

Fraud Detection in Online Assessments Using a Hybrid NLP–Ensemble Learning Framework for Secure and Reliable Evaluation Systems

1st Dulcie Suri

*Department of Animation
Chandigarh University
Mohali, Punjab, India
dulciesuri5@gmail.com*

4th Purnendu Bikash Acharjee

*Department of Computer Science
CHRIST University
Karnataka, India
pbacharya@gmail.com*

2nd Banda Divya

*Department of Computer Science and
Engineering
CMR Institute of Technology
Hyderabad, Telangana, India
divyabanda123@gmail.com*

5th Christina R

*Department of Social Work
St. Claret College Autonomous
Jalahalli, Karnataka, India
christina@claretcollege.edu.in*

3rd Tanuja Singh

*University Institute of Teachers
Training and Research
Chandigarh University
Punjab, India
tanuja.singh19@gmail.com*

6th P. Ravi Kumar

*Department of Electrical and
Electronics Engineering
KPR Institute of Engineering and
Technology
Coimbatore, Tamil Nadu, India
prkmephd@gmail.com*

Abstract – Teachers are wary of letting students take tests online because of the rising incidence of academic dishonesty, making fraud detection in online assessment an increasingly pressing issue. This became more of a problem during the coronavirus epidemic, when there was a pressing need to swiftly convert in-person classes, seminars, labs, and assessment assignments into online alternatives. Many students also find the altered activities more difficult, adding fuel to the fire of student-instructor friction about the prospect of fraudulent collaboration on independently administered assessments. This study employs AI techniques to overcome these obstacles; first, it uses data preprocessing and SMOTE to fix class imbalance; then, it uses feature engineering to make the model better at discriminating. In addition, social media posts are analysed using natural language processing (NLP) techniques to gain a better understanding of trends of online exam cheating. With this goal in mind, it create a new hybrid architecture called "FastText CNN with LSTM." This architecture combines CNN's global feature extraction capabilities with fastText embeddings' semantic text representation and LSTM's sequential dependency learning capabilities. This method outperforms previous models, providing a solid foundation for detecting online assessment fraud, with an accuracy rate of 94.50%.

Keywords—Fraud Detection (FD), Natural Language Preprocessing (NLP), Online Social Network (OSN).

I. INTRODUCTION

Nowadays, most Internet transactions are conducted using wireless mobile terminals, that can authenticate users using their passwords, fingerprints, noises, and photos. In order to commit fraud, the scammer can gather user data, including ID, password, age, occupation, and other details, and log in to several trading systems as an actual user. The fast evolution of information technology has made this type of fraudulent activity widespread, and it does enormous harm to consumers, companies, and society at large[1]. Further, conventional information security measures will not stop online

transaction theft after these personal details have been obtained. Internet platforms must consequently develop anti-fraud mechanisms to safeguard these new services. Since most popular online fraud detection systems use labelled data to build robust models, this data shortage could make current fraud detection algorithms useless[2]. Labelled data plays a big role in identifying service fraud or benign identities, but both parties need to collect enough business data to make a determination. Consequently, the newly-launched online services' fraud detection system experiences the infamous cold-start problem. Despite the scarcity of labelled data, millions of user behavioral records are continually amassed across all types of internet sites. They could be seen as unlabeled secondary data sources that could encourage us to use them to enhance fraud detection performance.

The rise of Internet-based financial services reliant on cutting-edge innovation like big data has been steadily gaining prominence, due to the expansion of the web. Unfortunately, people and companies have also fallen victim to related forms of online financial fraud, such as those involving credit cards, bank statements, insurance, etc. Online banking has made many people's lives easier, but unfortunately, there is a dark side to this convenience: the prevalence of transaction fraud[3]. The field of study devoted to developing models and technology for the identification of fraudulent transactions has been growing in the realm of online finance and electronic transactions. One of the biggest problems with financial services is the amount of money that can be lost due to fraudulent activity, especially when using credit cards. Finding fraudulent transactions quickly is critical for a number of reasons, including lowering financial losses, improving customer service, and protecting financial institutions' reputations. Finding a solution to this problem will require collaboration between academics and financial institutions to test and refine learning frameworks for reducing false positives in fraud detection systems and generalizing information from

complex patterns across datasets[4]. Any action with an underlying criminal intent, which is not always obvious, can be said to be fraudulent. Online fraud is becoming a more serious problem on a global scale. It is not always easy to tell when someone is being dishonest since fraud involves criminal intent. Trusting others becomes even more important when using online social network, despite the fact that they make interacting with numerous individuals incredibly convenient. The safe evolution of online social network is impossible without confidence[5]. As a result, cybersecurity and compliance reporting necessitate the monitoring and management of privileged user behaviors. User profile and model creation for anomaly monitoring and detection is one of the most applicable and powerful approaches.

It follows the outline for the remainder of the article. It lay out the context of this study in Section 2. The procedures and materials employed are detailed in the third section. The study's findings and analysis are detailed in Section 4. In the last section, they draw the conclusions of this investigation.

II. LITERATURE SURVEY

Fraudulent transactions undermine customer confidence and cause significant financial losses. High false-positive rates plague the machine learning-based fraud detection techniques used today, resulting in unnecessary transaction blocks and operational inefficiencies. An approach is to use deep architectures like as GRUs and VAEs, that can retain sequential dependencies in transaction data and enable unsupervised fraud detection using anomaly scores[6]. Detecting fraudulent activities with a high degree of accuracy is made even stronger by combining VAEs, that identify anomalies, with GRUs, when learn sequential data. Furthermore, attention techniques can be employed for feature importance learning, enabling the model to identify crucial transaction features for fraud detection[7]. There has been significant progress in improving fraud detection capabilities through the use of graph-based approaches, specifically GNNs. Automatic feature learning that takes into account interaction connections between transactions is made possible by GNNs, which aggregate information from nearby nodes to create representations for central nodes[8]. Learning spatial-temporal information successfully is a challenge for GNN-based techniques, notwithstanding their potential. Spatial aggregation is critical since fraudsters typically operate with restricted devices because of high equipment expenses[9]. Notifying cardholders might halt transactions, hence fraudulent actions often happen within restricted time constraints.

The Common solutions at the transaction level include sequence-based deep models that may detect patterns of user behavior. Among other things, they suggested a methodology for early fraud detection based on survival analysis that uses RNNs to process user activity sequences[10]. Using an RNN framework that effectively captures patterns on sequences of events, they demonstrated a real-time fraud detector. For the purpose of detecting fraudulent transactions, they transformed the sequential behavioral data into a tree-like structure that

represents user intentions, and they upgraded the sequence-based model to a tree-based LSTM network[11]. An online method for detecting fraudulent credit card operations using a neural classifier was presented. The system is dependent on the operation's data and its past performance; it was deployed in a transactional hub[12]. Overcoming the imbalance in the rate of normal and fraudulent activities was achieved using a nonlinear variant of Fisher's discriminant analysis. A model for detecting credit card fraud utilizing frequent item set mining was suggested by them[13]. An effective method for detecting fraudulent transactions was developed by them using a feed forward ANN.

Finding linguistic trends in the spoken languages of a population is next to impossible without word representations and other tools of natural language processing. Therefore, there is a great demand for learning strategies and better embedding methods to identify online assessment from text. This study proposes the NLP-CNN-LSTM technique to enhance word representations. To improve accuracy, the classifier extracts global features with local dependencies quickly using a CNN and LSTM combination. Additionally, this approach includes Fast Text Embedding.

III. PROPOSED SYSTEM

Due in large part to its low cost and ease of usage in reaching large groups, online survey research has grown in popularity in recent years. Online surveys allow participants to fill out questionnaires using a variety of digital tools that control the presentation of questions and the storage and retrieval of responses. When traditional survey methods failed to reach under-represented populations, researchers turned to these alternative approaches.

Fintech companies are increasingly confronted with the problem of fraudulent transactions. Based on real-world trends, this dataset labels 51,000+ transactions as either authentic or fraudulent. Data such as purchases, actions, payment methods, and gadget use are all part of it[14].

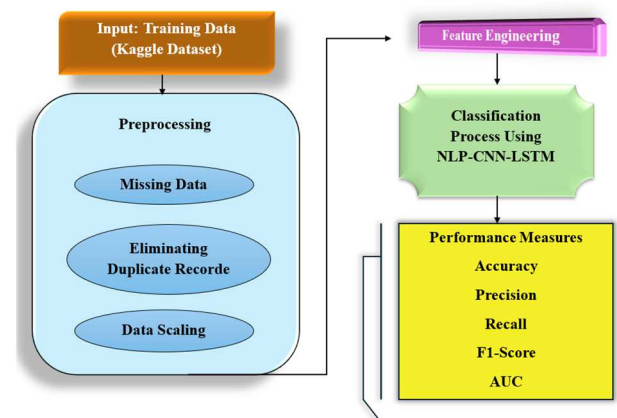


Fig. 1. Block Diagram Of Proposed Method

A new FD method for detecting online assessment is introduced in this research. The suggested FD method includes pre-processing and prediction based on NLP-

CNN-LSTM [15]. It is used to fine-tune the NLP-CNN-LSTM parameters, which greatly enhances the classification performance. The general procedure of FD in the online assessment model is shown in Figure 1.

A. Data Preprocessing:

The first and most important stage in getting datasets ready to analyze and analyze financial fraud is to do the necessary preparation. Data cleansing, including the elimination of duplicate records and the standardisation of data scale, are all part of these procedures. Our goal is to clean up and organise the dataset so that they can analyse it more accurately and use efficient models to find instances of financial fraud[16].

1) Addressing the Missing Data:

The process of managing missing values comprises filling in the blanks in our dataset that result from data being either missing or incomplete. For this aim, mean imputation is a commonly used approach. Mean imputation entails calculating the average of the available data points for a specific feature and subsequently utilising that average to substitute the missing values. The procedure can be encapsulated in Equation (1).

$$W_{imputed} = \frac{1}{M} \sum_{j=1}^M W_j \quad (1)$$

M is the total number of values that are not missing, $W_{imputed}$ is the set of values that are approximated to fill the gaps, and W_j is the set of values that were first observed.

2) Duplicate Records Eliminating:

To maintain the accuracy and reliability of our data, eliminate duplicate records to ensure our dataset contains unique data points and to prevent bias resulting from duplications. This procedure entails scrutinising each record and retaining just the unique entries by juxtaposing them with all other records as delineated in equation (2).

$$C_{unique} = \{c_i \in C : \text{No identical record in } C \text{ matches } c_i\} \quad (2)$$

The dataset that C_{unique} is referring to has had all duplicate records deleted. c_i stands for a unique record in this dataset, while C stands for the initial/original dataset, which might include duplicates.

3) Scaling the Data:

One of the most important parts of preparing data is scaling it so that all numerical features are on the same, similar scale. Standardisation and min-max scaling are two popular approaches to this problem. A feature, W , is transformed through standardisation so that it has a mean value of 0 and a SD of 1, indicating a consistent distribution. As shown in Equation (3), this transformation allows us to compare feature W with other features within a dataset effectively;

$$W_{scaled} = \frac{W - \tau}{\gamma} \quad (3)$$

In which W_{scaled} is the modified feature, W is the original feature, τ is the mean, and γ is the standard deviation. At the same time, as illustrated in equation (4), min-max scaling usually modifies feature W such that it falls inside a specified range [0,1].

$$W_{scaled} = \frac{W - W_{min}}{W_{max} - W_{min}} \quad (4)$$

In order to normalise or rescale a feature based on its size, it is necessary to consider the original feature W , its minimum value W_{min} , and its maximum value W_{max} . The result is represented by W_{scaled} .

B. Feature Engineering:

In order to make models better at predicting whether or not an online evaluation is fraudulent, domain expertise features have to be created. While online assessments do generate a mountain of data for institutions, not all of it is immediately useful for making predictions. They can extract a lot of useful information from the raw data. Based on our findings, it developed a handful of industry-standard predictors[17].

Enhanced features were developed for this research. Incorporating the following predictors into the dataset:

- Monthly transaction frequency—these factors provide light on how an account spends its money.
- Due to the fact that card thieves conceal information on card applications during purchases, missing data dummy variables play a significant role in predicting instances of credit card fraud.
- When making a larger purchase, fraudsters often use stolen cards at petrol stations and restaurants, so it's important to have dummy variables to signal when purchases occur at specific shops.
- Various account history characteristics, including the total number of transactions within the 8-month sample, the highest and average authorisation amounts, and more.
- A dummy variable that is set to true if the amount of an authorised transaction at a specific merchant is more than 10% of the standard deviation of the mean of the merchant's non-fraudulent transactions and the deviation of the transaction from the non-fraud mean is towards the fraudulent mean. In order to maximise classification accuracy, a 10% variation was utilised since it has been shown to be statistically significant.

C. Training the Model:

The main idea here is that a hybrid deep learning classifier, when used with effective word embedding, can FD in online assessments quite accurately.

To better depict the pre-processed text, word embedding techniques can be employed. Word2Vec and other prior word embedding systems perform an excellent job, but they have a few limitations, such adding more dimensions and not being able to build words that aren't in the user's vocabulary. The Fast text embedding method avoids this problem by creating new words outside of the corpus using the n-gram strategy. Words are best represented with Fast text embedding in this setting.

The research clearly shows that CNN and LSTM, two deep learning models for text analysis, are quite effective. In quick succession, CNN can glean features from text all across the world. The LSTM is capable of comprehending

traits that are dependent on one another over the long run. LSTM outperforms RNN because to its ability to resolve the vanishing gradient problem that RNN encounters. There is evidence in the literature that LSTM outperforms RNN. When used together, they can improve classifier performance, allowing for faster feature dependency extraction[18]. FastText integration with CNN and LSTM improves the accuracy of depression detection in the NLP technique.

The CS dataset stores data from tweets and posts along with their corresponding labels. An example of this is shown below. The equation $CS = \{C, K\}$ states that for every m tweets or posts, the sample C is represented by

With the corresponding labels $K = \{KZ_1, KZ_2, KZ_3, \dots, KZ_m\}$ the data set $C = \{CS_1, CS_2, CS_3, \dots, CS_m\}$. Additionally, every tweet or post CS_j consists of o words, represented as $CS_j = \{CX_1, CX_2, CX_3, \dots, CX_o\}$, with CX_j being a word in the tweet or post.

Equation (5) below provides a concise expression for each tweet or post.

$$CS_j = CX_{1:o} = CX_1 \oplus CX_2 \oplus CX_3 \oplus \dots \oplus CX_o \quad (5)$$

In order to remove any extraneous or undesired symbols, the data C from the tweets and posts is cleaned and pre-processed. Then, for every CS_j in C , every CS_{oj} becomes a mandatory one to check. The CS_{oj} can be any length you want it to be in a tweet or post. The inputs to any neural network must be uniform in size and shape. The model design typically receives text input of varying lengths, so that some sentences may be shorter than others. To ensure that sentences remain uniform in length, padding measures are employed. During creation, the current model design makes use of the padding="post" approach. In order to make sentences of the same length, this method will append zeroes to the end of the series. In order to keep the input side consistent, a padding technique is later applied to each CS_{oj} . Equation (6) represents the CS_{oj} scheme with the r word padding approach.

$$CS_{oj} = CX_{1:r} \quad (6)$$

In this case, the zero vectors $CX_{k+1}, CX_{k+2}, \dots, CX_r$ have m dimensions. Finding r value is what the PAD approach is all about. Equipped with these vectors, the rapid text technique constructs an embedding layer for words with dimensions cm . A dot product is created between the embedding vector and weight vector wv . This is achieved by applying 1D Convolution filters on each sentence. The non-linear activation is used to pipeline this outcome. With the help of the node's accumulated, weighted input, the activation function sets it ablaze, allowing it to generate superior outcomes. Because of its superior performance and ease of training, the ReLU is the recommended choice. For the following algorithms, Conv1D use ReLU as h .

The LSTM is then used to identify the output based on local dependencies. Equations (7) – (12) provide the calculations for the long short-term memory (LSTM) using conversion functions, which are defined as follows:

$$e_{hs} = \gamma(X_{se}[g_{s-1}, z_s] + c_e) \quad (7)$$

$$j_{hs} = \gamma(X_{se}[g_{s-1}, z_s] + c_j) \quad (8)$$

$$\tilde{d}_s = \text{tang}(X_{sd}[g_{s-1}, z_s] + c_d) \quad (9)$$

$$d_s = e_{hs} * d_{s-1} + j_{hs} * \tilde{d}_s \quad (10)$$

$$p_{hs} = \gamma(X_{sp}[g_{s-1}, z_s] + c_p) \quad (11)$$

$$g_s = p_{hs} * \text{tang}(d_s) \quad (12)$$

The whole process is laid out in a series of algorithms, starting with Algorithm 1:

Algorithm1: Performance Evaluation of NLP-CNN-LSTM Method	
1. Input:	NLP-CNN-LSTM Model.
2. Data:	bb = encoded index with pad, st = size of the test, Z = text label, fo = no. of epochs, at = size of the batch, $wtrain$ = data training, $wtest$ = data for testing, $ztrain$ = train labels, $ztest$ = testing labels, qt = random state;
3. Output:	Confusion matrix, Precision, Recall, F1-Score, Accuracy, TPR, FPR
4. Begin	
5.	$wtrain, wtest, ztrain, ztest = \text{train} - \text{test} - \text{split}(bb, Z, st, qt)$
6.	NLP-CNN-LSTM-model.fit($wtrain, ztrain, fo, at$);
7.	NLP-CNN-LSTM-model.fit($wtest, ztest$);
8.	$zpred = \text{NLP} - \text{CNN} - \text{LSTM} - \text{model.predict}(wtest)$;
9.	$\text{write}(\text{confusion matrix}(ztest, zpred))$;
10.	$\text{write}(\text{Precision})$;
11.	$\text{write}(\text{recall})$;
12.	$\text{write}(f1 - \text{score})$;
13.	$\text{write}(\text{accuracy})$;
14.	$\text{write}(\text{TPR})$;
15.	$\text{write}(\text{FPR})$;
16.	$\text{write}(\text{confusion matrix}(ztest, zpred))$;
17. End	

IV. RESULT AND DISCUSSION

These days, its common knowledge that online evaluation methods provide a greater danger of fraud. An innovative and reliable strategy for identifying possible instances of exam fraud in online multiple-choice question (MCQ) systems is presented in this study. Examinees' likelihood of communicating with one another is first statistically evaluated using the concordance of answers and answer time compared to null expectations; this assessment is then utilised to detect possible fraudulent activity.

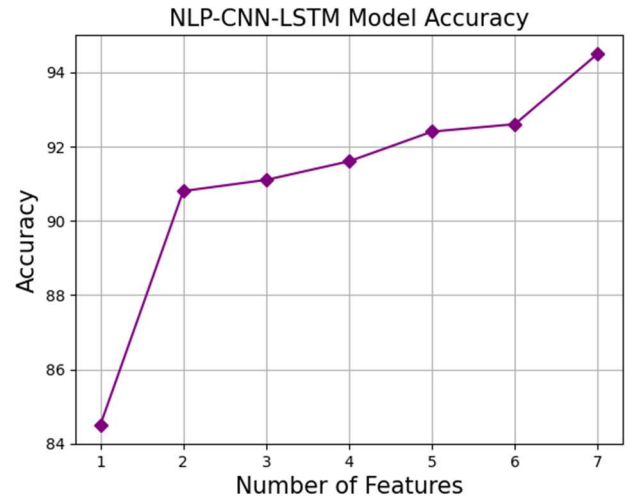


Fig. 2. Accuracy for NLP-CNN-LSTM Model

To create the graph, then iterated over each attribute, ranking them from most important to least, as shown in Figure 2. When using the full set of features, it found that the accuracy was the highest.

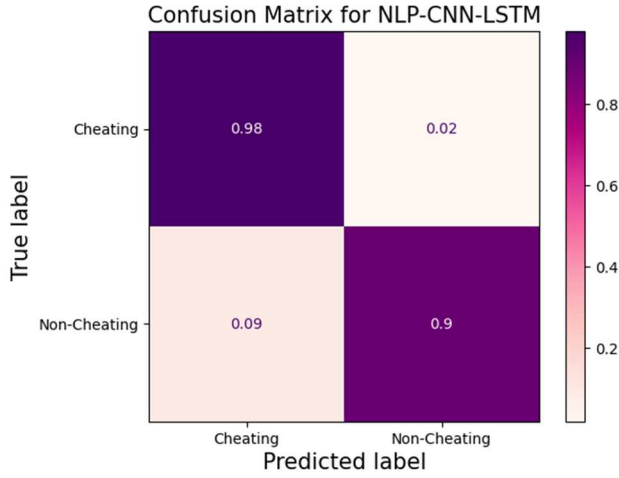


Fig. 3. Confusion Matrix

NLP-CNN-LSTM models for detecting cheating and non-cheating behaviours in online exams are compared in Figure 3, which displays confusion matrices. NLP-CNN-LSTM achieved 90% accuracy for cheating and 98% accuracy for non-cheating.

TABLE I. PERFORMANCE COMPARISON(%)

Models	Accuracy	F1-Score	AUC
NLP-CNN	88.16	88.84	87.23
BiLSTM	91.28	91.88	93.48
NLP-CNN-LSTM	94.50	94.96	98.14
CNN-LSTM	85.37	85.76	87.12
GRU-CNN	82.09	82.68	84.56
DQN	78.36	78.80	80.37

Table 1 shows the results of these evaluations in terms of accuracy, AUC, and F1 score. Using our full feature set regularly yields the best performance, according to the data.

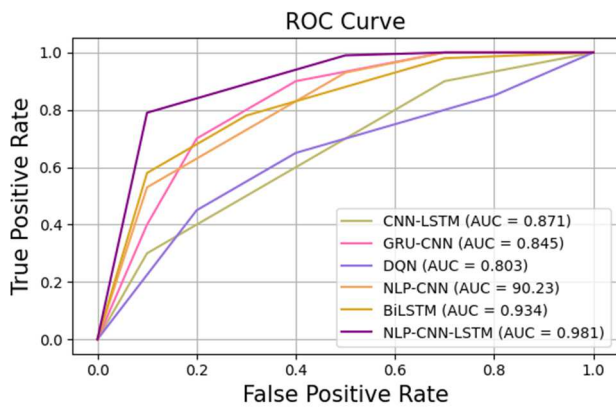


Fig. 4. Cross Validation Score for each Subsets

In order to make our system even better, they employed the metrics of ROC and AUC to determine how sensitive and specific our system. This makes it possible to

use various probability thresholds to summarise the model's TPR and FPR trade-off. The NLP-CNN-LSTM network utilised by us had the best AUC at 0.981, as demonstrated in Figure 4.

TABLE II. BEFORE AND AFTER PREPROCESSING COMPARISON

Models	Before Preprocessing		After Preprocessing	
	Precision	Recall	Precision	Recall
NLP-CNN	0.836	0.808	0.865	0.844
BiLSTM	0.862	0.843	0.897	0.878
NLP-CNN-LSTM	0.894	0.876	0.921	0.903
CNN-LSTM	0.799	0.782	0.834	0.815
GRU-CNN	0.771	0.758	0.807	0.786
DQN	0.735	0.716	0.766	0.744

Part one of Table 2 demonstrates that after the class imbalance problem was solved, the performance of the identical NLP-CNN-LSTM model increased dramatically across all four measures.

Performance Comparison

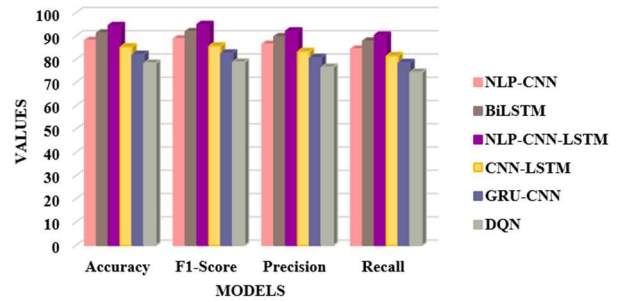


Fig. 5. Performance Comparison

Five classifiers are assessed: NLP-CNN, BiLSTM, CNN-LSTM, GRU-LSTM, DQN, and NLP-CNN-LSTM. The comparative criteria employed include recall, accuracy, precision, and F1 score. Figure 5 illustrates the performance comparison across various classifiers. The NLP-CNN-LSTM classifiers attain the best accuracy of 94.50%.

V. CONCLUSION

Among the many challenges posed by the global pandemic, EL has grown in importance as a means of continuing education on a global scale. Organisations in the field of education are increasingly worried about real-time online assessments. Exam invigilators face significant difficulties due to instances of fraudulent behaviour occurring during online exams. Educational institutions have tried numerous manual techniques in response to this major problem, but so far, none of them have been really innovative or effective. They start by inputting and preparing data using artificial intelligence, utilising the SMOTE to reduce data imbalance. The discriminative skills of the model are further improved through feature engineering. To better understand how to identify online assessments, researchers are analysing social media posts using NLP methods. Modern text analysis is facilitating the

disease diagnosis process by continuously monitoring the underlying trends in the shared text. Utilising word embedding techniques and deep learning models, social media posted content can be used for early depression diagnosis. FastText embedding and the CNN-LSTM hybrid classifier are the building blocks of the NLP-CNN-LSTM method. With a high accuracy of 94.50%, the suggested method outperforms the state-of-the-art in detecting depression.

REFERENCES

- [1] Z. Zhang, X. Zhou, X. Zhang, L. Wang, and P. Wang, "A Model Based on Convolutional Neural Network for Online Transaction Fraud Detection," *Secur. Commun. Networks*, vol. 2018, 2018, doi: 10.1155/2018/5680264.
- [2] C. Liu, Y. Gao, L. Sun, J. Feng, H. Yang, and X. Ao, "User Behavior Pre-training for Online Fraud Detection," *Proc. ACM SIGKDD Int. Conf. Knowl. Discov. Data Min.*, pp. 3357–3365, 2022, doi: 10.1145/3534678.3539126.
- [3] X. Wang, Z. Wan, and Y. Zhang, "A DQN-based Internet Financial Fraud Transaction Detection Method," *ACM Int. Conf. Proceeding Ser.*, 2021, doi: 10.1145/3487075.3487136.
- [4] E. Strelcenia and S. Prakoonwit, "A Survey on GAN Techniques for Data Augmentation to Address the Imbalanced Data Issues in Credit Card Fraud Detection," *Mach. Learn. Knowl. Extr.*, vol. 5, no. 1, pp. 304–329, 2023, doi: 10.3390/make5010019.
- [5] E. K. Boahen, W. Changda, and B. M. Brunel Elvire, "Detection of Compromised Online Social Network Account with an Enhanced Knn," *Appl. Artif. Intell.*, vol. 9514, pp. 777–791, 2020, doi: 10.1080/08839514.2020.1782002.
- [6] B. S. Jayaprakasam and C. Ubagaram, "VAE-GRU-Attention: Variational Autoencoder with Gated Recurrent Units for Fraudulent Transaction Detection," *International J. HRM Organ. Behav.*, 2022.
- [7] K. Prabhakar, M. S. Giridhar, A. Tatia, T. M. Joshi, S. Pal, and U. S. Aswal, "Comparative Evaluation of Fraud Detection in Online Payments Using CNN-BiGRU-A Approach," *Int. Conf. Self Sustain. Artif. Intell. Syst. ICSSAS 2023 - Proc.*, no. Iccsas, pp. 105–110, 2023, doi: 10.1109/ICSSAS57918.2023.10331745.
- [8] S. Khosravi, M. Kargari, B. Teimourpour, and M. Talebi, *Transaction fraud detection via attentional spatial-temporal GNN*, vol. 81, no. 4. Springer US, 2025. doi: 10.1007/s11227-025-06983-8.
- [9] M. Karthikeyan, S. Thota, K. Sailaja Kunar, R. Umanesan, S. Karunakaran, and M. Renuka, "Optimizing Financial Security with Cloud AI: Implementing Deep Q-Network and Transfer Learning for Risk Management and Fraud Detection," *2025 Int. Conf. Emerg. Smart Comput. Informatics, ESCI 2025*, pp. 1–6, 2025, doi: 10.1109/ESCI63694.2025.10988191.
- [10] M. Huang *et al.*, "AUC-oriented Graph Neural Network for Fraud Detection," *WWW 2022 - Proc. ACM Web Conf. 2022*, pp. 1311–1321, 2022, doi: 10.1145/3485447.3512178.
- [11] R. Rajkumar, N. Kogila, S. Rajesh, and A. R. Begum, "Intelligent System for Fraud Detection in Online Banking using Improved Particle Swarm Optimization and Support Vector Machine," *Proc. 8th Int. Conf. Commun. Electron. Syst. ICCES 2023*, no. Icces, pp. 644–649, 2023, doi: 10.1109/ICCES57224.2023.10192690.
- [12] Y. Abakarim, M. Lahby, and A. Attioui, "An efficient real time model for credit card fraud detection based on deep learning," *ACM Int. Conf. Proceeding Ser.*, no. July 2019, 2018, doi: 10.1145/3289402.3289530.
- [13] M. Mohammadi, S. Yazdani, and M. Khanmohammadi, "Presenting a Model for Financial Reporting Fraud Detection using Genetic Algorithm," *Adv. Math. Financ. Appl.*, vol. 6, no. 2, pp. 377–392, 2021, doi: 10.22034/amfa.2019.1872783.1252.
- [14] "Fraud Detection Dataset." Accessed: Oct. 02, 2025. [Online]. Available: <https://www.kaggle.com/datasets/ranjitmandal/fraud-detection-dataset-csv>
- [15] A. A. Albraikan *et al.*, "Optimal Deep Learning-based Cyberattack Detection and Classification Technique on Social Networks," *Comput. Mater. Contin.*, vol. 72, no. 1, pp. 907–923, 2022, doi: 10.32604/cmc.2022.024488.
- [16] A. A. Almazroi and N. Ayub, "Online Payment Fraud Detection Model Using Machine Learning Techniques," *IEEE Access*, vol. 11, no. November, pp. 137188–137203, 2023, doi: 10.1109/ACCESS.2023.3339226.
- [17] I. Karkaba, E. M. Adnani, and M. Erritali, "Deep Learning Detecting Fraud in Credit Card Transactions," *J. Theor. Appl. Inf. Technol.*, vol. 101, no. 9, pp. 3557–3565, 2023.
- [18] V. Tejaswini, K. S. Babu, and B. Sahoo, "Depression Detection from Social Media Text Analysis using Natural Language Processing Techniques and Hybrid Deep Learning Model," *ACM Trans. Asian Low-Resource Lang. Inf. Process.*, vol. 23, no. 1, 2024, doi: 10.1145/3569580.