



Tikrit Journal of Pure Science

ISSN: 1813 – 1662 (Print) --- E-ISSN: 2415 – 1726 (Online)

Journal Homepage: <http://tjps.tu.edu.iq/index.php/j>



A Hybrid Deep Learning Model to Accurately Detect Anomalies in Online Social Media

Darbaz M. Hussein , Hakem Beitollahi

Computer Science Department, Faculty of Science, Soran, Iraqi-Kurdistan Region

DOI: <http://dx.doi.org/10.25130/tjps.27.2022.074>

ARTICLE INFO.

Article history:

-Received: 22 / 6 / 2022

-Accepted: 28 / 8 / 2022

-Available online: / / 2022

Keywords: Anomalies, LSTM, CNN, ANN, Feature Selection

Corresponding Author:

Name: Darbaz Maarouf Hussein

E-mail:

dmh150h@cs.soran.edu.iq

Tel:

ABSTRACT

Online social media (OSM) generates a massive amount of data about human behavior based on their interactions. People express their opinions, comments and share information about variety of topics of their daily life through OSM. The majority of the comments are divided into three categories: Positive, negative, and natural. Regarding the negative comments, the OSM platform facilitates abnormal actions, such as unsolicited messages, misinformation, rumors, the dissemination of fake news and propaganda, as well as the dissemination of malicious links. Therefore, one of the most significant data analytic activities for identifying normal and deviant individuals on social networks is abnormality detection. This paper proposes a hybrid model based on three famous deep learning approaches to discover the behavioral abnormalities, and negative comments in the OSM platforms. The selected benchmark for our research is the airline sentiment in a Twitter dataset. In the proposed method, the dataset is fed to the LSTM network; next, the output of LSTM feeds the CNN network. CNN combines features and generates various feature maps. In the next step, we reduce and select the most important features. Finally, the selected features are given to ANN to classify the data. The proposed method (LSTM + CNN + ANN) is compared with various classical machine-learning (ML) techniques. The experimental results show that the proposed method enhances the accuracy and precision on average by %8.6 and %8.4, respectively in compared to the classical ML techniques.

1. Introduction

Nowadays, online social media (OSM) have become an integral part of people's lives, allowing them to communicate and interact on a large scale with others who share similar interests, values, and perspectives [1]. People utilize OSM to keep in touch with friends and relatives, as well as to share information, images, and videos. The most famous OSM platforms to share data between individuals are Facebook, MySpace, WhatsApp, Instagrams and Twitter. Twitter is one of the most common OSM platform in which users communicate via messages identified as tweets. Every day, more than 7, 00, 000 tweets are sent out by different people. Individual users can use these tweets to share their opinions, thoughts, feedback, and sentiments about a variety of topics and situations. Market analysis, topics of national significance, politics, well-known individuals, and so on could be the topic.

Millions of opinions and various comments are broadcasted through the OSM platforms in every day. Numerous of these comments are positive and cheerfulness that are out of the scope of this paper. However, several of them are upsetting, painful, disgusting and negative that we call them anomaly comments. Anomaly detection is the process of identifying data elements in a dataset that are distinct from all other data elements in the set such as rare events, abnormalities, deviants, and outliers are all terms [2]. Anomalies in a simple 2-dimensional dataset are depicted in (Fig. 1), where the normal and abnormal points have been explained. As depicted in the figure, G1 and G2 are normal regions; however, anomalies are points that are sufficiently far away from the regions, such as points F1 and F3, and points in area F2.

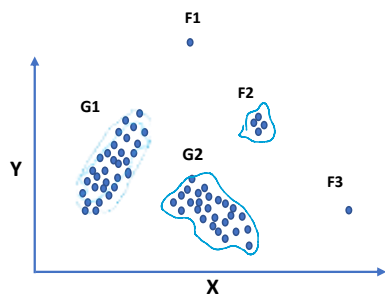


Fig. 1: (An example of anomalies. [5].)

Anomalies can be classified into the three types: Structural anomalies, content anomalies, and behavioral anomalies as shown in (Fig. 2).

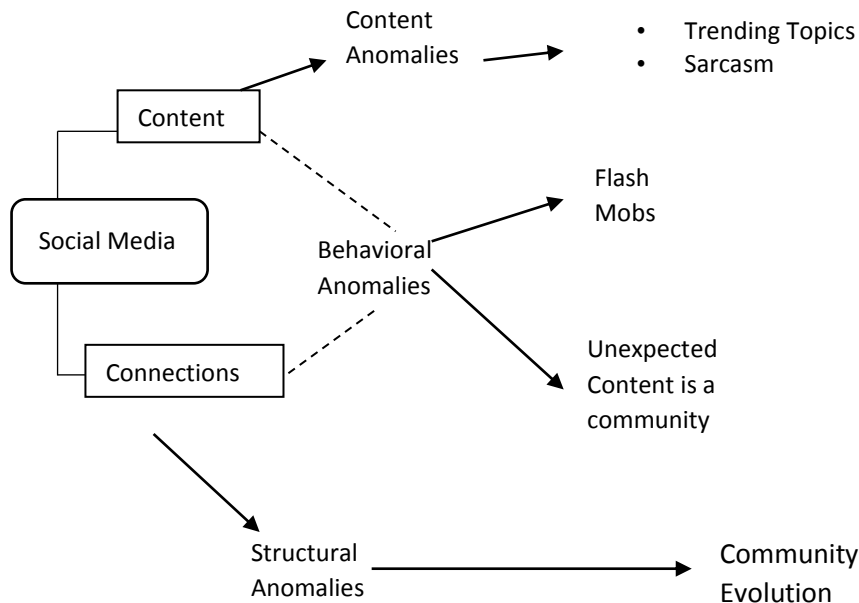


Fig. 2: (Various anomalies in OSM. [5].)

In the content anomalies, the abnormality data points (content) are happen anomalously, that displayed against the entire dataset. Trending topics, topic modeling, and sarcasm are all part of the identification process in this anomaly type.

In the behavioral anomalies, a something is related to do with particular human behavior. Flash mobs and unexpected content in a community are among the things that have been identified. The non-contextual properties of an instance are likewise defined by the behavioral attributes.

In the structural anomalies, anomalies can also be characterized according to graphical features when working with graphical structures, such as in social networks [3] [4] [5] [6].

Anomalies can occur for a variety of reasons, including popular subjects, bad thoughts (criticism), the evolution of new groups, hate speeches, flash mobs, and unexpected content. Anomaly detection is crucial for preventing hostile actions like bullying, terrorist attack preparation, and the dissemination of fraud information. New sorts of atypical behaviors have emerged because of the increasing popularity of social media, creating worry among a variety of groups.

Recently, many research communities in academia and industry have paid attention to discover anomalies in the OSM platforms. Techniques based on machine learning (ML) and deep-learning (DL) are attractive approaches to detect anomalies in the OSM platforms [7] [8] [9] [10]. This paper proposes a hybrid method by combining three DL techniques, namely long short-term memory (LSTM), convolutional neural network (CNN), and artificial neural network (ANN). The comments of the passengers of an airline on Twitter is selected as a benchmark dataset in this paper to detect the anomalies.

In the proposed method, first, a preprocessing is done on data and all text tweets are converted to appropriate numbers. To this end, the tokenizer class in the form of a sequence is used to express each text tweet by a sequence of numbers. Next, all tweet sequences are fed to an LSTM network to bring each tweet to a probability value between zero and one. The column of text tweets is removed from the dataset and instead, a new column of made probability values is added to the dataset. As the strength of each classifier technique highly depends on the features existence in the dataset, we decide to combine the features and acquire features that are

more powerful. The CNN network is a promising network for this goal. When several combined features are generated, we use the correlation operator to select the most important features. Finally, the selected features are fed to the ANN network to classify data to normal and anomaly data. The main contributions of our work are as follows:

- Proposing a hybrid model of three DL techniques (LSTM, CNN, and ANN) that is stronger from each DL technique, separately.
- Utilizing the LSTM network to convert the text tweets to a probability between zero and one.
- Utilizing the CNN network to combine features and generate several feature maps.
- Utilizing correlation operator to select the most important features

- Performing several experiments and comparing the proposed method with classical ML techniques including Decision Tree (Dtree), Support Vector Machine (SVM), Random Forest (RF), Logistic Regression (LR), and K-Nearest Neighbor (KNN). Experimental results show that the proposed method enhances the accuracy of detecting anomalies by 8.6% on average in compared to the classical ML techniques.

The rest of the paper is organized as follows. Section 2 presents the related work. Section 3 presents the proposed method. The experimental results are expressed in Section 4 and finally section 5 concludes the paper.

2. Related work

In this section, we discuss various ML algorithms models for detecting an anomaly in OSM.

2.1 Structural anomaly detection

According to Nicholas et al., [7] a two-stage approach for detecting anomalies in dynamic graphs is presented: To evaluate the normalcy of behavior, the first step employs basic, conjugate Bayesian models for discrete-time calculating processes to follow the pairwise linkages of all nodes in the graph. Then, on a much-decreased selection of possibly anomalous nodes, in the second step perform traditional network inference algorithms. The method's utility is shown using both simulated and actual data sets. For structural anomaly detection, a network embedding approach was utilized. Using a single parameter, a score was used to correlate structural inconsistencies with the embedding and to distinguish (segregation) anomalies from ordinary nodes [8]. To detect structural malware, an entropy-based methodology called the Shannon entropy method was utilized, followed by power law distribution [9]. Using point activity data from a user and pairwise communication data, techniques based on a variant of the Bayesian network Rose et al., [10], and Anshika et al., [11] have been utilized to discover group abnormalities. Anomalies are uncommon, unlawful, or harmful acts that can be classified in a variety of ways, such as abnormalities, outliers, the exceptions. The existence of anomalies is assessed based on a functional

structure that differs from that of the normal. It has been seen that the assaults have a large spread effect through engagements in social networking sites, as unlawful users utilize the sites differently than their friends, obeying patterns differently. Anomaly detection via graph embedding is previously worked to investigate graph embedding for anomaly detection [12] [13] [14]. While Agovic et al., [14] display that embedding may be used to discover outliers, an automated detection approach is missing. The commute time distance is used by Kumar et al., [13] to detect anomalies in dynamic graphs, with the eigenspace embedding serving solely to estimate the commute time distance. According to Jing et al., [12] recommends that spectral embedding be used to uncover abnormal community structure from several sources. The method combines graph partitioning to assess the abnormalities level and provides a unique dimension reduction technique to make the approach more effective and scalable for large networks [15]. Two graph metrics, ego, and egonet, were utilized and then to identify anomalous nodes in the provided graph using a power-law curve [16]. Anomaly scores were calculated by fitting the network's power-law curve: if a node was farther away from the curve, it was reflected as abnormal. The network's vertices were encoded to their vector form using a technique called "Clique Embedding" to detect structural anomalies. The pairwise distance of the network's vertex representation was lowered using this strategy. Vector representation was later computed using the reservoir sampling approach. Finally, anomalous vertices/nodes were detected using the K-Means clustering approach [17] [18]. The DBSCAN (Density-Based Spatial Clustering of Applications with Noise) method according to Ahmet [19] was utilized to monitor data points that didn't feasible in any of the clusters in a given network structure, resulting in the detection of anomalous network architectures [20].

2.2 Content anomaly detection

Statistical topic models based on Latent Dirichlet Allocation (LDA) were utilized to analyze the contents of textual logs in Amogh et al., [21], and Guixian et al., [22] to determine distinct themes/concepts within these logs at different theme levels. LDA endless by Words were also utilized to split the themes in Ying et al., [23] for detecting the trending topics, taking into account their biased role. TPMTM was utilized, which is a two-phase modeling approach that combines statistical modeling (LDA) with regular pattern mining to generate more thorough representations of rich themes and semantics [24]. The LDA algorithm and a mixture of Gaussian Mixture Models (GMM) may also be utilized to find anomalous themes [25]. According to Md. Shafiur et al., [1] classified spammers on Twitter using an SVM classifier method. To distinguish between spammers and real users, the authors used tweet text and user behavior to identify spammer

traits. The content-based model of machine learning algorithms is the subject of a lot of existing research. Anomaly classification models are learned using content-based attributes to categorize messages and profiles as abnormal [26]. The invention and development of a multinomial Nave Bayes classification using an upgraded Expectation Maximization (EM) method are used to detect textual anomalies. It employed a significant amount of unlisted data and applied an EM technique to increase the precision of the Naive Bayes classification. In a binary classification setting, this approach was utilized to discover abnormalities in the text [27]. For content anomaly detection, the Analytic Hierarchy Process (AHP) and Technique for Order of Preference by Similarity to Ideal Solution (TOPSIS) Ravneet et al., [28] were utilized. For the same purpose, C –Means and an optimized K –Means technique was applied [29]. [13] was developed a hybrid anomaly detection method called DT-SVMNB that cascades several machine learning algorithms, including support vector machine (SVM), the naive bayesian classifier (NBC), and decision tree (C5.0), to classify normal and anomalous users in social networks. It compiled a list of distinctive features using user profiles and content. The suggested machine learning model called DT-SVMNB is trained using two different types of datasets with the chosen features. In the social network, their model categorized people as depressed or suicidal. By using both synthetic and real social network datasets were tested. The performance analysis shows good accuracy, which solves the impact and efficiency of their suggested system.

2.3 Behavioral anomaly detection

Descriptive statistics were employed to classify respondent responses. The data set was then qualitatively analyzed, and the logistic regression model was utilized to identify characteristics associated with a change in relationship status on Facebook [30]. For a labeled dataset, the Logistic Regression technique linearly splits the dataset into normal and anomalous user behavior. However, the majority of the data taken from online social media in the form of user thoughts or feedback is unlabeled, making it incomplete to train the Logistic Regression model. Machine learning and statistical methods (Regression Model) were utilized to differentiate between depressed communities using mood, psycholinguistic processes, and content concerns derived from their members' posts [31]. The relationship between reason and influence between the variables is thought to remain unaltered when using a regression model. Because this assumption does not always true, assessing the values of a variable using the regression equation might lead to misleading and incorrect findings. For cleansing user activity logs in social media networks, restricted Local Differential Privacy (LDP) was used to detect behavioral anomalies. Outliers were recognized to the

duration of calls made between individuals using a Bayesian anomaly detection technique that was afterward applied to the remodeled stream of users [32]. When compared to the LDP model, the mistakes and ϵ are much greater. ϵ is also known as the privacy budget, or the privacy loss parameter, is a control parameter for the detection ratio of anomalous output from a dataset. Because the model's performance is susceptible to the skewed dataset, LDP with a Bayesian model yields poor accuracy. Social influence-based behavioral analysis is one of the techniques for detecting malicious behaviors. Suggested an unsupervised clustering to analyze Facebook user reactions. Due to the immediate nature of these reactions, studying them can help identify anomalies in Facebook accounts [33]. Sujatha and Balachandran's [6] efforts to give an insight into different OSM abnormalities and identify irrational user behavior in a dataset of tweets on India's demonetization policy. Also, the rule-based methods (TextBlob and VADER) were explored for converting an unlabeled dataset into a labeled one. Their proposed model was self-supervised learning. There are two steps to it. The labeling is given in the first phase based on the users' attitudes and semantic understanding. In the second stage, user behaviors are classified into normal and abnormal using a conventional supervised classifier.

3. Proposed method: The hybrid model of LSTM+CNN+ANN

Behavioral anomalies represent users' odd behavior regarding tweets posted in the OSM, which differs from ordinary tweets made by non-anomalous individuals. Following the detection of behavioral anomalies, any suspicious activities can be avoided. Furthermore, in most cases, the knowledge gained from the patterns in the tweets made by OSM users is crucial and valuable.

To discover an abnormality, it is required first to identify normal behavior. If the majority of the crowd-sourced data obtained via tweets for a specific worry, issue, or topic reveals similar patterns following analysis, it is said to as Normal. This paper proposes a hybrid structure of three deep neural networks, including LSTM, CNN, and ANN, to detect anomalies in OSM. The dataset we use for anomaly detection is the sentiment study of an American airline. Contributors were requested to classify positive, negative, and neutral tweets and then categorize unfavorable causes (such as "late flight" or "rude service"). People began expressing their ideas and opinions on the effects of the late flight or rude service in the airline sentiment program shortly after the announcement. The tweeting study aims to learn and understand how the general population feels about airline sentiment strategy.

3.1 The overall architecture of the model

First, the dataset is given to the LSTM network to convert the text tweet of the passengers to a probability value. The column of text tweets is

deleted, and the new column of probability values is added to the dataset. The updated dataset is given to the CNN network. The CNN network combines and mixtures the features of the dataset. As a result, several combined features are generated. Valuable features are selected and given to the ANN network

via the correlation operation. The ANN network, which includes the sigmoid activation function, classifies the data into normal and anomalies. (Fig. 3) shows the overall architecture of our proposed method. Below, we discuss the details of the architecture.

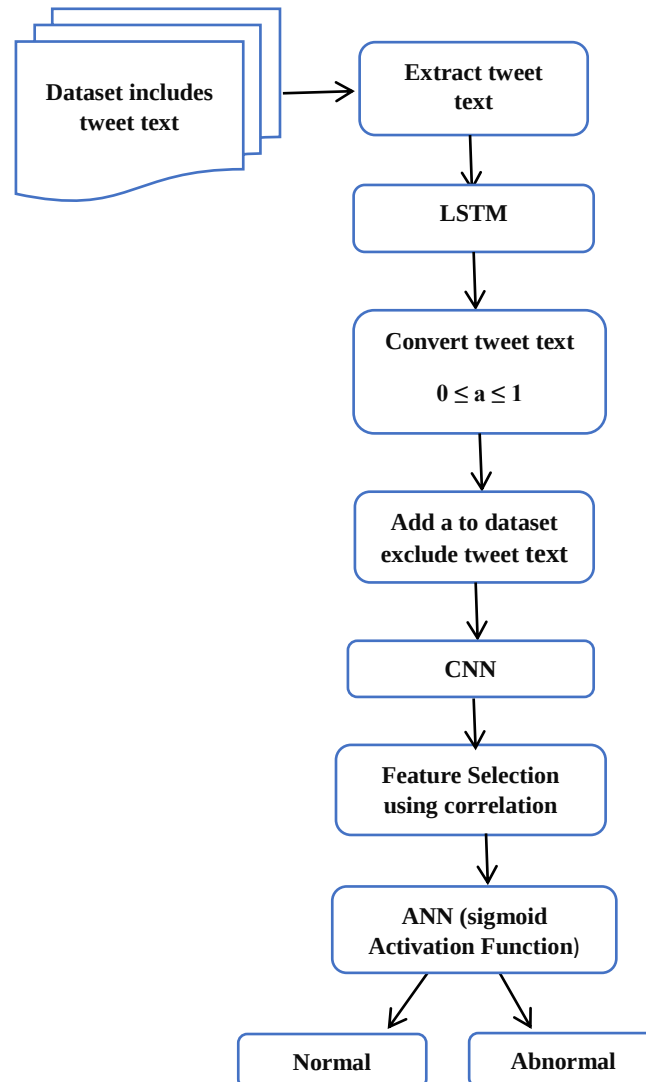


Figure 3: (Our proposed model architecture)

3.2 Dataset

The experiments are carried out on airline sentiment. A Twitter social network is made up of the comments of passengers. A sentiment study of each major American airline's concerns. Contributors were requested to classify positive, negative, and neutral tweets before categorizing unfavorable causes using twitter data from February 2015, followed by classifying negative reasons (such as "late flight" or "rude service"). This dataset has 14640 passengers' opinions on airline sentiment. The dataset is available on [Twitter US Airline Sentiment | Kaggle](#) [33].

3.3 Data Pre-Processing through the LSTM network

Data pre-processing is an essential step before any ML and DL technique.

Our dataset includes: _unit_id, _golden, _unit_state, _trusted_judgments, _last_judgment_, airline_sentimentairline_sentiment:confidence, negativereason, negativereason:confidence, airline airline_sentiment_gold, name, negativereason_gold, retweet_count, text, tweet_coord, tweet_created, tweet_id, tweet_location, user_timezone at preliminary features. One of the main important features is the tweet text of users. Moreover, features of the dataset have different and various text, including numeric, date, hour, text, and other categorical features. In the first step, we will form all features into numeric. To this end, the tweet text feature must be converted to a numeric value. Tweet text in the dataset has various lengths, from minimizing 10 characters to maximizing 24

characters, which is used to convert tweet texts into a numeric probability between zero and one. The architecture of the proposed LSTM network to

convert the text tweet to the numeric is given in (Fig. 4).

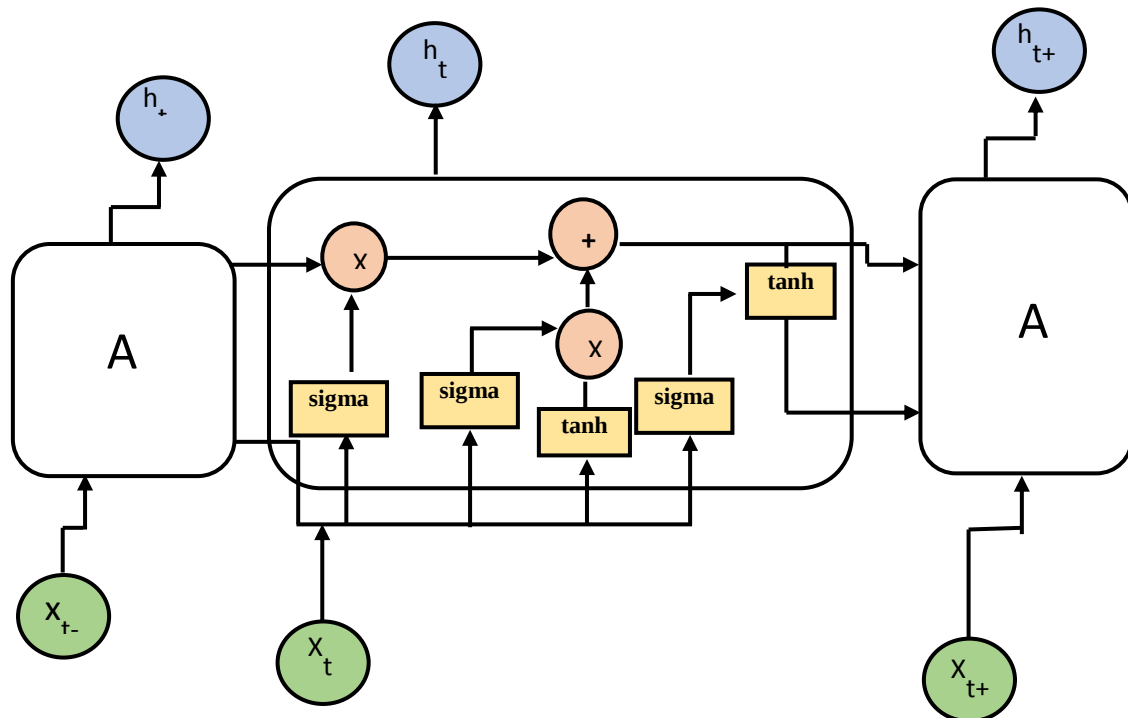


Figure 4: (The Model Architecture of LSTM)

Let us discuss the structure of the proposed LSTM network. The first Embedding layer is 32-neuron with 10000000 parameters. It uses a dense vector format to represent words. When a word is utilized, its position in the vector space is determined by the surrounding words. The next layer is a 64-neuron LSTM layer with 42240 parameters. "Embedding" is the input list of sentences, each of which is padd_sequence and has the same length. The activation function is a corrected linear activation function, which is extensively used. However, any other relevant activation function can be utilized. Next, a dense layer with 64 neurons and 4160 parameters with the activation function of 'Relu' is added. The last layer is a layer with 1-neuron, and 65 parameters with the activation function of 'sigmoid' that will be used for data classification. The total parameters of the dense layers are 10,046,465. Given that we are working with a binary dataset, we will use a generic Adam optimizer to regulate the learning rate and assign the loss as binary cross entropy. The LSTM network is ready and can be utilized in our proposed method.

Let us have a look at the column of text tweets. The texts have variable lengths with different characters (Maximum length has 24 characters and minimum length has two characters). We plan to convert any word to a sequence of integer numbers. To this end, a vocab_size of 100000 words is selected. However, if a word does not appear in the vocab_size it is filled through Oov_token=" <oov>." All the tweets, which are in the text feature, are put in the list. The

following steps are necessary and discussed to convert the texts to numeric.

- 1- Tokenizer classes were imported through TensorFlow.Keras which consists of a dictionary that recognizes anomaly words; because of that, it has been applied to all the sentences
- 2- The index is taken for each one of the words.
- 3- Each tweet is converted into a sequence, it means instead of each word a sequence of numbers is put.
- 4- The length of all tweets was set to the same size through pad sequence.

After setting the architecture, using the class of train_test_split, 80% and 20% of text tweets are used for training and testing, respectively. The LSTM network converts every text tweet to a probability value. The probability feature is made as a new feature and the text feature then is deleted from the dataset.

The Date and Time features are two features in the dataset. We created a new time-based feature through (hour*60+minute), and the date feature is separated and split into the day, month, and year. Next, the year's feature is deleted as all the samples have the same year. After all converting and splitting the features, totally 18 features are identified and generated. Next, all features are normalized. The specified attributes are trusted_judgments, airline_sentiment: confidence, retweet_count, New_Text, True, finalized, golden, in progress, American, Delta, Southwest, US Airways, United, Virgin America, hours, minutes, New_Time,

airline_sentiment. (Fig. 5) shows the correlation map of the 18 features of the dataset. As the correlation between the features are high in most of cases, if we

give these features to the classifier network (ANN), the good results are not achieved.

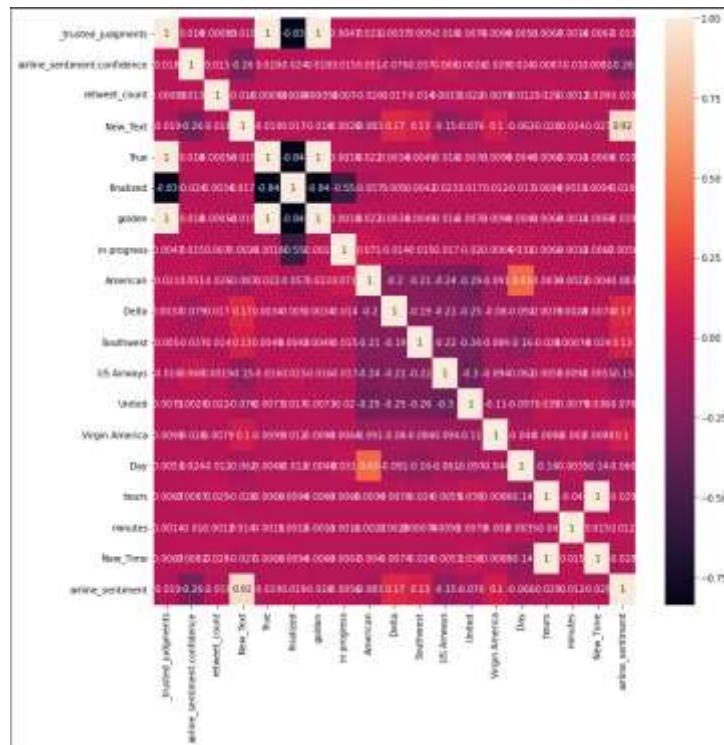


Fig. 5: (The correlation between features)

3.4 The CNN Network

The next step of our proposed method is to combine and merge features and generate new valuable features. The CNN network is normally used for image classification. A convolute operation runs on the image and scans the image via some filters and then new feature maps are generated. This operation is repeated several times with new filters and as a result many feature maps are generated from the image. Any image is three sets of numbers (between

0 and 255) for three channels of green, blue and red. However, for a black-white image, the image has exactly one set of numbers between zero and 255. If we put the numeric values of features in an image platform and then give it to the CNN network, CNN does the convolute operation and thereby, every vicinity features are combined and new feature maps are generated. That means we want to get several tens of new features. The architecture of the proposed CNN model is shown in (Fig. 6).

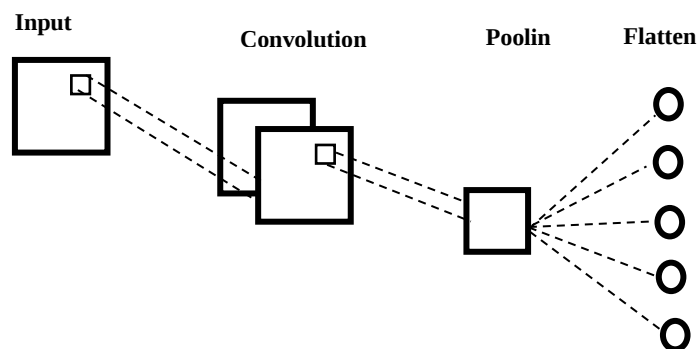


Fig. 6: (The Model Architecture of CNN)

The proposed CNN network has three layers including two convolutional layers and one pooling layer. The first convolutional layer has 12 filters of size (1,3), the activation function of 'Relu' and the input shape of (1,18,1). The second convolutional layer has 16 filters of size (1,3), and the activation function

of 'Relu'. The third layer is a pooling layer with the size of (1,2) which will be used to reduce number of parameters to a significant amount. The final step in the CNN network is to flatten the data to a 1D array to be the input of the ANN network.

3.5 Feature Selection

The CNN network combines 18 features and finally generates 25088 new features. The new features are the outcome of combination and mixture of 18 old features. However, not all generated features are useful to train the classifier network. Feature selection is a valuable step that is done to select the most important features. We generate the correlation map for all features. The features with low correlation amount are kept. Those features that have high correlations together, one of them is kept and the rest is deleted. Finally, after this procedure 102 features are selected and used to be applied to the classifier network of ANN.

3.6 ANN

The first step in preparing the ANN is to import the TensorFlow model. The Sequential class in Keras is then utilized to build our model, allowing us to group

a linear stack of layers into the model based on their arrival order. We generate a variety of dense layers within the model. Every dense layer is a neural network layer that accepts various arguments. We used three layers, two of them utilize the 'Relu' activation function, and the last layer uses the sigmoid activation function. The first dense layer has 6 units, and 114 parameters. The second dense layer has 6 units and 42 parameters. The third and final dense layer has one unit and seven parameters. We will use a generic Adam optimizer to regulate the learning rate and assign the loss as binary cross entropy. We can now fit the model to the dataset after it has been constructed. The training set will be used as input, and a batch size of 32 will be used to perform 100 epochs. The architecture for the ANN model is shown in (Fig. 7).

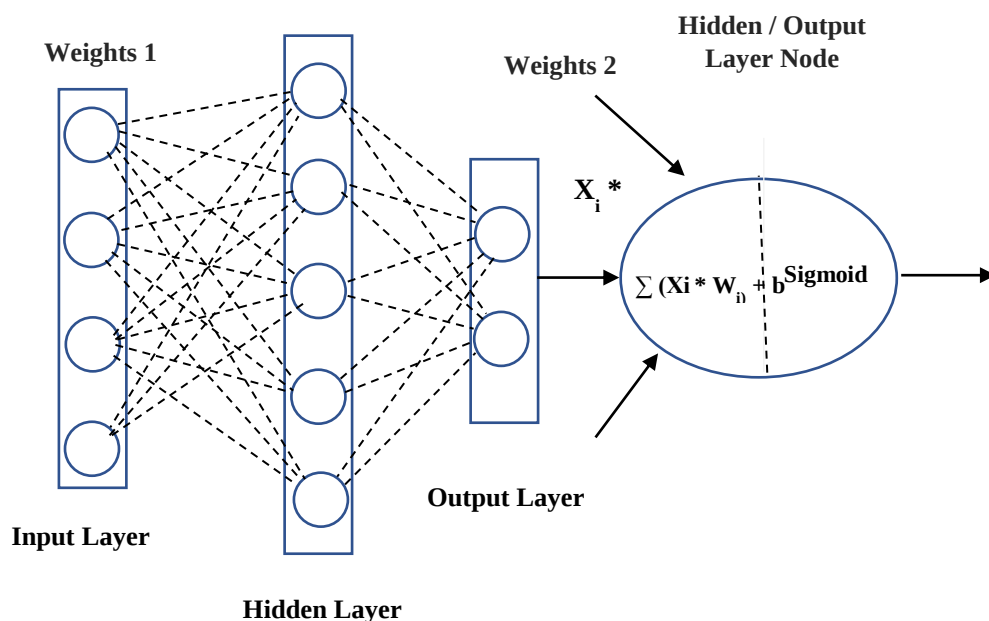


Fig. 7: (The Model Architecture of ANN)

3.7 Model training by supervised learning

Nowadays supervised learning is the foremost common sub-branch of machine learning. Ordinarily, modern machine learning professionals will start their travel with supervised learning calculations. It is supervised by utilizing of labeled datasets to train and prepare calculations that classify information or foresee results precisely. The label of the dataset consists of positive, neutral, and negative. The model training of dataset %70 is used and 30% is used for testing.

3.7.1 Logistic Regression

The logistic regression method is utilized when the dependent variable is double. The logistic regression is approximating the parameters of a logistic model and is a kind of distant regression. A logistic is utilized to deal with data that has two appropriate criteria and the relationship between criteria and forecast. A logistic is one of the forms of the regression analysis methods, which is utilized in the

place where the dependent variable is separated. Example: right or wrong, 0 or 1, etc. This means that the target variable can have only two values, and the relationship between the target variable and the independent variable is denoted by the sigmoid curve. The Eq. (1) contains the equation to calculate linear regression.

$$Y = \alpha_0 + \alpha_1 X_1 + \alpha_2 X_2 + \dots + \alpha_n X_n \dots (1)$$

While X_1 , X_2 , and so on are explanatory variables, Y is a dependent variable. In the Logistic Regression, the value α_0 causes the curve to move to the left, and the constants α_1 , α_2 , ..., α_n also known as the slope indicate how steep the curve should be. Later Sigmoid Function on the

Linear Regression using Eq. (2) is applied.

$$P = 1 / (1 + e^{-Y}) \dots (2)$$

The logistic function, commonly known as the sigmoid function, offers a "S"-shaped curve that can be utilized to map any real-world quantity to a number between 0 and 1. The sigmoid's value can be

classified as non-anomalous if it is greater than 0.5, and it can be classified as anomalous if it is lower than 0.5.

3.7.2 Random Forest

The random forest, as the name implies, is a collection of decision-making trees that work together as an ensemble. In the random forest, every tree sprays a class prediction, and the class with the most votes is the model's prediction. Individual training decision-making trees are developed by RFs. To generate a final prediction, all of the trees' predictions are pooled together. Ensemble approaches are used when they produce a final decision based on a set of data. They are referred to as ensemble approaches if they make a final decision based on a set of data. Gini Index is used to decide on how nodes on a decision tree branch and is given by Eq. (3).

$$G = 1 - \sum_{k=1}^c (p_i)^2 \dots (3)$$

While c is the total number of classes and p_i is the relative frequency of the class that is present in the dataset.

3.7.3 SVM

SVM is a supervised machine learning model that uses classification algorithms to separate between two classes in the training data. The SVM's goal is to identify the optimal line to segregate data. You can classify new text after SVM groups provide course information for each group. SVM is a speed and reliable classification method that works well with a restricted amount of data for analyze. When the data is trained SVM then classifies the data into normal and anomalies through a hyperplane.

3.7.4 KNN

The KNN approach saves whole usable input instances and distributes new information items depending on likeness measures. The KNN is also named lazy learning. Of course, the value of K refers to the number of neighbors. All the classification procedures depend on the KNN, the classification efficiency mainly be dependent on the K , to the calculate distance between points. There are many methods for this distance estimation, of which the most widely known methods are Euclidean distance, Manhattan distance, and Hamming distance. We were used Euclidian distance to find the distance between

points and then classify them into normal and anomalies.

4. Experimental and discussion

The proposed model has 14640 tweets in the dataset labeled and organized. The Python programming language is used to implement our proposed model. The Python packages used in our model are Numpy, Pandas, TensorFlow, and Sklearn libraries.

To evaluate the performance of the proposed model, we compare it with singular LSTM and singular ANN as deep-learning techniques. Moreover, the proposed model is compared with famous machine learning models, including logistic regression (LR), support vector machine (SVM), decision tree (Dtree), random forest (RF), and K-nearest neighbor (KNN).

Accuracy, Recall, F-score, and Precision are the most important metrics to compare the proposed technique with other classical ML and DL techniques. Relations (4) to (7) show Precision, Recall, F-score, and Accuracy, respectively. These measurements are used for each of the algorithms.

$$\text{Precision} = \frac{Tp}{Tp+Fp} \dots (4)$$

$$\text{Recall} = \frac{Tp}{Tp+Fn} \dots (5)$$

$$F - \text{score} = \frac{2 * \text{Precision} * \text{Recall}}{\text{Precision} + \text{Recall}} \dots (6)$$

$$\text{Accuracy} = \frac{Tp+Tn}{Tp+Tn+Fp+Fn} \dots (7)$$

As can be seen in (Table 1), our proposed method outperforms other classical techniques. The proposed method enhances the precision on average by 8.4% compared to previous classical ML and DL techniques. (Fig. 8) shows the detailed results of precision, recall, f-score, and accuracy for the proposed technique and other classical ML and DL techniques. Experimental results show that our proposed technique enhances the recall metric by 8% compared to previous classical ML and DL techniques. By investigating (Table 1) and (Fig. 8), we can see the proposed method enhances f1-score and accuracy on average by 9.4% and 8.6%, respectively. (Table 2) summarizes the improvement of our proposed technique over all other classical techniques, metric by metric.

Table 2: (Metrics Comparison for Anomaly Detection)

Algorithms	Recall	Precision	F1-Score	Accuracy
KNN(k=5)	%84	%84	%84	%88.28
SVM	%82	%86	%84	%88.68
DT	%86	%87	%86	%90.25
RF	%86	%87	%87	%90.27
LR	%83	%85	%84	%88.68
LSTM	%84	%89	%89	%92.95
ANN	%83	%88	%87	%91.82
LSTM+CNN+ANN	%91	%93	%93	%96.92

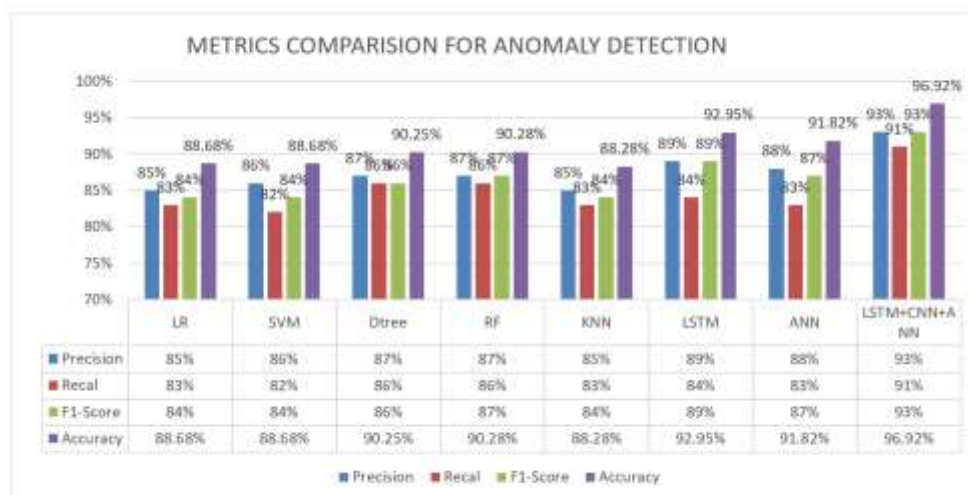


Fig. 8: (Comparison of anomaly detection by metrics)

Table 3: (Improvement over other algorithms by our proposed method)

Algorithms	Recall	Precision	F1-Score	Accuracy	Average Metrics
KNN(k=5)	%8.33	%10.71	%10.71	%9.78	%9.88
SVM	%10.97	%8.14	%10.71	%9.29	%9.78
DT	%5.81	%6.89	%8.14	%7.39	%7
RF	%5.81	%6.89	%6.89	%7.36	%6.74
LR	%9.64	%9.41	%10.71	%9.29	%9.76
LSTM	%8.33	%4.49	%4.49	%4.27	%5.39
ANN	%9.64	%5.68	%6.89	%5.55	%6.94

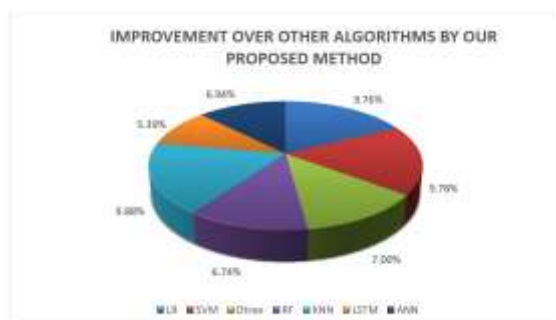


Fig. 9: (Improvement over other algorithms by our proposed method)

The average improvement (the average of all metrics) achieved by the proposed technique over all other classical techniques is shown graphically in (Fig.9). As can be seen in the figure, the proposed technique improves all performance metrics on average by 9.88%, 9.78%, 9.76%, 7%, 6.74%, 5.39%, and 6.94% in compared to KNN, SNM, LR, Dtree, RF, LSTM, and ANN, respectively.

5. Conclusion

While OSM has become a catalyst for millions of individuals worldwide to connect and benefit from various services, it may also provide fertile ground for harmful activities. The discovery of anomalies in the OSM platforms is the ultimate goal of many

academic and industrial researchers. This paper proposes a hybrid technique composed of LSTM, CNN, and ANN to detect anomalies in the OSM platform of Tweeter. The proposed method is compared with singular LSTM, CNN, and ANN. It also is compared with various machine-learning methods. The experimental results show that the proposed method outperforms all deep-learning and machine-learning methods in detecting anomalies in the Tweeter OSM platform. The results show that the proposed method enhances detection accuracy on average by 8.6% better than previous classical techniques. The proposed method improves the recall metric and precision on average by 8.3% and 8.4% compared to previous classical techniques. Our results show the superiority of our technique over classical deep-learning and machine-learning and encourage the managers of the OSM platforms to use our technique for detecting anomalies in their platforms. We decide to improve the technique's accuracy by utilizing better ideas for feature selections, such as genetic algorithms. Moreover, we intend to use optimization techniques (e.g., Cuckoo search) at the output of the ANN network for better classification.

References

- [1] Md. Shafiur, R. et al. (2021). An efficient hybrid system for anomaly detection in social networks. *springeropen*.
- [2] Keith, Henderson. et al. (2011). Graph Mining using Recursive Structural Features. San Diego, United States: p. 663–671.
- [3] Akoglu, L.; McGlohon, M. and Faloutsos, a. C. (2010). Spotting Anomalies in Weighted Graphs. *Advances in Knowledge Discovery and Data Mining*, Volume 6119: p. 410–421.
- [4] Abdolazim, R. et al. (2013). Anomaly Detection in Online Social Networks using Structure-Based Technique. London, UK, IEEE.
- [5] Sujatha, A. K. and Balachandran, K. (2020). Self-Supervised Learning Based Anomaly Detection in Online social media. *International Journal of Intelligent Engineering and Systems*, **13**(3).
- [6] Nicholas, A. H. et al. (2010). Bayesian anomaly detection methods for social networks. *The Annals of Applied Statistics*.
- [7] Renjun, H. et al. (2016). An embedding approach to anomaly detection. Helsinki, Finland, IEEE.
- [8] M.Sai, S. L. Y. et al. (2017). Identifying Malicious Data in social media. *International Research Journal of Engineering and Technology*.
- [9] Rose, Y.; Xinran, H. and Yan, L. (2015). Group Anomaly Detection in Social Media Analysis. *ACM Journals*, **10**(2): 1–22.
- [10] Anshika, C.; Himangi, M. and Anuja, A. (2019). Anomaly Detection using Graph Neural Networks. Faridabad, India, IEEE.
- [11] Jing, G. et al. (2011). A spectral framework for detecting inconsistency across multi-source object relationships. Vancouver, BC, Canada, IEEE.
- [12] Kumar, S. and Kamalika, D. (2014). Localizing anomalous changes in time evolving graphs. *ACM SIGMOD*: 1347-1358.
- [13] Md. Shafiur, R. et al. (2021). An efficient hybrid system for anomaly detection in social networks. *springeropen*.
- [14] Agovic, A. et al. (2009). Anomaly detection using manifold embedding and its applications in transportation corridors. *Intelligent Data Analysis*, **13**(3): 435-455.
- [15] Renjun, H. et al. (2016). An embedding approach to anomaly detection. Helsinki, Finland, IEEE.
- [16] Abdolazim, R. et al. (2013). Anomaly Detection in Online Social Networks using Structure-Based Technique. London, UK, IEEE.
- [17] Wenchao, Y. et al. (2018). A Flexible Deep Embedding Approach for Anomaly Detection in Dynamic Networks. London, United, Kingdom, *ACM SIGKDD*: p. 2672–2681.
- [18] Luan, T.; Liye, F. and Cyrus, S. (2016). Distance-based outlier detection in data streams. In: *Proc. of the VLDB Endowment*, **9**(12): 1089–1100.
- [19] Ahmet, Ş. D. (2019). ANOMALOUS ACTIVITY DETECTION FROM DAILY SOCIAL MEDIA USER MOBILITY DATA. *Omer Halisdemir University Journal of Engineering Sciences*, **8**(2): 638 - 651.
- [20] Sujatha, A. K. and Balachandran, K. (2020). Self-Supervised Learning Based Anomaly Detection in Online social media. *International Journal of Intelligent Engineering and Systems*, **13**(3).
- [21] Amogh, M.; Nisheeth, S. and Jaideep, S. (2012). Contextual Anomaly Detection in Text Data, **5**(4): 469-489.
- [22] Guixian, X. et al. (2019). Research on Topic Detection and Tracking for Online News Texts. *IEEE Access*, 30 April. Volume 7: 58407 - 58418.
- [23] Ying, F. et al. (2014). Self-Adaptive Topic Model: A Solution to the Problem of 'Rich Topics get Richer. *China Communications*, Dec, **11**(12): 35 - 43.
- [24] Than, T. W. and Sint, S. A. (2018). TPMTM: Topic Modeling over Papers' Abstract. *Advances in Science, Technology and Engineering Systems Journal*, Volume 3: 69-73.
- [25] Brejit, L. A. and Anjana.P.Nair. (2018). Anomalous Topic Discovery Based on Topic Modeling from Document Cluster. *International Research Journal of Engineering and Technology*, Feb, **5**(2): 966-972.
- [26] MUHAMMAD, A. and SAJID, Y. B., (2015). Classifier ensembles using structural features for spammer detection in online social networks. *Foundations of Computing and Decision Sciences*, **40**(2): 89-105.
- [27] Carl, S. and Alta, d. W. (2016). Semi-supervised machine learning for textual anomaly detection. Stellenbosch, South Africa, IEEE: 1-5.
- [28] Ravneet, K.; Sarbjeet, S. and Harish, K. (2018). AuthCom: Authorship Verification and Compromised Account Detection in Online Social Networks using AHP-TOPSIS Embedded Profiling based Technique. *ELSEVIER*, 15 December, Volume 113: 397-414.
- [29] Mei, M. et al. (2018). Using Semantic Clustering and Auto Encoders for Detecting Novelty in Corpora of Short Texts. Rio de Janeiro, Brazil, IEEE: 1-8.
- [30] Oliver, L. H. et al. (2017). Relationship Breakup Disclosures and Media Ideologies on Facebook. *New Media & Society*, **20**(5): 1931-1952.
- [31] Thin, N. et al. (2014). Affective and Content Analysis of Online Depression Communities. *IEEE Transactions on Affective Computing*, **5**(3): 217 - 226.
- [32] Randa, A. et al. (2019). Anomaly detection over differential preserved privacy in online social networks. Volume 14: 1-20.
- [33] Savyan, P. and S. Mary, S. B., (2017). Behaviour Profiling of Reactions in Facebook Posts for Anomaly Detection. Chennai, India, IEEE.
- [34] Anon. 2015. Kaggle. [Online] Available at: <https://www.kaggle.com/datasets/crowdfLOWER/twitter-airline-sentiment>

نموذج تعلم عميق hybrid لاكتشاف الحالات الشاذة بدقة في وسائل التواصل الاجتماعي عبر الإنترنت

دهرياز معروف حسين حاكم سيد بهيتولاها

قسم علوم الحاسوب ، جامعة سوران، اربيل، العراق

الملخص

تولد الوسائط الاجتماعية عبر الإنترنت (OSM) كمية هائلة من البيانات حول السلوك البشري بناءً على تفاعلاتها. يعبر الناس عن آرائهم وتعليقاتهم ويشاركون المعلومات حول مجموعة متنوعة من مواضيع حياتهم اليومية من خلال OSM. تنقسم غالبية التعليقات إلى ثلاث فئات: إيجابية وسلبية وطبيعية. فيما يتعلق بالتعليقات السلبية، تسهل منصة OSM الإجراءات غير الطبيعية، مثل الرسائل غير المرغوب فيها والمعلومات المضللة والشائعات ونشر الأخبار والدعاية المزيفة، فضلاً عن نشر الروابط الخبيثة. لذلك، فإن أحد أهم أنشطة تحليل البيانات لتحديد الأفراد العاديين والمنحرفين على الشبكات الاجتماعية هو اكتشاف الشذوذ. تقترح هذه الورقة نموذجاً هجيناً يعتمد على ثلاثة مناهج تعلم عميق شهيرة لاكتشاف التشوهات السلوكية والتعليقات السلبية في منصات OSM. المعيار المختار لبحثنا هو شعور شركات الطيران في مجموعة بيانات تويتر. في الطريقة المقترحة، يتم تغذية مجموعة البيانات إلى شبكة LSTM؛ بعد ذلك، يغذي ناتج LSTM شبكة CNN. تجمع CNN بين الميزات وتقوم بإنشاء خرائط ميزات متنوعة. في الخطوة التالية، نقوم بتقليل أهم الميزات واختيارها. أخيراً، يتم إعطاء الميزات المحددة لـ ANN لتصنيف البيانات. تتم مقارنة الطريقة المقترحة (LSTM + CNN + ANN) بمختلف تقنيات التعلم الآلي الكلاسيكية (ML). أظهرت النتائج التجريبية أن الطريقة المقترحة تعزز الدقة في المتوسط بنسبة 8.6% و 8.4% على التوالي مقارنة بتقنيات ML التقليدية.