

RESEARCH ARTICLE



Optimized Deep Learning Techniques to Identify Rumors and Fake News in Online Social Networks

Abu Sarwar Zamani^{1,2} , Aisha Hassan Abdalla Hashim², Sara Saadeldeen Ibrahim Mohamed¹ and Nasre Alam^{3,*}

¹Department of Computer and Self Development, Prince Sattam bin Abdulaziz University, Saudi Arabia

²Department of Electrical and Computer Engineering, International Islamic University Malaysia, Malaysia

³Department of Computer Science, Woldia University, Ethiopia

Abstract: The swift expansion of networking platforms has led to a significant proliferation of fake news on social media in recent years, posing a serious risk to public safety. This phenomenon carries various potential negative effects on society, including the erosion of public confidence in journalists and governmental institutions. Consequently, the identification of fake news has attracted considerable attention from researchers across various fields. As online and social media platforms have grown, it has become easier for false information to mix in with real or verified information. People who spread false information usually have some kind of political or social goal in mind when they spread their hoaxes. Because of this, it is of the utmost importance to come up with a trustworthy way to spot false information. This article describes a way to use deep learning to spot fake news. Methodology is made up of a set of input data. The information in this dataset comes from the social networking site Twitter. First, the raw data that is being used are preprocessed. Stop word removal, stemming, and tokenization are the main parts of data preprocessing. The NTLK library is used to get rid of stop words. Porter's Algorithm is used to do stemming. N-gram model is used to do tokenization. LSTM, CNN, and AdaBoost algorithms are used to build the model. Results have shown that LSTM is better than CNN and AdaBoost in terms of accuracy, specificity, and sensitivity. LSTM has achieved an accuracy of 99.24% for fake news detection. Specificity of LSTM is 99.2%. LSTM's sensitivity is 98.67%. LSTM has achieved an accuracy of 99.24% for fake news detection. Specificity of LSTM is 99.2% and sensitivity is 98.67%.

Keywords: fake news detection, deep learning, Long Short-Term Memory (LSTM), N-gram, porters stemming, social networks

1. Introduction

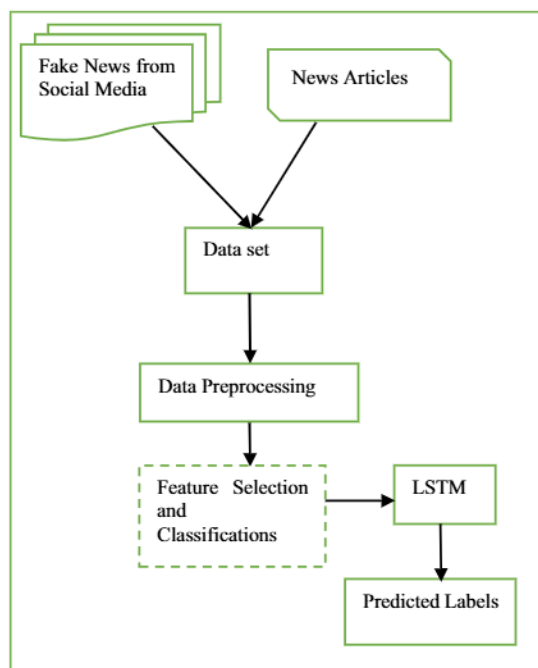
The proliferation of dishonest content on social media platforms such as Facebook, Twitter, and other similar platforms has led to an increase in the number of scholars engaged in identifying and debunking such misinformation [1]. Misinformation is being disseminated at an alarmingly rapid rate in an attempt to garner attention on social media and divert people's attention away from the most pressing issues. The vast majorities of people have a natural tendency to be skeptical about the reliability of information and choose to place their confidence in what they read on social media platforms instead. In addition, there is a phenomenon known as "confirmation bias," in which individuals have a tendency to assign more weight to information that confirms the worldviews they already hold. Research has demonstrated that, in general, people have a difficult time recognizing dishonesty in others. Misleading and manufactured news reports that are published in a wide variety of media venues are frequently disseminated with the goal of achieving a variety of unfavorable ends, such as the ridicule of the

general public or financial benefit. Multimodal signals can take a variety of forms, including but not limited to images, videos, articles, links, postings, and blogs. These and other forms of media are frequently seen on social networking sites. Rumors are harmful to social harmony because they are disruptive and damaging to social harmony. Because of the volume of transmission, it possesses and/or the harm it has already caused, locating its source or establishing whether or not such rumors are grounded in truth requires a significant amount of work. In addition, the dissemination of incorrect information regarding a significant news piece or occurrence can have a devastating effect on society in a relatively short amount of time [2].

The usage of social media is now so widespread that users may be misled into believing that the material they obtain on these platforms in the form of quotations, articles, images, posts, videos, and so on is accurate. This kind of information overload on social media platforms confuses consumers, sways beliefs, and moves political agendas forward. It may also be lucrative for online publishers. When making a fake website, it is essential to select a domain name that is phonetically comparable to that of a real website. The spread of stuff like this alters people's opinions and values by drawing their attention away from real news and onto fake information [3].

*Corresponding author: Nasre Alam, Department of Computer Science, Woldia University, Ethiopia. Email: nasarhi@wldu.edu.et

Figure 1
Block diagram of fake news detection



A general process of fake news detection is shown above in Figure 1. Input dataset is prepared by collecting tweets from twitter. This data is preprocessed. Preprocessing includes stop word removal, stemming, and tokenization. Appropriate features are selected to enhance the accuracy of classification model. Next, machine learning and deep learning methods are used to classify tokenized data [4, 5]. In this way, deep learning and machine learning are used to detect fake news and rumors.

The growth of online and social media platforms has made it easier for false information to blend in with genuine or confirmed content [6]. It is possible to utilize it to influence public opinion and, as a consequence, the perceptions, thoughts, and actions of individuals. Because of this, it is really straightforward for individuals who fabricate fake news to disseminate their links, posts, images, videos, and audio files throughout all of the social media platforms [7]. When they spread their hoaxes, the people that spread disinformation typically have some sort of political or societal agenda in mind. As a consequence of this, the development of a trustworthy system for recognizing misinformation is of the utmost importance [8].

A prevalent approach involves the utilization of Machine Learning (ML), which has demonstrated considerable efficacy in identifying and mitigating the spread of fake news across social media platforms. ML encompasses a variety of techniques and algorithms, such as support vector machines (SVM) and naïve Bayes classifiers. Both SVM and naïve Bayes classifiers are categorized as supervised models within the ML framework. They have been utilized to assess the veracity of news articles and have proven to be effective. Nonetheless, certain limitations exist regarding their application [9].

This article presents a methodology for fake news identification using deep learning. Methodology consists of input dataset. This dataset is created by collecting data from social networking site twitter. First, data preprocessing is performed on the raw input data. Data preprocessing mainly includes Stop Word Removal,

Stemming, and Tokenization. For stop word removal NTLK library is used. Stemming is performed using Porters Algorithm. Tokenization is accomplished by N-gram model. Model is build using LSTM, CNN and AdaBoost algorithms. Results have proven that the accuracy, specificity, and sensitivity of LSTM is better than that of CNN and AdaBoost methods.

2. Literature Review

Scholars introduced a framework for identifying fake news through n-gram analysis and Term Frequency-Inverse Document Frequency (Tf-Idf) as a method for extracting features. These two methods are used in conjunction with one another. In this study, a total of six different machine learning classifiers are utilized in order to compare and contrast two distinct feature extraction strategies. There are six supervised classifications [10] methods, and they are as follows: K-Nearest Neighbor, Support Vector Machine, Logistic Regression, Linear Space Vector Machine, and Decision Tree. Stochastic Gradient Descent is the sixth method (SGD). For the purpose of this experiment, we make use of a dataset that is open to the public and draws its contents from both fake and real news websites.

To the best of our knowledge, this research is the first research team to carry out an exhaustive examination of how members of public WhatsApp groups behave when debating contentious subjects. The dataset, which was compiled over the course of 28 days at three distinct levels, yielded the recovery of more than 270 thousand messages and seven thousand individuals (message, user, group). Lahby et al. [10] as well as Santos de Oliveira and Vaz de Melo [11] give a variety of metrics that are intended to characterize the activities that take place within WhatsApp groups. Among the several methods of information retrieval that have been studied, the Tf-Idf approach is the most effective one. The Tf-Idf method assigns a score to a term in any content based on the frequency with which it appears in that content. The TF-IDF feature extraction strategy is one that has the potential to be helpful when extracting textual features from a variety of sources, such as headlines, news stories, and textual posts. The following part will focus more attention on syntactic models and go into greater detail about them.

Yang et al. [12] dissected both false and correct reports using vector space modeling (VSM) in conjunction with rhetorical structure theory (RST). Multiple nodes in each dimension stand in for the rhetorical linkages that might be made between different pieces of news. These types of news stories are represented as vectors in an n-dimensional space. It is standard practice to divide news stories into two distinct categories: "fake" and "genuine," with the former being regarded as a subset of the latter. When a piece of news is included in a cluster, it is possible for that piece to be easily recognized as being a member of the cluster due to the fact that it shares certain features with the other items included in the cluster.

An unsupervised K-Nearest Neighbor (KNN) identification technique was proposed by Barhate et al. [13] as a means of locating users who engage in manipulative behavior. The individuals' neighborhoods are also determined in order to more accurately portray the similarities between them in terms of the number of expected followers [13]. When utilizing this model, it is possible to estimate the number of followers to within 8.42 percentage points of the actual figure.

Using an algorithm that was based on a support vector machine (SVM), Pennycook and Rand [14] proposed a satire identification model that could be utilized in four distinct fields: science, business, soft news, and civics. In this study, we make use of two different datasets: the first contains 240 news articles drawn from both satirical and serious sources, while the second contains 360 news

items drawn from the same sources. If you use a combination of five characteristics that has an accuracy of 90% and a recall of 84% for detecting satirical news, you can lessen the deceptive effect that satire has by reducing the likelihood that people will take it seriously. Machine learning techniques, like SVM and Nave Bayes, are commonly applied in text analysis of news articles, whether they are real or fake [15, 16].

The characteristics of fake news were outlined in a definition provided by Yang et al. [12] A piece of fabricated news is immediately identifiable by the fact that it presents material that is patently false and is written in such a way as to mislead its audience into believing that it is authentic. Second, the data that is produced by people who actively participate in social activities is enormous, unstructured, and loud, all of which provide substantial challenges. Additionally, the data is incomplete. The authors give a thorough analysis of the current state-of-the-art in recognizing fake news in OSM from the points of view of evaluation metrics, representative datasets, and data mining. In addition to this, the paper outlines unresolved difficulties and prospective future research fields [17] regarding the detection of fake news.

The novel approach that Bahad et al. [18] have shown, which continually learns to identify rumors, is an extremely interesting development. Recurrent neural networks (RNNs) are used in the suggested method to learn the hidden features required to capture the temporal context of posts [19]. The findings indicate that the RNN approach [18] is superior to the manually created rumor detection models that were used in the past and that it is possible to enhance its performance by including more hidden layers.

CSI (Capture, Score, and Integrate) is a model that was developed by Ruchansky et al. [20] for the purpose of making automated predictions that are more precise. This model incorporates all three of the following aspects: the content of articles; user comments received; and sources that people use to promote it. The response-and-text model makes use of an RNN in order to more accurately capture the rhythms of user behavior over time. In terms of accuracy on data taken from the real world, it has been demonstrated through experiments that CSI performs better than the most advanced models currently available [21].

Hsu et al. [22] introduce a common fake features network (CFFN) that is built on deep learning and uses contrastive loss to differentiate between false photos and real ones. In the first step of the process, high-performance GANs are put to use so that pairings of synthetic and real images can be generated. The subsequent step entails the formation of a dense network known as DenseNet, which is designed to take in paired data. The following stage consists of instructing the proposed network to differentiate between authentic and altered photographs. At the very end, a classification layer is added to the mix with the stages that came before it in order to identify whether or not the input image is faked. Using the CelebA dataset, the DCGAN, WGAP, WGAP-GP, LSGAN, and PGGAN networks were utilized to construct the training set of false photographs.

The research that was carried out by Jammal et al. [23] has the capability of revealing fake faces that were created by both humans and computers. A recommended ensemble-based neural network (NN) classifier with an AUC-ROC of 94–99% is offered in this article in order to recognize computer-generated images that were produced by a GAN. In addition, a neural network model that employs noise filtering and face cropping has been developed to recognize false face pictures that were created by humans and has achieved an AUC-ROC score of 74.9%. In order to verify the validity of the proposed model, the CelebA and PGGAN datasets are utilized.

Tan and Chan [24] proposed a phrase-based image captioning model to construct image description. To do so, they made use of a hierarchical LSTM architecture that they named phi-LSTM. The usual approaches that are considered state-of-the-art generate the image caption in a sequential fashion, whereas the phi-LSTM decodes the picture captioning from phrases to sentences. The suggested strategy beats other state-of-the-art models when tested on the MS-COCO dataset, as well as the Flickr8k and Flickr30k datasets.

In their generative model, Kumar and Verma [25] make use of a convolutional neural network (CNN), which is then followed by a deep recurrent architecture (a language generator). The model is able to produce complete sentences in natural language based on the input image. The model achieves a high level of performance when compared to other models that are considered to be state-of-the-art based on the BLEU-1 score.

Zhou et al. [26] suggest a two-stream network (face classification and patch triplet) for identifying face alterations. A convolutional neural network (CNN) model is trained in real time in order to determine whether or not a particular facial image has been modified. For the purpose of effective picture splicing recognition [27], steganographic elements are incorporated into the training of a second stream. Using the FaceSwap and SwapMe tools, an evaluation of the proposed model was carried out on the recently obtained dataset.

3. Research Methodology

3.1. Deep learning for fake news detection

This section presents a methodology for fake news identification using deep learning. Methodology consists of input dataset. This dataset is created by collecting data from social networking site twitter. First, data preprocessing is performed on the raw input data. Data preprocessing mainly includes Stop Word Removal, Stemming, and Tokenization. For stop word removal NTLK library is used [28]. Stemming is performed using Porters Algorithm. Tokenization is accomplished by N-gram model. Model is built using LSTM, CNN, and AdaBoost algorithms. Tokenization is accomplished by N-gram model. Single-occurring words in a document are called uni grams. Bigram is a combination of two words. Trigram is a combination of three words.

Because it is an RNN network that incorporates gates such as I/P, forget, and O/P in addition to an additional memory cell, the technique known as LSTM (Long Short-Term Memory) [29] has gained a lot of popularity in recent years. Since LSTM networks are capable of remembering data for very long periods of time, gradient descent may be replaced with their utilization. When this occurs, the networks are able to recognize patterns and sequences with increased precision. In order to guarantee that the data is accurate, the input and output gates are put to use in the intervals between the time steps. The input gate I in Equation (2) determines the contents of the memory cell that it is associated with. The Forget Gate (f) is a device that works by changing the current state of the cell so that it is in line with Equation (1). This causes all memories to be rewritten. Everything that came before this is completely unimportant at this point. The final step, which is referred to as the output gate (y), is responsible for controlling the information that is passed on to the subsequent stage. Due to the fact that all three gates are attached to the memory cell, it is possible to trace the timing of the output. LSTM is shown in Figure 2. A schematic representation of the LSTM network is provided in Figure 3. At each time step, an embedding x_i is given to the network as an input, and the network's

Figure 2
LSTM architecture

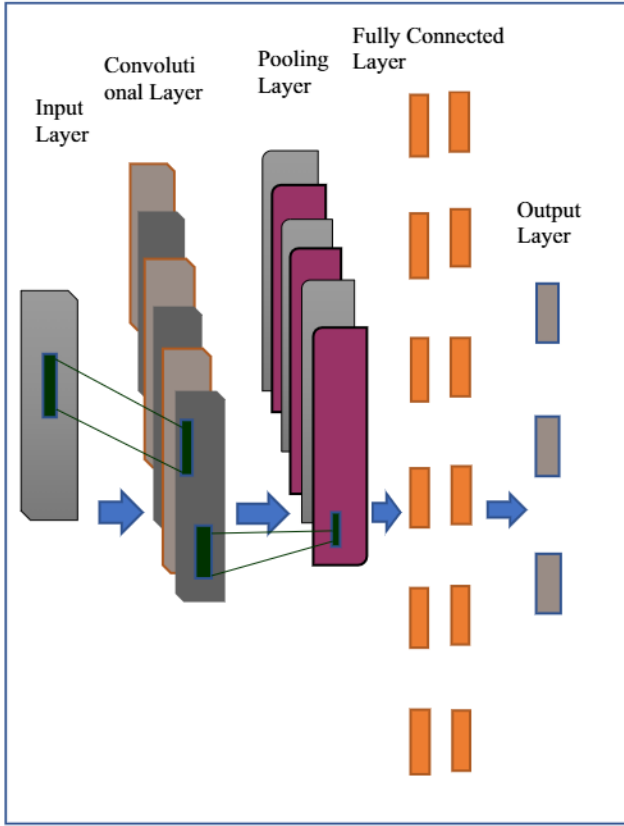
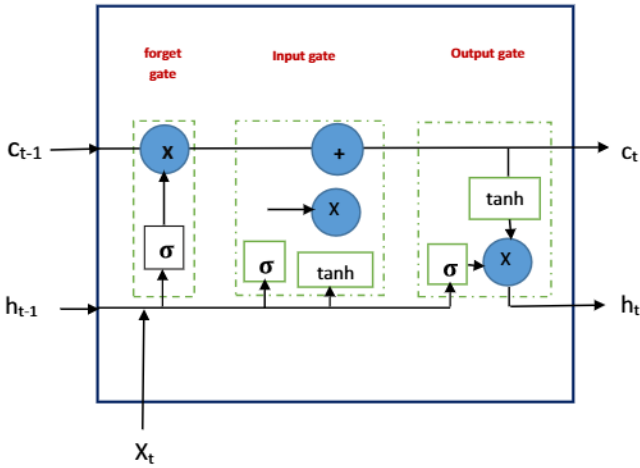


Figure 3
Convolution neural network architecture



output h_i is computed by making use of the most recent embedding x_i in addition to the most recent cell state c_{i-1} and the output h_{i-1} . Depending on the current status of the cell, the data can either be added or erased.

$$f_t = \sigma(w_f \cdot [y_{t-1}, x_t] + b) \quad (1)$$

$$i_t = \sigma(w_i \cdot [y_{t-1}, x_t] + b) \quad (2)$$

Two of the components that go into the construction of a CNN are known as the convolution and pooling layers [30]. Image processing is just one of its many strong suits, and it also has the ability to recognize interdependencies. The data that is input into the system is processed by the convolution layer in order to extract features from the data. Convolutional operations on the embedding matrix can be carried out with the assistance of this approach. It is possible to recover the word embedding vectors that were generated by the many different ways of word embedding from their storage position in the embedding matrix. Following the convolution layer in a CNN comes the pooling layer, which is responsible for a variety of tasks, including the decrease of dimensionality and the selection of features through pooling. Each of the three potential strategies for pooling resources—namely, maximum pooling, minimum pooling, and average pooling—is an option that can be implemented.

After the collection of the information, those features are then processed by a neural network that has all of its connections complete. Utilizing activation functions is required in order to produce an output. The convolutional neural network is illustrated here in Figure 3. These results were achieved by utilizing two layers of convolution and then pooling the averages together. A convolutional neural network, also referred to as a CNN, is composed of several layers including the convolutional layer, pooling layer, activation layer, and connected convolutional layer. A deep CNN is comprised of a number of convolutional layers that are coupled to one another. The primary function of the convolutional layer, as may be deduced from its own name, is to filter the input data. Only a small portion of the original image's pixels are processed by the filter that is located in the convolutional layer (say, 3×3).

Researchers at the University of Michigan have created and given the name AdaBoost to an innovative gradient-boosting technique for binary classification. Following the construction of the initial tree decision tree, the accuracy of the tree is evaluated in comparison to the training set. In this section, the usefulness of the tree decision tree will be evaluated. The objective of this method is to develop a unified, comprehensive categorization plan by combining various classification techniques. The training data serve as the foundation for building the initial model, which is subsequently refined through the creation of additional models to address any deficiencies in the original model. Upon successful prediction of all data in the training set, the model creation process is deemed complete. If not, it will persist until the maximum number of potential models is reached. A multitude of categorization models were merged into one in order to produce the most effective model possible. The AdaBoost sensor is a favorite among a lot of people when it comes to sensors that are able to detect pedestrians. Prior to establishing the feature values, the images are divided into rectangular sections through the process of cropping. Drivers will have an easier time recognizing people if their windows are marked in any order they choose. The identical strategies are utilized, with the exception that the windows in the illustrated image are selected in a different order this time around. Aside from that, nothing has changed from before. Windows are judged to be pedestrian if they are not rejected by any of the models [31]. It is possible to keep repeating this approach until a hierarchy of classification standards has been established.

3.2. Dataset description

The information has been gathered from two distinct sources in order to distinguish between legitimate and fraudulent news. The dataset of fake and authentic news is sourced from Kaggle¹. There are over 40,000 articles in the collection overall, including

¹<https://www.kaggle.com/>

both legitimate and fraudulent news. The real news and fake news are divided into two distinct datasets, each containing roughly 20,000 articles. On the other hand, Lara-Navarra et al. [7] identify another pretrained dataset called the glove Twitter data.

3.3. Data visualization and preprocessing

The two classes that make up the dataset are designated as a true category and a false category, respectively. By presenting data in a visual context, such as map or graph, data visualization aids in our understanding of what relative data implies. This makes the data easier for the human mind to understand, making it simpler to find trends, patterns, and outliers in vast datasets. The dataset is divided into two groups: original news and fake news. Class “1” represents the first category, which is factual news, while class “0” represents the second category, which is fake news.

In order to increase efficiency, data preprocessing is an essential step that entails modifying data before it is executed. It entails data transformation and cleaning, as will be seen in the part that follows. Figure 4 illustrates the number of articles in relation to the associated class labels, fake and real, respectively, to begin the preprocessing of the data. It may be observed that the two datasets don’t differ all that much. The dataset is obviously stable. The fake news category is represented by the “0” class (blue bar) in the picture, whereas the true news category is represented by the “1” class (orange bar). Since the contents of the topic section differ across the two categories, only the primary text needs to be processed; the matching subject, title, and date of each news item can be removed.

Figure 5 provides an explanation of the several topics that combine to form the dataset. The count of each subject indicates how news has diffused throughout society. Whereas the blue bar indicates phony or unreliable news, the orange bar indicates reliable and accurate news. Global and political news are among the subjects covered by actual news. On the other hand, fake news is expressed in the fields of politics, news, the left, US news, and news from the Middle East.

3.4. Deep learning LSTM model: A neural network

We have employed neural networks to construct our model. A sequential model is utilized for a simple stack of layers

Figure 4
Articles in alignment with fabricated and authentic categories

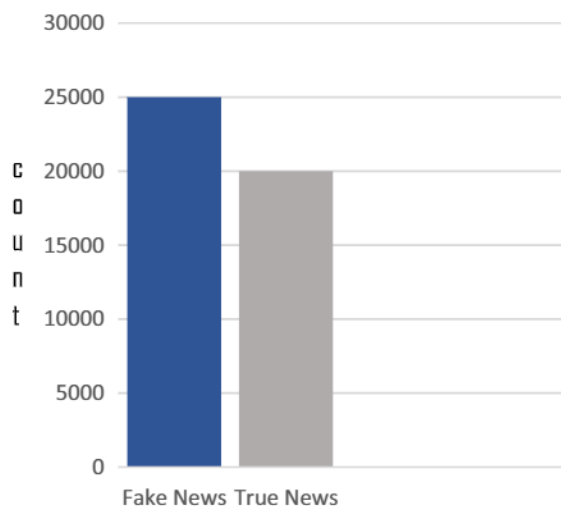
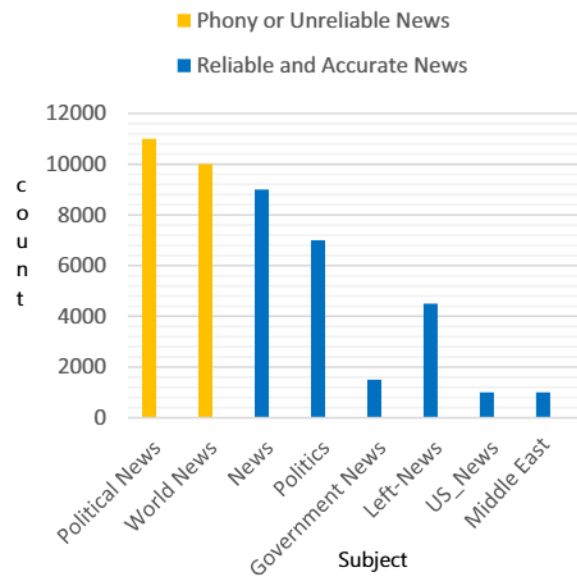


Figure 5
Current trends in various categories of fabricated information and authentic news



consisting of one input tensor and one output tensor for each layer. The input shape that the model should expect must be known. Therefore, in a sequential model, the input shape information needs to be sent to the first layer. The goal of sequence classification, a predictive modeling issue, is to forecast a category given a series of inputs that span time or space. The challenge is compounded by the potential variations in sequence lengths, a diverse set of input symbols, and the requirement for the model to grasp the intricate connections or patterns among symbols within the input sequence over an extended period. In order to identify false information, a long short-term memory model has been applied.

An LSTM network is characterized by its utilization of LSTM cell blocks instead of traditional neural network layers. These cells consist of three parts, the input gate, neglect gate, and output gate. The new sequence value x_t is appended to the previous output from column $ht-1$ on the left-hand side, as shown in Figure 2. Using a tanh layer is the initial step in dampening this combined input.

- 1) During the second stage, the input is then channeled through an input gate. The input of an input gate, consisting of a layer of sigmoid-activated nodes, is then scaled by a compressed input. These input gate sigmoids possess the capability to eliminate any unnecessary parts of the input vector.
- 2) Since a sigmoid function yields values between 0 and 1, it is possible to “switch off” specific input values by training the weights that connect the input to these nodes to create values that are close to zero.
- 3) The data flow through this cell continues with the forget gate loop. In LSTM cells, st is a forget gate variable.
- 4) To create an effective layer of recurrence, the input data is supplemented by $st-1$. Gradient disappearance is less likely when an addition operation is used instead of a multiplication procedure.
- 5) A forget gate, on the other hand, controls this repeated loop. It works similarly to an input gate but helps the network figure out which state variables should be “kept” or “lost.”
- 6) Lastly, an output gate controls the output of the tanh squashing function, which is located at the output layer. Which values from

the cell ht are genuinely permitted as an output are determined by this gate.

The input is compressed between -1 and 1 using a tanh activation function. Here's how this can be shown:

$$g = \tanh(b^g + x_t U^g + h_{t-1} V^g) \quad (3)$$

where the input and prior cell output weights are represented by the symbols V_g and b_g , respectively, and b_g is the input bias.

The forgot gate is represented by:

$$f = \sigma(b^f + x_t U^f + h_{t-1} V^f) \quad (4)$$

The element-wise product of the previous state and the forget gate yields $st-1$.

Output gates, however, can be expressed as:

$$o = \sigma(b^o + x_t U^o + h_{t-1} V^o) \quad (5)$$

As seen in Figure 2, the network's final output will be ht . To improve our system, we have employed the stacked LSTM model. The return sequence is set to true when utilizing the stacked version of the LSTM algorithm. If the return sequence is enabled, the LSTM layer that follows utilizes the output of each neuron's hidden state as an input. A complicated LSTM model with many LSTM and dense layers is needed to categorize a specific item of content as authentic news or fake news.

- 1) Each word is represented by 32 length vectors in the first layer, which is called the embedding layer.
- 2) The LSTM layer comprises the following two layers, with 128 and 64 memory units, respectively.
- 3) Two Dense output layers are present. The reLu activation function and 32 memory units make up the first Dense layer.
- 4) The output layer, which has a sigmoid activation function and a single neuron, is the next dense layer.

A regular layer of neurons is all that constitutes a dense layer in a neural network. Tight links exist because every neuron in the layer above it receives information from every other neuron in the layer below it. This layer is composed of the activations from the previous layer, a bias vector b , and a weight matrix W . In their recommended networks, Tan et al. [32] have typically selected one or two thick layers to prevent over-fitting.

The most popular activation function for CNN neurons' outputs is the Rectifier Unit (ReLU) [33]. The ReLU function has one main advantage over other activation functions: it does not excite every neuron at once. ReLU is a post-convolution nonlinear activation function, similar to sigmoid or tanh. ReLU is represented by:

$$\sigma = \max(0, z) \quad (6)$$

The probability that the word j will occur in relation to the word i by presuming that Q_{ij} is sigmoid. The global objective function that was inferred can be stated as follows:

$$J = - \sum_{i \in corpus, j \in context(i)} \text{Log} Q_{ij} \quad (7)$$

Since the classification problem is fake news, we employ the Dense output layer to predict two classes: actual news and fake news, which are assigned the numbers 0 and 1, respectively. The loss function, metrics, and optimizer are intended for the whole model compilation.

Ten training epochs were used to train the model. Additionally, the Binary Cross-Entropy loss function has been put into practice, and the weights are upgraded using an Adam optimizer. The choice for the learning rate is 0.01. There is a batch size of 256. In contrast, 100 is the embedding size. We have lowered the batch size in order to improve our model's accuracy.

A study was conducted using a random sampling method. The sample included 22 first-year in-service postgraduate science teachers from one of the Colleges of Education in Bhutan. Out of these teachers, 13 (59%) were male. Therefore, it is essential to select a model's hyperparameters in a manner that ensures training is efficient regarding both time and accuracy of fit. The primary distinction is found in the internal processes of the LSTM network's cells. Tables 1 and 2 outline all relevant hyperparameters for an LSTM model that are essential for enhancing performance, along with the recommended values considered best practice.

Table 1
LSTM layered architecture

Layer (type)	Output size	Param Number
Embedding_1 (embedding)	300 x 100	1,000,000
Lstm_1 (LSTM)	300 x 128	117,248
Lstm_2 (LSTM)	64	49,408
Dense_1 (Dense)	32	2080
Dense_2 (Dense)	1	33

Table 2
Hyperparameters for proposed model

Hyperparameters	Value
Layer for embedding	1
LSTM layer	2
Layer with high concentration	2
Loss Function	Binary cross entropy
Function for activation	ReLU
Optimizer	Adam
Learning rate	0.01
Epoch count	10
Size of embedding	100
Group quantity	256

4. Result and Discussion

This set of data is made up of information from Twitter, a social networking website. It is a pretrained dataset and it consists of four attributes: title, main text, subject, and date. Pretrained word vectors are available in this dataset. For experimental analysis, 20000 features are used. 16000 features are used for testing, 2000 for testing and remaining 2000 features for the validation purpose. Vectorization consists of word and its frequency count. The first step is to preprocess the raw data that will be used. The three most important parts of data preprocessing are getting rid of stop words, figuring out where words come from, and turning words into tokens. The NTLK

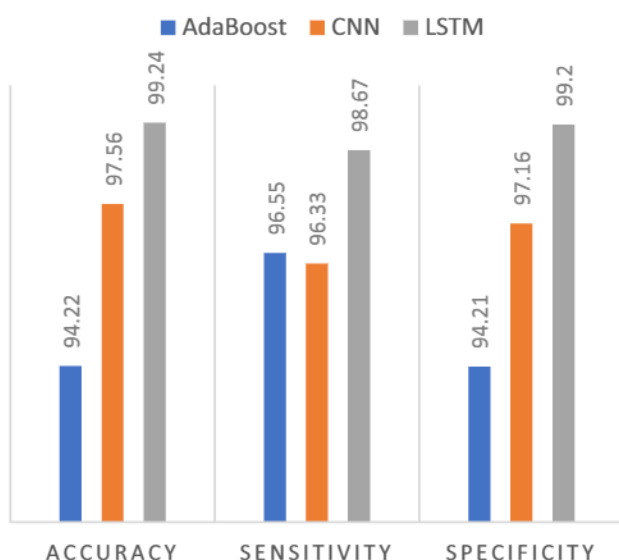
library is used to get rid of these “stop words.” Porter’s Algorithm is used to find the stems of words. Tokenization is achievable with the N-gram model’s assistance. The model was constructed using the CNN, AdaBoost, and LSTM algorithms. Figure 6 and Table 3 present the results. LSTM has achieved an accuracy of 99.24% for fake news detection. Accuracy of LSTM is 1.68% higher than accuracy of CNN and 5.02% higher than accuracy of AdaBoost algorithm. Specificity of LSTM is 99.2%. Specificity of LSTM is 2.04% more than CNN and 4.99% more than AdaBoost technique. LSTM’s sensitivity is 98.67%. LSTM has outclassed CNN and AdaBoost for the identification of fake news and rumors.

In conclusion, there exists a strong correlation between Western rhetoric and English writing. Rhetoric encompasses more than just the mere use of persuasive language. Nevertheless, incorporating figures of speech can greatly enhance one’s writing abilities in English. The art of persuasion, known as rhetoric, encompasses various rhetorical devices that English writers should focus on in order to enhance their understanding of rhetoric and the concepts it entails. By incorporating a rhetorical perspective into their writing, writers can improve their proficiency in utilizing devices like contrast and exaggeration.

Table 3
Accuracy, specificity, and sensitivity comparison of different classifiers

Parameter/Algorithm name	AdaBoost (%)	CNN (%)	LSTM (%)
Accuracy	94.22	96.55	94.21
Sensitivity	97.56	96.33	97.16
Specificity	99.24	98.67	99.2

Figure 6
Result comparison of classifiers for fake news detection in social media data



5. Conclusions and Future Work

This article outlines one technique for identifying fake news that makes use of deep learning. A collection of data is an essential component of any methodology. This dataset contains information

that was obtained from the social networking website known as Twitter. In the beginning, the raw data that are going to be used will be preprocessed. The three most important aspects of data preprocessing are known as stop word removal, stemming, and tokenization. In order to eliminate stop words, the NTLK library is utilized. Porter’s Algorithm is used to complete the stemming process. The N-gram model is used to help with tokenization. The CNN, AdaBoost, and LSTM algorithms were utilized in the construction of the model. In terms of accuracy, specificity, and sensitivity, the results have demonstrated that LSTM performs significantly better than both CNN and AdaBoost. LSTM has achieved an accuracy of 99.24% for fake news detection. Specificity of LSTM is 99.2% and sensitivity is 98.67%. LSTM has outclassed CNN and AdaBoost for the identification of fake news and rumors. The future path is to expand and improvise on the current work in order to automate the process of creating an automated system for e-commerce websites, where the identification of false news has taken on equal importance.

Funding Support

This study is supported via funding from Prince Sattam bin Abdulaziz University project number (PSAU/2024/R/1445).

Ethical Statement

This study does not contain any studies with human or animal subjects performed by any of the authors.

Conflicts of Interest

The authors declare that they have no conflicts of interest to this work.

Data Availability Statement

The Twitter Sentiment Analysis datasets that support the findings of this study are openly available at <https://www.kaggle.com/competitions/tweet-sentiment-extraction/discussion/142142>. The glove Twitter data that support the findings of this study are openly available at <https://www.kaggle.com/datasets/bertcarremans/glovetwitter27b100d.txt>.

Author Contribution Statement

Abu Sarwar Zamani: Conceptualization, Writing – original draft, Project administration, Funding acquisition. **Aisha Hassan Abdalla Hashim:** Writing – review & editing, Supervision. **Sara Saadeldeen Ibrahim Mohamed:** Investigation, Visualization. **Nasre Alam:** Software, Writing – review & editing, Data curation.

References

- [1] t'Serstevens, F., Piccillo, G., & Grigoriev, A. (2022). Fake news zealots: Effect of perception of news on online sharing behavior. *Frontiers in Psychology*, 13, 859534. <https://doi.org/10.3389/fpsyg.2022.859534>
- [2] Igwebuike, E. E., & Chimunya, L. (2021). Legitimizing falsehood in social media: A discourse analysis of political fake news. *Discourse & Communication*, 15(1), 42–58. <https://doi.org/10.1177/1750481320961659>

- [3] Bertoni, V. B., Saurin, T. A., Fogliatto, F. S., Falegnami, A., & Patriarca, R. (2021). Monitor, anticipate, respond, and learn: Developing and interpreting a multilayer social network of resilience abilities. *Safety Science*, 136, 105148. <https://doi.org/10.1016/j.ssci.2020.105148>
- [4] Sheikhi, S. (2020). An efficient method for detection of fake accounts on the Instagram platform. *Revue d'Intelligence Artificielle*, 34(4), 429–436. <https://doi.org/10.18280/ria.340407>
- [5] Raghuvanshi, A., Singh, U. K., Sajja, G. S., Pallathadka, H., Asenso, E., Kamal, M., ..., & Phasinam, K. (2022). Intrusion detection using machine learning for risk mitigation in IoT-enabled smart irrigation in smart farming. *Journal of Food Quality*, 2022(1), 3955514. <https://doi.org/10.1155/2022/3955514>
- [6] Veluri, R. K., Patra, I., Naved, M., Prasad, V. V., Arcinas, M. M., Beram, S. M., & Raghuvanshi, A. (2022). Learning analytics using deep learning techniques for efficiently managing educational institutes. *Materials Today: Proceedings*, 51, 2317–2320. <https://doi.org/10.1016/j.matpr.2021.11.416>
- [7] Lara-Navarra, P., Falciani, H., Sánchez-Pérez, E. A., & Ferrer-Sapena, A. (2020). Information management in healthcare and environment: Towards an automatic system for fake news detection. *International Journal of Environmental Research and Public Health*, 17(3), 1066. <https://doi.org/10.3390/ijerph17031066>
- [8] Kaliyar, R. K., Goswami, A., & Narang, P. (2021). DeepFakE: Improving fake news detection using tensor decomposition-based deep neural network. *The Journal of Supercomputing*, 77(2), 1015–1037. <https://doi.org/10.1007/s11227-020-03294-y>
- [9] Jasti, V. D. P., Zamani, A. S., Arumugam, K., Naved, M., Pallathadka, H., Sammy, F., ..., & Kaliyaperumal, K. (2022). Computational technique based on machine learning and image processing for medical image analysis of breast cancer diagnosis. *Security and Communication Networks*, 2022(1), 1918379. <https://doi.org/10.1155/2022/1918379>
- [10] Lahby, M., Aqil, S., Yafooz, W. M. S., & Abakarim, Y. (2022). Online fake news detection using machine learning techniques: A systematic mapping study. In M. Lahby, A. S. K. Pathan, Y. Maleh, & W. M. S. Yafooz (Eds.), *Combating fake news with computational intelligence techniques* (pp. 3–37). Springer. https://doi.org/10.1007/978-3-030-90087-8_1
- [11] Santos de Oliveira, L., & Vaz de Melo, P. O. (2022). Large-scale and long-term characterization of political communications on social media. In 2022: *Companion Proceedings of the 28th Brazilian Symposium on Multimedia and the Web*, 31–34. https://doi.org/10.5753/webmedia_estendido.2022.225803
- [12] Yang, S., Shu, K., Wang, S., Gu, R., Wu, F., & Liu, H. (2019). Unsupervised fake news detection on social media: A generative approach. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 33(1), 5644–5651. <https://doi.org/10.1609/aaai.v33i01.33015644>
- [13] Barhate, S., Mangla, R., Panjwani, D., Gatkai, S., & Kazi, F. (2020). Twitter bot detection and their influence in hashtag manipulation. In 2020 *IEEE 17th India Council International Conference*, 1–7. <https://doi.org/10.1109/INDICON49873.2020.9342152>
- [14] Pennycook, G., & Rand, D. G. (2021). The psychology of fake news. *Trends in Cognitive Sciences*, 25(5), 388–402. <https://doi.org/10.1016/j.tics.2021.02.007>
- [15] Jin, B., & Xu, X. (2024). Machine learning predictions of regional steel price indices for east China. *Ironmaking & Steelmaking*. Advance online publication. <https://doi.org/10.1177/03019233241254891>
- [16] Jin, B., & Xu, X. (2024). Palladium price predictions via machine learning. *Materials Circular Economy*, 6(1), 32. <https://doi.org/10.1007/s42824-024-00123-y>
- [17] Kaliyar, R. K., Goswami, A., & Narang, P. (2021). FakeBERT: Fake news detection in social media with a BERT-based deep learning approach. *Multimedia Tools and Applications*, 80(8), 11765–11788. <https://doi.org/10.1007/s11042-020-10183-2>
- [18] Bahad, P., Saxena, P., & Kamal, R. (2019). Fake news detection using Bi-directional LSTM-recurrent neural network. *Procedia Computer Science*, 165, 74–82. <https://doi.org/10.1016/j.procs.2020.01.072>
- [19] Jin, B., & Xu, X. (2024). Price forecasting through neural networks for crude oil, heating oil, and natural gas. *Measurement: Energy*, 1, 100001. <https://doi.org/10.1016/j.meaeene.2024.100001>
- [20] Ruchansky, N., Seo, S., & Liu, Y. (2017). CSI: A hybrid deep model for fake news detection. In *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*, 797–806. <https://doi.org/10.1145/3132847.3132877>
- [21] Davoudi, M., Moosavi, M. R., & Sadreddini, M. H. (2022). DSS: A hybrid deep model for fake news detection using propagation tree and stance network. *Expert Systems with Applications*, 198, 116635. <https://doi.org/10.1016/j.eswa.2022.116635>
- [22] Hsu, C. C., Zhuang, Y. X., & Lee, C. Y. (2020). Deep fake image detection based on pairwise learning. *Applied Sciences*, 10(1), 370. <https://doi.org/10.3390/AP10010370>
- [23] Jammal, A. A., Thompson, A. C., Mariottoni, E. B., Berchuck, S. I., Urata, C. N., Estrela, T., ..., & Medeiros, F. A. (2020). Human versus machine: Comparing a deep learning algorithm to human gradings for detecting glaucoma on fundus photographs. *American Journal of Ophthalmology*, 211, 123–131. <https://doi.org/10.1016/j.ajo.2019.11.006>
- [24] Tan, Y. H., & Chan, C. S. (2019). Phrase-based image caption generator with hierarchical LSTM network. *Neurocomputing*, 333, 86–100. <https://doi.org/10.1016/j.neucom.2018.12.026>
- [25] Kumar, A., & Verma, S. (2021). CapGen: A neural image caption generator with speech synthesis. In *Data Analytics and Management: Proceedings of ICDAM*, 605–616. https://doi.org/10.1007/978-981-15-8335-3_46
- [26] Zhou, P., Han, X., Morariu, V. I., & Davis, L. S. (2017). Two-stream neural networks for tampered face detection. In 2017 *IEEE Conference on Computer Vision and Pattern Recognition Workshops*, 1831–1839. <https://doi.org/10.1109/CVPRW.2017.229>
- [27] Asim, M. N., Ghani Khan, M. U., Malik, M. I., Razzaque, K., Dengel, A., & Ahmed, S. (2019). Two stream deep network for document image classification. In 2019 *International Conference on Document Analysis and Recognition*, 1410–1416. <https://doi.org/10.1109/ICDAR.2019.00227>
- [28] Bird, S., Klein, E., & Loper, E. (2009). *Natural language processing with Python: Analyzing text with the natural language toolkit*. USA: O'Reilly Media.
- [29] Huan, J. L., Sekh, A. A., Quek, C., & Prasad, D. K. (2022). Emotionally charged text classification with deep learning and sentiment semantic. *Neural Computing and Applications*, 34(3), 2341–2351. <https://doi.org/10.1007/s00521-021-06542-1>
- [30] Umer, M., Imtiaz, Z., Ahmad, M., Nappi, M., Medaglia, C., Choi, G. S., & Mehmood, A. (2023). Impact of convolutional

- neural network and FastText embedding on text classification. *Multimedia Tools and Applications*, 82(4), 5569–5585. <https://doi.org/10.1007/s11042-022-13459-x>
- [31] Ibrahim, Y., Okafor, E., Yahaya, B., Yusuf, S. M., Abubakar, Z. M., & Bagaye, U. Y. (2021). Comparative study of ensemble learning techniques for text classification. In *2021 1st International Conference on Multidisciplinary Engineering and Applied Science*, 1–5. <https://doi.org/10.1109/ICMEAS52683.2021.9692306>
- [32] Tan, M., Chen, B., Pang, R., Vasudevan, V., Sandler, M., Howard, A., & Le, Q. V. (2019). MnasNet: Platform-aware neural architecture search for mobile. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2820–2828. <https://doi.org/10.1109/CVPR.2019.00293>
- [33] Li, Y., & Yuan, Y. (2017). Convergence analysis of two-layer neural networks with ReLU activation. In *Proceedings of the 31st International Conference on Neural Information Processing Systems*, 597–607. <https://dl.acm.org/doi/abs/10.5555/3294771.3294828>

How to Cite: Zamani, A. S., Hashim, A. H. A., Mohamed, S. S. I., & Alam, N. (2025). Optimized Deep Learning Techniques to Identify Rumors and Fake News in Online Social Networks. *Journal of Computational and Cognitive Engineering*, 4(2), 142–150. <https://doi.org/10.47852/bonviewJCCES2023348>