

The Interpretable Artificial Neural Network in Vehicle Insurance Claim Fraud Detection Based on Shapley Additive Explanations

Alan Wilson ¹, Kaixian Xu ², Zhaoyan Zhang ³ and Yu Qiao ^{4,*}

¹ Intact Financial Corporation, Toronto, ON M5G 0A1, Canada

² Risk & Quant Analytics, BlackRock, Jersey City, NJ 08540, USA

³ Zhongke Zhidao (Beijing) Technology Co., Ltd., Beijing 102627, China

⁴ Meta Platforms, Inc., Bellevue, WA 98005, USA

Abstract: Vehicle insurance claim fraud presents a major challenge in the insurance industry, leading to financial losses and increased premiums for policyholders. Traditional fraud detection methods, such as rule-based systems and manual claim assessment, struggle to handle the complexity and growing volume of fraudulent claims. With the advancement of Machine Learning (ML), models such as Artificial Neural Networks (ANNs) have significantly improved fraud detection accuracy. However, a key limitation of existing ML-based methods is their lack of interpretability, making it difficult for insurers to justify fraud detection decisions. To address this issue, this study proposes an interpretable fraud detection framework based on an ANN integrated with Shapley Additive Explanations (SHAP). The framework involves preprocessing insurance claim data, training an ANN for fraud prediction, and applying SHAP to analyze feature importance and provide interpretability. Experimental results demonstrate that the proposed model achieves high accuracy in fraud detection while offering insights into influential features affecting claim decisions. The findings highlight the importance of incorporating explainability into ML-based fraud detection, ensuring transparency and trustworthiness in the insurance industry.

Keywords: vehicle insurance claim fraud detection; shapley additive explanations; machine learning; neural network

1. Introduction

Vehicle insurance claim fraud has become a significant challenge in the insurance industry [1–3], leading to substantial financial losses and increased premiums for policyholders. It encompasses various deceptive activities, such as exaggerated claims, staged accidents, false documentation, and misrepresentation of damages or injuries. Fraudulent claims not only strain the financial stability of insurance companies but also impact genuine policyholders by increasing overall insurance costs [4, 5]. Consequently, developing effective and reliable fraud detection methods has become a critical necessity for insurance companies to mitigate risks and ensure fairness in claim settlements.

The global insurance industry faces significant financial burdens due to fraudulent claims. Studies suggest

that insurance fraud accounts for billions of dollars in losses annually, with a considerable portion attributed to vehicle insurance claims. Traditional methods of fraud detection, such as manual claim assessment and rule-based systems, are often inefficient and time-consuming [6–8], making them inadequate for handling the increasing volume of claims. Additionally, sophisticated fraudulent activities have evolved over time, making it challenging for insurers to detect fraud using conventional approaches. The necessity for more advanced and automated detection methods has become evident [9–11], as timely identification of fraudulent claims can save insurance companies from considerable financial losses and improve overall operational efficiency.

Over the last decades, traditional fraud detection methods have relied on rule-based systems, expert-driven heuristics, and statistical models which have been widely used in many domains [12,13]. These approaches primarily involve predefined rules, such as flagging claims that exceed a certain threshold or exhibit unusual patterns. While rule-based systems can effectively detect simple fraud cases, they suffer from limited adaptability and high false-positive rates. Moreover, fraudsters often find ways to circumvent these predefined rules, making them less effective in handling complex fraud schemes. Statistical models, such as logistic regression [14,15], have been explored for fraud detection by analyzing claim characteristics and identifying anomalies. However, these methods struggle to capture the intricate patterns and nonlinear relationships present in fraudulent claims.

With the advent of Machine Learning (ML) [16–19], fraud detection has witnessed significant progress. Supervised and unsupervised ML algorithms have been widely employed to improve the accuracy and efficiency of fraud detection models. Supervised learning methods, including decision trees, random forests, Support Vector Machines (SVMs), and Artificial Neural Networks (ANNs) [20–22], leverage labeled datasets to identify fraudulent claims based on historical patterns. Unsupervised learning techniques, such as clustering and anomaly detection, have also been utilized to uncover hidden fraud patterns in claim data. The ability of machine learning models to process large volumes of data [23–25] and identify subtle fraudulent behaviors has greatly enhanced the detection of insurance fraud. However, despite these advancements, the interpretability of ML-based fraud detection models remains a challenge.

One of the primary limitations of existing machine learning models in fraud detection is their lack of interpretability. Many high-performing models, especially deep learning-based approaches like ANNs [26–28], operate as “black boxes,” making it difficult to understand the reasoning behind their predictions. In the insurance industry, where transparency and accountability are crucial, decision-makers require interpretable models to justify claim rejections and comply with regulatory standards. Without proper explanations, insurers may face challenges in defending their fraud detection decisions, potentially leading to legal and ethical concerns. Therefore, there is an urgent need for interpretable fraud detection frameworks that not only achieve high accuracy but also provide clear insights into the factors influencing fraud predictions.

To address the interpretability challenge in vehicle insurance fraud detection, this study proposes an explainable ANN framework based on Shapley Additive Explanations (SHAP). The proposed framework integrates multiple stages, including data preprocessing, ANN-based fraud detection, and SHAP-based interpretability analysis. As illustrated in Figure 1, the framework begins with data preprocessing, which involves string feature transformation, normalization, data balancing, and data splitting. The processed data is then fed into an ANN model for fraud prediction, leveraging its ability to capture complex patterns in insurance claim data. Finally, SHAP analysis is employed to interpret the ANN predictions by generating feature importance visualizations, such as summary plots, dependence plots, force plots, and waterfall plots. This interpretability enables insurers to gain deeper insights into fraudulent claim patterns, make informed decisions, and build trust with policyholders.

2. Literature Review

2.1. Insurance Fraud Detection

Insurance fraud detection has attracted considerable attention, prompting the development of diverse machine learning approaches aimed at improving detection accuracy and efficiency due to their strong

performance across various tasks. For example, Gangadhar et al. introduced a Chaotic Variational Autoencoder-based one-class classifier specifically designed for insurance fraud detection, demonstrating superior performance in identifying fraudulent transactions [29]. Additionally, Asgarian et al. developed AutoFraudNet, a multimodal network that leverages multiple data modalities to improve fraud detection in the auto insurance sector [30]. In another study, Gupta et al. applied a Markov model combined with machine learning techniques to detect fraud in health insurance, achieving high accuracy and F1-scores, underscoring the model's effectiveness [31].

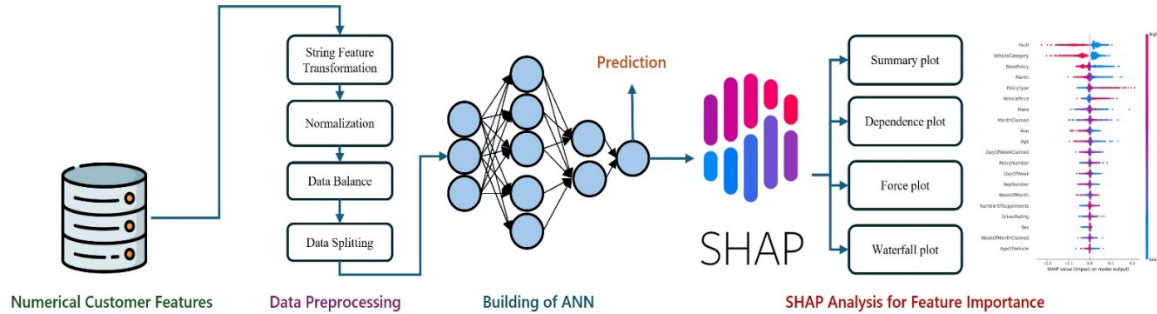


Figure 1. The process of the proposed SHAP-based interpretable ANN in vehicle insurance claim fraud detection.

2.2. The Advancements Related to the Interpretability of Machine Learning Models

As ML models become increasingly complex, the need for interpretability has grown across various domains. Numerous studies have proposed interpretability methods to enhance model transparency and trustworthiness. Ribeiro et al. introduced Local Interpretable Model-agnostic Explanations (LIME), which approximates black-box model predictions using interpretable local surrogate models, making it widely used for feature importance analysis [32]. Lundberg and Lee proposed Shapley Additive Explanations (SHAP), which utilizes cooperative game theory to assign feature importance scores and has become a standard in ML interpretability [33]. In deep learning, Selvaraju et al. developed Grad-CAM, a method that visualizes convolutional neural network (CNN) activations to highlight regions influencing predictions, improving interpretability in image classification tasks [34]. Additionally, Shrikumar et al. introduced DeepLIFT, a technique that tracks contributions of input features in deep networks, making it useful for medical and financial applications [35]. For tree-based models, Lundberg et al. extended SHAP for gradient-boosting models like XGBoost, offering a more robust feature attribution framework [36]. Caruana et al. proposed Generalized Additive Models with Pairwise Interactions (GA2Ms) to balance interpretability and predictive power in structured data analysis [37]. Although interpretability methods have been extensively studied in domains such as healthcare, finance, and image recognition, their application to vehicle insurance claim fraud detection remains limited. The lack of interpretable fraud detection frameworks hinders transparency and regulatory compliance in the insurance industry. This gap underscores the necessity of integrating interpretable machine learning models into fraud detection systems to enhance decision-making and accountability.

3. Method

3.1. Dataset Description and Preprocessing

Our study employs a publicly available dataset from Kaggle, consisting of 32 features designed to identify fraudulent vehicle insurance claims. The primary goal is to develop an effective fraud detection model, using the “FraudFound_P” feature as the target variable. The dataset encompasses a range of attributes related to insurance claims, including policyholder details, claim characteristics, and accident-specific information. Notable features include Month, WeekOfMonth, and DayOfWeek, among others. The target variable, “FraudFound_P,” classifies claims as either fraudulent (1) or legitimate (0). Within the dataset, fraudulent claims constitute approximately 5.99% of the total, while legitimate claims account for 94.01%.

To improve model performance, we applied several preprocessing steps. First, categorical string features were converted into numerical representations. Next, normalization was performed to ensure that all numerical

features were on a consistent scale. Given the class imbalance in the dataset, we employed the Synthetic Minority Over-sampling Technique (SMOTE) [38–40] to balance the distribution of fraudulent and legitimate claims. Finally, the dataset was split into training and testing sets in a 7:3 ratio, ensuring a reliable evaluation of the model's effectiveness.

3.2. Artificial Neural Network

Artificial Neural Networks (ANNs) are a class of machine learning models inspired by the structure and functionality of biological neural networks, which are widely used in many domains [41–43]. They consist of interconnected layers of artificial neurons that process information through weighted connections and activation functions. Each neuron receives inputs, applies a transformation using a weight and bias system, and passes the result through an activation function to determine the output. ANNs are widely used in classification and regression tasks due to their ability to capture complex patterns in data [44,45].

In this study, we designed an ANN to detect fraudulent vehicle insurance claims. The model architecture consists of an input layer, multiple hidden layers, and an output layer. The input layer takes numerical features from the preprocessed dataset. Given that the dataset contains 32 features, the input layer is designed to accommodate these variables. The hidden layers play a crucial role in feature extraction and representation learning. Our model includes three hidden layers, each containing a different number of neurons to effectively capture intricate relationships within the data. The first hidden layer consists of 32 neurons, followed by two additional layers with 16 neurons each. These layers are designed to progressively extract meaningful patterns from the input features. To introduce non-linearity and enhance the learning capability of the network, an activation function is applied to each neuron in the hidden layers. This activation function helps the model learn complex representations by introducing non-linear transformations to the data. Additionally, to prevent overfitting and ensure efficient learning, appropriate weight initialization and regularization techniques are incorporated. The output layer consists of a single neuron, responsible for producing the final classification result—whether a claim is fraudulent or legitimate. Given that this is a binary classification task, a suitable activation function is applied in the output layer to ensure the predicted value falls within the expected range. The network is trained using a binary cross-entropy loss function, which measures the difference between predicted probabilities and actual labels. For optimization, the model utilizes an adaptive gradient-based optimization algorithm to update weights and minimize the loss function efficiently. The training process is conducted for 20 epochs with a batch size of 32, ensuring a balanced trade-off between computational efficiency and model convergence.

3.3. Shapley Additive Explanations

SHAP is a widely used interpretability method for machine learning models [46–48], based on Shapley values from cooperative game theory. It provides a principled way to quantify the contribution of each feature to a model's prediction by distributing the total prediction difference among input features fairly. SHAP is particularly valuable for complex models, such as artificial neural networks, where interpretability is often a challenge. By assigning importance scores to individual features, SHAP helps in understanding how different inputs influence the model's output.

A key advantage of SHAP is its consistency and ability to provide both global and local interpretability. Global interpretability allows us to analyze the overall impact of features across all predictions, while local interpretability explains individual predictions, making it possible to understand why a specific claim was classified as fraudulent or legitimate. The SHAP framework offers various visualization tools, such as summary plots, dependence plots, force plots, and waterfall plots, to enhance the interpretability of machine learning models.

In this study, we employed SHAP to analyze the output of our artificial neural network model for Vehicle Insurance Claim Fraud Detection. By computing SHAP values, we identify the most influential features contributing to fraud classification. This helps in understanding which factors drive fraudulent claims and provides transparency in decision-making. Specifically, SHAP allows us to explore the impact of properties such as claim-related attributes, policyholder information, and accident details on fraud detection. The insights gained from SHAP analysis not only improve model interpretability but also assist insurance companies in better

assessing risk factors and enhancing fraud detection strategies.

4. Results and Discussion

4.1. The Prediction Performance of the ANN

To assess the effectiveness of the proposed ANN in vehicle insurance claim fraud detection, we evaluated its performance using multiple metrics, including accuracy, precision, recall, F1-score, and Area Under the Curve (AUC). The results are illustrated in Figures 2–5, detailing the training process, performance metrics, ROC curve, and confusion matrix.

Figure 2 presents the accuracy and loss curves during model training. The left plot displays the accuracy of both the training and validation datasets over 20 epochs. The training accuracy steadily increases, reaching over 90%, while the validation accuracy fluctuates around 82%, indicating some generalization challenges. The right plot depicts the training and validation loss trends. The training loss consistently decreases, demonstrating effective learning, whereas the validation loss increases after a few epochs, suggesting possible overfitting. These trends indicate that while the model learns well on the training set, further tuning may be required to enhance generalization.

Figure 3 shows the key performance metrics of the trained model. The ANN achieved an accuracy of 83.26%, an F1-score of 84.10%, a recall of 88.53%, and a precision of 80.09%. These results highlight the model's effectiveness in detecting fraudulent claims, with a strong recall score indicating that the majority of fraud cases were correctly identified. However, the slightly lower precision suggests that some legitimate claims were misclassified as fraudulent, which may require further refinement of the model.

Figure 4 illustrates the Receiver Operating Characteristic (ROC) curve, which is a graphical representation of the model's ability to distinguish between fraudulent and legitimate claims at various classification thresholds. The curve plots the True Positive Rate (TPR) against the False Positive Rate (FPR), with the diagonal line representing a random classifier. The proposed ANN achieves an AUC score of 0.91, demonstrating a high level of discrimination between the two classes. A higher AUC value indicates better overall performance, confirming the model's reliability in fraud detection.

Figure 5 displays the confusion matrix, which provides a detailed breakdown of the model's classification results. The matrix shows that the model correctly classified 3392 legitimate claims and 3850 fraudulent claims. However, 957 legitimate claims were misclassified as fraudulent (false positives), and 499 fraudulent claims were missed (false negatives). While the high recall score indicates that most fraud cases were detected, the presence of false positives highlights the need for further tuning to reduce unnecessary fraud alerts.

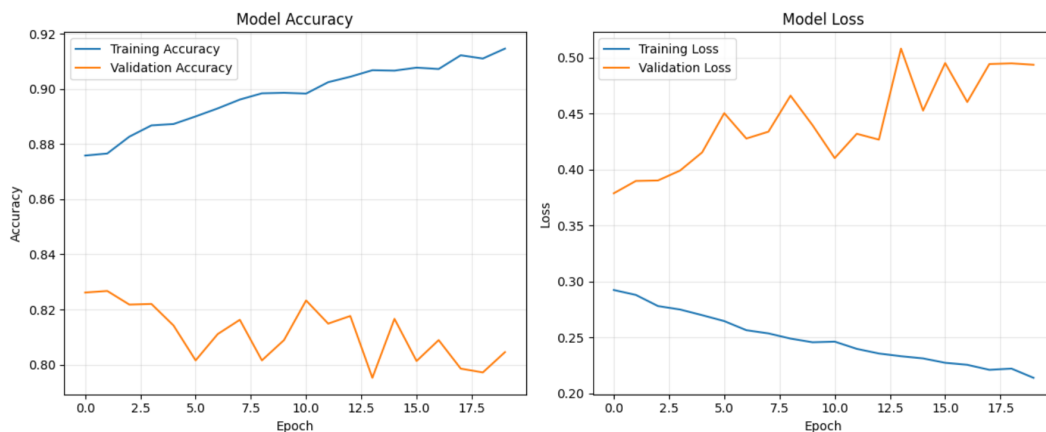


Figure 2. The accuracy and loss curve during the model training process.

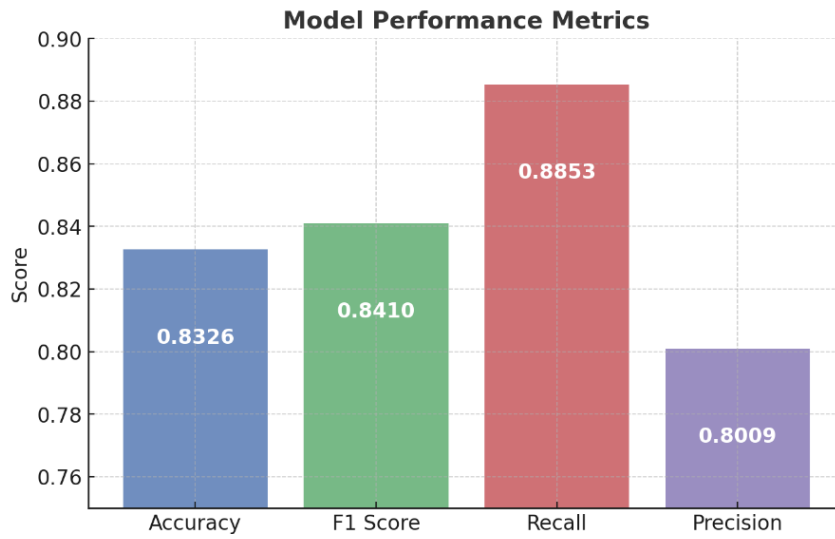


Figure 3. The model performance evaluated by different metrics.

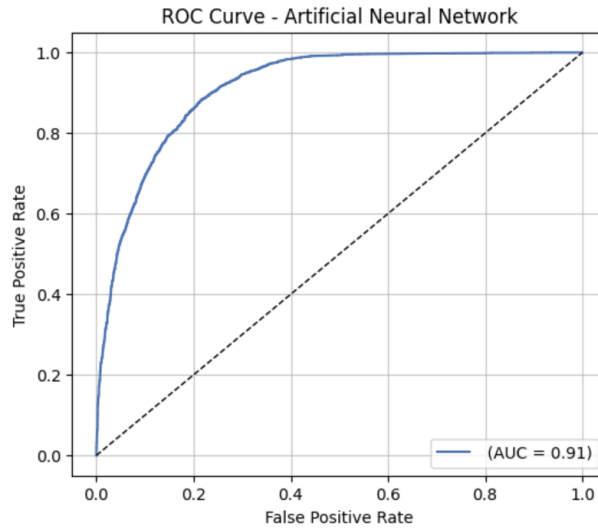


Figure 4. The ROC curve of the proposed ANN.

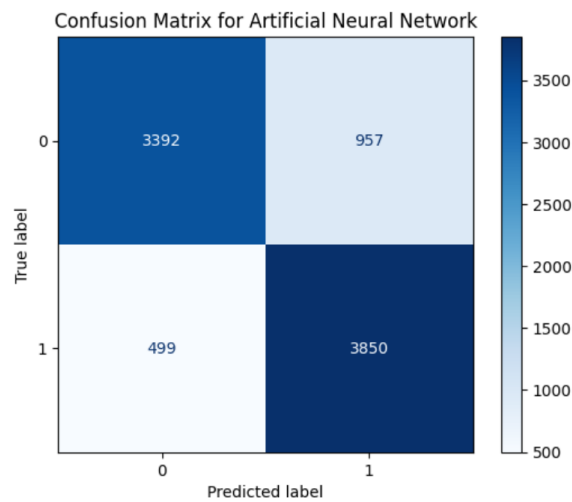


Figure 5. The confusion matrix of the proposed ANN.

4.2. The Interpretability Analysis of the Proposed ANN Model Based on SHAP

Figure 6 presents the SHAP-based feature importance analysis for some crucial properties, which provides insights into how different attributes influence the prediction of fraudulent vehicle insurance claims. The left panel illustrates the distribution of SHAP values for each feature, where the color gradient represents feature

values (blue for low values and red for high values). The right panel ranks the features based on their average absolute SHAP values, highlighting their overall impact on the model's output.

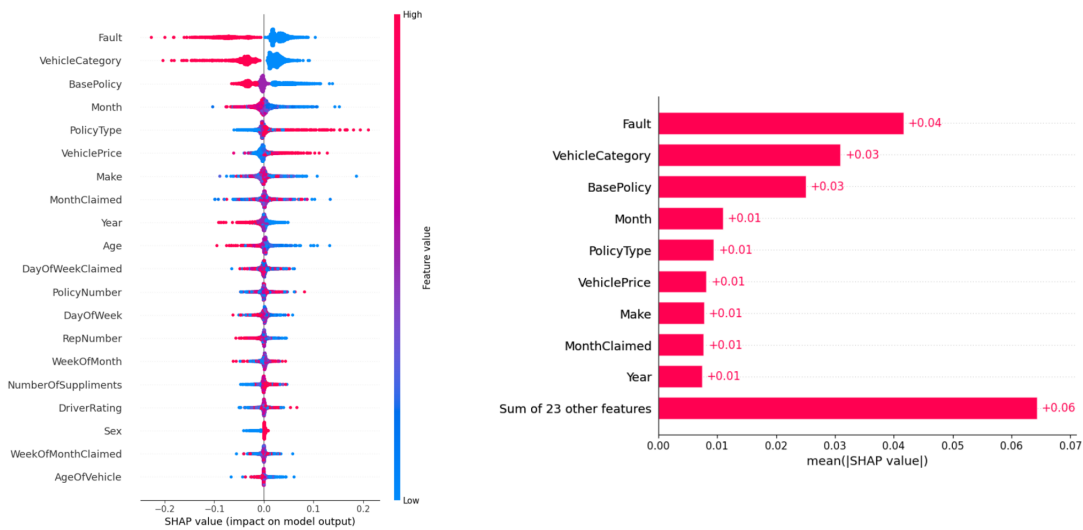


Figure 6. The importance of some crucial features (summary plot) based on SHAP.

Among all the features, “Fault” exhibits the highest contribution to fraud prediction, suggesting that whether the policyholder is at fault in an accident significantly affects the likelihood of a fraudulent claim. “VehicleCategory” follows as the second most influential factor, indicating that certain vehicle categories may have a higher fraud risk. “BasePolicy”, which represents the type of insurance policy, also plays a critical role, suggesting that different coverage plans may influence fraud patterns.

Other important features include “Month”, “PolicyType”, and “VehiclePrice”, which contribute to fraud classification, though with slightly lower impact. The presence of “Make”, referring to the car manufacturer, implies that fraud trends may vary across vehicle brands. Additionally, “MonthClaimed” and “Year” are also among the top contributing factors, indicating that temporal patterns in claim filings may be relevant for fraud detection.

Figure 7 presents the SHAP dependence plots for some selected features, illustrating how individual features influence the model's predictions. Each plot shows the SHAP values (impact on the model output) for a given feature, with color gradients representing another related feature. These visualizations provide deeper insights into the relationships between features and their effect on fraud detection. The first plot shows the Month feature, where SHAP values tend to increase in later months of the year, suggesting that claims filed in certain months may have a higher likelihood of being fraudulent. A similar trend is observed in the DayOfWeek and DayOfWeekClaimed features, where certain days exhibit a stronger impact on fraud detection. The WeekOfMonth plot indicates a slight negative correlation between the feature and its SHAP values, meaning that claims filed earlier in the month might have a different fraud risk than those filed later. Additionally, the Make plot suggests that certain vehicle brands influence the fraud likelihood differently, with varying SHAP distributions across different car manufacturers. Other notable dependencies include AccidentArea, where claims in urban areas appear to have slightly higher SHAP values compared to rural areas, and VehicleCategory, which exhibits variations in impact across different vehicle classes.

Figures 8 and 9 illustrate the force and waterfall plots generated by SHAP for four individual predictions. These visualizations help explain how different features contribute to each specific classification decision made by the model.

In Figure 8, the force plots depict the influence of features on the final prediction score. The blue segments represent factors that decrease the fraud probability, while the red segments indicate features that increase it. For example, in the first case, features such as VehicleCategory, BasePolicy, and AgeOfPolicyHolder contribute negatively (blue) to reducing the likelihood of fraud, whereas factors like Fault and PolicyType push the prediction towards a higher fraud probability. These plots provide a clear breakdown of how individual features

impact each specific case.

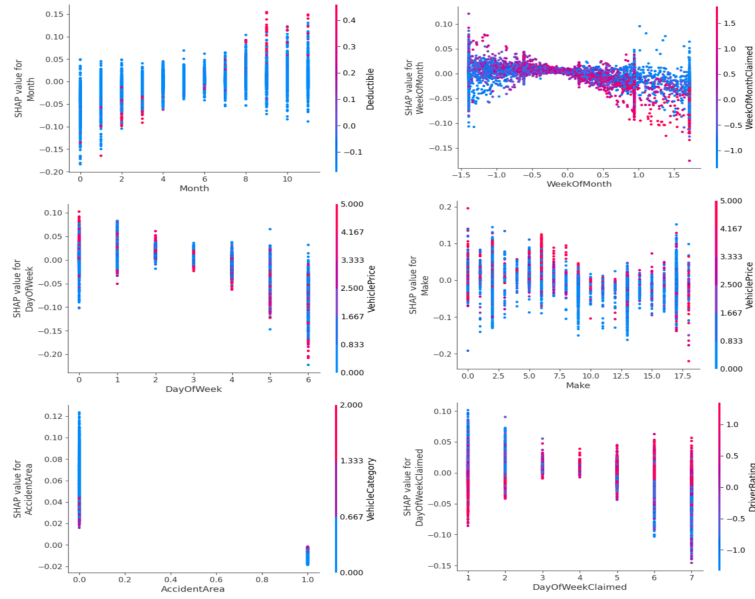


Figure 7. The relationship among some features (dependence plot) based on SHAP.

Figure 9 presents the waterfall plots, which further break down the contributions of different features toward the model's output for each instance. Here, features are ranked based on their impact on the final prediction, with positive contributions shown in red and negative contributions in blue. For instance, VehicleCategory and BasePolicy consistently appear as significant factors across multiple cases, emphasizing their importance in fraud detection. Other features, such as PolicyType, MaritalStatus, and AgeOfPolicyHolder, also show varying levels of influence, highlighting the complexity of fraud prediction.

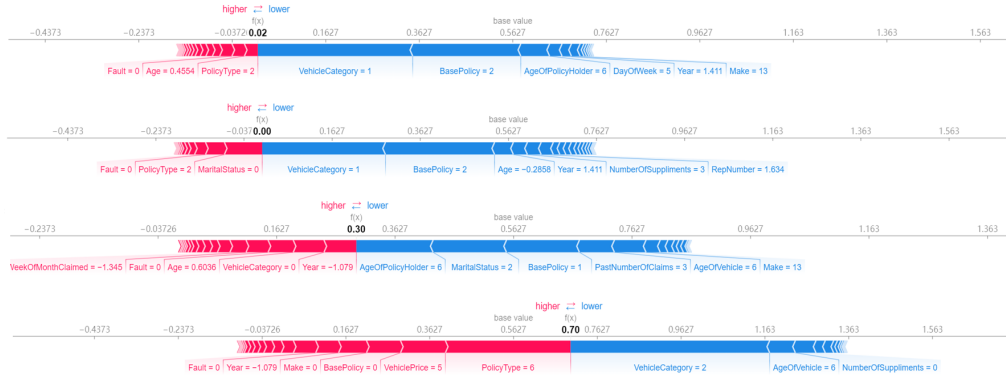


Figure 8. The force plot for the first four examples based on SHAP.

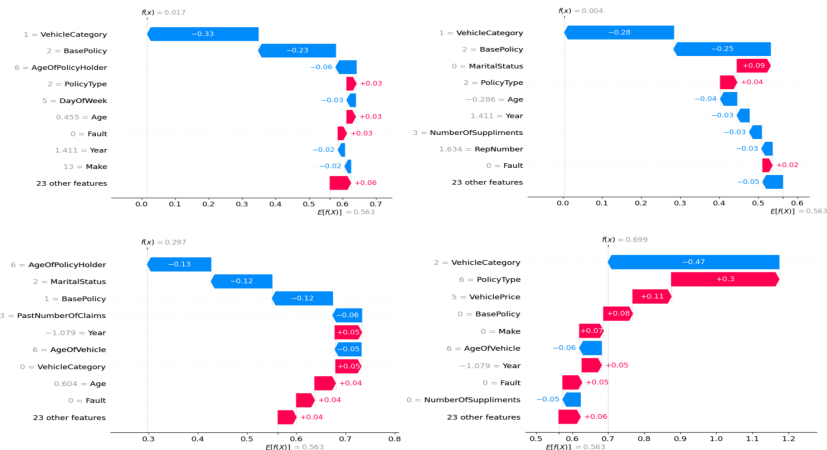


Figure 9. The waterfall plot for the first four examples based on SHAP.

5. Discussion

While the proposed ANN model demonstrates strong prediction and interpretability performance in detecting fraudulent vehicle insurance claims, several limitations remain that could be addressed in future research: (1) One notable limitation is the class imbalance in the dataset. Despite using SMOTE to balance the training data, synthetic data generation may introduce biases that do not fully represent real-world fraud patterns. Future work could explore cost-sensitive learning or anomaly detection methods that do not require artificial oversampling. (2) Another challenge is feature engineering. Although SHAP analysis provides insights into important features, some categorical variables (e. g., vehicle make, policy type) were converted into numerical representations without exploring potential interactions. More advanced techniques, such as feature embedding or graph-based representations, could enhance the model's ability to capture relationships among categorical variables. (3) Additionally, the model interpretability can be further improved. While SHAP provides useful explanations, decision-makers in the insurance industry might require more intuitive explanations. Future work could explore rule-based models or hybrid approaches that combine deep learning with explainable models, such as decision trees or case-based reasoning systems.

6. Conclusions

This study presents an interpretable vehicle insurance fraud detection framework that combines an ANN with SHAP analysis to enhance fraud prediction and model transparency. The proposed model effectively detects fraudulent claims, achieving strong classification performance while providing insights into feature importance. SHAP-based interpretability allows insurers to understand the key factors influencing fraud predictions, aiding in better decision-making and regulatory compliance. Despite its effectiveness, the study identifies limitations such as dataset imbalance, feature representation challenges, and generalizability concerns. Future work could explore cost-sensitive learning, advanced feature engineering techniques, and domain adaptation to improve fraud detection accuracy and adaptability. By integrating interpretability into ML-driven fraud detection, this research contributes to developing more transparent and reliable fraud prevention systems for the insurance industry.

Funding

This research was supported by the U.S. National Science Foundation under Grant No. 1563372 and by the National Natural Science Foundation of China under Grant No. 719740361.

Author Contributions

Conceptualization and methodology, A. W. and Y. Q.; writ-ing—original draft preparation and writing—review and editing, A. W. K. X., Z. Z., Y. Q. and Y. J. All authors have read and agreed to the published version of the manuscript.

Institutional Review Board Statement

Not applicable.

Informed Consent Statement

Not applicable.

Data Availability Statement

Data is available upon request from the corresponding author.

Conflicts of Interest

The authors declare no conflict of interest.

References

- 1 Roy R, George KT. Detecting Insurance Claims Fraud Using Machine Learning Techniques. In Proceedings of the 2017 International Conference on Circuit, Power and Computing Technologies (ICCPCT), Kollam, India, 20–21 April 2017; pp. 1–6.
- 2 Li P, Shen B, Dong W. An Anti-Fraud System for Car Insurance Claim Based on Visual Evidence. *arXiv* 2018; arXiv:1804.11207.
- 3 Roriz R, Pereira JL. Avoiding Insurance Fraud: A Blockchain-Based Solution for the Vehicle Sector. *Procedia Computer Science* 2019; **164**: 211–218.
- 4 Viaene S, Ayuso M, Guillen M, *et al.* Strategies for Detecting Fraudulent Claims in the Automobile Insurance Industry. *European Journal of Operational Research* 2007; **176**(1): 565–583.
- 5 Emerson RW. Insurance Claims Fraud Problems and Remedies. *UMLR* 1991; **46**: 907.
- 6 Zhang Z. RAG for Personalized Medicine: A Framework for Integrating Patient Data and Pharmaceutical Knowledge for Treatment Recommendations. *Optimizations in Applied Machine Learning* 2024; **4**(1).
- 7 Xu K, Gan Y, Wilson A. Stacked Generalization for Robust Prediction of Trust and Private Equity on Financial Performances. *Innovations in Applied Engineering and Technology* 2024; **3**(1): 1–12.
- 8 Zhou T, Zhang G, Cai Y. Residual Self-Attention-Based Temporal Deep Model for Predicting Aircraft Engine Failure within a Specific Cycle. *Optimizations in Applied Machine Learning* 2023; **3**(1).
- 9 Huang W, Ma J. Analysis of Vehicle Fault Diagnosis Model Based on Causal Sequence-to-Sequence in Embedded Systems. *Optimizations in Applied Machine Learning* 2023; **3**(1).
- 10 Huang W, Cai Y, Zhang G. Battery Degradation Analysis through Sparse Ridge Regression. *Energy & System* 2024; **4**(1).
- 11 Ma J, Xu K, Qiao Y, *et al.* An Integrated Model for Social Media Toxic Comments Detection: Fusion of High-Dimensional Neural Network Representations and Multiple Traditional Machine Learning Algorithms. *Journal of Computational Methods in Engineering Applications* 2022; **2**(1): 1–12.
- 12 Ma J, Zhang Z, Xu, K, *et al.* Improving the Applicability of Social Media Toxic Comments Prediction Across Diverse Data Platforms Using Residual Self-Attention-Based LSTM Combined with Transfer Learning. *Optimizations in Applied Machine Learning* 2022; **2**(1).
- 13 Ma J, Chen X. Fingerprint Image Generation Based on Attention-Based Deep Generative Adversarial Networks and Its Application in Deep Siamese Matching Model Security Validation. *Journal of Computational Methods in Engineering Applications* 2024; **4**(1): 1–13.
- 14 LaValley MP. Logistic regression. *Circulation* 2008; **117**(18): 2395–2399.
- 15 Hosmer DW Jr, Lemeshow S, Sturdivant RX. *Applied Logistic Regression*; John Wiley & Sons: Hoboken, NJ, USA, 2013.
- 16 Zhou Z, Wu J, Cao, Z, *et al.* On-Demand Trajectory Prediction Based on Adaptive Interaction Car Following Model with Decreasing Tolerance. In Proceedings of the 2021 International Conference on Computers and Automation (CompAuto), Virtual, 7–9 September 2021; pp. 67–72.
- 17 Zhang H, Zhu D, Gan Y, *et al.* End-to-End Learning-Based Study on the Mamba-ECANet Model for Data Security Intrusion Detection. *Journal of Information, Technology and Policy* 2024; **2**(1): 1–17.
- 18 Zhang G, Zhou T, Cai Y. Coral-Based Domain Adaptation Algorithm for Improving the Applicability of Machine Learning Models in Detecting Motor Bearing Failures. *Journal of Computational Methods in Engineering Applications* 2023; **3**(1): 1–17.
- 19 Zhang G, Zhou T. Finite Element Model Calibration with Surrogate Model-Based Bayesian Updating: A Case Study of Motor FEM Model. *Innovations in Applied Engineering and Technology* 2024; **3**(1): 1–13.
- 20 Gan Y, Chen X. The Research on End-to-end Stock Recommendation Algorithm Based on Time-frequency Consistency. *Advances in Computer and Communication* 2024; **5**(4).
- 21 Gan Y, Ma J, Xu K. Enhanced E-Commerce Sales Forecasting Using EEMD-Integrated LSTM Deep Learning Model. *Journal of Computational Methods in Engineering Applications* 2023; **3**(1): 1–11.
- 22 Chen X, Gan Y, Xiong S. Optimization of Mobile Robot Delivery System Based on Deep Learning. *Journal*

- of *Computer Science Research* 2024; **6**(4): 51–65.
- 23 Chen X, Wang M, Zhang H. Machine Learning-Based Fault Prediction and Diagnosis of Brushless Motors. *Engineering Advances* 2024; **4**(3).
- 24 Wang Z, Zhao Y, Song C, *et al.* A New Interpretation on Structural Reliability Updating with Adaptive Batch Sampling-Based Subset Simulation. *Structural and Multidisciplinary Optimization* 2024; **67**(1): 7.
- 25 Ye X, Luo K, Wang H, *et al.* An Advanced AI-Based Lightweight Two-Stage Underwater Structural Damage Detection Model. *Advanced Engineering Informatics* 2024; **62**: 102553.
- 26 Wang X, Zhao Y, Wang Z, *et al.* An Ultrafast and Robust Structural Damage Identification Framework Enabled by an Optimized Extreme Learning Machine. *Mechanical Systems and Signal Processing* 2024; **216**: 111509.
- 27 Zhao Y, Dai W, Wang Z, *et al.* Application of Computer Simulation to Model Transient Vibration Responses of GPLs Reinforced Doubly Curved Concrete Panel under Instantaneous Heating. *Materials Today Communications* 2024; **38**: 107949.
- 28 Hao Y, Chen Z, Sun X, *et al.* Planning of Truck Platooning for Road–Network Capacitated Vehicle Routing Problem. *arXiv* 2024; arXiv:2404.13512.
- 29 Gangadhar KSNVK, Kumar BA, Vivek Y, *et al.* Chaotic Variational Auto Encoder Based One Class Classifier for Insurance Fraud Detection. *arXiv* 2022; arXiv:2212.07802.
- 30 Asgarian A, Saha R, Jakubovitz D, *et al.* AutoFraudNet: A Multimodal Network to Detect Fraud in the Auto Insurance Industry. *arXiv* 2023; arXiv:2301.07526.
- 31 Gupta RY, Mudigonda SS, Baruah PK, *et al.* Markov Model with Machine Learning Integration for Fraud Detection in Health Insurance. *arXiv* 2021; arXiv:2102.10978.
- 32 Ribeiro MT, Singh S, Guestrin C. “Why Should I Trust You?” Explaining the Predictions of Any Classifier. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining San Francisco, CA, USA, 13–17 August 2016; pp. 1135–1144.
- 33 Lundberg SM, Lee SI. A Unified Approach to Interpreting Model Predictions. *Advances in Neural Information Processing Systems* 2017; **30**.
- 34 Selvaraju RR, Cogswell M, Das A, *et al.* Grad-Cam: Visual Explanations from Deep Networks via Gradient-Based Localization. In Proceedings of the IEEE International Conference on Computer Vision, Venice, Italy, 22–29 October 2017; pp. 618–626.
- 35 Shrikumar A, Greenside P, Kundaje A. Learning Important Features through Propagating Activation Differences. In Proceedings of the International Conference on Machine Learning, Centre, Sydney, Australia, 6–11 August 2017; pp. 3145–3153.
- 36 Lundberg SM, Erion G, Chen H, *et al.* From Local Explanations to Global Understanding with Explainable AI for Trees. *Nature Machine Intelligence* 2020; **2**(1): 56–67.
- 37 Caruana R, Lou Y, Gehrke J, *et al.* Intelligible Models for Healthcare: Predicting Pneumonia Risk and Hospital 30-Day Readmission. In Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, Sydney, NSW, Australia, 10–13 August 2015; pp. 1721–1730.
- 38 Bunkhumpornpat C, Sinapiromsaran K, Lursinsap C. DBSMOTE: Density-Based Synthetic Minority Over-Sampling Technique. *Applied Intelligence* 2012; **36**: 664–684.
- 39 Chawla NV, Bowyer KW, Hall LO, *et al.* SMOTE: Synthetic Minority Over-Sampling Technique. *Journal of Artificial Intelligence Research* 2002; **16**: 321–357.
- 40 Mansourifar H, Shi W. Deep Synthetic Minority Over-Sampling Technique. *arXiv* 2020; arXiv:2003.09788.
- 41 Dai W. Evaluation and Improvement of Carrying Capacity of a Traffic System. *Innovations in Applied Engineering and Technology* 2022; **1**(1): 1–9.
- 42 Dai W. Safety Evaluation of Traffic System with Historical Data Based on Markov Process and Deep-Reinforcement Learning. *Journal of Computational Methods in Engineering Applications* 2021; **1**(1): 1–14.
- 43 Dai W. Design of Traffic Improvement Plan for Line 1 Baijiahu Station of Nanjing Metro. *Innovations in Applied Engineering and Technology* 2023; **10**.
- 44 Agatonovic-Kustrin S, Beresford R. Basic Concepts of Artificial Neural Network (ANN) Modeling and Its

- Application in Pharmaceutical Research. *Journal of Pharmaceutical and Biomedical Analysis* 2000; **22**(5): 717–727.
- 45 Wu W, Dandy GC, Maier HR. Protocol for Developing ANN Models and ITS application to the Assessment of the Quality of the ANN Model Development Process in Drinking Water Quality Modelling. *Environmental Modelling & Software* 2014; **54**: 108–127.
 - 46 Bordt S, von Luxburg U. From Shapley Values to Generalized Additive MODELS and back. In Proceedings of the International Conference on Artificial Intelligence and Statistics, Valencia, Spain, 25–27 April 2023; pp. 709–745.
 - 47 Nohara Y, Matsumoto K, Soejima H, *et al.* Explanation of Machine Learning Models Using Shapley Additive Explanation and Application for Real Data in Hospital. *Computer Methods and Programs in Biomedicine* 2022; **214**: 106584.
 - 48 Movsessian A, Cava DG, Tcherniak D. Interpretable Machine Learning in Damage Detection Using Shapley Additive Explanations. *ASCE-ASME Journal of Risk and Uncertainty in Engineering Systems, Part B: Mechanical Engineering* 2022; **8**(2): 021101.

