

Leveraging AI for Fraud Detection and Prevention in Insurance Claims

Thulasiram Prasad Pasam

NTT DATA, Inc. USA

ABSTRACT

The research examines the way artificial intelligence can be used to identify fraud in insurance. It assesses the accuracy of various machine learning models with specific emphasis on their capability of recognizing fake claims. It is discovered that the model is successful in determining authentic claims. It is not particularly excellent at detecting fraudulent activity, indicating the need for development. Mitigating data imbalance and investigating hybrid AI-rule-based systems are recommended to enhance model accuracy. It also highlights the need for developing comprehensible models that can be trusted and relied on by all stakeholders in the time of using AI in fraud detection for the insurance sector.

Keywords: Fraud Detection, Insurance Fraud, Artificial Intelligence, Hybrid AI Systems, Anomaly Detection, Model Interpretability, Machine Learning, Data Imbalance

INTRODUCTION

Insurance fraud poses a large threat to the insurance sector because it results in significant financial losses alongside diminished customer confidence. Digital platforms created modernized fraudulent claims that were sophisticated for detection purposes. Fraud detection systems analysts observe powerful capabilities to overcome this issue through the utilization of artificial intelligence (AI). Insurance businesses utilize natural language processing and machine learning using AI technology to detect fraudulent claims using quick and effective identification techniques. The lack of detection capabilities permits AI models to identify hidden data patterns along with unusual occurrences by processing large volumes of data. The research document confirms the way artificial intelligence assists insurance companies to identify fraud activities and minimize financial damages.

Aim

The primary goal of this project is to investigate the use of artificial intelligence in identifying and preventing fraud in insurance claims using modern data analysis techniques.

Objectives

- To determine the efficacy of AI approaches in detecting fraudulent insurance claims across many datasets
- To evaluate the effectiveness of machine learning models in fraud detection systems
- To discover the obstacles of applying AI technologies in insurance fraud prevention processes
- To recommend the most effective AI approaches for improving fraud detection in the insurance business

Research Questions

- What are the top AI algorithms for identifying fraudulent insurance claims across several datasets?
- What variables influence the success of machine learning models in enhancing fraud detection inside insurance systems?
- What are the primary hurdles in using AI technology for fraud prevention in the insurance industry?
- What AI approaches are advised to improve fraud detection and prevention in the insurance industry?

Research Rationale

Insurance fraud stands as a major concern that generates heavy financial problems for insurance companies and their policyholders. The current established methods that fraud detection operates lack efficiency in the time it comes to

detecting complicated fraudulent schemes. The rising number of fraudulent insurance claims leads insurance companies to deal with higher operational expenses while facing decreased policyholder trust [1]. The large volume of complex data requires insufficient manual detection methods because these methods lead to errors and are inefficient. The fraud identification capabilities of machine learning technologies offer promise in improving an insurance company's detection accuracy.

LITERATURE REVIEW

Efficacy of AI Approaches in Detecting Fraudulent Insurance Claims

AI techniques applied to detect fraudulent insurance claims have received major cultural attention throughout the last few years. AI detection systems demonstrate success by using machine learning technologies for their operations. The combination of fast data assessment with pattern detection systems can now be achieved through methods that outperform standard analytical techniques [2]. The combination of Random Forests and Support Vector Machines is very effective in classifying fraud claims during supervised machine learning algorithm training with labeled data. The unsupervised learning method of anomaly detection allows for detecting irregular patterns in claims data by working without needing prior instance labeling.



Fig 1: Searching Financial Fraud with AI

AI becomes most effective against changing fraudulent techniques because its learning abilities allow it to adapt to shifting patterns as it develops. AI models outperform traditional rule-based techniques in terms of precision and recall, as well as accuracy [3]. The level of success AI achieves in fraud detection strictly rests on the quality of training data provided to the models. Model generalization together with accurate predictions require high-quality datasets that are properly labeled.

Effectiveness of Machine Learning Models in Fraud Detection Systems

Multiple research studies suggest that machine learning algorithms work well in insurance fraud detection systems. The models process a large volume of data for fraud detection by identifying signs hidden within the data. Unsupervised learning techniques, like anomaly detection can find anomalous patterns in data without the need for labeled information [4]. Random Forests and Gradient Boosting constitute the supervised learning techniques that classify fraudulent claims by studying historical data. The models achieve excellent prediction accuracy after receiving training with labeled datasets.

Machine learning models gain better performance through improved ability to resist fraud scheme changes because they process more data points. These models prove efficient due to their ability to process complex high-dimensional data to detect unsuspected kinds of fraud. Hybrid machine learning models use many strategies that have proved effective for improving detection accuracy [5]. The performance of machine learning models depends entirely on the quality along with quantitative aspects of the training data used.

Challenges in Implementing AI Technologies for Fraud Prevention

The implementation of AI-based fraud prevention systems in insurance demands solutions to multiple barriers. The main problem about AI model training stems from low-quality data input. AI system functionality suffers from negative effects due to inaccurate information together with missing or biased data [6]. AI models encounter operational difficulties because fraudulent schemes change at a high pace forcing them to maintain continuous adaptation. The lack of appropriate labeled training data has a detrimental influence on the supervised model's operational capabilities.



Fig 2: Fraud Detection Methods using Machine learning

Stakeholders face challenges understanding AI decision-making processes because these models normally operate without clear explanation capabilities. There arises suspicion about their systems among end-users in the time of AI systems operating without transparency. The adoption of AI technologies within existing operational frameworks proves to be expensive and demanding because expertise and time and funding are necessary [7]. For example, companies to use AI-driven fraud detection systems need to ensure compliance with legal and ethical standards.

Recommended AI Approaches for Enhancing Fraud Detection in Insurance

Several AI techniques need implementation for improving insurance sector fraud detection. Random Forest and XGBoost represent supervised learning models that show strong effectiveness in identifying insurance claim fraudulence through classification processes. The supervised learning models generate accurate fraud predictions in the time of receiving big high-quality datasets for training purposes [8]. Unsupervised learning models, specifically anomaly detection, provide organizations with substantial value because they detect fraudulent patterns that lack labeled data.

Neural networks have become the recommended analytical method in the time of processing complex and large-scale data using deep learning models. Claims description texts benefit from natural language processing (NLP) for analysing their contents that improves fraud detection performance. Rules-based algorithms offer enterprises with a balanced means of detecting fraud in the time of combined with machine learning models in hybrid AI systems [9]. Typical effectiveness requires continuous model updating along with fraud patterns recognition for maintaining constant operational performance.

Literature gap

The literature has not adequately explored two areas. The first one is about the way hybrid AI systems that combine rule-based algorithms and machine learning can be used in fighting against fraud. Secondly, there is little attention paid towards continuously updating models that can adapt with emerging tricks of fraud over time. This is very important in the time of fraud detection systems are to remain effective and accurate over extended periods of operation.

METHODOLOGY

Research Methods

The research uses a *positivist philosophy* to study objective outcomes in fraud detection by implementing AI detection methods. Research tests existing theories about AI's ability to detect fraudulent insurance claims through a *deductive approach* [10]. The research method provides grounds for hypothesis testing that refers to previous scientific work. The research studies AI technologies applied in insurance through an exploratory method to find new findings about the industry. The *collection of secondary data* enables researchers to analyze established datasets for better efficiency because it draws information from previously performed studies thus increasing both credibility and scope of the investigation.

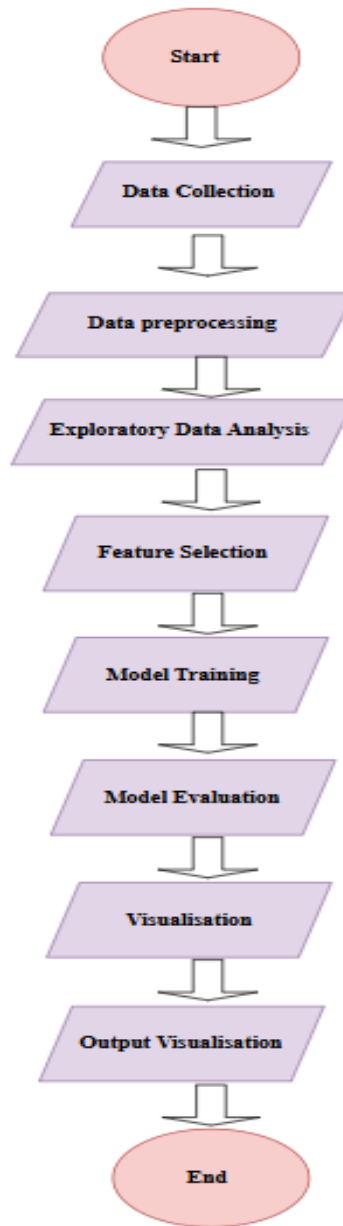


Fig 3: Flowchart

Data Description

This study relies on historical insurance claims data that consists of different features including policy information with associated claim costs as well as incident descriptions along with customer records. The essential aspects of this dataset consist of policy_state together with insured_sex and incident_type and policy_annual_premium and fraud_reported that represents the fraud detection outcomes. The dataset's several types of variables allow for substantial AI analysis using both numerical and categorical characteristics. The information gathered from trusted sources contains comprehensive documentation regarding legitimate and fraudulent claims.

Tools and Techniques

The research uses Python as its primary programming language, combining Pandas for data administration, Scikit-learn for machine learning techniques, and Seaborn for data representation. The main classification algorithms used for this analysis consist of Random Forest and XGBoost. Machine learning models demonstrate effectiveness in discovering fraudulent claims throughout big and intricate data systems [11]. The preprocessing of categorical variables utilizes OneHotEncoder alongside LabelEncoder but the normalization of numerical features depends on MinMaxScaler.

DATA ANALYSIS

▼ Loading the dataset

```
[121] Data_insurance= pd.read_csv("insurance_claims.csv")
```

```
[122] Data_insurance.head()
```

	months_as_customer	age	policy_number	policy_bind_date	policy_state	policy_csl	policy_deductable
0	328	48	521585	2014-10-17	OH	250/500	1000
1	228	42	342868	2006-06-27	IN	250/500	2000
2	134	29	687698	2000-09-06	OH	100/300	2000
3	256	41	227811	1990-05-25	IL	250/500	2000
4	228	44	367455	2014-06-06	IL	500/1000	1000

Result: 40 columns

Fig 4: Loading the dataset

The data contains 40 columns that present different features about insurance claims. The dataset contains three major groups of information consisting of customer profile information and policy specifics and claims facts. The collection of variables provides understanding about customer actions and different claim categories along with policy aspects.

Data preprocessing

▼ Checking missing values

```
[123] print(Data_insurance.isnull().sum())
```

Fig 5: Checking missing values

The code examines data missingness through the `isnull().sum()` function. The operation enables the detection of columns containing missing values so preprocessing tasks can be properly managed. The maintenance of data integrity heavily depends on proper handling of missing values.

▼ Checking the duplicates value

```
[124] print(Data_insurance.duplicated().sum())
```

0

▼ Dropping irrelevant columns

```
[125] Data_insurance.drop(columns=["policy_number", "policy_bind_date", "policy_state"], inplace=True)
```

Fig 6: Checking the duplicate values

The `duplicate().sum()` function ensures evaluation of dataset duplications which reveals no duplicates exist. Verification of data uniqueness completes before the procedure continues to eliminate unnecessary columns. The analysis does not require columns “policy_number”, “policy_bind_date”, “policy_state” making them suitable for removal. This data cleaning procedure results in a streamlined data collection that prepares the dataset for model training purposes.

▼ Converting categorical columns to numerical

```
[126] label_encoder = LabelEncoder()
```

```
[127] Data_insurance["insured_sex"] = label_encoder.fit_transform(Data_insurance["insured_sex"])
```

Fig 7: Converting categorical columns to numerical

The code performs conversion of the categorical “insured_sex” column through LabelEncoder to create numerical values. Machine learning models need numerical data for their operations and require this transformation to achieve processing capability. The “insured_sex” column receives encoding that transforms text values “Male” and “Female” into numeric labels through fit_transform().

▼ Encoding incident type as a categorical variable

```
[128] Data_insurance["incident_type"] = label_encoder.fit_transform(Data_insurance["incident_type"])

[129] numeric_columns = Data_insurance.select_dtypes(include=["number"]).columns
categorical_columns = Data_insurance.select_dtypes(exclude=["number"]).columns
Data_insurance[numeric_columns] = Data_insurance[numeric_columns].fillna(Data_insurance[numeric_columns].median())
for col in categorical_columns:
    Data_insurance[col] = Data_insurance[col].fillna(Data_insurance[col].mode()[0])
print(Data_insurance.isnull().sum())
```

Fig 8: Encoding incident type with the categorical variable

Through LabelEncoder the “incident_type” column became numerical values that can work with machine learning models. The code addresses missing values through filling numerical columns with the median and categorical columns with the mode. The completion of the dataset through this method enables precise training as it eliminates all missing values from the data.

▼ Scaling the numerical features

```
[130] scaler = MinMaxScaler()
scale_columns = ["age", "policy_annual_premium", "capital-gains", "capital-loss", "total_claim_amount"]
Data_insurance[scale_columns] = scaler.fit_transform(Data_insurance[scale_columns])
```

Fig 9: Scaling the numerical features

The code makes use of MinMaxScaler to normalize numerical features that occupy defined numerical boundaries. The model performance benefits from column scaling that includes the variables “age”, “policy_annual_premium”, and “capital-gains”. Data standardization through scaling remains critical because it prevents models from showing preference for features whose values are substantial.

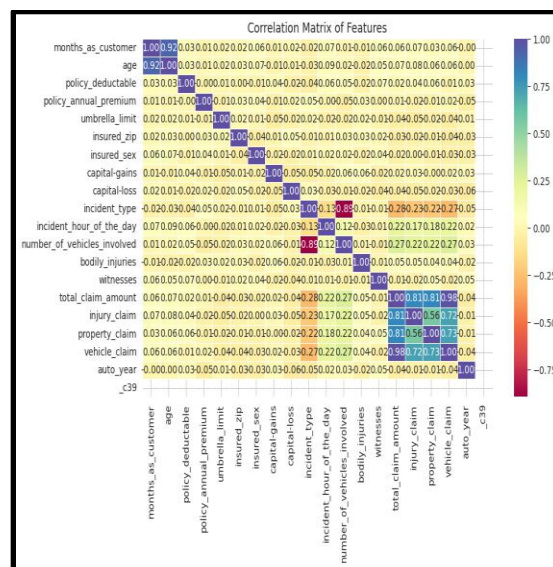


Fig 10: Correlation Matrix

The correlation heatmap shows the way numerical variables are related in the dataset. There is a strong positive correlation for some pairs of data like “months_as_customer” and “age”. A few features have low correlation with others such as “capital-gains” or/and “capital-loss”. The heatmap assists in recognizing important variables for fraud detection models. It reveals possible multicollinearity between some independent variables.

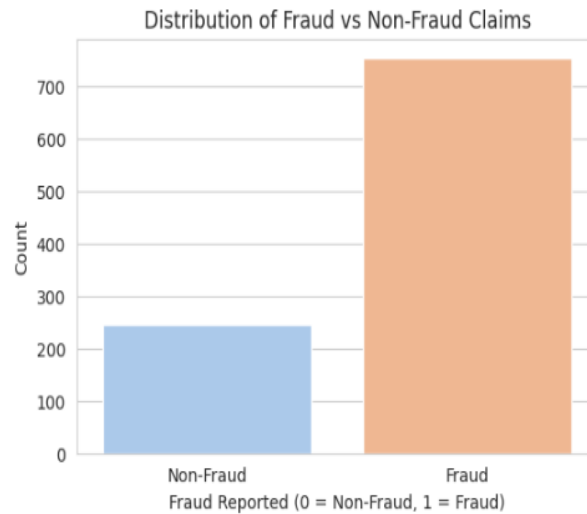


Fig 11: Distribution of Fraud vs Non-fraud Claims

One can display the way fake and real claims are spread throughout the data in this bar plot. Most of the claims are true with very few being false and there is an evident disproportion in the distribution of the two classes as indicated by the plot that makes fraud detection difficult. Discrepancies can need to be taken into account at the model-building stage, especially with respect to class weighting or resampling techniques that can be employed.

✓ Split the data into features and target variable

```
[133] selected_features = ["capital-gains", "incident_type", "incident_hour_of_the_day", "policy_annual_premium",
X = Data_insurance[selected_features]
y = Data_insurance["fraud_reported"]
```

✓ Splitting data into training and test sets

```
[134] X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2, random_state=42)
```

Fig 12: Splitting data into training and test sets

The code splits the dataset by taking the feature selections as well as the target variable called fraud_reported. Some of the selected features are "capitalgains", "incident_type", "incident_hour_of_the_day" while others are related to fraud detection. The data is divided 80/20 into two sets such as training set used to train model and test set used to assess model's performance.

✓ Training a prediction model using Random Forest

```
[135] model = RandomForestClassifier(n_estimators=100, random_state=42)
model.fit(X_train, y_train)
```

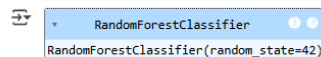


Fig 13: Training the Random Forest Model

A RandomForestClassifier is set up with 100 estimators and random state of 42 to ensure reproducibility. The model is trained using the training data, X_train and y_train. It is decided to use Random Forest because it can easily manage complex data and identify complex patterns. Consequently, the model has been prepared to enhance fraud detection accuracy by predicting on unseen data.

✓ Classification report the model

```
✓ [136] y_pred = model.predict(X_test)
0s accuracy = accuracy_score(y_test, y_pred)
classification_rep = classification_report(y_test, y_pred)
conf_matrix = confusion_matrix(y_test, y_pred)
```

Fig 14: Classification report for the model

The code above assesses the way well the trained model performs on the test set. One utilized `accuracy_score()` to determine the accuracy of the model. It also gives precision, recall as well as F1-score in details through `classification_report()`. The `confusion_matrix()` that accounts for true positives, false positives, true negatives and false negatives helps in evaluating model performance.

✓ Output the evaluation metrics

```
[137] print(f"Model Accuracy: {accuracy:.2f}\n")
print("Classification Report:\n", classification_rep)
print("Confusion Matrix:\n", conf_matrix)
```

➡ Model Accuracy: 0.72

Classification Report:

	precision	recall	f1-score	support
N	0.74	0.94	0.83	145
Y	0.47	0.13	0.20	55
accuracy			0.72	200
macro avg	0.60	0.54	0.52	200
weighted avg	0.67	0.72	0.66	200

Confusion Matrix:

```
[[137  8]
 [ 48  7]]
```

Fig 15: Displaying the Evaluation metrics

The model can describe that a claim is fake or real with a 72% accuracy. The model has a recall of 0.94 for non-fraudulent claims, indicating that it is effective at detecting real claims. However, it detects fraud on average, with precision and recall for false claims at 0.47. This implies that many fraudulent claims are overlooked, whereas the bulk of false negatives are successfully discovered. This suggests that there is still plenty of space for improvement in fraud detection for the most part. The confusion matrix reveals 137 true negatives, 8 false positives, 48 false negatives, and 7 real positives, indicating areas for improvement.

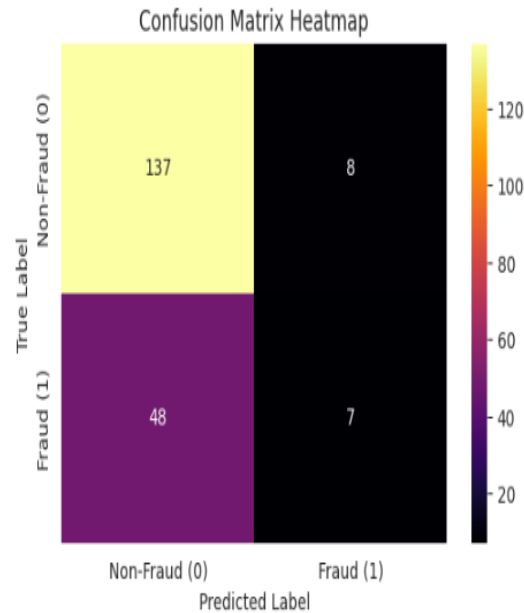


Fig 16: Displaying the Confusion matrix heatmap

One can see that the model identified 137 cases correctly as negatives, but missed 8 cases (positives) and wrongly identified 48 cases as positives while getting only 7 true positives by looking at the confusion matrix heatmap. The intensity of colors used indicates the way many predictions fall under every category. It can be observed from the heatmap that the model tends to misclassify non-fraudulent claims as fraudulent since there are many false negatives than false positives. This clearly shows that there is a great necessity for improving the model.

FUTURE DIRECTIONS

Researchers can consider using complex artificial intelligence (AI) systems, for example deep learning and neural networks to enhance fraud detection in insurance claims. The improvement can also be achieved through integration of better quality and more diverse data into the models [12]. Integrated systems can be able to identify any anomalies in a matter of seconds. Dealing with imbalanced data through oversampling, class weighting techniques etc. can increase the detection ability particularly in relation to fraud detection. It is important to continue researching the way models can be communicated and understood, allowing stakeholders to rely on them [13]. Finally, hybrid models that combine rule-based systems with artificial intelligence can give more effective fraud prevention solutions.

CONCLUSION

The above data concludes AI models have a lot of promise for making fraud detection better in the insurance sector. This study points out where current machine learning techniques are strong but also emphasizes their weaknesses especially in the time it comes to identifying fake claims. The model is effective in identifying non-fraudulent claims that it is not very accurate in the time of it comes to fraudulent ones. Research can be geared towards improving the interpretability of models, dealing with data imbalance as well as incorporating real-time detection in the future. It can be advisable to adopt hybrid approaches that combine rule-based systems with artificial intelligence for enhanced fraud detection. Further studies help in creating trustworthy, adaptable and transparent fraud prevention systems within the insurance industry.

REFERENCES

- [1] Agarwal, S., 2023. An intelligent machine learning approach for fraud detection in medical claim insurance: A comprehensive study. *Scholars Journal of Engineering and Technology*, 11(9), pp.191-200.
- [2] Aboah, A., Wang, B., Bagci, U. and Adu-Gyamfi, Y., 2023. Real-time multi-class helmet violation detection using few-shot data sampling technique and yolov8. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition* (pp. 5350-5358).
- [3] Goriparthi, R.G., 2023. AI-Enhanced Data Mining Techniques for Large-Scale Financial Fraud Detection. *International Journal of Machine Learning Research in Cybersecurity and Artificial Intelligence*, 14(1), pp.674-699.

- [4] Usmani, U.A., Happonen, A. and Watada, J., 2022, July. A review of unsupervised machine learning frameworks for anomaly detection in industrial applications. In *Science and Information Conference* (pp. 158-189). Cham: Springer International Publishing.
- [5] Chiericato, M., Frangiamore, F., Morassi, M., Baresi, C., Nici, S., Bassetti, C., Bnà, C. and Galelli, M., 2022. A hybrid machine learning/deep learning COVID-19 severity predictive model from CT images and clinical data. *Scientific reports*, 12(1), p.4329.
- [6] Rana, N.P., Chatterjee, S., Dwivedi, Y.K. and Akter, S., 2022. Understanding dark side of artificial intelligence (AI) integrated business analytics: assessing firm's operational inefficiency and competitiveness. *European Journal of Information Systems*, 31(3), pp.364-387.
- [7] Issa, H., Jabbouri, R. and Palmer, M., 2022. An artificial intelligence (AI)-readiness and adoption framework for AgriTech firms. *Technological Forecasting and Social Change*, 182, p.121874.
- [8] Hanafy, M.O.H.A.M.E.D. and Ming, R., 2021. Using machine learning models to compare various resampling methods in predicting insurance fraud. *Journal of Theoretical and Applied Information Technology*, 99(12), pp.2819-2833.
- [9] Balcerek, S., Karovič, V. and Karovič, V., 2021. Application of business rules mechanism in IT system projects. *Developments in Information & Knowledge Management for Business Applications: Volume 2*, pp.33-112.
- [10] Pala, S.K., 2022. Investigating fraud detection in insurance claims using data science. *International Journal of Enhanced Research in Science, Technology & Engineering* ISSN, pp.2319-7463.
- [11] Ashtiani, M.N. and Raahemi, B., 2021. Intelligent fraud detection in financial statements using machine learning and data mining: a systematic literature review. *Ieee Access*, 10, pp.72504-72525.
- [12] Sacchi, R., Terlouw, T., Siala, K., Dirnaichner, A., Bauer, C., Cox, B., Mutel, C., Daioglou, V. and Luderer, G., 2022. PProspective EnvironMental Impact asSEment (premise): A streamlined approach to producing databases for prospective life cycle assessment using integrated assessment models. *Renewable and sustainable energy reviews*, 160, p.112311.
- [13] Michel-Villarreal, R., Vilalta-Perdomo, E., Salinas-Navarro, D.E., Thierry-Aguilera, R. and Gerardou, F.S., 2023. Challenges and opportunities of generative AI for higher education as explained by ChatGPT. *Education Sciences*, 13(9), p.856.