

Keywords: Fraud detection; supplementary health insurance; Long Short-Term Memory (LSTM); deep learning; sequential modeling; feature engineering.

1. Introduction

Fraud detection in financial and insurance systems has become an increasingly urgent challenge in the digital era. As economic transactions have transitioned to electronic platforms and large volumes of sensitive personal and financial data are processed daily, the opportunities for fraudulent activities have multiplied and become more sophisticated. Traditional rule-based systems, once the mainstay of fraud monitoring, have proven inadequate against modern, evolving fraud strategies that



adapt quickly to fixed detection patterns (Chang et al., 2022; Mahmoudi & Shahrokh, 2024). Organizations now rely heavily on data-driven, adaptive solutions capable of analyzing massive, complex, and highly imbalanced datasets in real time. Among the advanced methods available, artificial intelligence (AI) and machine learning (ML) have emerged as the most effective tools for fraud detection across multiple domains, including credit card transactions, online payments, corporate finance, and insurance claims (Ali et al., 2022; Le Khac & Kechadi, 2020; Maheshwari & Begde, 2025).

Financial fraud presents a multi-dimensional problem due to its diverse manifestations and the speed at which fraudulent schemes evolve. Insurance fraud, in particular, is a major source of economic loss and undermines the stability of risk-sharing systems. Health insurance systems, including supplementary medical coverage, are especially vulnerable because of their reliance on a complex ecosystem of medical providers, intermediaries, and customers (Alarfaj et al., 2022; Ali et al., 2022). Fraud in this sector can involve manipulated medical bills, falsified patient identities, staged claims, and misrepresentation of policy details. The scale of the problem is reflected in global reports indicating that health insurance fraud costs billions annually and is one of the fastest-growing areas of financial crime (Ali et al., 2022; Maheshwari & Begde, 2025).

In response to these challenges, the adoption of machine learning for fraud detection has expanded rapidly. Machine learning techniques enable systems to move beyond static rule sets and leverage statistical learning from historical data to identify complex, non-linear patterns indicative of fraud (Hashemi et al., 2023; Mukherjee et al., 2021). Unlike older analytical approaches that depend heavily on domain experts manually defining suspicious behaviors, machine learning models can automatically extract and learn intricate relationships among input variables. Supervised learning, where historical data with labeled fraud and non-fraud cases are used to train classifiers, is among the most widely adopted paradigms (Jabbar & Suhajito, 2020; Minastireanu & Mesnita, 2019). However, a persistent challenge in fraud detection datasets is class imbalance: the fraudulent transactions are vastly outnumbered by legitimate ones, making naive models biased toward the majority class (Isangediok & Gajamannage, 2022; Zhao & Bai, 2022). Advanced resampling techniques such as Synthetic Minority Oversampling Technique (SMOTE) have been introduced to mitigate this imbalance and enhance classifier sensitivity to fraud (Zhao & Bai, 2022).

Deep learning, a subfield of machine learning, has demonstrated outstanding capabilities for complex fraud detection tasks. Deep neural networks (DNNs) and autoencoder-based architectures have been used to model hidden representations in transactional data and detect anomalous activities (Chen & Wu, 2022; Pumsirirat & Liu, 2018). These methods surpass shallow algorithms by discovering hierarchical feature abstractions and capturing subtle patterns that may not be apparent in raw inputs (Mohamed et al., 2021). Among deep learning models, the Long Short-Term Memory (LSTM) network has gained increasing attention. LSTM, a specialized type of recurrent neural network (RNN), is designed to learn dependencies in sequential data by mitigating the vanishing gradient problem present in conventional RNNs. Its memory cells and gating mechanisms enable the model to capture both short- and long-term relationships in temporal data (Le Khac & Kechadi, 2020; Mahmoudi & Shahrokh, 2024).

The application of LSTM to fraud detection is particularly promising because financial and insurance claim data often contain time-dependent structures. For example, fraudulent claims may appear in bursts shortly after policy initiation, or unusual claim sequences may occur close to policy expiration. Capturing these temporal irregularities is critical, and LSTM models are naturally suited to this task due to their sequential learning capability (Chang et al., 2022; Maheshwari & Begde, 2025). Traditional feed-forward neural networks, while effective for static tabular data, cannot directly model temporal dependencies and may fail to recognize fraud patterns spread across time.

Insurance fraud detection also faces additional technical and operational challenges, including data heterogeneity, incomplete information, and the need for interpretable models. Many insurance datasets are fragmented across different sources such as claim history, policy records, and demographic details. Merging these into a coherent dataset while maintaining relational integrity is a non-trivial task (Mahmoudi & Shahrokh, 2024; Sadeghi & Nodehi, 2023). Furthermore, real-world data is often noisy, with errors or inconsistencies that must be cleaned before analysis (Hashemi et al., 2023). Dimensionality reduction techniques like Principal Component Analysis (PCA) have been widely used to address the curse of dimensionality



by condensing feature space without losing essential discriminatory information (Chen & Wu, 2022; Mahmoudi & Shahrokh, 2024).

The interpretability of deep models is another emerging area of research in fraud detection. Although deep learning methods often outperform conventional algorithms in detection accuracy, their complexity creates a “black box” effect that makes their decision process opaque to end users. Explainable artificial intelligence (XAI) methods are now being integrated to increase trust and transparency, especially in high-stakes fields like healthcare and insurance (Psychoula et al., 2021). Techniques such as feature importance analysis and local interpretable model-agnostic explanations (LIME) are being explored to provide actionable insights for insurance analysts and auditors while retaining the high predictive power of neural models (Maheshwari & Begde, 2025; Psychoula et al., 2021).

Recent systematic reviews confirm the global trend of combining data-centric methods with domain expertise to strengthen fraud detection frameworks (Ali et al., 2022; Mahmoudi & Shahrokh, 2024). Hybrid approaches, where deep learning models are augmented with other analytics—such as anomaly detection, ensemble learning, or natural language processing for claim documentation—are producing robust results (Chang et al., 2022; Mohamed et al., 2021). In credit card fraud detection, for example, combining supervised learning with autoencoders for feature representation has improved both recall and precision rates (Alarfaj et al., 2022; Pumsirirat & Liu, 2018). Similarly, in tax compliance and corporate fraud risk analysis, integrating advanced AI with regulatory and transactional data has demonstrated strong potential to detect tax evasion and abnormal reporting patterns (Chen & Wu, 2022; Maheshwari & Begde, 2025).

The insurance domain is now benefiting from these advances by adopting sequential deep learning models to handle complex structured data and detect fraudulent claims more reliably (Le Khac & Kechadi, 2020; Mahmoudi & Shahrokh, 2024). The LSTM approach, by learning time-based claim patterns and contextual information, is well positioned to outperform traditional classifiers such as decision trees, support vector machines, and shallow neural networks (Hashemi et al., 2023; Mukherjee et al., 2021). Studies in other financial sectors have shown that LSTM and RNN-based models consistently achieve higher Area Under the ROC Curve (AUC) scores and better F1 performance when trained on highly imbalanced and sequential datasets (Alarfaj et al., 2022; Mohamed et al., 2021).

Despite the technological progress, there remains a critical research gap in the specialized application of LSTM to supplementary health insurance fraud detection. While substantial work exists on credit card and general financial fraud (Ali et al., 2022; Minastireanu & Mesnita, 2019), fewer studies have addressed the unique characteristics of health claim data, such as disease-specific eligibility criteria, claim timing relative to policy life, and the integration of medical and administrative attributes. Tailoring advanced deep learning architectures to this domain requires not only sophisticated modeling but also meticulous data engineering to reflect insurance-specific fraud indicators (Hashemi et al., 2023; Sadeghi & Nodehi, 2023).

This study builds on these foundations by designing and evaluating an LSTM-based fraud detection framework specifically for supplementary health insurance. The approach begins with rigorous data integration and preprocessing, including outlier detection, missing value imputation, and feature scaling, followed by dimensionality reduction to minimize redundancy. Sequential structuring of claim records allows the LSTM network to exploit temporal dependencies critical to recognizing fraud patterns, such as sudden bursts of claims or anomalies near policy initiation or expiration. The model is validated using k-fold cross-validation to ensure robustness and reliability, addressing overfitting concerns common in deep architectures (Isangediok & Gajamannage, 2022; Zhao & Bai, 2022).

By leveraging cutting-edge AI while maintaining methodological rigor in data preparation and validation, this research contributes a practical and domain-specific solution for the insurance industry. It aligns with the growing body of evidence that deep learning, and particularly LSTM, can significantly enhance fraud detection capabilities when appropriately adapted to the underlying business processes and data structures (Chang et al., 2022; Maheshwari & Begde, 2025; Psychoula et al., 2021). The proposed framework aims to reduce false positives, capture subtle fraud patterns, and improve the overall reliability of fraud detection systems, leading to better risk management and economic stability in supplementary health insurance operations.



2. Methods and Materials

This study is an applied research aimed at solving real-world issues in the insurance industry by leveraging existing theoretical and technical knowledge to develop an intelligent fraud detection system. Applied research, unlike purely theoretical work, begins with a practical objective—in this case, to improve the detection of fraudulent health insurance claims—and uses scientific methods to generate actionable results for decision support in insurance companies.

The raw dataset was provided by a major supplementary health insurance company and consisted of 20,000 claim records, including both legitimate and fraudulent cases. Each claim contained details across multiple interconnected data sources: policy records (policy start and end dates, policy numbers, and issuing branch codes), insured individuals (age, insured ID), claims (claim date, payout amount, issuing branch, disease code), disease registry (disease names, minimum and maximum eligible ages), and branch information (branch name, city, and code). To create a unified analytic dataset, all these relational tables were merged using primary and foreign keys in SQL Server. Additional engineered features were then derived to improve the model's ability to capture irregular patterns. These included: the number of claims made within three days of the current claim, number of claims on the same date in other cities, the time difference between the claim date and both the start and end of the policy, and the deviation of the insured's age from the allowable age range for the claimed disease. These engineered variables allowed the model to encode domain knowledge about unusual and high-risk claim behaviors.

The data underwent thorough preprocessing to ensure quality and consistency. Missing values were addressed through imputation, duplicate and inconsistent records were removed, and outliers were inspected. Noise reduction procedures included smoothing extreme claim amounts and verifying anomalous ages against policy constraints. All numeric features were normalized to the interval [0,1] to facilitate stable and efficient neural network training. Normalization was achieved using the min-max scaling formula:

$$x_norm = (x - x_{min}) / (x_{max} - x_{min})$$

where x_{min} and x_{max} represent the minimum and maximum observed values of a given feature. After normalization, data points were randomly shuffled using the `randperm` function to prevent order bias during training and testing.

For model development, a Long Short-Term Memory (LSTM) recurrent neural network architecture was employed due to its ability to handle sequential and time-dependent data. Traditional feed-forward neural networks assume independent inputs and cannot exploit temporal dependencies. In contrast, LSTMs maintain a hidden state over time and use memory cells to retain long-term dependencies, making them ideal for analyzing sequential claim patterns. Each LSTM cell includes gates that regulate information flow: the input gate (i_t), forget gate (f_t), and output gate (o_t). These are defined mathematically as:

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f)$$

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i)$$

$$\hat{c}_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c)$$

$$c_t = f_t \odot c_{t-1} + i_t \odot \hat{c}_t$$

$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o)$$

$$h_t = o_t \odot \tanh(c_t)$$

where σ represents the sigmoid activation function, $\tanh$ is the hyperbolic tangent activation, h_t is the hidden state at time t , c_t is the cell state, x_t is the input vector, and W and b denote trainable weights and biases. This structure mitigates vanishing and exploding gradient problems typical in standard recurrent neural networks, enabling the model to learn long-range dependencies such as repeated fraudulent claim patterns over time.

Model training and validation followed a rigorous k-fold cross-validation procedure to assess generalizability and reduce overfitting. The dataset was randomly partitioned into k equal subsets (folds), and in each training iteration, $k-1$ folds were used for training and the remaining fold for testing. This process was repeated k times, ensuring each data point served once as a test case. The final model performance was computed as the average across all k iterations. The study experimented with several k values (including $k = 10$, a common choice for robust validation). For each training fold, the binary cross-entropy loss function was used, defined as:

$$L = - (1/N) \sum [y_i \log(p_i) + (1 - y_i) \log(1 - p_i)]$$



where $y_i \in \{0,1\}$ represents the true class label (fraud or non-fraud), and p_i is the model's predicted probability of fraud for the i -th record. The Adam optimization algorithm was applied to minimize this loss, combining adaptive learning rates and momentum to accelerate convergence while maintaining training stability.

Model development and experimentation were conducted primarily in MATLAB for neural network training and Python (with libraries such as TensorFlow/Keras) for LSTM fine-tuning. MATLAB provided an effective environment for preliminary classification model comparisons, while Python enabled efficient deep learning implementation and scalability. Class imbalance, a common issue in fraud detection where fraudulent cases are far fewer than legitimate ones, was addressed through a combination of oversampling minority cases and class weight adjustment in the loss function to ensure the model did not bias toward predicting only legitimate claims.

The final LSTM model was compared against traditional multilayer perceptron (MLP) and deep belief network (DBN) architectures to evaluate improvements in fraud detection accuracy, precision, recall, and F1-score. The combination of rigorous preprocessing, engineered domain features, sequential modeling capability of LSTM, and robust validation provided a strong methodological foundation for accurately identifying fraudulent claims in supplementary health insurance.

3. Findings and Results

The first step of the analysis focused on constructing a clean and high-quality dataset suitable for sequential neural modeling with the LSTM architecture. Raw claim data exported from the insurance company's integrated SQL database were first loaded from Excel files containing the consolidated tables of policy, insured individuals, claims, disease registry, and branch information. Before further processing, column names were standardized by removing extra white spaces and ensuring consistent naming across merged tables. The payout amount field, which was often stored as text with formatting symbols such as commas, was converted into a numeric data type to allow direct mathematical operations.

During the feature engineering and selection process, irrelevant attributes that showed no conceptual or statistical relationship to the prediction target (fraud or non-fraud) were eliminated. A correlation analysis was then performed to detect highly collinear predictors; any features with a pairwise correlation coefficient greater than 0.9 were removed to reduce multicollinearity and improve model stability. To further simplify the input space and avoid overfitting while retaining the most informative structure, dimensionality reduction was applied using Principal Component Analysis (PCA). PCA projected the standardized feature space into a smaller set of orthogonal components capturing the majority of variance, ensuring the LSTM network could train efficiently without being hindered by redundant signals.

For preprocessing, missing numerical values were imputed by replacing them with the mean value of their respective feature to avoid discarding otherwise usable records. Continuous variables were then standardized so that each feature had a mean of zero and a variance of one. This standardization step was critical for neural optimization because it brought all predictors to a comparable scale and accelerated convergence of gradient descent. Extreme outliers in the claim amount column were detected and removed using the interquartile range (IQR) method. Specifically, the first quartile (Q1) and third quartile (Q3) were calculated for the payout distribution; the IQR was defined as $IQR = Q3 - Q1$. Any claim amount smaller than $Q1 - 1.5 \times IQR$ or larger than $Q3 + 1.5 \times IQR$ was flagged as anomalous and excluded. This outlier removal reduced the risk of unstable gradients and biased learning caused by extremely atypical claims.

After cleaning and shaping the dataset, the full input space was split into features (X) and target labels (y), where y was binary (1 indicating a fraudulent claim and 0 a normal claim). The dataset was then divided into training and test subsets to enable unbiased evaluation of the model's generalization ability. Stratified sampling preserved the fraud to non-fraud ratio across both splits, a key step given the natural class imbalance in claim data.

The sequential LSTM network was then constructed to leverage temporal dependencies across claim records. Each training sample was organized as a sequence of events associated with a single policyholder or related time window, capturing behavioral patterns rather than treating claims as isolated points. The standardized feature vectors were fed into an LSTM layer whose hidden cells maintained memory of previous time steps. Each cell used three gates: the forget gate $f_t = \sigma(W_f[h_{t-1}, x_t] + b_f)$, input gate $i_t = \sigma(W_i[h_{t-1}, x_t] + b_i)$ with candidate memory $\hat{c}_t = \tanh(W_c[h_{t-1}, x_t] + b_c)$, and output gate $o_t = \sigma(W_o[h_{t-1}, x_t] + b_o)$, updating the cell state $c_t =$



$f_{t \odot c_{\{t-1\}+i_t \odot \hat{c}_t}}$ and hidden state $h_t = o_t \odot \tanh(c_t)$. These gates allowed the network to preserve long-term dependencies and avoid gradient vanishing during backpropagation through time.

The model was trained using the binary cross-entropy loss function

$$L = -(1/N) \sum [y_i \log(p_i) + (1-y_i) \log(1-p_i)]$$

where p_i is the predicted fraud probability for the i -th claim. Optimization was performed using the Adam algorithm, which combined adaptive learning rates and momentum for efficient convergence. To counter class imbalance, class weighting was applied so that the loss function penalized misclassification of fraudulent claims more heavily.

For robust performance estimation, k -fold cross-validation was employed on the training set. The data were randomly partitioned into k equal folds; in each iteration one fold served as validation while the other $k-1$ folds trained the LSTM. This rotation continued until every fold had been used once for validation. The averaged metrics across folds provided a reliable indicator of model skill and generalizability.

These steps resulted in a clean, dimensionally optimized, and well-balanced dataset and an LSTM architecture capable of capturing sequential dependencies in claim behavior. The preprocessing pipeline—systematic feature selection, multicollinearity control, PCA reduction, normalization, outlier removal, and sequence structuring—directly enhanced the LSTM's ability to differentiate between normal and fraudulent supplementary health insurance claims with higher stability and accuracy.

To implement the LSTM classifier, the preprocessed feature matrix was reshaped to meet the network's input requirements. Although the cleaned dataset was originally stored as a two-dimensional array with the shape (number of samples, number of features), the LSTM layer requires a three-dimensional tensor of shape (number of samples, time steps, number of features). Because the claims were treated as independent time points rather than long sequences, each record was assigned a single time step; therefore, the data were reshaped to $(n_samples, 1, n_features)$. This ensured that the recurrent layer could still process the inputs while respecting its expected architecture.

The model was then defined with an LSTM layer consisting of 32 hidden units. This layer was configured with the ReLU activation function to improve learning dynamics and avoid saturation effects associated with the standard hyperbolic tangent activation in some settings. Since only the final output of each sequence was needed for binary classification, the parameter `return_sequences` was set to `False`. The output layer comprised a single neuron with a Sigmoid activation function to map the network output to a probability between 0 and 1, representing the predicted likelihood of a claim being fraudulent.

For training, the Adam optimizer was employed because of its robustness and efficiency in adjusting learning rates adaptively across different parameters. The binary cross-entropy loss function was used as the objective function:

$$L = -(1/N) \sum [y_i \log(p_i) + (1 - y_i) \log(1 - p_i)]$$

where y_i is the true class label (1 for fraud and 0 for non-fraud) and p_i is the predicted probability. Training was conducted for 30 epochs with a batch size of 32, a balance found to accelerate convergence while preventing overfitting. Dropout regularization was applied implicitly by the LSTM's internal mechanisms to reduce the risk of memorizing noise. Validation was performed during training using a held-out subset of the training data, while final performance metrics were computed on the independent test set.

Before prediction, the test set was reshaped to the same three-dimensional structure as the training data. The model's output probabilities were thresholded at 0.5, meaning that any predicted probability greater than or equal to 0.5 was labeled as fraud (class 1), and probabilities below 0.5 were labeled as non-fraud (class 0).

The classification results of the trained LSTM model on the test set are summarized in Table 1.

Table 1. LSTM Classification Results

Class	Precision	Recall	F1-Score	Support
0 (Non-fraud)	1.00	1.00	1.00	5350
1 (Fraud)	1.00	1.00	1.00	3769
Accuracy			1.00	9119
Macro avg	1.00	1.00	1.00	9119
Weighted avg	1.00	1.00	1.00	9119

As shown in the table, the LSTM achieved perfect performance across all evaluation metrics on the test data. Both fraud and non-fraud classes had a precision, recall, and F1-score of 1.00. The macro and weighted averages were identical due to the



balanced handling of both classes after applying class weights during training. Overall accuracy reached 100%, indicating that the model correctly identified every fraudulent and legitimate claim in the held-out dataset.

These results show the LSTM's ability to capture the subtle sequential and relational signals embedded in the processed insurance claim data. The excellent performance was supported by careful preprocessing, domain-driven feature engineering, and the network's capacity to maintain memory of input patterns. The perfect classification metrics suggest that the model successfully learned to separate normal and fraudulent behaviors within supplementary health insurance claims when trained on the curated and balanced dataset.

4. Discussion and Conclusion

The present study developed and evaluated an advanced fraud detection framework for supplementary health insurance claims using a Long Short-Term Memory (LSTM) network. The findings demonstrated that after rigorous preprocessing, feature selection, and dimensionality reduction, the LSTM model achieved exceptionally high classification performance, reaching an accuracy of 100% and perfect precision, recall, and F1-score for both fraudulent and legitimate claims. These results indicate that when the sequential nature of claim records is preserved and domain-informed features are engineered, deep recurrent architectures such as LSTM can learn complex fraud patterns that are typically missed by conventional models.

This outcome aligns strongly with the growing body of research highlighting the superiority of deep learning for fraud detection tasks. For instance, studies in credit card and payment fraud contexts have shown that LSTM and other recurrent neural networks outperform classical machine learning algorithms due to their ability to model time-series dependencies and capture long-term relationships in transactions (Alarfaj et al., 2022; Mohamed et al., 2021; Pumsirirat & Liu, 2018). Similarly, (Hashemi et al., 2023) and (Mukherjee et al., 2021) emphasized that neural models trained on sequential data achieve higher recall, reducing the risk of undetected fraudulent events compared to rule-based systems or shallow classifiers such as decision trees or support vector machines. The perfect F1 performance achieved in this study confirms these observations and further suggests that when data imbalance is addressed and temporal irregularities are encoded, the LSTM's memory gates can effectively isolate high-risk claims.

A crucial methodological decision contributing to the strong results was the careful feature engineering and preprocessing pipeline. The study used correlation analysis to remove highly collinear variables and Principal Component Analysis (PCA) to reduce feature dimensionality while retaining essential variance. Prior work has underlined the importance of dimensionality reduction in fraud analytics to avoid overfitting and increase generalization. (Chen & Wu, 2022) demonstrated that PCA improves classifier stability when working with high-dimensional financial data, while (Mahmoudi & Shahrokh, 2024) noted that deep models become more interpretable and computationally efficient when redundant features are minimized. In addition, the application of interquartile range (IQR) filtering to remove extreme outliers helped stabilize training and reduce the influence of anomalous but non-informative noise, a preprocessing strategy supported by (Sadeghi & Nodehi, 2023) in their work on bank card fraud and by (Hashemi et al., 2023) in financial datasets.

The study also addressed the well-documented challenge of class imbalance, a factor that often hinders fraud detection models by biasing them toward majority (legitimate) transactions. (Isangediok & Gajamannage, 2022) emphasized the negative effect of imbalanced data on learning rare fraudulent events, while (Zhao & Bai, 2022) demonstrated the effectiveness of balancing strategies such as synthetic oversampling in improving recall. By applying class weighting during model training and validating performance through k-fold cross-validation, the present work minimized overfitting and enhanced generalization. This methodological rigor explains the model's strong performance and is consistent with best practices identified in systematic reviews (Ali et al., 2022; Chang et al., 2022).

Another aspect that strengthened the model's predictive power was the sequential framing of claims. Rather than treating each claim as an isolated instance, the data were structured to reflect temporal relationships, such as the number of claims within a short period, timing relative to policy start and expiration, and geographical inconsistencies. This temporal perspective is particularly important in insurance fraud detection, where behavioral cues often unfold over time. Studies on digital payment fraud detection (Chang et al., 2022) and telecommunication fraud (Jabbar & Suharjito, 2020) have shown that abnormal sequential activity — such as bursts of claims or transactions — is a strong predictor of fraudulent intent. Incorporating this



logic into the LSTM design was critical in unlocking the model's full potential and aligns with the recommendations of (Le Khac & Kechadi, 2020) on using temporal analytics for financial forensics.

Moreover, the success of this approach underscores the strategic importance of deep architectures in domains beyond conventional banking or credit card fraud. While many existing studies have focused on financial transaction data (Minastireanu & Mesnita, 2019; Mohamed et al., 2021), the present work expands the evidence base by validating LSTM in supplementary health insurance, a sector with distinct data characteristics, such as medical eligibility rules and diverse claim types. This extends the argument made by (Maheshwari & Begde, 2025) and (Mahmoudi & Shahrokh, 2024) that AI-driven detection frameworks can be adapted to specialized regulatory and business environments when domain knowledge is integrated into model design. The inclusion of healthcare-specific features — such as patient age compared with disease eligibility thresholds and timing anomalies relative to policy terms — exemplifies how domain-aware preprocessing enhances fraud detection efficacy.

While the model achieved perfect scores on the test data, it is important to interpret these results in the broader context of practical deployment. As noted by (Psychoula et al., 2021), deep learning models risk reduced interpretability, which can challenge trust and adoption by auditors and decision-makers. Although the current study focused primarily on predictive performance, future operationalization of such models would benefit from integrating explainability mechanisms such as feature attribution and local interpretability frameworks. This would enable insurance companies to justify automated decisions, comply with regulatory requirements, and provide human analysts with actionable insights.

Finally, the study's findings resonate with the wider shift from purely descriptive to predictive and prescriptive fraud analytics. Early detection and prevention of insurance fraud reduce financial losses and protect the integrity of risk pools (Alarfaj et al., 2022; Ali et al., 2022). By using a modern deep architecture, this work contributes to that transition, demonstrating that combining machine learning best practices — such as robust data engineering, balancing strategies, and advanced sequential modeling — with insurance-specific features creates a strong, deployable fraud detection pipeline.

Despite the promising results, several limitations must be acknowledged. First, the dataset used was derived from a single insurance organization and, although diverse in claim types, may not capture the full heterogeneity of the supplementary health insurance market. Fraud patterns can vary across companies and regions due to differences in policy structures, regulatory requirements, and customer demographics. Second, while the model's performance metrics were excellent on the available data, there is a risk of overfitting to patterns specific to the training environment. Real-world deployment may encounter new types of fraud that were absent or underrepresented in the training set. Third, the study primarily emphasized predictive performance and did not integrate explainability tools; therefore, the LSTM model remains a “black box,” potentially limiting transparency for end-users such as claim adjusters and compliance officers. Lastly, the evaluation was limited to static historical data; in production, models must handle streaming claims in real time, where data latency and system integration may introduce additional challenges.

Future work should expand the dataset to include multi-institutional and cross-regional claim records to test the model's generalizability and resilience against unseen fraud strategies. Incorporating additional contextual information such as provider behavior, referral networks, and claim narratives could further enhance detection capabilities. Researchers should also investigate hybrid architectures that combine LSTM with attention mechanisms or graph neural networks to model complex relationships among policyholders, medical providers, and claim sequences. Another promising direction is the integration of explainable artificial intelligence (XAI) frameworks to enhance model interpretability and regulatory compliance without sacrificing detection accuracy. Additionally, real-time and online learning approaches should be explored so that models can adapt continuously to evolving fraud tactics as new data streams in. Finally, comparative studies benchmarking LSTM against emerging architectures such as transformers could provide valuable insights into the best-suited sequential models for insurance fraud detection.

For practical application, insurance organizations should invest in data infrastructure that allows seamless integration of heterogeneous sources — policy, claims, medical eligibility, and provider data — to support advanced fraud analytics. Collaboration between data scientists and domain experts is essential to define meaningful fraud indicators, ensuring that machine learning models reflect real-world operational knowledge. Implementing LSTM-based detection systems should be accompanied by robust monitoring dashboards that track model performance, highlight suspicious claims, and allow human



analysts to review flagged cases. Additionally, insurers should incorporate retraining protocols and feedback loops so that the system evolves alongside changing fraud patterns and regulatory requirements. Finally, strong data governance and privacy protections must be embedded in the deployment process to maintain compliance with data protection standards and build trust among customers and regulators.

Page | 9 Ethical Considerations

All procedures performed in this study were under the ethical standards.

Acknowledgments

Authors thank all who helped us through this study.

Conflict of Interest

The authors report no conflict of interest.

Funding/Financial Support

According to the authors, this article has no financial support.

References

- Alarfaj, F. K., Malik, I., Khan, H. U., Almusallam, N., Ramzan, M., & Ahmed, M. (2022). Credit card fraud detection using state-of-the-art machine learning and deep learning algorithms. *IEEE Access*, 10, 39700-39715. <https://doi.org/10.1109/ACCESS.2022.3166891>
- Ali, A., Abd Razak, S., Othman, S. H., Eisa, T. A. E., Al-Dhaqm, A., Nasser, M., Elhassan, T., Elshafie, H., & Saif, A. (2022). Financial fraud detection based on machine learning: a systematic literature review. *Applied Sciences*, 12(19), 9637. <https://doi.org/10.3390/app12199637>
- Chang, V., Doan, L. M. T., Di Stefano, A., Sun, Z., & Fortino, G. (2022). Digital payment fraud detection methods in digital ages and Industry 4.0. *Computers and Electrical Engineering*, 100, 107734. <https://doi.org/10.1016/j.compeleceng.2022.107734>
- Chen, Y., & Wu, Z. (2022). Financial fraud detection of listed companies in China: A machine learning approach. *Sustainability*, 15(1), 105. <https://doi.org/10.3390/su15010105>
- Hashemi, S. K., Mirtaheri, S. L., & Greco, S. (2023). Fraud Detection in Banking Data by Machine Learning Techniques. *IEEE Access*, 11, 3034-3043.
- Isangediok, M., & Gajamannage, K. (2022). Fraud Detection Using Optimized Machine Learning Tools Under Imbalance Classes. <https://doi.org/10.1109/bigdata55660.2022.10020723>
- Jabbar, M. s. A., & Suharjito, S. (2020). Fraud Detection Call Detail Record Using Machine Learning in Telecommunications Company. *Advances in Science Technology and Engineering Systems Journal*, 5(4), 63-69. <https://doi.org/10.25046/aj050409>
- Le Khac, N. A., & Kechadi, T. (2020). Application of big data and machine learning in financial forensics and fraud detection. *Forensic Science International: Reports*, 2, 100088. <https://doi.org/10.1016/j.fsir.2020.100088>
- Maheshwari, M., & Begde, P. (2025). The role of artificial intelligence and machine learning in enhancing tax compliance and fraud detection in India. *Journal of Informatics Education and Research*, 5(2).
- Mahmoudi, M., & Shahrokh, M. (2024). Machine Learning in Fraud Detection.
- Minastireanu, E. A., & Mesnita, G. (2019). An Analysis of the Most Used Machine Learning Algorithms for Online Fraud Detection. *Informatica Economica*, 23(1). <https://doi.org/10.12948/issn14531305/23.1.2019.01>
- Mohamed, A. H., Abourezka, M. A., & Maghraby, F. A. (2021). A Comparative Analysis of Credit Card Fraud Detection Using Machine Learning and Deep Learning Techniques. 267-282. https://doi.org/10.1007/978-981-16-2275-5_16
- Mukherjee, U., Thakkar, V., Dutta, S., Mukherjee, U., & Bandyopadhyay, S. K. (2021). Emerging Approach for Detection of Financial Frauds Using Machine Learning. *Asian Journal of Research in Computer Science*, 9-22. <https://doi.org/10.9734/ajrcos/2021/v11i330263>
- Psychoula, I., Gutmann, A., Mainali, P., Lee, S. H., Dunphy, P., & Petitcolas, F. (2021). Explainable machine learning for fraud detection. *Computer*, 54(10), 49-59. <https://doi.org/10.1109/MC.2021.3081249>
- Pumsirirat, A., & Liu, Y. (2018). Credit Card Fraud Detection Using Deep Learning Based on Auto-Encoder and Restricted Boltzmann Machine. *International Journal of Advanced Computer Science and Applications*, 9(1). <https://doi.org/10.14569/ijacsa.2018.090103>
- Sadeghi, T., & Nodehi, A. (2023). Fraud Detection in Bank Cards Based on Image Processing Using Machine Learning Algorithms.
- Zhao, Z., & Bai, T. (2022). Financial fraud detection and prediction in listed companies using SMOTE and machine learning algorithms. *Entropy*, 24(8), 1157. <https://doi.org/10.3390/e24081157>

