

A Hybrid Deep Neural Network for Arabic Fake News Classification based on Temporal CNN and BiLSTM with Attention

Azhar A. Hadi^{1*}, Abbas F. Hommadi¹, Hussein A. Ismael²

¹.Electrical Engineering Department, Engineering College, University of Babylon, Babylon, Iraq.

².Computer Center, University of Babylon, Babylon, Iraq.

³.Information Technology College, University of Babylon, Babylon, Iraq.

Received: 18 Feb 2025/ Revised: 04 Jun 2026/ Accepted: 25 Jul 2026

Abstract

The proliferation of fake news on social media poses a significant threat to societal harmony, especially in languages like Arabic, which is distinguished by its complexity. Thus, Artificial intelligence techniques are necessary to mitigate the impact of misinformation. Many researchers have conducted this challenge by introducing machine learning and deep learning approaches. This study presents a new hybrid deep learning network for enhancing the classification of Arabic fake news. This network combines three mechanisms: Temporal Convolutional Networks, Bidirectional Long Short-Term Memory, and Attention Mechanism. This mixture model produces a strong feature extraction process with temporal awareness. Moreover, the attention layer enables the model to focus only on relevant features and ignore irrelevant ones to concentrate on salient features. Also, several stages of pre-processing Arabic text, representing words using a pre-trained word embedding model (Glove, Ara2Vec), and extracting features through advanced layers (BiLSTM, TCN, and Attention). Extensive experimentation on large-scale and well-known Arabic fake news datasets (AFND and AraNews) demonstrates the efficacy of the proposed model. The hybrid model shows superior performance compared to baseline and previous state-of-the-art studies by achieving an accuracy of 88% and 95% on the AFND and AraNews datasets, respectively. Our results show that the suggested model can be used as a powerful technique to restrain the widespread online misinformation in the Arabic language.

Keywords: Arabic Fake News; Temporal Convolution Network; BiLSTM; Attention Mechanism; Deep Learning.

1- Introduction

Nowadays, with the internet age, people strongly rely on online sources to get fed with the news and events regarding their interests. As the number of internet users and online news sources has increased dramatically over the last two decades, users can easily access and engage in events locally and globally. Moreover, individuals can share, re-post, or even create such events on the internet with ease. [1]. All these capabilities make people a gear part of the online information era. To this end, all kinds of information can be exposed to the citizens of the internet, including false information, with no obvious hint of which should be trusted. Thus, building a trustworthy environment to exchange online information among people and organizations interchangeably become a need in the modern world of the online age [2]. Therefore, this

introduction will describe what false information is and the damage it could cause.

According to [3], [4] Fake news is described as nonexistent or skewed facts that can spread widely on an online platform to change facts and people's thoughts about something. Moreover, P. Bahad et al. point out that misinformation is designed to be forged information with a deceiving legal appearance [5]. Both definitions indicate the numerous potential harmful effects that misinformation can cause, including political, financial, and social damage. This includes planting inaccurate news on the internet to manipulate people's opinions. Trump's win in 2016 during the US presidential election demonstrates a good example of how fake news could impact society to gain political power [6]. According to [7] Social network data is used to mine people's political opinions and sentiments, and then use their political favour to propagate inaccurate information to get some political gain. In addition, fake or untrusted news can mislead others intentionally and

✉ Azhar A. Hadi
engbabil@gmail.com

unintentionally. Thus, it helps spread misunderstanding about something. Another example is during COVID-19, false health information was spread, which can cause harm to public health, mentally and physically [8].

Detection of online false information is necessary for two reasons. First, the rapid distribution of information through online platforms such as Meta (Facebook) and X (Twitter) makes it difficult to stop or even mimic the propagation. Second, it can be challenging to distinguish between false and true information [5], [9], [10]. One way to tackle this challenge is by manually categorizing online information using experts and journalists. This approach is an unrealistic solution due to the vast amount of online information, which makes it costly and time-consuming. Another approach is by automatically classifying online data based on Artificial Intelligence (AI). This approach becomes a more handy and feasible solution due to the fast and reliable outcome. Thus, many researchers have developed and introduced deep learning approaches to conquer this challenge. For example, authors in [11] employ advance transformer technique called MARBERTv2 to identify fake news in multi-dialectal Arabic tweets. Another research [12] merges traditional 2D-CNN with a transformer-based model, namely AraBERT, to achieve this task on a variety of Arabic fake news datasets

The Arabic language occupies the 4th rank among world internet users. However, there is a significant insufficiency of research dedicated to the Arabic language compared to other languages, like English [13]. This research seeks to reduce the gap by making the following contributions:

1. Introduce a hybrid deep learning model that integrates a Temporal Convolutional Network (TCN) for feature extraction and Bidirectional Long Short-Term Memory (BiLSTM) with an Attention mechanism for capturing contextual and sequential dependencies to improve the classification task.
2. Utilize different word embeddings, GloVe and Ara2Vec, to deliver rich and various semantic representations of Arabic text.
3. Conduct extensive experiments and comparative analysis on two popular datasets (AFND and AraNews) to demonstrate the superiority of our study.

The rest of this paper is outlined as follows: Section II summarizes the related works, Section III introduces DL techniques, Section IV demonstrates the proposed methodology, Section V presents a description of the experimental studies and discusses the results, and Section VI presents the paper's conclusion.

2- Related Work

Due to the importance of detecting online misleading information, many studies have been done in this regard.

These studies tackle the problem based on ML and DL approaches. The following presents a summary overview of such studies.

The authors in [3] have presented an ML (Support Vector Classification (SVC), linear SVC, KMeans, and variation Bayesian estimation of a Gaussian mixture (Bayesian Gaussian mixture)) and DL (Capsule networks, deep double Q-learning (DDQN), Deep pyramid CNN (DPCNN), and Information Distilled LSTM (ID-LSTM)) models to classify Arabic fake news using the AFND dataset. According to their results, DL methods were found to be superior to the ML methods based on accuracy metrics. The capsule network achieved 79% accuracy for a binary classifier.

The study in [14] presented a comprehensive machine-learning approach for Arabic fake news detection. This notable approach is based on joining content-related and user-related features with sentiment analysis. Random Forest, Decision Tree, AdaBoost, and Logistic Regression algorithms are used in this study to find effective ML algorithms. The study collects and annotates 10K Arabic tweets as fake or real, where the best result showed 76% accuracy by applying the Logistic Regression model. Although the outstanding feature engineering approach is used, the dataset is considered small. Thus, their work cannot be generalized to Arabic content due to the language variations.

The researchers in [15] Investigate machine learning approaches to classify Arabic online false information. This research compared three different ML methods: Naive Bayes, Logistic Regression, and Random Forest, where the AraNews dataset [16] is used. The TF-IDF is used to extract features from Arabic text. The study highlights that Random Forest outperforms the other two methods by scoring 86% accuracy due to its ensembling behaviour. The paper [17] addresses the need to identify misleading information during the COVID-19 pandemic in Arabic. While there is overwhelming content published online throughout the pandemic, this research conducts a developed machine learning classifier (Logistic Regression) to tackle the misinformation detection surrounding COVID-19. Moreover, a new benchmark dataset has been launched to serve the research thirst for detecting false Arabic information. One major limitation of this work is trust in only three manual annotators, which may result in potential subjectivity and biased decisions.

The studies in [18], and [19] present a hybrid deep learning approach to tackle the information credibility in the online Arabic context. Their presented models combine the robustness of CNN and LSTM methods to improve the detection accuracy. The findings in [18] were able to present a model with outstanding results by achieving 81.6% accuracy using the AFND dataset [13]. The other study [19] demonstrates that the hybrid model is superior to CNN alone and LSTM alone by achieving 91.4%

accuracy compared to 85.9% and 87.8% for CNN and LSTM using the AraNews dataset [16], which contains Arabic news articles from multiple domains. These studies introduced a significant hybrid DL approach that surpasses traditional ML methods. However, reasoning about other evaluation metrics, such as false positives and false negatives, and their impacts on identifying misinformation will give a more general assessment that is general.

The research [20] utilized an advanced deep learning approach to improve the performance of checking fake news. The attention mechanism with BiLSTM is leveraged in their study to identify important features in the text. Moreover, ANN is used as a classifier to reveal fake information from the corpus. The dataset called AFND [13] was used in the research, which comprised more than 600,000 articles from different domains. The proposed work exceeds other deep learning techniques, such as GRU, LSTM, and CNN techniques, by recording 81% accuracy.

The study in [18] points out the difficulties of uncovering online false information in the Arabic language by inspecting textual and semantic meaning without relying on external features. An ensemble deep-learning model has been suggested in this study. By ensembling CNN and LSTM schemes, the authors were able to present a model with outstanding results by achieving 81.6% accuracy.

The authors in [21] Conduct the challenges associated with Arabic sarcastic unreal news by utilizing language computing and machine learning approaches. They gathered 3,185 fake articles from two ironic websites and 3,710 real articles from BBC-Arabic, CNN-Arabic and Al-Jazeera news to examine the best model for predicting unreal news stories. The research discloses the outperformance of the CNN network with pre-trained word embedding compared to other machine learning techniques, such as XGBoost. Regardless of the findings, the research does not explore diverse data and digging more fake news ironic attributes to boost the performance. Albalawi et al. [22] proposed the integration of textual and visual features to address the problem of detecting fake Arabic social media content. Thus, the detection of fake news does not depend only on textual features, as previous studies suggested, but also on visual information. This study collected and annotated about ~4k tweets from the X platform with text and image content. To extract textual features, the study employed pre-trained BERT, while they used a mix of VGG-19 and ResNet50 to obtain visual features. Their findings highlight that the model based on textual features solely gives a better performance than integrating textual with visual features.

3- Background

3-1- Deep Learning

Deep Learning (DL) is an innovative approach to the Machine Learning (ML) process, which builds on top of Neural Networks to perform learning at different levels of abstraction [23][24]. It has had a positive impact on diverse fields that entail feature identification, including image recognition, drug discovery, object identification, and speech recognition. DL applies the back-propagation method, which involves tweaking or optimizing neural network parameters. DL techniques are found to be superior to traditional ML, which provides scale and very high accuracy, especially with large training datasets. Some of the common DL techniques that are used in various fields are the Convolution Neural Networks (CNN), which can be used for image, video, and speech data. The other DL technique is Recurrent Neural Networks (RNN), which are used for sequential data such as text and speech data [25][26].

3-2- Recurrent Neural Network (RNN)

One of the most promising deep learning models is the Recurrent Neural Network (RNN). It is a good learning model for sequential data processing, such as language processing and speech recognition. Through the internal state of the neural network's recollection of prior inputs, it acquires characteristics for time-series data. On the basis of historical and current data, an RNN can also forecast future information. However, learning stored data over an extended period is challenging in the RNN structure due to the gradient disappearing and exploding. In addition, RNN has weakened in identifying long-term dependencies. Thus, Long Short-Term Memory (LSTM) is proposed to resolve traditional RNN issues [27], [28].

3-3- Long Short-Term Memory (LSTM)

LSTM networks are recurrent neural networks of a higher order intended for processing sequential data. The LSTM scheme was proposed by Hochreiter and Schmidhuber [29], adding memory cells and gating functions. There are different gate designs in LSTM, including the input, forget, and output gates that control the information flow. This controlling mechanism guarantees the preservation of crucial information in long sequences while fading irrelevant information [30]. The principal equations that define the functioning of an LSTM unit are formalised as in Eq.1 to 6:

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (1)$$

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (2)$$

$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, x_t] + b_C) \quad (3)$$

$$C_t = f_t * C_{t-1} + i_t * \tilde{C}_t \quad (4)$$

$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (5)$$

$$h_t = o_t * \tanh(C_t) \quad (6)$$

At each time step t , the LSTM unit processes the input x_t , previous hidden state h_{t-1} , and cell state C_{t-1} via multiple gates to control information flow. To decide which segment of C_{t-1} to ignore, the forget gate f_t is used. To specify the new information to be included in the cell state C_t , the input gate i_t is applied. Candidate cell state $\tilde{C}_t$ is calculated to offer a new cell state. To produce the updated cell state, C_t is used to combine the old and new information. While the output gate chooses what part of C_t to output, the hidden state h_t is generated from C_t [7], [31]. BiLSTM networks are an extension of traditional LSTM networks that encode information from both the past and the future states of a sequence. BiLSTM enhances the capacity of a network to capture context and dependencies. BiLSTM processes the input sequence in two directions. In other words, the network gains context from preceding and succeeding elements of the sequence. The network gains a more inclusive context of the input sequence. At each time step, the output from the two directions gets concatenated to be used in the outputs. BiLSTM have many applications, including sequence labelling tasks, quality classification tasks, and the machine translation task [29], [30].

3-4- Attention Mechanism

The attention mechanism is considered one of the evolutionary architectures in DL, especially for addressing sequential data challenges. It was presented by Bahdanau, Cho, and Bengio [32] to mitigate significant flaws in standard RNN and its enhancements, like LSTM and gated recurrent unit GRU. This mechanism supports the model to emphasize important parts of the sequential data rather than considering the whole sequence equally, which enables the model to prioritize relevant features in the sequence. The emphasizing selectivity empowers the model performance and makes it applicable to different tasks such as text classification, text summarization and speech recognition [32].

The most well-known use of attention is the transformer design, presented by Vaswani et al. in 2017. A transformer significantly improves the efficiency and effectiveness by adapting the parallel processing. This parallelism design allows processing the whole sequence, while the transformers' self-attention mechanism computes attention levels for each part in the input sequence [33].

To enable the transformer model to pay attention to particular parts of the input sequence, a scaled dot-product attention layer is facilitated. The scaled dot-product attention mechanism plays a cornerstone component in transformer architecture, where its job is to weigh the

sequence parts by computing the attention score for each part in the input sequence, as shown in Eq.7.

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}} + \text{mask}\right) V \quad (7)$$

Where Q represent queries to search for relevant data, K represents keys that are examined about the queries, V represents values that are accumulated based on the attention scores, and d_k is the key dimension. This equation can be described as follows. First, it obtains the attention scores by applying the dot product between matrices Q and K . This resulting matrix of attention scores is scaled by d_k to avoid extremely large values that can affect the gradient flow while training. To dismiss particular parts, like padding tokens, a mask can be applied optionally. Second, the softmax function is utilized to result in normalized attention weights, which indicate the importance of each part in the sequence. Finally, the weighted sum is computed by multiplying normalized attention weights by the value matrix. This step accumulates the value matrix information relevant to the attention weight that each part receives [34].

Moreover, many state-of-the-art architectures, such as BERT (Bidirectional Encoder Representations from Transformers) and GPT (Generative Pre-trained Transformer), have been derived from the attention mechanism. Acting attention mechanism as corn stone in such cutting-edge techniques revolutionizes sequential data processing and results in more accurate and context-aware decisions.

3-5- Temporal Convolutional Networks TCN

TCN was first introduced by [35] for video action segmentation. In traditional techniques, the combination of CNN and RNN is used to extract features, which leads to missing long-term relationships and is more complicated to learn. On the other hand, TCN integrates feature extraction methods by utilising 1D convolution to acquire low, medium, and high dependencies with sequential data. TCN utilized four types of convolutions. First, 1D Convolutions are used to obtain temporal relationships in sequential data. Second, Causal Convolutions are used to guarantee there is no information leak from the future to the current state. Third, Dilated Convolutions helps to acquire long-term relationships. Fourth, Residual Connections are used to avoid vanishing gradients in more complex networks [35].

3-6- Dataset Description

This study uses two well-known datasets for detecting Arabic misinformation: the Arabic Fake News Dataset (AFND) [13] and Arabic News (AraNews) [16] Dataset, as

shown in Fig.1. Both benchmarks are available on Kaggle¹ and GitHub² websites.

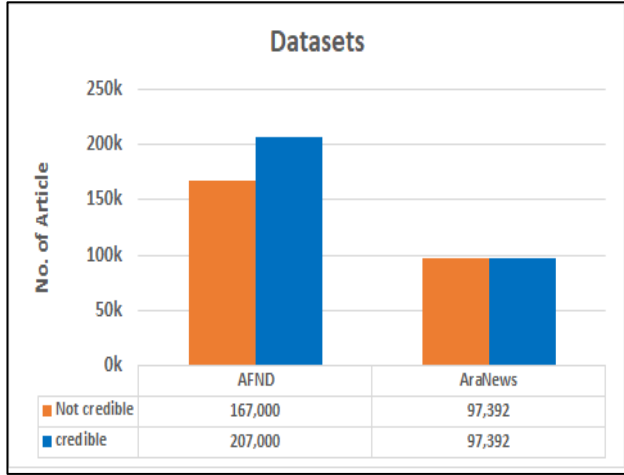


Fig.1: Datasets Summary

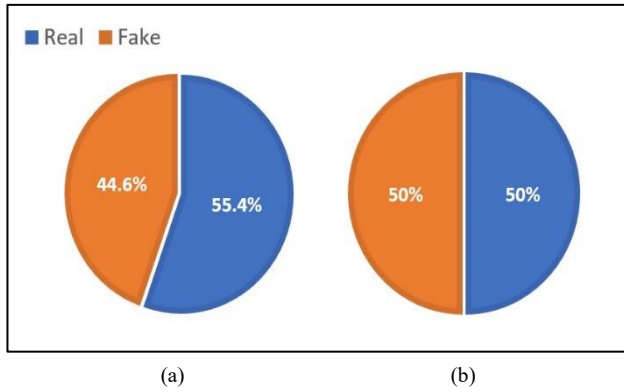


Fig.2: Datasets Distribution (a) AFND (b) AraNews

First, the AFND [13] The dataset consists of about 607K articles. These articles were gathered during 6 month period from 134 Arabic news websites in 19 nations. By using the Arabic fact-checking service Misbar, the articles are labelled as credible, not credible, or undecided. This study considers only credible and not credible articles and ignores the undecided ones. Thus, the task becomes a binary classification where the ratio of each category is 55.4% and 44.6% for fake and real news, respectively. Furthermore, articles with fewer than 10 words are discarded [13]. The AFND distribution can be seen in Fig.2 (a).

Second, the AraNews dataset is a large-scale dataset for Arabic false news. It consists of more than 5 million Arabic articles that are assembled from 50 newspapers.

¹ <https://www.kaggle.com/datasets/murtadhayaseen/arabic-fake-newsdataset-afnd>

² <https://github.com/UBC-NLP/wanlp2020-Arabic-fake-news-detection>

This data covers a wide range of topics, varying from politics to entertainment topics, published in 15 Arabic countries in addition to the United States of America and the United Kingdom [16]. This dataset is perfectly balanced with binary classes, as shown in Fig.2 (b).

There are several challenges associated with these datasets. Including the noisiness, different Arabic dialects, and specific language difficulties. Therefore, several steps in pre-processing data are taken to overcome these challenges, as described in Section IV.

4- Methodology

This part describes the stages of the proposed work. These stages are: (i) Pre-processing data, (ii) Text representation, (iii) Feature extraction, (iv) Classifying text. Fig.3 demonstrates the proposed model.

4-1- Pre-processing Data

Preprocessing text is a vital step in text mining due to its ability to transform raw text into cleaned and structured data that can be used for further analysis [36]. It contains several techniques such as cleaning, normalizing, and tokenizing input text. First, cleaning and normalizing are two of the most important pre-processes in the preparation of data for NLP. Both techniques are used to eliminate noise, error, ambiguity, and unnecessary content from an original text document. The Arabic text is cleaned by eliminating non-Arabic words, stop words, white spaces, and punctuation marks.

Table 1: Arabic Text Preprocessing

Preprocessing steps	Before	After
Removing Punctuation	كيف جاءت الفكرة؟	كيف جاءت الفكرة
Removing Numbers and Special Characters	شعار # نحب الحياة 22	شعار نحب الحياة
Removing Stop Words	شارك في هذه الندوة	شارك الندوة
Tokenization	وذكر البنك المركزي	[وذكر ، البنك ، المركزي]
Normalization	شركتان فرنسية وبريطانية	شركتان فرنسية وبريطانية

In addition, encountered some problems with word embedding because some Arabic characters, like "ة" for example, in this sentence "شركتان فرنسية وبريطانية تعلنان نتائج ايجابي" were processed by normalize the text by using the Tashaphyne library and converting the character "ة" to "ه", the example above became "شركتان فرنسية وبريطانية تعلنان نتائج". The second, tokenising technique is a fundamental step in NLP as it converts the continuous flow of text into discrete pieces to make it easier for further processing by considering the splitting criteria such as commas, semicolons and full stops. Tokenization is a relatively

straightforward process and can be done after cleaning and normalizing, more detailed examples of Arabic text pre-processing are shown in Table 1.

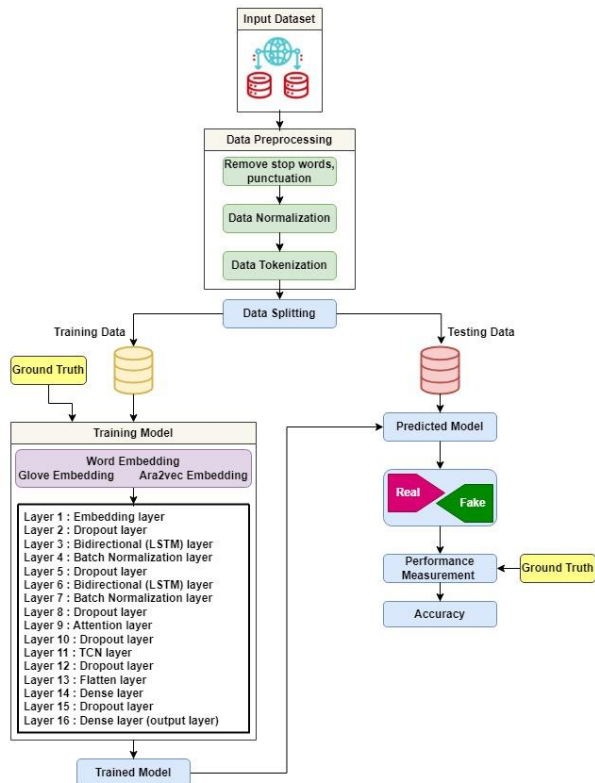


Fig.3: Proposed Model

4-2- Text Representation

To make input text understandable by a computational model, there is a need to convert the text into numerical form. There are different techniques used to achieve this task, including bag-of-words (BoW), term frequency-inverse document frequency (TF-IDF), and more sophisticated techniques such as word embeddings (Word2Vec, GloVe) and contextual embeddings (BERT, GPT [37], [38]). In this work, the input text is represented by word embedding (Word2Vec, GloVe) due to its simplicity compared with contextual embedding, being more representative than traditional BoW and TF-IDF, and the availability of an Arabic version for Word2Vec and GloVe.

4-2-1-Word Embedding

Word Embedding is trained to produce numerical representations for text in which words with a similar meaning have close vectors. Dense vectors are utilized by using some attributes in the sentences to produce a more representative vector than the high, sparse vectors of the

classical TF-IDF method. It is difficult and time-consuming to build an embedding matrix from scratch. Therefore, this work is based on a pre-trained model, which includes Word2Vec and GloVe Embedding.

First, AraVec is a specific version of Word2Vec for the Arabic language. This model is trained on diverse Arabic NLP corpus (Tweets, World Wide Web pages and Wikipedia Arabic articles), Soliman et al. This diversity guarantees wide applicability across different NLP tasks, which vary from text classification to machine translation and sentiment analysis [39], [40]. AraVec2 is used in this work, where it has been applied to 66,900,000 tweets to produce a 300-dimensional embedding vector [39].

Second, GloVe is presented by Stanford scientists as a type of word embedding [41]. It is formed to learn general statistical features over a corpus, as well as the nearby context around a specific text. The main concept of Glove is the co-occurrence matrix of words, which defines the word frequency. This model yields low-dimensional embedding vectors where semantically similar words are assigned to vectors that are near each other. (GloVe v.1.2 glove. Wikipedia.6B) The GloVe-Arabic version has been used in this work.

4-3- Feature Extraction

This stage is the most vital phase to understanding raw text and eventually building an accurate model [42]. However, the Arabic language has hardness because of its plentiful morphology and various dialects [43]. To overcome these challenges, an extensive feature extraction scheme has been proposed by employing the benefit of multiple advanced neural network layers. To effectively capture dependencies among sequence parts from both directions, BiLSTM is utilised for this purpose. This behaviour is highly advantageous in Arabic text. Moreover, an attention layer has been included to ensure the most significant tokens get focused dynamically from the input sequence. Prioritising features is important to realize an Arabic article. To reveal temporal patterns in lengthy Arabic sequences, TCN is utilized. Finally, the integration of multiple neural network schemes, namely BiLSTM, attention, and TCN, has promoted the proposed architecture and discovered crucial patterns to distinguish Arabic fake content.

4-4- Classifying Text

In order to categorise the input article as legitimate or fake, the final block in the proposed work carries out the classification process. This process takes place after multiple hidden layers that are essential to obtain significant features from raw text [44]. In our work, the final block consists of three layers: dense, dropout, and sigmoid layers. To compute the likelihood of a certain

class, a sigmoid function shown in Eq.8 is implemented[45].

$$\sigma(x) = \frac{1}{1 + e^{-x}} \quad (8)$$

Where x is a vector of input features.

5- Experiments and Results

5-1- Experiment Settings

Many experiments were carried out to enhance the effectiveness of the proposed deep learning model and achieve the best results, as all results were conducted on the Google Colab platform. Google Colab is based on Jupyter Notebook, which provides many core Python libraries to implement and develop learning algorithms. It is a web-based work environment, with the following free resources: Disk 107.7 GB, System RAM 12.7 GB, and a T4 GPU. Moreover, the following libraries have been used: Keras Library, Scikit-learn Library, and Genism Library. Two datasets The AFND and AraNews, have been used to evaluate our proposed model.

Table 2: Hyperparameters of the proposed model

Name	Size
Max-Length	200
Embedding Dim.	300
Batch Size	64
Dropout Rate	0.2 and 0.5
Optimizer	Adam
Activation Functions	ReLU and Sigmoid
Epochs	20
Early Stopping	Patience=3 verbose=1, and monitor=validation loss
Regularization	L2 (0.01)

In text representation, two different word embeddings have been studied and compared, Glove and Ara2Vec, separately, to get rich semantic relationships for words. The proposed model architecture includes a two-layer Bidirectional LSTM (256 and 128 units, respectively), one attention layer with a size of 128 of the hidden state, a TCN layer (128 units), and finally a classification process. The sigmoid activation has been applied at the output layer, while the RELU layer (Rectified Linear Unit Activation Function) has been used with the other layers. To mitigate overfitting, dropout, and L2-regularisation are used.

Finally, batch normalisation is used to improve the optimization process efficiency and stability by speeding up convergence [46]. Table 3 summarizes the proposed model layers, number of nodes, and number of parameters. Hyperparameters are essential to achieve a high-

performance model. In our experiments, hyperparameters are chosen as shown in Table 2.

Table 3: Summary of proposed model architecture

Layers	Output	Parameters
Embedding	(None, 200, 300)	15898800
Dropout	(None, 200, 300)	0
Bidirectional-LSTM	(None, 200, 300)	1140736
BatchNormalization	(None, 200, 512)	2048
Dropout	(None, 200, 512)	0
Bidirectional-LSTM	(None, 200, 512)	656384
BatchNormalization	(None, 200, 512)	1024
Dropout	(None, 200, 512)	0
Attention	(None, 200, 256)	456
Dropout	(None, 200, 256)	0
TCN	(None, 200, 128)	476288
Dropout	(None, 200, 128)	0
Flatten	(None, 25600)	0
Dense	(None, 64)	1638464
Dropout	(None, 64)	0
Dense	(None, 1)	65

Furthermore, both datasets (AFND and AraNews) are partitioned into 80% for the training set, 10% for the validation set, and 10% for the test set.

5-2- Results and Discussion

The evaluation metric used to assess the performance of our proposed model is accuracy, as shown in Eq. 9.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (9)$$

Where TP (True Positive) refers to the number of fake articles predicted as fake, TN (True Negative) refers to the number of credible articles predicted as real, FP (False Positive) refers to the number of credible articles predicted as fake, and FN (False Negative) refers to the number of fake articles predicted as real [47].

The experiments are conducted on two well-known datasets (AFND and AraNews), where two-word embeddings (Glove and Ara2Vec) are used to represent words. Moreover, the proposed model is extensively compared against several baseline models (CNN, LSTM, BiLSTM, CNN with LSTM, and CNN with BiLSTM) as well as the state of artwork in previous studies.

Table 4: Result comparison based on AFND Dataset

Models	Embedding	Acc.	Embedding	Acc.
CNN	Arb2vec	78%	Glove	73%
LSTM	Arb2vec	86%	Glove	85%
BiLSTM	Arb2vec	84%	Glove	86%

CNN+LSTM	Arb2vec	80%	Glove	81%
CNN+BiLSTM	Arb2vec	81%	Glove	86%
Previous Study[3]	-	-	-	79%
Previous Study[20]	-	-	Glove	81%
Previous Study[18]	-	-	Glove	81%
Proposed Model	Arb2vec	87%	Glove	88%

Table 5: Result comparison based on the AraNews 16k Dataset

Models	Embedding	Acc.	Embedding	Acc.
CNN	Arb2vec	85%	Glove	91%
LSTM	Arb2vec	86%	Glove	92%
BiLSTM	Arb2vec	87%	Glove	92%
CNN+LSTM	Arb2vec	89%	Glove	93%
CNN+BiLSTM	Arb2vec	88%	Glove	93%
Previous Study[19]	-	-	-	91%
Proposed Model	Arb2vec	89%	Glove	95%

First case, the models are compared based on the AFND dataset, as shown in Table 4. The result demonstrates that the Glove approach is the best word embedding by achieving a higher accuracy compared to the AraVec word embedding. In terms of accuracy, the proposed model records 88%, achieving +9%, +7%, and +7% higher than [3], [18], [20], respectively. This outperformance stems from the complex feature-extraction layers designed into our model. Where feature acquisition plays an important role in the classification process.

The second case is to study the proposed work using the AraNews dataset with 16K articles, as shown in Table 5. This dataset is balanced and compromised randomly. Our proposed model attains an accuracy of 95% with Glove representation. It surpasses the previous study [19], which achieves 91% and the other baseline models.

Finally, 50k articles of the AraNews dataset are conducted and the results are shown in Table 6. Our proposed model attained an accuracy of 92% with Glove and 90% with

Ara2vec, compared with a previous study [19] that achieved 91%. In comparison with 16K AraNews, the proposed accuracy declined by 3%. This is due to the noisiness of AraNews articles. However, the proposed model keeps its superiority by at least 3% more than the other studied models. In conclusion, Table 7 shows the superiority of our proposed model against previous studies based on multiple datasets.

Fig 4 and Fig 5 show the confusion matrix for our proposed model against the AraNews and AFND datasets, respectively. The following Fig. 6 demonstrates the learning curves of our proposed model and baseline models. It found that the convergence of the proposed model is better than the others.

Table 6: Proposed model result for the AraNews dataset (50k)

Models	Embedding	Acc.	Embedding	Acc.
CNN	Arb2vec	84%	Glove	84%
LSTM	Arb2vec	87%	Glove	86%
BiLSTM	Arb2vec	88%	Glove	85%
CNN+LSTM	Arb2vec	89%	Glove	88%
CNN+BiLSTM	Arb2vec	89%	Glove	89%
Proposed Model	Arb2vec	90%	Glove	92%

Table 7: Previous studies vs Proposed model comparison

Study	Models	Datasets	Acc.
[3]	-	AFND	79%
[20]	CNN+LSTM	AFND	81%
[19]	CNN+LSTM	AraNews	91%
[18]	Bi-LSTM with Atten.	AFND	81%
[11]	AraBERT	AraNews	78%
[12]	MARBERTv2	ArabFake	89.96%
Proposed model	CNN + BiLSTM + Atten.	AraNews	95%
		AFND	88%

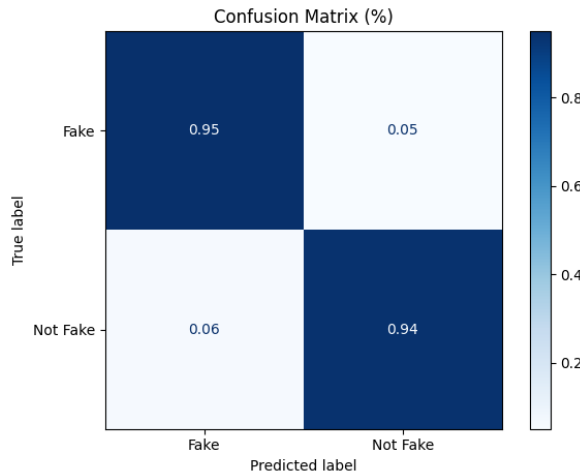


Fig.4: A confusion matrix of the AraNews Dataset

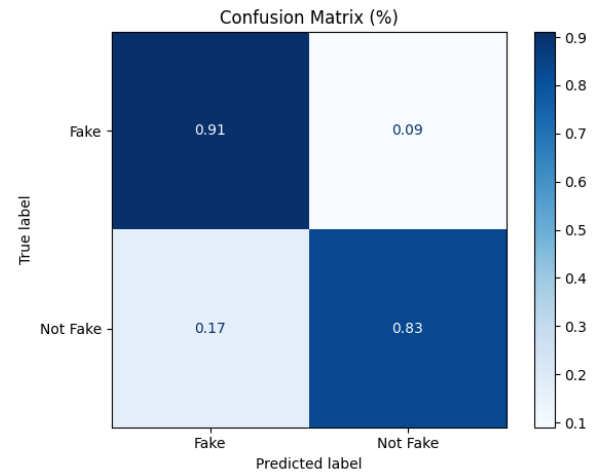


Fig.5: A confusion matrix of AFND Dataset

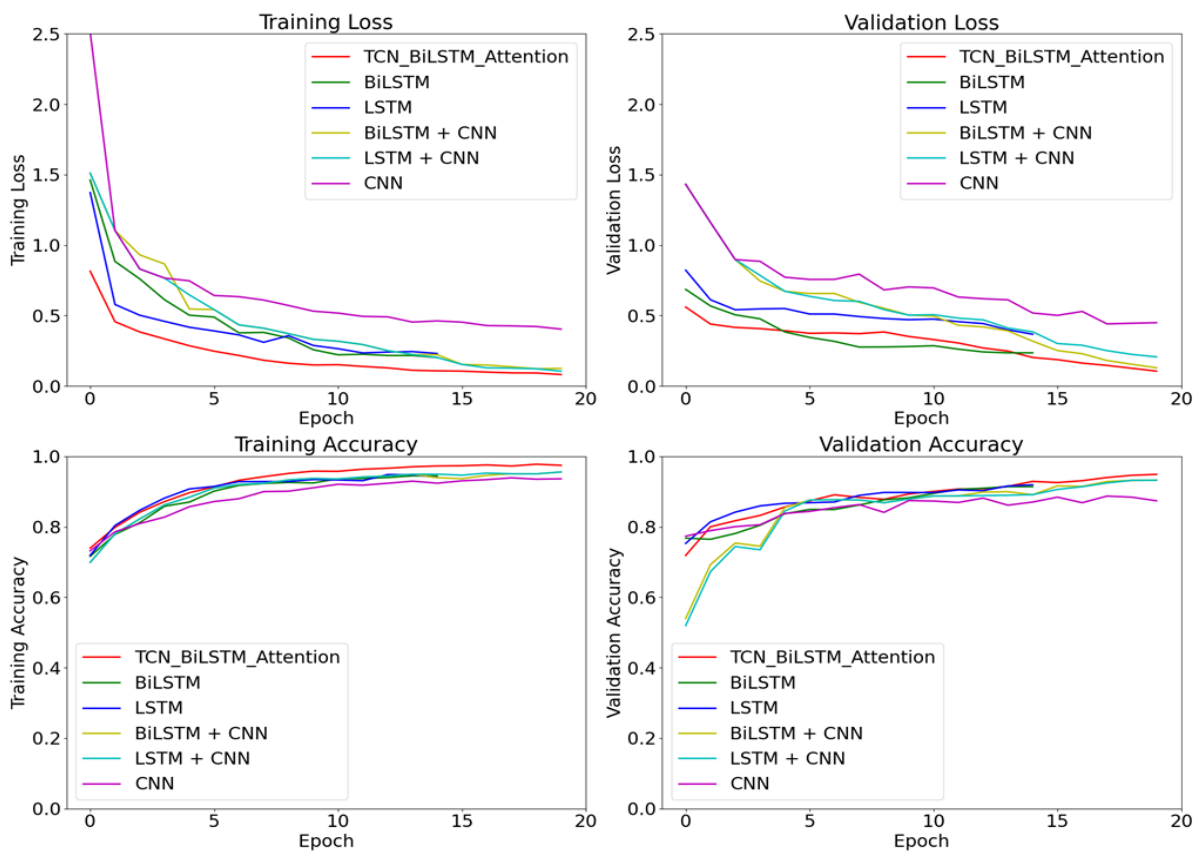


Fig.6: Learning Curve

6- Conclusions

In this study, a robust hybrid deep learning model has been developed by incorporating TCN for feature extraction and BiLSTM with an attention mechanism network for

detecting fake news in the Arabic context. The proposed model achieves a high-performance accuracy of 88% and 95% with AFND and AraNews datasets, respectively. These results exceed the baseline methods (CNN, LSTM, BiLSTM, CNN with LSTM, and CNN with BiLSTM) by at least (4%, 3%, 3%, 2%, and 2%), respectively, and

state-of-the-art methods by at least 7%. The hybrid methodology succeeds in tackling the complexity of Arabic content involving different dialects and morphological abundance. According to our extensive experiments, Glove provides the best word representation compared to AraVec, where the model accuracy improved by 1% (AFND) and 6% (AraNews). In future work, several directions can advance this research. First, advanced pre-trained transformer models could be applied, such as BERT and GPT, to get the best representation of Arabic content. Second, examining our work against more diverse datasets is much preferable since Arabic has many dialects, ranging from standard to region-based accents.

References

- [1] M. Alkhair, K. Meftouh, K. Smaïli, and N. Othman, "An Arabic Corpus of Fake News: Collection, Analysis and Classification," in *Communications in Computer and Information Science*, 2019. doi: 10.1007/978-3-030-32959-4_21.
- [2] M. Al-Yahya, H. Al-Khalifa, H. Al-Baity, D. Alsaeed, and A. Essam, "Arabic Fake News Detection: Comparative Study of Neural Networks and Transformer-Based Approaches," *Complexity*, vol. 2021, 2021, doi: 10.1155/2021/5516945.
- [3] A. Khalil, M. Jarrah, M. Aldwairi, and Y. Jararweh, "Detecting Arabic Fake News Using Machine Learning," in *2021 2nd International Conference on Intelligent Data Science Technologies and Applications, IDSTA 2021*, 2021. doi: 10.1109/IDSTA53674.2021.9660811.
- [4] Y. Almsrahad and N. Moghaddam Charkari, "Image Fake News Detection using Efficient NetB0 Model," *Journal of Information Systems and Telecommunication (JIST)*, vol. 12, no. 45, pp. 41–48, Mar. 2024, doi: 10.61186/jist.40976.12.45.41.
- [5] P. Bahad, P. Saxena, and R. Kamal, "Fake News Detection using Bi-directional LSTM-Recurrent Neural Network," in *Procedia Computer Science*, 2019. doi: 10.1016/j.procs.2020.01.072.
- [6] A. Monnier, "Narratives of the Fake News Debate in France," *IAFOR Journal of Arts & Humanities*, vol. 5, no. 2, 2018, doi: 10.22492/ijah.5.2.01.
- [7] S. Senhadji and R. A. S. Ahmed, "Fake news detection using naïve Bayes and long short term memory algorithms," *IAES International Journal of Artificial Intelligence*, vol. 11, no. 2, 2022, doi: 10.11591/ijai.v11.i2.pp746-752.
- [8] Y. Bang, E. Ishii, S. Cahyawijaya, Z. Ji, and P. Fung, "Model Generalization on COVID-19 Fake News Detection," in *Communications in Computer and Information Science*, 2021. doi: 10.1007/978-3-030-73696-5_13.
- [9] J. A. Nasir, O. S. Khan, and I. Varlamis, "Fake news detection: A hybrid CNN-RNN based deep learning approach," *International Journal of Information Management Data Insights*, vol. 1, no. 1, 2021, doi: 10.1016/j.jjimei.2020.100007.
- [10] S. R. Sahoo and B. B. Gupta, "Multiple features based approach for automatic fake news detection on social networks using deep learning," *Appl. Soft Comput.*, vol. 100, 2021, doi: 10.1016/j.asoc.2020.106983.
- [11] N. A. Othman, D. S. Elzanfaly, and M. M. M. Elhawary, "Arabic Fake News Detection Using Deep Learning," *IEEE Access*, vol. 12, no. September, pp. 122363–122376, 2024, doi: 10.1109/ACCESS.2024.3451128.
- [12] A. Maher et al., "ArabFake: A Multitask Deep Learning Framework for Arabic Fake News Detection, Categorization, and Risk Prediction," *IEEE Access*, vol. 12, no. November, pp. 191345–191360, 2024, doi: 10.1109/ACCESS.2024.3518204.
- [13] A. Khalil, M. Jarrah, M. Aldwairi, and M. Jaradat, "AFND: Arabic fake news dataset for the detection and classification of articles credibility," *Data Brief*, vol. 42, 2022, doi: 10.1016/j.dib.2022.108141.
- [14] G. Jardaneh, H. Abdelhaq, M. Buzz, and D. Johnson, "Classifying Arabic tweets based on credibility using content and user features," in *2019 IEEE Jordan International Joint Conference on Electrical Engineering and Information Technology, JEEIT 2019 - Proceedings*, 2019. doi: 10.1109/JEEIT.2019.8717386.
- [15] F. Aljwari, W. Alkaber, A. Alshutayri, E. Aldahri, N. Aljojo, and O. Abouola, "Multi-scale Machine Learning Prediction of the Spread of Arabic Online Fake News," *Postmodern Openings*, vol. 13, no. 1 Supl, 2022, doi: 10.18662/po/13.1supl/411.
- [16] E. M. B. Nagoudi, A. Elmadany, M. Abdul-Mageed, T. Alhindi, and H. Cavusoglu, "Machine Generation and Detection of Arabic Manipulated and Fake News," Nov. 2020, [Online]. Available: <http://arxiv.org/abs/2011.03092>
- [17] A. R. Mahlous and A. Al-Laith, "Fake News Detection in Arabic Tweets during the COVID-19 Pandemic," *International Journal of Advanced Computer Science and Applications*, vol. 12, no. 6, 2021, doi: 10.14569/IJACSA.2021.0120691.
- [18] S. E. Sorour and H. E. Abdelkader, "AFND: ARABIC FAKE NEWS DETECTION WITH AN ENSEMBLE DEEP CNN-LSTM MODEL," *J. Theor. Appl. Inf. Technol.*, vol. 100, no. 14, 2022.
- [19] T. A. Wotaifi and B. N. Dhannoon, "An Effective Hybrid Deep Neural Network for Arabic Fake News Detection," *Baghdad Science Journal*, vol. 20, no. 4, 2023, doi: 10.21123/bsj.2023.7427.
- [20] T. A. Wotaifi and B. N. Dhannoon, "Attention Mechanism Based on a Pre-trained Model for Improving Arabic Fake News Predictions," *Iraqi Journal of Science*, vol. 64, no. 11, 2023, doi: 10.24996/ij.s.2023.64.11.45.
- [21] H. Saadany, E. Mohamed, and C. Orasan, "Fake or Real? A Study of Arabic Satirical Fake News," Nov. 2020, [Online]. Available: <http://arxiv.org/abs/2011.00452>
- [22] R. M. Albalawi, A. T. Jamal, A. O. Khadidos, and A. M. Alhothali, "Multimodal Arabic Rumors Detection," *IEEE Access*, vol. 11, 2023, doi: 10.1109/ACCESS.2023.3240373.
- [23] S. F. Ahmed et al., "Deep learning modelling techniques: current progress, applications, advantages, and challenges," *Artif. Intell. Rev.*, vol. 56, no. 11, 2023, doi: 10.1007/s10462-023-10466-8.
- [24] J. Kasubi, M. D. Huchaiah, I. Gad, and M. K. Hooshmand, "A Comparison Analysis of Conventional Classifiers and Deep Learning Model for Activity Recognition in Smart

- Homes based on Multi-label Classification,” *Journal of Information Systems and Telecommunication (JIST)*, vol. 12, no. 46, pp. 127–137, Jun. 2024, doi: 10.61186/jist.36294.12.46.127.
- [25] P. Kim, *MATLAB Deep Learning: With Machine Learning, Neural Networks and Artificial Intelligence*. 2017.
- [26] K. Jindal and R. Aron, “A Hybrid Machine Learning Approach for Sentiment Analysis of Beauty Products Reviews,” *Journal of Information Systems and Telecommunication*, vol. 10, no. 37, 2022, doi: 10.52547/jist.15586.10.37.1.
- [27] Q. Zhang, L. T. Yang, Z. Chen, and P. Li, “A survey on deep learning for big data,” 2018. doi: 10.1016/j.inffus.2017.10.006.
- [28] J. Yu, A. de Antonio, and E. Villalba-Mora, “Deep Learning (CNN, RNN) Applications for Smart Homes: A Systematic Review,” 2022. doi: 10.3390/computers11020026.
- [29] M. Schuster and K. K. Paliwal, “Bidirectional recurrent neural networks,” *IEEE Transactions on Signal Processing*, vol. 45, no. 11, 1997, doi: 10.1109/78.650093.
- [30] Z. Huang, W. Xu, and K. Yu, “Bidirectional LSTM-CRF Models for Sequence Tagging,” Aug. 2015, [Online]. Available: <http://arxiv.org/abs/1508.01991>
- [31] A. Graves and J. Schmidhuber, “Framewise phoneme classification with bidirectional LSTM and other neural network architectures,” *Neural Networks*, vol. 18, no. 5–6, pp. 602–610, 2005, doi: 10.1016/j.neunet.2005.06.042
- [32] D. Bahdanau, K. H. Cho, and Y. Bengio, “Neural machine translation by jointly learning to align and translate,” in *3rd International Conference on Learning Representations, ICLR 2015 - Conference Track Proceedings*, 2015.
- [33] P. Zhou et al., “Attention-based bidirectional long short-term memory networks for relation classification,” in *54th Annual Meeting of the Association for Computational Linguistics, ACL 2016 - Short Papers*, 2016. doi: 10.18653/v1/p16-2034.
- [34] A. Vaswani et al., “Attention is all you need,” in *Advances in Neural Information Processing Systems*, 2017.
- [35] C. Lea, R. Vidal, A. Reiter, and G. D. Hager, “Temporal convolutional networks: A unified approach to action segmentation,” in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 2016. doi: 10.1007/978-3-319-49409-8_7.
- [36] S. Vijayarani and M. P. Research Scholar, “Preprocessing Techniques for Text Mining-An Overview.”
- [37] S. Akuma, T. Lubem, and I. T. Adom, “Comparing Bag of Words and TF-IDF with different models for hate speech detection from live tweets,” *International Journal of Information Technology (Singapore)*, vol. 14, no. 7, 2022, doi: 10.1007/s41870-022-01096-4.
- [38] M. SATHVIK, “Enhancing Machine Learning Algorithms using GPT Embeddings for Binary Classification,” Mar. 27, 2023. doi: 10.36227/techrxiv.22331053.v1.
- [39] A. B. Soliman, K. Eissa, and S. R. El-Beltagy, “AraVec: A set of Arabic Word Embedding Models for use in Arabic NLP,” in *Procedia Computer Science*, 2017. doi: 10.1016/j.procs.2017.10.117.
- [40] S. F. Sabbeh and H. A. Fasihuddin, “A Comparative Analysis of Word Embedding and Deep Learning for Arabic Sentiment Classification,” *Electronics (Switzerland)*, vol. 12, no. 6, 2023, doi: 10.3390/electronics12061425.
- [41] HANSON ER, “GloVe: Global Vectors for Word Representation Jeffrey,” *AES: Journal of the Audio Engineering Society*, vol. 19, no. 5, 1971.
- [42] H. A. Ismael, N. H. Al-A’araji, and B. K. Shukur, “An Enhanced Diabetic Foot Ulcer Classification Approach Using GLCM and Deep Convolution Neural Network,” *Karbala International Journal of Modern Science*, vol. 8, no. 4, 2022, doi: 10.33640/2405-609X.3268.
- [43] H. Himdi, G. Weir, F. Assiri, and H. Al-Barhamtoshy, “Arabic Fake News Detection Based on Textual Analysis,” *Arab. J. Sci. Eng.*, vol. 47, no. 8, 2022, doi: 10.1007/s13369-021-06449-y.
- [44] S. Sharma, S. Sharma, and A. Athaiya, “ACTIVATION FUNCTIONS IN NEURAL NETWORKS,” *International Journal of Engineering Applied Sciences and Technology*, vol. 04, no. 12, 2020, doi: 10.33564/ijeast.2020.v04i12.054.
- [45] H. Pratiwi et al., “Sigmoid Activation Function in Selecting the Best Model of Artificial Neural Networks,” in *Journal of Physics: Conference Series*, 2020. doi: 10.1088/1742-6596/1471/1/012010.
- [46] . A. Ismael, N. H. Al-A’araji, and B. K. Shukur, “Enhanced the prediction approach of diabetes using an autoencoder with regularization and deep neural network,” *Periodicals of Engineering and Natural Sciences*, vol. 10, no. 6, 2022, doi: 10.21533/pen.v10i6.3394
- [47] S. R. Durugkar, R. Raja, K. K. Nagwanshi, and S. Kumar, “Introduction to data mining,” 2022. doi: 10.1002/9781119792529.ch1.