

Exploring a Hybrid Deep Learning Framework to Automatically Discover Topic and Sentiment in COVID-19 Tweets

Khandaker Tayef Shahriar*^{1,2} and Iqbal H. Sarker*³

¹Department of Computer Science and Engineering, Chittagong University of Engineering & Technology, Chittagong, 4349, Bangladesh.

²International Islamic University Chittagong, Chattogram, 4318, Bangladesh, Google Scholar: scholar.google.com.

³Center for Securing Digital Futures, Edith Cowan University, Perth, WA 6027, Australia, Google Scholar: scholar.google.com.

Contributing authors: tayef@iiuc.ac.bd; m.sarker@ecu.edu.au;

Abstract

COVID-19 has created a major public health problem worldwide and other problems such as economic crisis, unemployment, mental distress, etc. The pandemic is deadly in the world and involves many people not only with infection but also with problems, stress, wonder, fear, resentment, and hatred. Twitter is a highly influential social media platform and a significant source of health-related information, news, opinion and public sentiment where information is shared by both citizens and government sources. Therefore, an effective analysis of COVID-19 tweets is essential for policymakers to make wise decisions. However, it is challenging to identify interesting and useful content from major streams of text to understand people's feelings about the important topics of the COVID-19 tweets. In this paper, we propose a deep learning framework for analyzing topic-based sentiments by extracting key topics with significant labels and classifying positive, negative, or neutral tweets on each topic to quickly find common topics of public opinion and COVID-19-related attitudes. While building our model, we take into account hybridization of Bidirectional Long Short-Term Memory (BiLSTM) and Gated Recurrent Unit (GRU) structures for sentiment analysis to achieve our goal. The experimental results show that our topic identification method extracts better topic labels and the sentiment analysis approach using our proposed hybrid deep

learning model achieves the highest accuracy compared to traditional models.

Keywords: COVID-19, Deep Learning, Sentiment Analysis, Topic Identification.

1 Introduction

Social networks play an important role in a major crisis, as people provide feedback and share ideas with others using these online channels that generate information and new ideas related to disaster response [1]. Twitter is regarded as one of the most important social media platforms to explain the behavior and predict the outcomes of the pandemic [2]. At the end of December 2019, a novel coronavirus outbreak that caused COVID-19 was reported [3]. Due to the rapid spread of the virus, the World Health Organization (WHO) declared a state of emergency [4]. The impact of the COVID-19 pandemic on the Twitter platform is huge. Thus, it is very necessary to know the topics related to the COVID-19 tweets and the associating sentiment polarities users are producing in this platform on a regular basis to get an overview of the pandemic situation, human needs, and steps to reduce the harmful impact of the pandemic. The Twitter platform can be considered as a source of useful information to highlight user conversations and better understand public perception towards the COVID-19 pandemic situation. Therefore, by focusing on the context of the COVID-19 pandemic, in this work, we analyze people's tweets on the basis of extracting meaningful topics in an unsupervised manner and predicting sentiments classes in a supervised manner.

Let's consider a dataset of Twitter that contains user-generated tweets having the sentiment polarities as the target label about COVID-19 related issues. It is very difficult to parse all the tweets and express the internal context manually [5]. Therefore, finding the internal topics automatically of all tweets will help the policymakers in the relevant departments to implement mandatory measures to reduce the negative impact of the pandemic. However, sentiment analysis is a popular and significant current research paradigm that reflects public perceptions of the event. Sentiment analysis of tweets helps to determine whether a given tweet is neutral, positive, or negative [6]. To achieve our goal, we consider the following research questions:

RQ 1: How to automatically discover the topics associated with the significant labels through analyzing COVID-19 tweets?

RQ 2: How to design an effective framework for predicting the sentiments of tweets with associated topic labels?

To answer these questions, we develop a hybrid deep learning framework by considering the sentiment terms and aspect terms of tweets. However, it is a challenging task to extract meaningful topic labels automatically by machine instead of following the manually annotated topic labeling approach with the diverse human interpretations [7]. Therefore, in this paper, we use Latent Dirichlet Allocation (LDA) [8], which is an unsupervised probabilistic algorithm to discover topics from tweets. Thus for the topic identification purpose, we do not need any labeled dataset. A collection of topics found in the tweets is generated by LDA. In our earlier paper [9], we proposed

“SATLabel” for identifying topic labels comparing the effectiveness with the manual labeling approach only. In this paper, we expand our previous work by changing the LDA-based optimal topic selection process and evaluating the effectiveness with respect to other topic labeling techniques. We also integrate the sentiment detection approach to build a complete framework for topic-based sentiment analysis. Our experimental results show that the label produced by the approach in the proposed framework has the highest soft cosine similarity score with tweets of the same topic compared to other topic labeling methods. In recent years, deep learning models have become very successful in many fields. Deep learning models automatically extract features from various neural network models and learn from their errors [10]. Deep learning models are extensively used in the field of sentiment analysis by Convolutional Neural Network (CNN), Recurrent Neural Network (RNN), Long Short Term Memory (LSTM), BiLSTM, and GRU [11]. GRU is good at capturing order details and long-distance dependency in a sequence from aspect terms and sentiment terms [12]. On the other hand, BiLSTM can read the input sequences in the forward and the backward direction to increase the amount of information and improve context successfully [13]. To improve the performance of sentiment prediction, we propose a hybrid deep learning model based on the combination of the above GRU and BiLSTM features with the use of the Global Average Pooling layer.

Conventional sentiment analysis methods directly follow the classification of raw text, while our proposed framework extracts latent topics through the implementation of LDA which are considered as additional relevant features to provide deeper insights to topic-aware sentiment analysis. Our proposed hybrid model incorporates both long-term contextual patterns (via BiLSTM) and short-term dependencies (via GRU) to improve sentiment classification performance. Our proposed method also provides structural analysis by integrating topic modeling with hybrid deep learning models to identify sentiment variations on pandemic-related themes. The primary contributions of this paper can be summarized as follows:

- We effectively use sentiment terms and aspect terms to identify topic labels in COVID-19 tweets.
- We propose a hybrid deep learning model combining GRU and BiLSTM features with the use of the Global Average Pooling layer.
- We have conducted extensive experiments utilizing real-world COVID-19 tweets to show the effectiveness of our proposed framework.

The whole paper is formed as follows. Section 2 reviews work on related topic modeling and sentiment analysis. The methodology and architecture of the proposed framework are depicted in section 3. After that, the results of the experiments are analyzed in section 4. Next, we present the discussion and the conclusion section of our work that provide directions for future work.

2 Related Work

The suffering caused by COVID-19 is severe [14]. COVID-19 news has been propagated widely on social media surpassing other news. The Tweeter remains on top to broadcast pandemic-related news [15]. Tweets related to COVID-19 can be helpful in identifying meaningful topic labels to highlight user conversation and understand ideas about people's needs and interests. Many researchers use the LDA algorithm to extract hidden themes from texts. Hingmire et al. [16] proposed a paper to create an LDA-based topic model but expert annotation is needed to assign the topic to class labels. Wang et al. [17] presented a paper that minimizes the problem of data sparsity without labeling the main topics particularly. Hourani et al. [18] presented a technique of text classification according to their topics that required a labeled dataset. Asmussen et al. [19] presented a paper where labeling the topic depends on the researcher's point of view without having any machine-dependent automatic method. Maier et al. [20] introduced the applicability and accessibility of communication researchers using LDA based topic labeling system based on extensive knowledge of the context. Guo et al. [21] compared LDA analysis with dictionary-based analysis by implementing a manual topic labeling technique. Ordun et al. [22] labeled each LDA-generated topic with the first three terms having the highest associated probabilities in the corresponding single document of each topic to interpret. Since the Mallet wrapper of the LDA algorithm tends to offer a better division of topics than the inbuilt Gensim version, the LDA Mallet models were implemented in the study [23] to generate topics where each topic is an integration of keywords and each keyword provides a certain weighted percentage to the topic and the topic labels were inferred from keywords.

Twitter sentiment analysis has become a major source of interest for over a decade [24]. In history, many works were performed in the field of sentiment analysis using deep learning models [11, 25]. Behera et al. [26] presented a hybrid model called Co-LSTM, which highly fits with big social data. Poria et al. [27] considered different aspects of analysis in order to investigate multi-modal sentiment analysis implementing three deep learning-based structures. Kaur et al. [28] examined emotions and sentiments using TextBlob and IBM Tone analyzer by translating 16,138 tweets into English. Nemes and Kiss [29] scrutinized tweets using TextBlob and RNN. Jiang et al. [30] proposed a paper based on aspect-level sentiment conversion modules by extracting aspect-specific sentiment words. Imran et al. [31] proposed a multi-layer LSTM based model for predicting both emotions and sentiment polarities. Authors in [32]

explored sentiments of Indians in COVID-19 related tweets after the lockdown. For analysis purposes, they collected 24,000 tweets from March 25th to March 28th, 2020 using the keywords: #IndiaLockdown and #IndiafightsCorona. The results showed that although there were negative sentiments about the closure, tweets possessing positive sentiments were actually present. Rustam et al. [33] created the feature set by combining the term frequency-inverse document frequency and bag-of-words for performance comparison of different machine learning techniques by classifying Twitter data as positive, neutral or negative sentiment. They investigated the dataset using the LSTM deep learning model. Naseem et al. [34] compared conventional machine learning techniques and deep learning methods using the application of various word embeddings such as Glove [35], fastText [36], and Word2Vec [37] by categorizing COVID-19 related tweets into negative, positive, and neutral classes. Their results suggested that deep learning (DL) methods performed better than traditional ML techniques.

Baired et al. [38] used SenticNet for fine-tuning various features from the dataset. They investigated short microblogging (e.g., Twitter), messaging, and sentiment analysis by combining knowledge-based and statistical methods. A study [39] suggested a way to classify sentiments using the deep learning model. It used NLP to model the topic to identify key issues related to COVID-19 expressed on social media. The classification was performed by applying the LSTM Recurrent Neural Networks (LSTM RNN) model. Ahmed et al. [40] used the opinion mining approach to identify sets of users with similar attitudes. At the beginning of the COVID-19 pandemic, Xue et al. [41] performed sentiment analysis on Twitter data expressing that fear of the unknown environment was prevalent on Twitter. Medford et al. [42] explored sentiments and investigated topics using LDA by collecting the Twitter data from January 14th to 28th, 2020. Xiang et al. [43] collected 82,893 tweets for performing sentiment analysis and topic modeling by applying the NRC Lexicon and LDA respectively. Satu et al. [44] presented "TClustVID" that performs a cluster-based division of tweets and topic modeling. They produced themes of positive, negative, and neutral topics for the cluster investigation and the realization of the COVID-19 pandemic status. Chandrasekaran et al. [45] analyzed key topics and corresponding sentiments by identifying the patterns and trends of COVID-19 related tweets before and after the outbreak. They used LDA to extract 26 topics from the COVID-19 related tweets and also clustered them into 10 broad themes to analyze the effect of COVID-19 on various domains. By analyzing the above works we can reach the conclusion that most of the works do not consider a complete framework for the automatic topic identification and sentiment analysis of COVID-19 related tweets. Moreover, most of the works followed a way of labeling topics manually which is expensive, time-consuming, and dictates difficult human interpretations of the tweets. However, a topic-based analysis of sentiments of the COVID-19 tweets categorizes the predicted multiclass sentiment polarities according to the extracted topics in order to gain an overview of broader public opinion. Therefore, in this paper, we propose a framework for automatically identifying key topics associated with specific topic labels and predicting

sentiment polarities to present a topic-based sentiment analysis of the COVID-19 related tweets.

3 Methodology

Topic identification and sentiment analysis (also called opinion mining) play an important role in obtaining information from social media platforms. In this paper, we implement two methods for in-depth analysis of the COVID-19 related tweets on Twitter. The first one is the LDA-based automatic topic modeling and labeling approach, and the second one is multiclass sentiment analysis based on a hybrid deep learning model. Our proposed framework aims to illustrate public sentiments in tweets towards some specific topic labels. In this section, we present our proposed framework for conducting a topic-based sentiment analysis of COVID-19 related tweets as shown in Fig. 1. In Fig. 1, the “topic identification” module generates a topic label, and the “deep learning model” simultaneously predicts the sentiment polarity of each tweet. Therefore, we can perform topic-based sentiment analysis to classify the sentiment polarities of all tweets into corresponding topic labels to provide insight into the COVID-19 pandemic situation. Text pre-processing is one of the most important steps to analyze and extract features from textual data for further processing. After performing normalization of noisy, non-grammatical, and unstructured tweets, we follow a set of processing steps to produce the expected result. By following the topic identification step, our framework separates positive, neutral, or negative sentiments from COVID-19 tweets by implementing a GRU-BiLSTM based hybrid deep learning model. The final summary generated by the proposed framework provides decision support to the policymakers based on the topic-based sentiment analysis which can be presented in any format, either text or charts.

3.1 Preparing Dataset

TextBlob, VADER and other sentiment analysis tools are often used to detect sentiments from COVID-19 related tweets [29, 30, 33, 40]. But in this work, we only consider the publicly available manually tagged dataset with sentiment polarities of the COVID-19 related tweets from the Kaggle repository¹. There are two CSV files in the dataset. One is Corona.NLP.train.CSV and another is Corona.NLP.test.CSV. There are 41,157 tweets in Corona.NLP.train.CSV file and 3,798 tweets in Corona.NLP.test.CSV file. There are five categories of sentiments as Positive, Extremely Positive, Neutral, Negative, and Extremely Negative which are transformed into three classes namely Positive, Neutral and Negative as a target label to achieve better accuracy. There are 7,713 Neutral, 15,398 Negative, and 18,046 Positive tweets available in the dataset which are marked with 0, 1, and 2 respectively. From training dataset, we have kept 20%

¹<https://www.kaggle.com/datatattle/covid-19-nlp-text-classification>

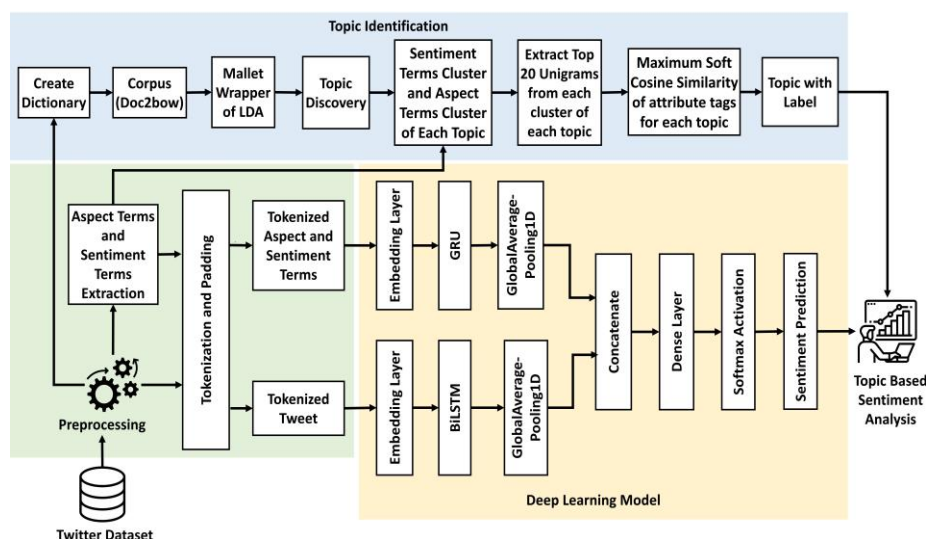


Fig. 1 Architecture of the proposed framework

of data as validation set. We prepare the Twitter dataset to get the standard form by using several text processing functions.

- 1) *Conversion of words into lowercase:* We transform all words in the tweets into lowercase.
- 2) *Deletion of URLs and links:* We replace all the URLs and hyperlinks in the tweets with empty strings.
- 3) *Elimination of mentions:* Normally people use @ in tweets to mention users. In our work, we eliminate the mentions.
- 4) *Removal of hashtags:* We identify hashtags and remove them with empty strings because hashtag words do not carry useful meaning in sentiment analysis [46].
- 5) *Handling contractions:* Generally, people use common English contractions in social media. We apply the regular expression to handle common contractions. For example, we replace won't with would not.
- 6) *Elimination of punctuations:* There is no importance of punctuations in tweets. We withdraw punctuation symbols such as &, -, etc.
- 7) *Managing words less than two characters:* We remove the words whose length is less than two characters because the word such as 'I' does not really carry much influence in determining sentiment.
- 8) *Stop words elimination:* Stop words in the English language such as pronouns, articles, prepositions, etc. have no emotional impact on the sentiment classification process. We remove the stop words.
- 9) *Dealing with Unicode and Non-English words:* To develop a clean and noise-free dataset, we remove the tweets that are not in English, and Unicode like "\u018e".

Table 1 Examples of Sentiment Terms and Aspect Terms

Sample Tweet	Sentiment Terms	Aspect Terms
Please read the thread.	read	thread
To enjoy and relax for your dinner it is a great place.	enjoy, relax, great	dinner, place
Links with info on communicating with children regarding COVID-19.	communicate, covid	links, info, children
The retail store owners right now	retail, right	owners, store

10) *Tokenization*: We vectorize the corpus by transforming each text into a sequence of integers based on word count.

3.2 Extracting Sentiment Terms and Aspect Terms

Sentiment terms capture the tone or perception of a sentence. Usually, adjectives and verbs are regarded as sentiment terms in a sentence that reflects the expressed opinion of the text. Noun phrases and nouns are regarded as aspect terms. Aspect terms are often considered as features that describe the event, product, or entity [47]. We follow the precise parts of speech tagging which is an effective way to extract sentiment terms and aspect terms from texts. Examples of sentiment terms and aspect terms are shown in Table 1.

3.3 Feature Extraction

To implement the machine learning and deep learning models, feature extraction is essential from the tweets [48]. Two methods of extracting features, a bag of words (BoW) and Word2vec are utilized in the proposed framework.

The boW is an easy way to extract features and it is used to retrieve information and perform NLP tasks [49]. BoW generates features to train the ML models on the basis of the presence of the individual words in the text and is widely used in text classification tasks. BoW produces vocabulary for all unique words and their frequency of occurrences in the documents to train models. In our work, we implement Doc2bow to extract numerical features from tweets to train the LDA model, which is a simple technique in Gensim [50] that converts documents into a representation of the bag of words as shown in the topic identification section of Fig. 1. In this paper, we use word2vec (skip-gram version) as embedding layers to extract numerical features from tokenized

texts and develop our framework for sentiment analysis purposes [37, 51].

Algorithm 1: Automatic Topic Labeling

Input: T : Number of Tweets in the dataset
Output: Topic Label (T_{Label})
 // Preprocess Tweets
 1 **for each** $t \in \{1, 2, \dots, T\}$ **do**
 2 $T_p \leftarrow \text{Preprocess}(t)$
 // Initialize Corpus
 3 Initialize an empty list *Corpus*
 // Corpus Development, Sentiment, and Aspect Terms
 Extraction
 4 **for each** $t_p \in \{1, 2, \dots, T_p\}$ **do**
 5 Append dictionary representation of Doc2Bow(t_p) to *Corpus*
 6 $S_{T_p} \leftarrow \text{Sentiment-Terms}(t_p)$
 7 $A_{T_p} \leftarrow \text{Aspect-Terms}(t_p)$
 // Topic Discovery using LDA
 8 $K \sim \text{Mallet}(\text{LDA}(\text{Corpus}))$
 // Cluster Creation Based on Dominant Topics
 9 **for each** $k \in \{1, 2, \dots, K\}$ **do**
 10 **for each** $t_p \in \{1, 2, \dots, T_p\}$ **do**
 11 $k_{\text{dominant}, t_p} \sim \text{dominant_topic}(t_p, k)$
 // Create Clusters from Topics
 12 $C_S \sim \text{Cluster}(S_{T_p} \rightarrow k_{\text{dominant}, t_p})$
 13 $C_A \sim \text{Cluster}(A_{T_p} \rightarrow k_{\text{dominant}, t_p})$
 // Extract Unigrams and Generate Topic Labels
 14 **for each** $k \in \{1, 2, \dots, K\}$ **do**
 // Extract Top 20 Unigrams from Sentiment Terms Cluster
 15 $U_S \sim \text{max-count}(\text{Top-Unigrams}(C_S \rightarrow k))$
 // Extract Top 20 Unigrams from Aspect Terms Cluster
 16 $U_A \sim \text{max-count}(\text{Top-Unigrams}(C_A \rightarrow k))$
 // Create Attribute Tags
 17 $A_{\text{tag}} \leftarrow U_S + U_A$
 // Assign Topic Label
 18 $T_{Label} \leftarrow \text{max-soft-cosine-similarity}(A_{\text{tag}}, k)$

3.4 Topic Identification

LDA is a well-performed topic modeling algorithm for finding hidden themes available in the corpus on an unlabeled dataset. But the challenge is how to assign important labels on topics generated by the LDA. The working principle for generating automatic topic labels from the Twitter dataset is presented in Algorithm 1. The steps for the identification of topics and the generation of significant topic labels automatically as output in the proposed framework are discussed below:

3.4.1 Creating Dictionary and Corpus

The systematic way of generating a number of language lexicons is supported by a dictionary and the corpus usually illustrates an arbitrary sample of that language. The corpus of a document is made up of words or phrases. In the Natural Language Processing (NLP) paradigm, the language corpus plays an important role in creating a knowledge-based system and text mining. In the proposed framework, we construct a dictionary from the preprocessed text and then create a corpus. Documents in the dictionary are converted to the Bag of Words (BoW) format using Doc2bow embedding [50] for the corpus development. Corpus comprises the words id and its frequency in all documents. Each word is considered as a normalized and tokenized string.

3.4.2 Topic Discovery

The BoW corpus is transferred to the LDA Mallet wrapper. The presence of a set of topics on the corpus is discovered by the LDA. LDA Mallet wrapper works fast and provides a better classification of topics using the Gibbs Sampling [52] method. LDA produces the most prominent words in the themes. So by using the weightage of keywords one can infer prevalent topics in texts. In order to overcome the difficult manual labeling method, our framework produces automated topic labels using sentiment terms and aspect terms without human interpretation. Based on the result of the topic coherence score, we select an LDA model that gets a total of 14 topics itself. Then we get the main topic as dominant topic in each tweet from LDA model to identify the distribution of topics across all tweets in the dataset.

3.4.3 Labeling Topic Using Top Unigrams

We generate sentiment terms cluster and aspect terms cluster by experimenting with tweets corresponding to each LDA-generated topic. Therefore, we find 14 sentiment terms cluster and 14 aspect terms cluster from LDA-generated topics. Unigram is a single n-gram word sequence. The use of unigrams can be seen in NLP, cryptography, and mathematical analysis. However, the soft cosine similarity takes into account the similarity of the features in the vector space model [53]. For each topic, we select top 20 unigrams as threshold value because of high frequencies of these unigrams and extract the top 20 unigrams from the sentiment terms cluster and the top 20 unigrams from the aspect terms cluster. Then we create all the possible combinations of the top 20 unigrams of sentiment terms with the top 20 unigrams of aspect terms for each topic respectively to generate an appropriate attribute tag. We choose an attribute tag that has the highest soft cosine similarity value in relation to the tweets of that topic to assign with a significant topic label. We use sentiment term and aspect term blended attribute tag to label each topic because that feature tag describes the topic of the tweets. To classify a tweet into categories with a specific topic label from test data, we detect the topic number that contains a significant percentage impact on that tweet.

3.5 Building a Hybrid Deep Learning Model for Sentiment Analysis

In our proposed framework we implement a hybrid deep learning model for multiclass sentiment classification using two powerful feature extractors. Gated Recurrent Unit (GRU) and Bidirectional Long Short Term Memory (BiLSTM) are used for feature extraction in two branches of the hybrid model in the proposed framework as shown in Fig. 1. Both of the extractors are followed by an individual Global Average Pooling layer. In embedding layers, we apply the word embeddings generated by Word2vec (skip-gram version) model to measure the semantic relationship between each word pair to represent. The GRU branch takes tokenized sentiment terms and aspect terms as input, and the BiLSTM branch takes the tokenized tweet as input. The concatenation layer merges the features coming from two branches followed by the fully connected dense layer with softmax activation function for multiclass sentiment classification i.e. positive, neutral, or negative as shown in Fig. 1. Below we describe the main models used to build the GRU-BiLSTM based hybrid model in the proposed framework:

Long Short Term Memory (LSTM) is a kind of RNN architecture that handles the vanishing gradient problem using exclusive units [54]. Data is stored for a long time in the memory cell of the LSTM unit and three gates control the access to information inside and outside the cell. For instance, which information from the previous state cell will be memorized and which information will be deleted is determined by “Forget Gate”, while which information should enter the cell state is determined by “Input Gate”, and the output is controlled by the “Output Gate”. Bidirectional LSTM, more commonly known as BiLSTM, is a standard LSTM extension that enhances the model performance in the order of sequence [55]. It is a type of sequence processing model combining two LSTMs: one receives the input approaching in the forward direction and the other receives it in the backward direction. In natural language processing activities, Bidirectional LSTM is especially considered a popular option to carry on.

Gated Recurrent Unit (GRU) is another kind of recurrent network where information flow is managed and controlled between cells by implementing gating techniques in the neural network. GRU contains an update gate and a reset gate without having an output gate. The GRU structure makes it possible to take on the dependence of large sequences of data adaptively, without losing any data from earlier segments of the sequence. Moreover, GRU is a relatively simple type and much faster to calculate which often provides comparable functionality. The performance of the GRUs is better on certain smaller and less frequent datasets.

Global Average Pooling (GAP) does not slide in the structure of a small window but is measured over the entire output feature map in the previous layer. GAP typically regularizes the whole network structure, and each output channel corresponds to the features of each class, making the relationship between the output and the feature class more intuitive. Moreover, the GAP

Table 2 The model hyper-parameters

Hyper-parameter	Initial Value	Hyper-parameter Space	Optimal Value
Word embedding size for word2vec	100	50, 100, 150	100
The number of hidden neurons in BiLSTM	128	64, 128, 256, 512	256
The number of hidden neurons in GRU	128	64, 128, 256	128
Batch Size	32	16, 32, 64	32

layer contains no data parameter. Therefore, using GAP helps to improve network performance and increase cognitive accuracy [56]. In addition, the GAP layer avoids overfitting because it does not need parameter optimization. It serves as a flatten layer for transforming a multi-dimensional feature space into a one-dimensional feature map. In addition, it takes less time for computation. Our hybrid deep learning-based model in the proposed framework effectively utilizes the advantages of BiLSTM, GRU, and GAP models for multi- class sentiment classification. To predict the sentiment polarities, we apply a fully connected dense layer to extract the features from the feature map of the concatenation layer as shown in Fig. 1. The Softmax classifier acts as its input and takes the output to the final step. Table 2 shows the hyper-parameters used in the proposed hybrid model. As summarized in Table 2, the hyper-parameters characteristics include word embedding size of word2vec, number of hidden neurons in BiLSTM, GRU, and batch size. Due to the different lengths of tweets in the dataset, we take the maximum length of the tweets when we input tweets to the model. RMSProp optimizer is used because RMSProp incorporates the best properties to adjust learning rate adaptively and optimize model parameters [57]. The pseudo-code of the proposed hybrid deep learning model in the framework is provided in Algorithm 2. The model summary for

the hybrid deep learning model is as shown in Table 3. The shape of the output of the hybrid deep learning model is 3 as seen in the last row of Table 3.

Algorithm 2: Pseudo code of the proposed hybrid model

Input: COVID-19 Training Set T_1 , Testing Set T_2

Output: Predicted sentiment classes: S_{pos} , S_{neu} , S_{neg}

// Preprocess the data

- 1 Preprocess Tweets
 - 2 Extract tokenized sentiment terms and aspect terms as tok_{sa}
 - 3 Extract tokenized tweets as tok_{tweets}
 - // Train Word2Vec model
 - 4 Initialize hyperparameters for the Word2Vec model
 - 5 Train the Word2Vec embedding model using T_1
 - // Build the hybrid model
 - 6 Initialize hyperparameters for GRU and BiLSTM layers
 - // Branch 1 for tokenized sentiment terms
 - 7 Create an embedding layer using the trained Word2Vec model for tok_{sa}
 - 8 Add a GRU layer
 - 9 Add a GlobalAveragePooling layer
 - // Branch 2 for tokenized tweets
 - 10 Create an embedding layer using the trained Word2Vec model for tok_{tweets}
 - 11 Add a BiLSTM layer
 - 12 Add a GlobalAveragePooling layer
 - // Concatenate features from both branches
 - 13 Concatenate the output features from the two branches
 - // Add the output layer
 - 14 Add a dense layer with softmax activation for final classification
 - // Train the hybrid model
 - 15 Train the hybrid model using T_1 to compute the initial weights
 - 16 **for** $epoch = 1$ to $max\ epochs$ **do**
 - // Select a mini-batch for training
 - 17 Select a mini-batch from T_1
 - // Forward and backward propagation
 - 18 Perform forward propagation and compute the loss
 - 19 Perform backpropagation and update weights using RMSProp optimization
 - // Evaluate the model
 - 20 Use the trained hybrid model to predict sentiment classes for the testing set T_2
 - 21 Output the predicted sentiment classes: S_{pos} , S_{neu} , S_{neg}
-

3.6 Topic Based Sentiment Analysis

After predicting the sentiments of the tweets, the final step of our proposed framework is to classify the sentiment polarities for each of the 14 topics. We

Table 3 Model Summary for Hybrid Deep Learning Model

Layer (Type)	Output Shape	Param #	Connected to
embedding 1 input (InputLayer)	(None, 23)	0	[]
embedding input (InputLayer)	(None, 37)	0	[]
embedding 1 (Embedding)	(None, 23, 100)	3,565,900	['embedding 1 input[0][0]']
embedding (Embedding)	(None, 37, 100)	3,565,900	['embedding input[0][0]']
gru (GRU)	(None, 23, 128)	88,320	['embedding 1[0][0]']
bidirectional (Bidirectional)	(None, 37, 512)	731,136	['embedding[0][0]']
global average pooling1d 1	(None, 128)	0	['gru[0][0]']
global average pooling1d	(None, 512)	0	['bidirectional[0][0]']
concatenate (Concatenate)	(None, 640)	0	['global average pooling1d 1[0][0]', 'global average pooling1d[0][0]']
dense (Dense)	(None, 3)	1,923	['concatenate[0][0]']
activation (Activation)	(None, 3)	0	['dense[0][0]']

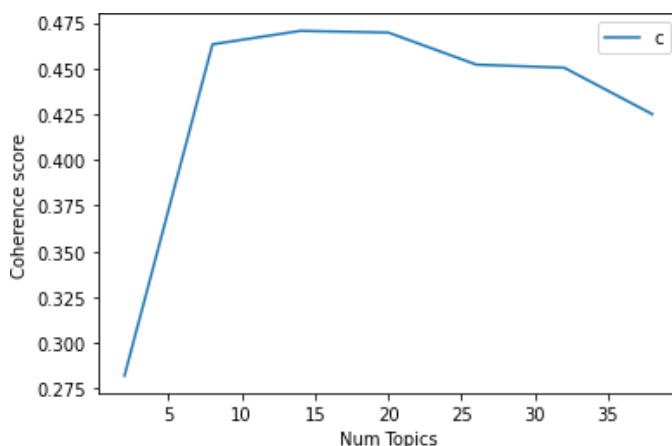


Fig. 2 Coherence score for the number of topics.

visualize and count the number of positive, negative, and neutral tweets to provide a summary to support decision-making for the policymakers. Our proposed framework effectively conducts topic-based sentiment analysis for providing a comprehensive overview of the broader public opinion. In the experimental section, we demonstrate the performance and effectiveness of our framework with respect to other benchmark models.

4 Experimental Results

In this section, we present the outcomes and findings of our experiments. First, we search trending topics with specific labels. Then we perform sentiment analysis of the tweets using a hybrid deep learning model. We also make the comparison of the results of the proposed framework with the corresponding benchmark models.

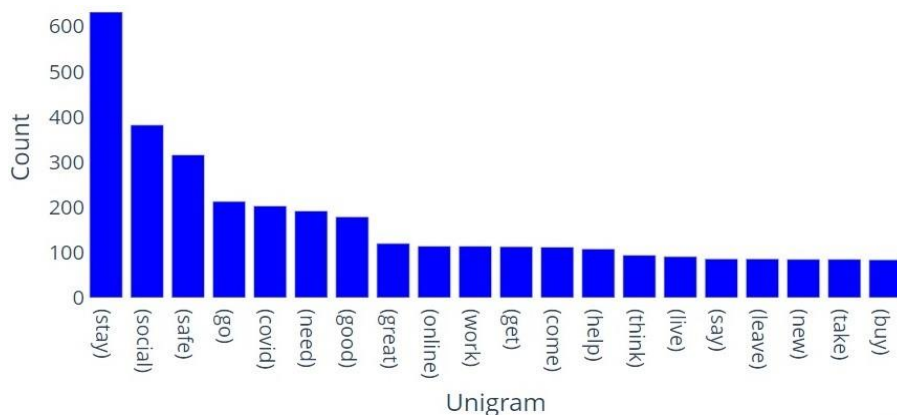


Fig. 3 Top 20 unigrams from sentiment terms cluster of topic no. 3

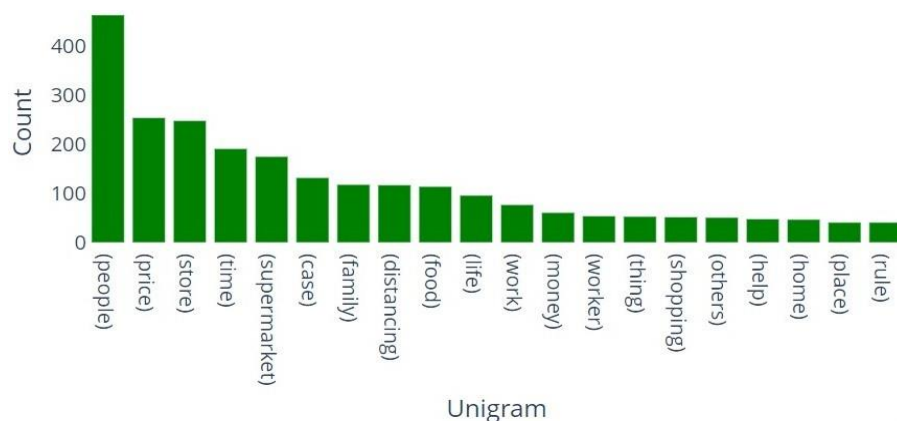


Fig. 4 Top 20 unigrams from aspect terms cluster of topic no. 3

4.1 LDA Based Topic Identification

4.1.1 Finding The Optimal Number of Topics for LDA

We build a function to extract a number of LDA models with multiple coherence values for the number of topics in order to obtain the optimal number of topics. LDA produces common co-occurred words and builds them into separate themes. We choose the optimal number of topics on the basis of the gensim coherence model [41]. We select the LDA model that returns the number of 14 topics for this dataset because it generates the highest coherence value. This is illustrated in Fig. 2, which presents the coherence scores for various topic numbers across multiple LDA models.

4.1.2 Selection of Top Unigram Features from Clusters

We generate sentiment terms cluster and aspect terms cluster for each topic. We find at the top count 20 unigrams in each cluster. Fig. 3 and Fig. 4 show the top 20 unigrams from sentiment terms cluster and aspect terms cluster of the topic no. 3 respectively. Then we get the topic label depending on the highest soft cosine similarity value of the attribute tag in relation to the tweets for that particular topic, which is generated by the combination of top unigrams of sentiment terms and aspect terms.

4.1.3 Qualitative Evaluation of Topic Labels

In table 4, we present part of a set of tweets assigned by the proposed framework generated topic labels. Table 4 presents that the detected topic labels are well-aligned and closely coherent with the descriptions of the tweets. We can categorize tweets and extract useful information related to a topic, by simply separating the tweets using the key label generated by the proposed method for that topic.

4.1.4 Effectiveness Analysis

In this experiment section, we compare the Soft Cosine Similarity (SCS) values of topic labels obtained by our proposed method with other two approaches as shown in Fig. 5. One is the approach followed by the first three important keywords of each topic [22] for topic labeling and another is the manual topic label inferred from keywords for LDA-generated 14 topics [23] as LDA generates keyword importance percentages for each topic. SCS is used to find semantic text similarities between two documents. The high value of SCS provides a high degree of similarity and the low SCS value provides less similarity to unrelated documents.

From Fig. 5, we find that SCS values generated by the proposed method for all topics are higher than other approaches except topics no. 1, 10, and 12. In the case of these three topics, the approach followed by Ordun et al. [22] finds slightly higher SCS values. For topics no. 9 and 13 we find the same SCS values for our proposed method and method followed by Ordun et al. [22]. In case of topic no. 3 we find the same SCS value of the method followed by Prabhakaran [23] with our proposed method. From Fig. 5, we can see that the topic labels produced by the proposed method in the framework illustrate higher SCS values for a maximum number of topics compared to other labeling methods. Therefore, our proposed framework works better and traces better topic labels from the dataset in an unsupervised manner to reduce the complex workload of human labeling.

Table 4 Example of Topics Detected on Tweets.

Sample Tweet	Topic No	Topic Label
create contact list with phone numbers of neighbours schools employer chemist GP set up online shopping accounts if poss adequate supplies of regular meds but not over order	5	online slot
Coronavirus Australia: Woolworths to give elderly, disabled dedicated shopping hours amid COVID-19 outbreak	1	thank store
My food stock is not the only one which is empty... PLEASE, don't panic, THERE WILL BE ENOUGH FOOD FOR EVERYONE if you do not take more than you need. Stay calm, stay safe.	2	hoard food
Dear Coronavirus, I've been following social distancing rules and staying home to prevent the spread of you. However, now I've spent an alarming amount of money shopping online. Where can I submit my expenses to for reimbursement? Let me know. #coronapocalypse #coronavirus	5	safe distancing
Cashier at grocery store was sharing his insights on #Covid 19 To prove his credibility he commented "I'm in Civics class so I know what I'm talking about".	10	learn behavior
Due to COVID-19 our retail store and classroom in Atlanta will not be open for walk-in business or classes for the next two weeks, beginning Monday, March 16. We will continue to process online and phone orders as normal! Thank you for your understanding!	6	small event
Airline pilots offering to stock supermarket shelves in #NZ lockdown #COVID-19	12	covid supermarket
@TartiiCat Well new/used Rift S are going for \$700.00 on Amazon rn although the normal market price is usually \$400.00 . Prices are really crazy right now for vr headsets since HL Alex was announced and it's only been worse with COVID-19. Up to you whethe	11	crude market
Response to complaint not provided citing COVID-19 related delays. Yet prompt in rejecting policy before consumer TAT is over. Way to go ?	7	learn scam

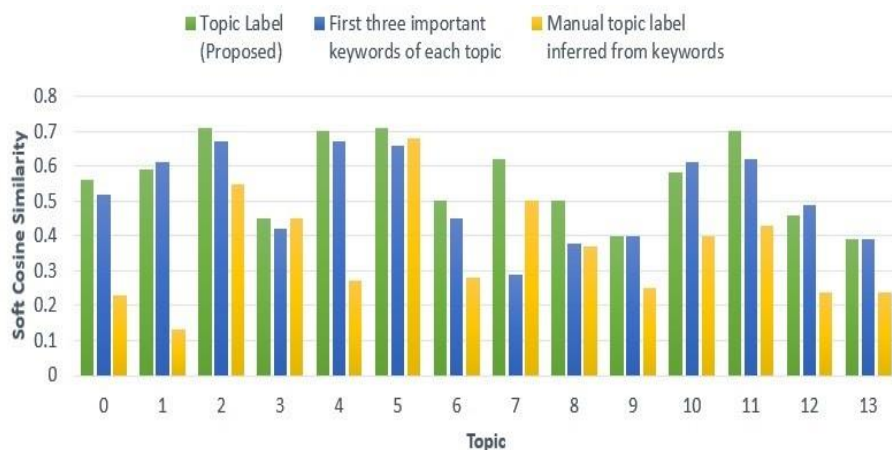


Fig. 5 Comparison of proposed topic label approach with other approaches

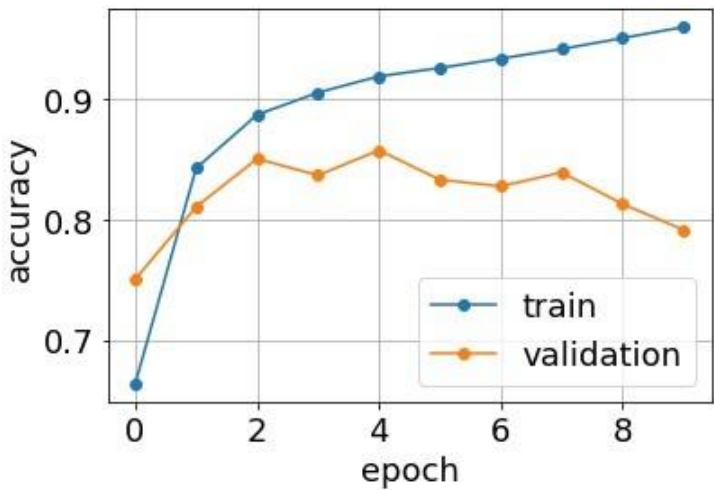


Fig. 6 Model accuracy graph (adopted from [58]).

Table 5 Evaluation of metrics of different models.

Model	Precision	Recall	F1-score
Simple RNN	0.82	0.82	0.82
CNN	0.82	0.81	0.81
GRU	0.85	0.85	0.85
BiLSTM	0.85	0.85	0.85
LSTM	0.85	0.85	0.85
Hybrid Model (Proposed)	0.86	0.86	0.86

4.2 Sentiment Analysis Using Hybrid Deep Learning Model

4.2.1 Effect of Model Iterations

Multiple iterations of model training affect model performance [58]. The validation accuracy of the model first rises and then falls with increasing model iterations. It can be seen in Fig. 6 and Fig. 7, when the number of iterations of the model increases, the model gradually overfits, and performance decreases. These figures illustrate the accuracy and loss graph of the proposed hybrid deep learning model in the framework respectively. Based on these plots, we find the best model at epoch no. 4. The fourth epoch is selected as the best epoch for model selection due to the model’s performance on validation data without exposure to test data. To mitigate the overfitting problem, we have saved the best model at epoch 4, which is similar to the early stopping strategy, and then we have loaded the best model for further performance analysis.

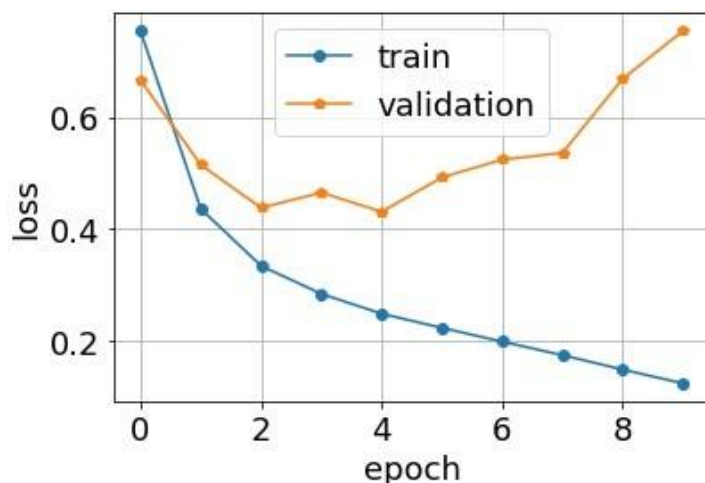


Fig. 7 Model loss graph (adopted from [58]).

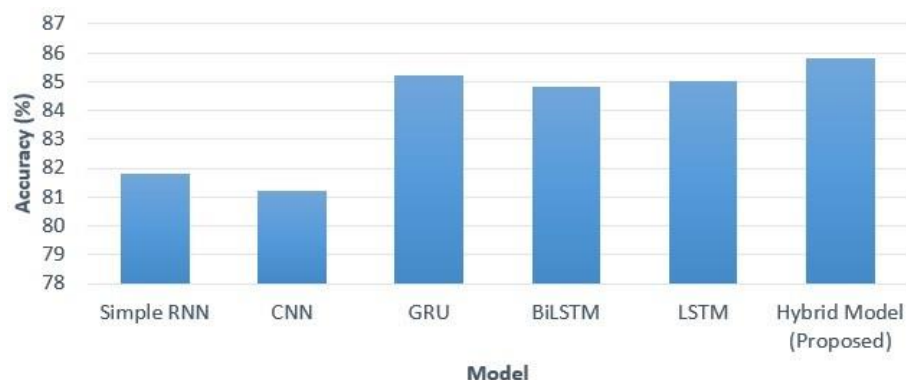


Fig. 8 Performance comparison of different models.

4.2.2 Performance Comparison of Different Models

We compare the sentiment analysis method in the proposed framework with the basic deep learning models such as SimpleRNN, CNN, GRU, BiLSTM, and LSTM using standard parameters in the dataset [58]. By implementing the word2vec (skip-gram version), we also train different deep learning models for the dataset and compare their results with the proposed hybrid model in the framework. Fig 8 represents the results of accuracy comparisons and Table 5 shows the precision, recall, and f1-score comparisons of different models. For simplicity, we have rounded the metrics to two decimal places to gain clarity on the differences in model performance. These differences show meaningful improvements in the context of the proposed hybrid model. We have performed multiple training runs to generate results to ensure robustness. From Fig. 8, it can be seen that almost all deep deep learning models provide good results

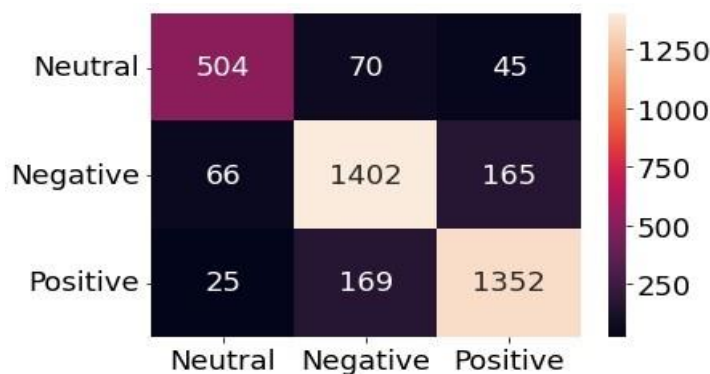


Fig. 9 Confusion matrix of proposed hybrid deep learning model.

with an accuracy level of more than 80%. The values reveal that our proposed approach outperforms other models and provides about 86% of accuracy. From the results of Table 5, it can be seen that the performance of the proposed hybrid model is higher compared to other basic deep learning models in terms of precision, recall, and f1-score for the weighted average of negative, neutral, and positive sentiment predictions. The experimental results show that the classification performance of the proposed hybrid deep learning classifier is better with the input of aspect terms and sentiment terms in the GRU branch of the model.

4.2.3 Error Analysis

The confusion matrix is used to perform the detailed error analysis as shown in Fig. 9. It is clear from the matrix that few data are incorrectly classified. For example, 45 out of 619 cases of the neutral class were predicted positive. In the negative class, 66 of the 1633 data were mistakenly classified as neutral. Error analysis shows that the positive class achieved the highest rate of accurate classification (87.45 %) while the neutral class achieved the lowest (81.42 %). A possible reason for the high incorrect prediction in the neutral class may be the presence of a small number of neutral class tweets in the training dataset. Here, we consider natural data distributions to evaluate the performance of hybrid deep learning models, without introducing any data augmentation techniques that may disrupt the natural properties of the dataset and lead to a lack of interpretability of the model. However, sentiment analysis is widely subjective, depends on individual perceptions, and people may consider a tweet in many ways. Therefore, by creating a balanced dataset with a variety of high amounts of data, the incorrect predictions can be reduced to some degree.

5 Discussion

We present topic-based sentiment analysis to extract useful information from COVID-19-related tweets. We plot and count the number of positive, neutral,

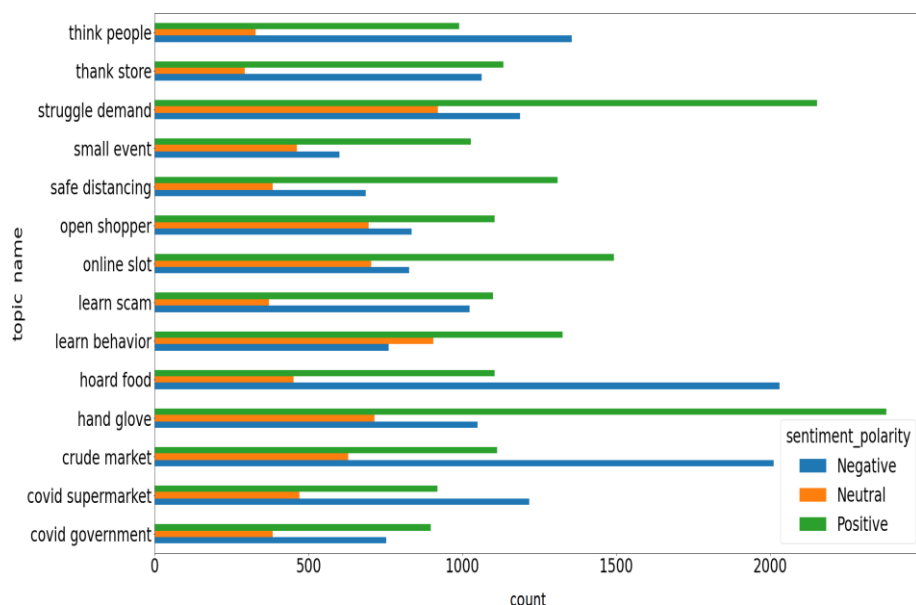


Fig. 10 Topic based sentiment analysis.

and negative tweets for each of the fourteen topics as shown in Fig. 10. The chart reports that in each topic most of the tweets are positive. In topics like covid supermarket, crude market, hoard food and think people we find that number of tweets that carry negative sentiment is higher. The topic-based sentiment analysis related to the COVID-19 tweets helps to understand people's perceptions and emotions for taking necessary steps to reduce suffering and compensate for the loss of resources. Government and policymakers can therefore take the necessary steps to understand the details of the human problems and their needs on the basis of public sentiments that reflect tweets on social media in order to minimize the detrimental effects of the COVID-19 pandemic.

It is difficult for the sentiment classifier to infer the correct polarity due to the limited textual context present in tweets. An interpretable topic distribution provided by LDA helps with the larger thematic structure of the data along with the sequential word dependencies. Thus domain-specific context changes for crises like COVID-19 are ignored by traditional sentiment models. Our proposed framework ensures alignment of sentiment classifications with thematic diversity in the dataset. Overall, in this paper, we present a data-driven [7] framework for the topic-based sentiment analysis to extract beneficial information from the COVID-19 related Twitter dataset. We use a laptop with a 2.0 GHz Core i3 processor and 4 GB RAM to execute our proposed technique. We write codes in the Python programming language and run our codes within the Google Colaboratory platform.

We

firmly

believe

that the proposed framework can be used effectively in other areas of applications such as agriculture, health, education, business, cyber security, etc. It requires a very large dataset and computational resources to train a transformer [59] from scratch like BERT, ALBERT, and RoBERTA to achieve better results, which we are not aligned in this study due to infrastructure constraints and dataset size. However, our objective is to measure the performance of the hybrid configuration of BiLSTM and GRU rather than focusing on the benchmark architecture. Hence, in this paper, we use classical deep learning classifiers such as BiLSTM and GRU to build the hybrid model [60]. To better understand the decision-making process, in the future we will incorporate explainable AI (XAI) techniques [48]. While our goal is to develop a methodological framework to balance interpretability and performance, we will analyze the potential and awareness of recent popular methods such as large language models (LLMs) [59] in our future work.

6 Conclusions

In this paper, we have presented a hybrid deep learning framework that effectively and automatically identifies key topic labels and corresponding sentiments from COVID-19 tweets. We have used the popular topic modeling algorithm LDA to extract hidden topics from tweets. Then we have used unigrams of the sentiment terms and aspect terms of the tweets to produce significant and meaningful topic labels on the basis of soft cosine similarity (SCS) values. Our proposed topic labeling method performs better and helps to categorize a huge amount of tweets corresponding to semantically similar topic labels to highlight user conversations. In our framework, we utilize the advantages of GRU, BiLSTM, and GAP models explained in the methodology section of this paper. Sequential dependencies are well captured by the BLSTM and GRU architectures but are critical for determining distant contextual information and encounters challenges handling ambiguous text without context modeling. We have also categorized sentiment polarities towards each topic. Overall, our proposed framework effectively facilitates topic-based sentiment analysis based on COVID-19 tweets and detects various issues related to the COVID-19 pandemic. We believe that the proposed framework will help policymakers to identify the problems associated with COVID-19 by analyzing tweets.

List of Abbreviations

BiLSTM: Bidirectional Long Short Term Memory;

BoW: Bag of Words;

CNN: Convolutional Neural Network;

GAP: Global Average Pooling;

GRU: Gated Recurrent Unit;

LDA: Latent Dirichlet Allocation;

LSTM: Long Short Term Memory;

RNN: Recurrent Neural Network;

SCS: Soft Cosine Similarity;

WHO: World Health Organization;

XAI: Explainable AI

Author Contributions: **K.T.S.** performed extensive work of literature review, framework construction, experimentation, and writing. **I.H.S.** provided the necessary guidance and supervision throughout this work. All authors have read and agreed to the published version of the manuscript.

Availability of Data and Material: Data used in this work can be made available upon reasonable request.

Consent for Publication

Not applicable.

Conflicts of Interests: The authors declare no conflict of interest.

Funding: None

Acknowledgments: This work has been done as a part of Master's thesis at Chittagong University of Engineering and Technology.

References

- [1] O'brien, Michael, Kathleen Moore, and Fiona McNicholas. "Social media spread during Covid-19: the pros and cons of likes and shares." *Ir Med J* 113, no. 4 (2020): 52.
- [2] Jahanbin, Kia, and Vahid Rahmanian. "Using twitter and web news mining to predict COVID-19 outbreak." *Asian Pacific journal of tropical medicine* 13, no. 8 (2020): 378.
- [3] Malta, Monica, Anne W. Rimoin, and Steffanie A. Strathdee. "The coronavirus 2019-nCoV epidemic: Is hindsight 20/20?." *EClinicalMedicine* 20 (2020).
- [4] Sohrabi, Catrin, Zaid Alsafi, Niamh O'Neill, Mehdi Khan, Ahmed Kerwan, Ahmed Al-Jabir, Christos Iosifidis, and Riaz Agha. "World Health Organization declares global emergency: A review of the 2019 novel coronavirus (COVID-19)." *International journal of surgery* 76 (2020): 71-76.
- [5] El Rahman, Sahar A., Feddah Alhumaidi AlOtaibi, and Wejdan Abdullah AlShehri. "Sentiment analysis of twitter data." In *2019 international conference on computer and information sciences (ICCIS)*, pp. 1-4. IEEE, 2019.
- [6] Mehta, Pooja, and Sharnil Pandya. "A review on sentiment analysis methodologies, practices and applications." *International Journal of Scientific and Technology Research* 9, no. 2 (2020): 601-609.
- [7] Sarker, Iqbal H. "Data science and analytics: an overview from data-driven smart computing, decision-making and applications perspective." *SN Computer Science* 2, no. 5 (2021): 1-22.
- [8] Blei, David M., Andrew Y. Ng, and Michael I. Jordan. "Latent dirichlet allocation." *Journal of machine Learning research* 3, no. Jan (2003): 993-1022.
- [9] Shahriar, K.T., Moni, M.A., Hoque, M.M., Islam, M.N., Sarker, I.H. (2022). SATLabel: A Framework for Sentiment and Aspect Terms Based Automatic

- Topic Labelling. In: Skala, V., Singh, T.P., Choudhury, T., Tomar, R., Abul Bashar, M. (eds) Machine Intelligence and Data Science Applications. Lecture Notes on Data Engineering and Communications Technologies, vol 132. Springer, Singapore.
- [10] Zeng, Daojian, Yuan Dai, Feng Li, R. Simon Sherratt, and Jin Wang. "Adversarial learning for distant supervised relation extraction." *Computers, Materials & Continua* 55, no. 1 (2018): 121-136.
- [11] Zhang, Lei, Shuai Wang, and Bing Liu. "Deep learning for sentiment analysis: A survey." *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery* 8, no. 4 (2018): e1253.
- [12] Zhang, Yuteng, Wenpeng Lu, Weihua Ou, Guoqiang Zhang, Xu Zhang, Jinyong Cheng, and Weiyu Zhang. "Chinese medical question answer selection via hybrid models based on CNN and GRU." *Multimedia tools and applications* 79, no. 21 (2020): 14751-14776.
- [13] Sharfuddin, Abdullah Aziz, Md Nafis Tihami, and Md Saiful Islam. "A deep recurrent neural network with bilstm model for sentiment classification." In 2018 International conference on Bangla speech and language processing (ICBSLP), pp. 1-4. IEEE, 2018.
- [14] Wang, Huwen, Zezhou Wang, Yinqiao Dong, Ruijie Chang, Chen Xu, Xiaoyue Yu, Shuxian Zhang et al. "Phase-adjusted estimation of the number of coronavirus disease 2019 cases in Wuhan, China." *Cell discovery* 6, no. 1 (2020): 1-8.
- [15] Abd-Alrazaq, Alaa, Dari Alhuwail, Mowafa Househ, Mounir Hamdi, and Zubair Shah. "Top concerns of tweeters during the COVID-19 pandemic: infoveillance study." *Journal of medical Internet research* 22, no. 4 (2020): e19016.
- [16] Hingmire, Swapnil, Sandeep Chougule, Girish K. Palshikar, and Sutanu Chakraborti. "Document classification by topic labeling." In *Proceedings of the 36th international ACM SIGIR conference on Research and development in information retrieval*, pp. 877-880. 2013.
- [17] Wang, Bo, Maria Liakata, Arkaitz Zubiaga, and Rob Procter. "A hierarchical topic modelling approach for tweet clustering." In *International Conference on Social Informatics*, pp. 378-390. Springer, Cham, 2017.
- [18] Hourani, Asseel S. "Arabic Topic Labeling using Naïve Bayes (NB)." In 2021 12th International Conference on Information and Communication Systems (ICICS), pp. 478-479. IEEE, 2021.
- [19] Asmussen, Claus Boye, and Charles Møller. "Smart literature review: a practical topic modelling approach to exploratory literature review." *Journal of Big Data* 6, no. 1 (2019): 1-18.
- [20] Maier, Daniel, Annie Waldherr, Peter Miltner, Gregor Wiedemann, Andreas Niekler, Alexa Keinert, Barbara Pfetsch et al. "Applying LDA topic modeling in communication research: Toward a valid and reliable methodology." *Communication Methods and Measures* 12, no. 2-3 (2018): 93-118.
- [21] Guo, Lei, Chris J. Vargo, Zixuan Pan, Weicong Ding, and Prakash Ishwar. "Big social data analytics in journalism and mass communication: Comparing dictionary-based text analysis and unsupervised topic modeling." *Journalism & Mass Communication Quarterly* 93, no. 2 (2016): 332-359.

- [22] Ordun, Catherine, Sanjay Purushotham, and Edward Raff. "Exploratory analysis of covid-19 tweets using topic modeling, umap, and digraphs." *arXiv preprint arXiv:2005.03082* (2020).
- [23] S. Prabhakaran, Topic Modeling with Gensim (Python), Machine Learning Plus (2018).
- [24] Khurana, Himja, and Sanjib Kumar Sahu. "Bat inspired sentiment analysis of Twitter data." In *Progress in Advanced Computing and Intelligent Engineering*, pp. 639-650. Springer, Singapore, 2018.
- [25] Prabha, M. Indhraom, and G. Umarani Srikanth. "Survey of sentiment analysis using deep learning techniques." In *2019 1st International Conference on Innovations in Information and Communication Technology (ICIICT)*, pp. 1-9. IEEE, 2019.
- [26] Behera, Ranjan Kumar, Monalisa Jena, Santanu Kumar Rath, and Sanjay Misra. "Co-LSTM: Convolutional LSTM model for sentiment analysis in social big data." *Information Processing & Management* 58, no. 1 (2021): 102435.
- [27] Poria, Soujanya, Navonil Majumder, Devamanyu Hazarika, Erik Cambria, Alexander Gelbukh, and Amir Hussain. "Multimodal sentiment analysis: Addressing key issues and setting up the baselines." *IEEE Intelligent Systems* 33, no. 6 (2018): 17-25.
- [28] Kaur, Simranpreet, Pallavi Kaul, and Pooya Moradian Zadeh. "Monitoring the dynamics of emotions during COVID-19 using Twitter data." *Procedia Computer Science* 177 (2020): 423-430.
- [29] Nemes, László, and Attila Kiss. "Social media sentiment analysis based on COVID-19." *Journal of Information and Telecommunication* 5, no. 1 (2021): 1-15.
- [30] Jiang, Qingnan, Lei Chen, Wei Zhao, and Min Yang. "Toward aspect-level sentiment modification without parallel data." *IEEE Intelligent Systems* 36, no. 1 (2021): 75-81.
- [31] Imran, Ali Shariq, Sher Muhammad Daudpota, Zenun Kastrati, and Rakhi Batra. "Cross-cultural polarity and emotion detection using sentiment analysis and deep learning on COVID-19 related tweets." *Ieee Access* 8 (2020): 181074-181090.
- [32] Barkur, Gopalkrishna, and Giridhar B. Kamath. "Sentiment analysis of nationwide lockdown due to COVID 19 outbreak: Evidence from India." *Asian journal of psychiatry* 51 (2020): 102089.
- [33] Rustam, Furqan, Madiha Khalid, Waqar Aslam, Vaibhav Rupapara, Arif Mehmood, and Gyu Sang Choi. "A performance comparison of supervised machine learning models for Covid-19 tweets sentiment analysis." *Plos one* 16, no. 2 (2021): e0245909.
- [34] Naseem, Usman, Imran Razzak, Matloob Khushi, Peter W. Eklund, and Jinman Kim. "COVIDSenti: A large-scale benchmark Twitter data set for COVID-19 sentiment analysis." *IEEE Transactions on Computational Social Systems* 8, no. 4 (2021): 1003-1015.
- [35] Pennington, Jeffrey, Richard Socher, and Christopher D. Manning.

- "Glove: Global vectors for word representation." In Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP), pp. 1532-1543. 2014.
- [36] Bojanowski, Piotr, Edouard Grave, Armand Joulin, and Tomas Mikolov. "Enriching word vectors with subword information." *Transactions of the association for computational linguistics* 5 (2017): 135-146.
- [37] Mikolov, Tomas, Kai Chen, Greg Corrado, and Jeffrey Dean. "Efficient estimation of word representations in vector space." *arXiv preprint arXiv:1301.3781* (2013).
- [38] Stappen, Lukas, Alice Baird, Erik Cambria, and Björn W. Schuller. "Sentiment analysis and topic recognition in video transcriptions." *IEEE Intelligent Systems* 36, no. 2 (2021): 88-95.
- [39] Jelodar, Hamed, Yongli Wang, Rita Orji, and Shucheng Huang. "Deep sentiment classification and topic discovery on novel coronavirus or COVID-19 online discussions: NLP using LSTM recurrent neural network approach." *IEEE Journal of Biomedical and Health Informatics* 24, no. 10 (2020): 2733-2742.
- [40] Ahmed, Md Shoaib, Tanjim Taharat Aurpa, and Md Musfique Anwar. "Detecting sentiment dynamics and clusters of Twitter users for trending topics in COVID-19 pandemic." *Plos one* 16, no. 8 (2021): e0253300.
- [41] Xue, Jia, Junxiang Chen, Chen Chen, Chengda Zheng, Sijia Li, and Ting-shao Zhu. "Public discourse and sentiment during the COVID 19 pandemic: Using Latent Dirichlet Allocation for topic modeling on Twitter." *PloS one* 15, no. 9 (2020): e0239441.
- [42] Medford, Richard J., Sameh N. Saleh, Andrew Sumarsono, Trish M. Perl, and Christoph U. Lehmann. "An "infodemic": leveraging high-volume Twitter data to understand early public sentiment for the coronavirus disease 2019 outbreak." In *Open forum infectious diseases*, vol. 7, no. 7, p. ofaa258. US: Oxford University Press, 2020.
- [43] Xiang, Xiaoling, Xuan Lu, Alex Halavanau, Jia Xue, Yihang Sun, Patrick Ho Lam Lai, and Zhenke Wu. "Modern senicide in the face of a pandemic: An examination of public discourse and sentiment about older adults and COVID-19 using machine learning." *The Journals of Gerontology: Series B* 76, no. 4 (2021): e190-e200.
- [44] Satu, Md Shahriare, Md Imran Khan, Mufti Mahmud, Shahadat Uddin, Matthew A. Summers, Julian MW Quinn, and Mohammad Ali Moni. "TClustVID: A novel machine learning classification model to investigate topics and sentiment in COVID-19 tweets." *Knowledge-Based Systems* 226 (2021): 107126.
- [45] Chandrasekaran, Ranganathan, Vikalp Mehta, Tejali Valkunde, and Evangelos Moustakas. "Topics, trends, and sentiments of tweets about the COVID-19 pandemic: Temporal infoveillance study." *Journal of medical Internet research* 22, no. 10 (2020): e22624.
- [46] Performing Sentiment Analysis Using Twitter Data!, <https://www.analyticsvidhya.com/blog/2021/07/>

- [performing-sentiment-analysis-using-twitter-data/](#), Accessed: 2022-09-28.
- [47] Wang, Wenya, Sinno Jialin Pan, Daniel Dahlmeier, and Xiaokui Xiao. "Coupled multi-layer attentions for co-extraction of aspect and opinion terms." In Proceedings of the AAAI conference on artificial intelligence, vol. 31, no. 1. 2017.
 - [48] Sarker, Iqbal H. "AI-driven cybersecurity and threat intelligence: cyber automation, intelligent decision-making and explainability." Springer Nature, 2024.
 - [49] Rustam, Furqan, Arif Mehmood, Muhammad Ahmad, Saleem Ullah, Dost Muhammad Khan, and Gyu Sang Choi. "Classification of shopify app user reviews using novel multi text features." IEEE Access 8 (2020): 30234-30244.
 - [50] Habibabadi, Sedigheh Khademi, and Pari Delir Haghighi. "Topic modelling for identification of vaccine reactions in twitter." In Proceedings of the Australasian Computer Science Week Multiconference, pp. 1-10. 2019.
 - [51] Jang, Beakcheol, Inhwon Kim, and Jong Wook Kim. "Word2vec convolutional neural networks for classification of news articles and tweets." PloS one 14, no. 8 (2019): e0220976.
 - [52] Boussaadi, Smail, Hassina Aliane, and Pr Ouahabi Abdeldjalil. "The researchers profile with topic modeling." In 2020 IEEE 2nd International Conference on Electronics, Control, Optimization and Computer Science (ICECOCS), pp. 1-6. IEEE, 2020.
 - [53] Sidorov, Grigori, Alexander Gelbukh, Helena Gómez-Adorno, and David Pinto. "Soft similarity and soft cosine measure: Similarity of features in vector space model." Computación y Sistemas 18, no. 3 (2014): 491-504.
 - [54] Hochreiter, Sepp, and Jürgen Schmidhuber. "Long short-term memory." Neural computation 9, no. 8 (1997): 1735-1780.
 - [55] Siامي-Namini, Sima, Neda Tavakoli, and Akbar Siامي Namin. "The performance of LSTM and BiLSTM in forecasting time series." In 2019 IEEE International Conference on Big Data (Big Data), pp. 3285-3292. IEEE, 2019.
 - [56] Li, Jianjun, Yu Han, Ming Zhang, Gang Li, and Baohua Zhang. "Multi-scale residual network model combined with Global Average Pooling for action recognition." Multimedia Tools and Applications 81, no. 1 (2022): 1375-1393.
 - [57] Kamila, Sabyasachi, Mohammad Hasanuzzaman, Asif Ekbal, and Pushpak Bhattacharyya. "Resolution of grammatical tense into actual time, and its application in Time Perspective study in the tweet space." PloS one 14, no. 2 (2019): e0211872.
 - [58] Shahriar, Khandaker Tayef, Muhammad Nazrul Islam, Md Musfique Anwar, and Iqbal H. Sarker. "COVID-19 analytics: Towards the effect of vaccine brands through analyzing public sentiment of tweets." Informatics in medicine unlocked 31 (2022): 100969.
 - [59] Sarker, Iqbal H. "LLM potentiality and awareness: a position paper from the perspective of trustworthy and responsible AI modeling." Discover Artificial Intelligence 4, no. 1 (2024): 40.



- [60] Ezen-Can, Aysu. "A Comparison of LSTM and BERT for Small Corpus." arXiv preprint arXiv:2009.05451 (2020).