

Article

Exploring the Evolution of Sentiment in Spanish Pandemic Tweets: A Data Analysis Based on a Fine-Tuned BERT Architecture

Carlos Henríquez Miranda ^{1,*}, German Sanchez-Torres ¹  and Dixon Salcedo ² 

¹ Facultad de Ingeniería, Universidad del Magdalena; Santa Marta 470001, Colombia; gsanchez@unimagdalena.edu.co

² Department of Computer Science and Electronics, University of the Coast, Barranquilla 080020, Colombia

* Correspondence: chenriquezm@unimagdalena.edu.co

Abstract: The COVID-19 pandemic has had a significant impact on various aspects of society, including economic, health, political, and work-related domains. The pandemic has also caused an emotional effect on individuals, reflected in their opinions and comments on social media platforms, such as Twitter. This study explores the evolution of sentiment in Spanish pandemic tweets through a data analysis based on a fine-tuned BERT architecture. A total of six million tweets were collected using web scraping techniques, and pre-processing was applied to filter and clean the data. The fine-tuned BERT architecture was utilized to perform sentiment analysis, which allowed for a deep-learning approach to sentiment classification. The analysis results were graphically represented based on search criteria, such as “COVID-19” and “coronavirus”. This study reveals sentiment trends, significant concerns, relationship with announced news, public reactions, and information dissemination, among other aspects. These findings provide insight into the emotional impact of the COVID-19 pandemic on individuals and the corresponding impact on social media platforms.

Keywords: deep learning; fine-tuning; natural language processing; evolution of feelings



Citation: Miranda, C.H.;

Sanchez-Torres, G.; Salcedo, D.

Exploring the Evolution of Sentiment in Spanish Pandemic Tweets: A Data Analysis Based on a Fine-Tuned BERT Architecture. *Data* **2023**, *8*, 96. <https://doi.org/10.3390/data8060096>

Academic Editor: Joaquín Torres-Sospedra

Received: 15 March 2023

Revised: 26 April 2023

Accepted: 10 May 2023

Published: 29 May 2023



Copyright: © 2023 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

During the beginning of the COVID-19 pandemic, humanity was isolated from work-places, schools, recreational centers, parks, and sports venues, resulting in consequences for people’s mental health [1]. Problems, such as stress, anxiety, fear, bad mood, irritability, frustration, and boredom, among others, have arisen with the emergence of the pandemic [2]. At the start of the pandemic, there was very little information, leading many researchers to access data and information from various sources during the COVID-19 outbreak. All of this data reflected the point of view of individuals, organizations, and government agencies on the coronavirus [3].

In parallel, there was a progressive increase worldwide in the use of the internet and its various tools for sharing information in recent years [4]. Common people have shifted from consumers to producers of data through different sources, such as blogs, websites, social networks, and apps, among others [5]. Nowadays, social media platforms are the most widely-used tools for spreading data, for example, Twitter, a social media platform with over 500 million users worldwide, is where opinions and comments about any topic, ranging from criticisms of or congratulations for products and services to sports and political issues, are most widely spread [6].

During the pandemic, social media platforms began to play an important role in times of confinement, as they were tools that kept people connected with the world without leaving home [7]. In the case of Twitter, it was used to share opinions about Covid, show points of view, and in some cases, write comments based on the sentiments people were feeling at that moment [3]. For this reason, sentiment analysis techniques (SA) began to be applied to decipher and discover the emotions of different Twitter users [8] during the

pandemic. The techniques of SA have progressed in recent years, moving from dictionary strategies [9], lexicon [10], semantic similarity [11], and machine learning [12], to deep learning [13].

Within deep learning, in the context of opinion classification, BERT architecture has been used. BERT (Bidirectional Encoder Representations from Transformers) is a state-of-the-art language model architecture developed by Google in 2018 [14]. It is composed of a set of self-attention cells, similar to transformers, that enable the prediction of missing words in a sentence by using a bidirectional semantic representation of words, considering both the preceding and following context of the word. BERT uses a dynamic word representation system that adjusts to the context and produces a richer representation of words and their relationships.

As part of the fine-tuning process, the pre-trained BERT model is adapted to a specific dataset, in this case, a sentiment exploration dataset. During the fine-tuning process, the model is adjusted to predict the correct output for the training dataset. Once the fine-tuning process is complete, the model can be used to make predictions on new data. Fine-tuning BERT on a specific dataset is highly effective for managing many natural language processing tasks.

BERT with fine-tuning enables sentiment exploration because it allows the model to better understand the context in which language is used. The pre-trained model already has a deep understanding of language. Furthermore, the use of a bidirectional semantic representation allows the model to consider the context of words, which significantly improves the accuracy of sentiment exploration.

Regarding the language used in sentiment analysis systems, most current research has focused on performing analyses on COVID-19-related tweets in English using different techniques and methodologies [15–19]. Few studies have been conducted in Spanish and particularly in Latin-American dialects [15]. Therefore, this article seeks to show the results obtained from a sentiment analysis of six million comments obtained from Twitter between 2020 and 2021 by users based in Colombia. Throughout the process, web scraping extraction techniques, natural language processing, and deep learning were applied for sentiment classification. This study makes an important contribution to understanding the evolution of sentiments during the pandemic and how these sentiments are changing over time, which has implications for policymakers, health professionals, and social media users alike.

The rest of the article is organized as follows. Section 2 addresses the background and similar work, Section 3 describes the materials and methods, Section 4 presents the experiments and results, and finally, the conclusions are provided in Section 5.

2. Background

Currently, a large amount of data is produced worldwide that is attractive to various governmental, commercial, and industrial sectors. However, the manual extraction and processing of the information make this process very complex. Therefore, there have been efforts for some time to work on systems that allow for the automatic analysis of large amounts of data, based on advances in disciplines such as natural language processing (NLP), machine learning, data mining, and cloud computing, among others. Within NLP is sentiment analysis (SA), an area that seeks to analyze the opinions, sentiments, values, attitudes, and emotions of people towards entities, such as products, services, organizations, individuals, problems, events, themes, and their attributes [16]. SA has shown a great research trend in recent years, mostly in the English language [11,17]. However, recent contributions have been made in other languages, such as Spanish [18], French [19], and Chinese [13], among others.

The majority of SA approaches focus on detecting the general polarity (positive or negative) of a paragraph or complete text [20]. Other approaches include sentence-level classification of the expressed sentiments in each sentence [21] and aspect-based classification of sentiments to specific characteristics of an entity found in each sentence [11].

The most-used technique to address SA is machine learning (ML), which depends on the existence of pre-labeled training documents, i.e., those that have already been assigned a polarity. Within the literature related to ML, there are works such as the one by De Freitas et al. [22], which performed an ontology-supported analysis in the domain of movies and hotels in Portuguese. Other examples include the study by Steinberger et al. [20], which presented a supervised approach to restaurant reviews in Czech, and the study by Manek et al. [23], which proposed a system in the English language based on the GINI index of movies.

In recent years, more advanced and precise approaches have used deep learning [24–26]. Deep learning consists of algorithms based on neural networks that represent a high-level of generalization regarding data processing; this is achieved via the multiple layers accumulated between themselves by alternating linear and non-linear transformations [27].

Various social networks have been studied. The work by Bonifazi et al. [28], using information based on Reddit data, the concept of scope in social network analysis was addressed, proposing a multidimensional view that included temporal scope. It introduced the concept of sentiment scope and developed a general framework for analyzing it on any social network platform. Bonifazi et al. [29], presented three novel approaches to analyzing COVID-19 Reddit posts: dynamic classification, virtual subreddit creation, and user community identification, addressing gaps in the previous research.

Several works are found in the literature related to sentiment analysis on Twitter that revolves around COVID-19. There are works on English tweets [30], where 500,000 tweets were analyzed in April 2020, finding that 36% of people had optimistic opinions, while only around 14% were negative. Likewise, an analysis was carried out in March 2020 on 50,000 tweets from several countries, resulting in most people having a positive and hopeful approach [31]. However, there were cases of fear, sadness, and disgust worldwide. Another contribution is found in the study by Boon-Itt et al. [32], where over 100,000 tweets were analyzed between December 2019 and March 2020, resulting in people having a negative perspective towards COVID-19, and the findings of predominant themes, such as pandemic emergency and learning how to control the disease.

In other languages, there are contributions, where more than 6 million comments in English and Portuguese were analyzed, resulting in finding sentiment trends and relationships with announced news, as well as a comparison of human behavior in two different geographical locations affected by the pandemic [33]. In another study [3], data were processed from users who shared their location as being “Nepal” between 21 May 2020 and 31 May 2020. The results of the study concluded that while most people in Nepal adopted a positive and hopeful approach, there were cases of fear, sadness, and disgust. Another contribution is found in a study [34] where Twitter messages collected during the first months of the COVID-19 pandemic in Europe were analyzed, finding that lockdown announcements correlated with a deterioration of mood in almost all surveyed countries, recovering in a short time.

Other contributions have focused on a Twitter analysis considering several aspects, such as the impact of vaccines on the pandemic [35,36], the impact of COVID on the financial market [37], the use or non-use of masks [38], the effects of confinement [39], and the relationship of COVID with announcements and news [34], among others. In the case of Colombia, some contributions can be found, such as a study [40] where 38,000 Twitter posts on COVID-19 vaccination were analyzed, highlighting opposition to the government with feelings of anger. In addition, a study analyzed the sentiments of 72,000 Twitter comments related to isolation, identifying the most frequent themes and words, resulting in fear as the predominant sentiment during the entire confinement period [41].

In another study [42], the researcher proposed a new approach to sentiment analysis of COVID-19 news headlines using deep neural networks, with high accuracy and an analysis of over 73,000 pandemic-related tweets from six global channels. In study by Kumari et al. [43], the researcher proposed a deep learning mechanism for the sentiment analysis of COVID-19-related Twitter data, utilizing an intelligent lead-based BiLSTM

and intelligent lead optimization to eliminate loss and improve accuracy. The proposed mechanism outperformed the baseline KNN technique, as assessed by metrics such as accuracy, sensitivity, and specificity.

Using pre-trained models for Spanish text has seen significant advancements in recent years. A BERT-based pre-trained language model (BETO) was introduced, which was trained exclusively on Spanish data from various sources, including Wikipedia [44]. This model has been utilized in different applications for the sentiment analysis of Spanish tweets [45], achieving accuracy levels of approximately 65%. Furthermore, another researcher presented the CaTrBETO model, which combines a caption transformer (CaTr) and BETO to understand the sentiments in tweets by jointly analyzing images and text [46]. The impact of COVID-19 on mental and physical health was discussed in [25], with social media found to be inducing fear and anxiety. The study analyzed Twitter data using a Hybrid Heterogeneous Support Vector Machine (H-SVM) for sentiment classification, outperforming RNN and SVM. In another paper [47], the sentiment analysis of Twitter posts related to COVID-19 from March to mid-April 2020 was addressed, using natural language processing techniques and seven different LSTM-based deep learning models for analysis. The models were trained to classify tweets into three classes (negative, neutral, and positive), which is advanced compared with traditional machine learning classifiers. Similar works proposed different deep learning models. Jojoa et al. [48], proposed a distilBERT transformer model for detecting positive and negative sentiments in open-text survey responses during the COVID-19 pandemic. The study aimed to contribute to understanding people's sentiments during the pandemic and addressing future challenges related to lockdown. In another study [49], a new approach to the sentiment analysis of Moroccan tweets related to COVID-19 was proposed. The model achieved high accuracy (86%) and outperformed well-known machine learning algorithms. The study showed that users' sentiments change over time and are affected by the evolution of the epidemiological situation in Morocco.

In general, efforts have been made to develop systems that automatically analyze large amounts of data through natural language processing, machine learning, data mining, and cloud computing. The majority of the approaches focus on detecting polarity, while the more advanced techniques involve deep learning. There have been several contributions in various languages, including English, Spanish, French, and Chinese. Sentiment analysis has been applied to analyze tweets about COVID-19 in different languages, finding various sentiments such as fear, optimism, and disgust, among others.

3. Materials and Methods

Sentiment analysis has become an essential tool for understanding human behavior and its relation to relevant social events or phenomena, such as the COVID-19 pandemic. In this sense, the methodology described in this section is a significant contribution to the analysis of sentiments expressed in Spanish tweets collected on Twitter during the period of 2020–2021.

The methodology consisted of several phases (see Figure 1), starting with data extraction through web-scraping techniques on the Twitter platform. Next, the pre-processing stage took place, which involved cleaning and filtering the collected data, and eliminating irrelevant or duplicated words or terms. Fine-tuning of the BERT architecture was then applied, a deep learning technique that allows for training a language model to identify and classify sentiments expressed in Spanish tweets. Once sentiment extraction had been carried out, the data analysis phase proceeded, which involved interpreting and visualizing the obtained results. The methodology described allowed for the identification of trends and patterns in the evolution of sentiments over time.

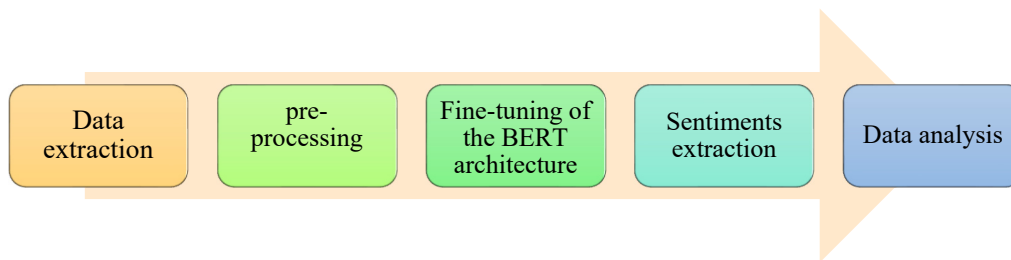


Figure 1. The general methodology applied to exploring the evolution of sentiment.

3.1. Data Extraction

The data extraction process was carried out using data scraping on Twitter. Scraping is a technique for data gathering that transforms the unstructured data found into structured data, which is then stored locally for further analysis. We used the Python language. The general steps for data extraction are shown in Figure 2. Algorithm 1 describes the general procedure applied in this phase.

Algorithm 1 Data extraction

```

1: fields = [id, data, user, text, retuits, likes, reply, quote, retweeted, lang, ... ]
2: Keywords = [COVID, COVID19, COVID-19, Coronavirus, SARS-CoV-2]
3: max_tuits = 10.000 each day
4: for each day from 01/04/2020 to 30/12/2021 do:
5:     search_str = "{} lang:{} since:{} until:{}",keywords,'es', day, day + 1
6:     tuits_generator = sntwitter.TwitterSearchScrapper(search_str).get_items()
7:     for count in range (0,max_tuits) do:
8:         tuit = tuits_generator.next()
9:         tuits_db = tuits.append(tuit[fields])
10: save_on_file (tuits_db)

```

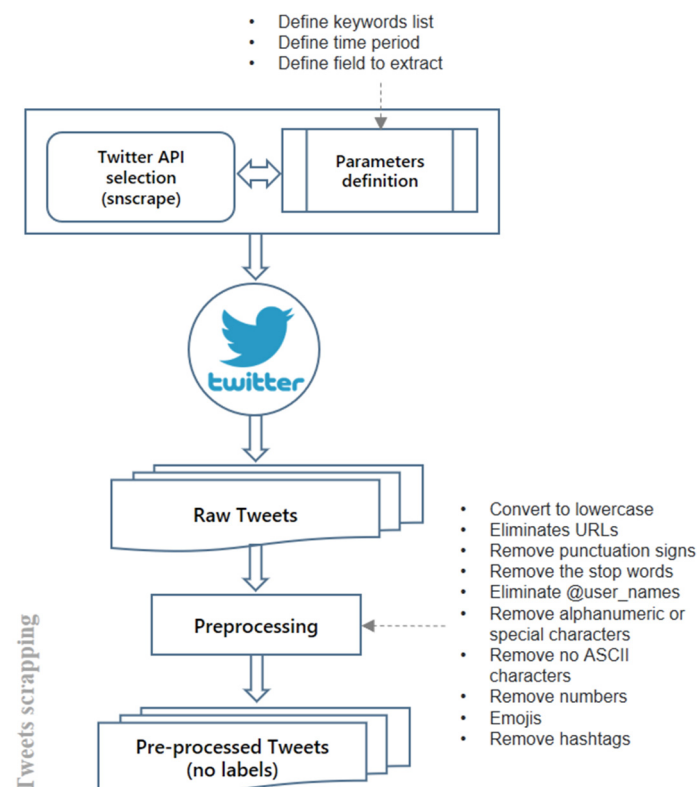


Figure 2. General description of data extraction process.

3.2. Data Pre-Processing

In this phase, a series of regular expressions were used to pre-process the text, removing irrelevant elements, such as mentions, links, Emojis, and punctuation marks. In addition, the hashtags were expanded and replaced by their corresponding text, removing the “#” symbol and separating the words that compose them. The hashtag expansion was carried out using a dictionary of 637,000 Spanish words. This pre-processing process permitted the analysis of sentiments in social networks since it allowed for reducing the complexity of the data by eliminating irrelevant elements and standardizing the information. We considered that expanding the hashtags and replacing them with their corresponding expanded text would permit us to hold key elements for sentiment identification. In this way, the application of sentiment analysis techniques, such as machine learning, natural language processing, and data mining, was facilitated.

3.3. Fine-Tuning of BERT Architecture

BERT has achieved state-of-the-art performance in various NLP tasks by employing a transformer-based architecture [50] (see Figure 3). The model is pre-trained on a large amount of text data and fine-tuned on a smaller labeled dataset for specific tasks [14].

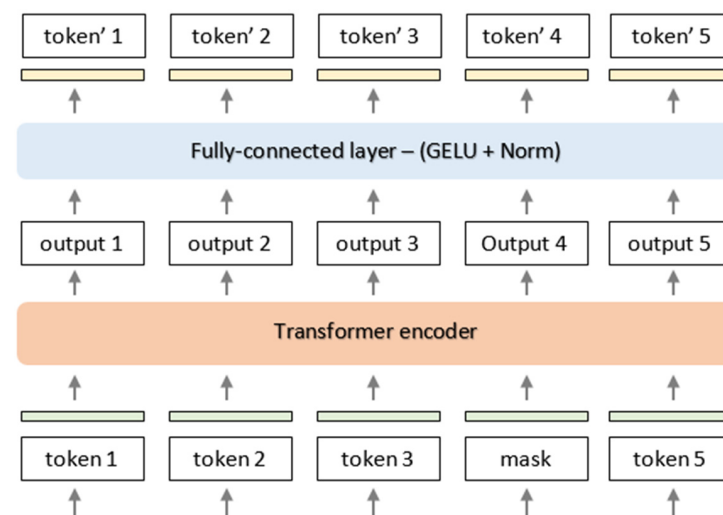


Figure 3. The Transformer based BERT base architecture.

One of the most important aspects of the BERT architecture is that it allows researchers to capture a vast amount of contextual information from a large corpus of text, including bidirectional language models and masked language models. This permitted us to understand the context and relationships between the words and phrases better than what would be possible via other architectures. Unlike context-free models, such as Word2Vec and GloVe, BERT offers contextual embeddings for every word within a text, which is a crucial aspect of its architecture. Additionally, BERT’s ability to handle long-term dependencies and its bidirectional nature allows it to perform well on tasks that require more complex language understanding, such as sentiment analysis. Finally, BERT’s architecture is highly modular and can be fine-tuned for specific tasks with relative ease, making it a versatile tool for natural language processing applications. The BERT architecture comprises three primary components, including an embedding layer, several transformer layers [51], and a task-specific output layer. During the embedding layer phase, an input word token would be transformed into a continuous vector space through an embedding matrix. In the subsequent Transformer layer, each word token would interact with all of the other word tokens via an attention mechanism, allowing for the exchange of information between them [52,53].

The main weakness of BERT is its inability to handle very long text sequences, as it considers only up to 512 tokens. This means that long sequences must be split into multiple

short sequences of 512 tokens, which can limit the technique's effectiveness for certain applications [50]. However, in the case of Twitter analysis, the length of the sequence to be analyzed is not a significant limitation due to the inherent characteristics of tweets.

Fine-tuning BERT with tweets in Spanish was necessary to improve its performance in analyzing Spanish text, which has its own unique linguistic characteristics, such as differences in syntax, morphology, and vocabulary, compared to other languages. By fine-tuning BERT on Spanish tweets, the model could better capture the nuances of the sentiments in Spanish text, and therefore produce more accurate results.

The general methodology for carrying out the fine-tuning process of the BERT-based architecture followed the classical guidelines of this type of process in the area of machine learning. It began by selecting a pre-trained architecture. Next, the dataset that was to be used in the model adjustment phase was defined in order to train and validate the model. The dataset was selected based on criteria such as language, origin (Twitter), and the amount of available data.

3.4. Sentiment Extraction

Each tweet was fed into a classification model based on the BERT architecture for sentiment analysis. The input data was preprocessed using techniques such as tokenization, attention masks, and padding to create suitable input features for the model. To perform this task efficiently on large volumes of data, the processing was carried out using the CuDF library, which leverages the power of GPU acceleration. The resulting sentiment scores were then used to generate insights into the attitudes and opinions expressed in the tweets, which could be used for a variety of applications in fields such as social media analytics, market research, and public opinion monitoring.

3.5. Data Analysis

In the data analysis phase, we used the CuDF library to preprocess and manipulate the data. In this phase, sentiments were analyzed over time, discriminating by sentiment type. We performed an analysis of relevant events during the pandemic and searched for patterns to determine if they influenced the generation of sentiments analyzed in the collected data. To show and analyze the evolution of the sentiments, we created visualizations, such as time-series graphs and heatmaps, highlighting changes and trends in the sentiments over time. We also performed statistical analysis and hypothesis testing to identify significant differences in the sentiments between time periods, sentiment types, and other relevant variables.

4. Results

The experiments were conducted on a computer equipped with an Intel Xeon Gold 6230R 2.10G 2933 MHz 26-core processor, 256 GB DDR4 2933 memory, and 2 NVIDIA RTX A6000 48 GB DH.

4.1. Data Extraction

This phase involved the collection of 6,306,621 tweets written in the Spanish language, with an approximate frequency of 10,000 tweets per day, during the analysis period. The tweet collection process was consistent throughout most of the analysis days, with a few exceptions where the API returned errors, leading to some missing data. The total number of tweets per day obtained is presented in Figure 4. The storage of this data was conducted using The Hierarchical Data Format version 5 (HDF5), an open-source file format that enables efficient storage of large and complex data, making it an ideal choice for the handling of the extensive dataset in this study.

Different search terms, such as COVID, COVID19, COVID-19, Coronavirus, and SARS-CoV-2, were used. The data was collected day by day, from the date 1 April 2020 to the date 30 December 2021, until 10,000 comments written in the Spanish language were obtained. The used parameters and the extracted fields are described in Table 1.

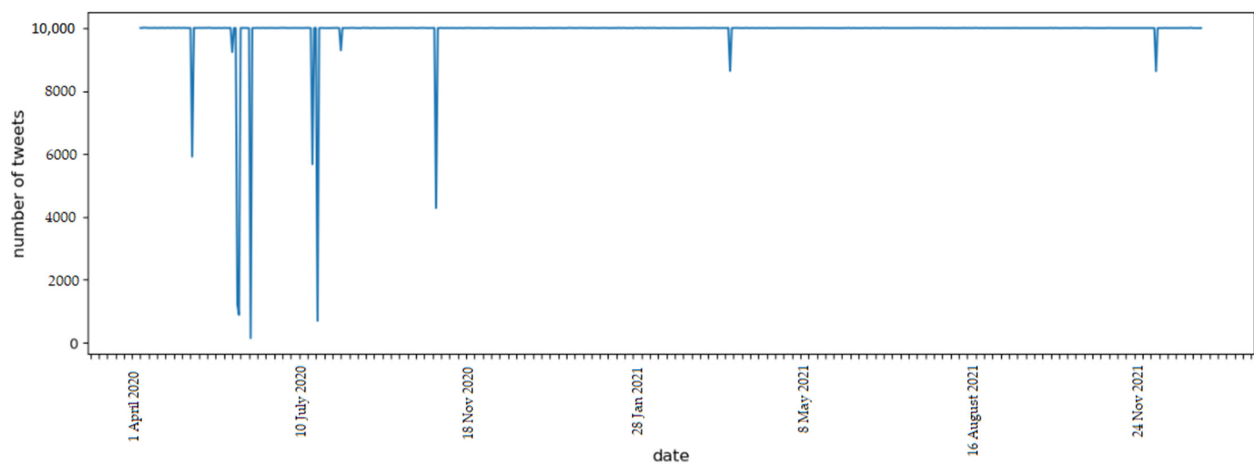


Figure 4. Total number of recollected tweets by day.

Table 1. This is a table. Tables should be placed in the main text near the first time they are cited.

Parameters	Description
keywords	COVID, COVID19, COVID-19, Coronavirus, SARS-CoV-2
Time period	1 April 2020 to 30 December 2021
Number of tuits	10.000 each day
Extracted fields	id, fecha, user, text, retuits, likes, reply, quote, retuiteado, lang, source, coordinates, user_location, user_verified

4.2. Data Pre-Processing

The Algorithm 2 procedure was applied to each of the collected tweets, using the CuDF library for GPU acceleration. In this work, the replacement of emojis with words associated with the sentiment they reflect was not addressed. Table 2 shows an example of the result of applying this procedure. Some hashtags were not properly expanded even though a dictionary of 637,000 words in Spanish was used. This study did not consider the analysis of possible errors in the expansion of hashtags and what impact this could have on the classification of the sentiments they represent.

Algorithm 2 Tuit preprocessing (parameter: text_tuit)

```

1: mentions = re.compile(r'@[^\s]+')
2: links = re.compile(r'https://[^\s]+')
3: emoji = re.compile(unicode for emojis)
4: punctuation = re.compile(r'[^\w\s\#]')
5: spaces = re.compile(r'\s+')
6: hashtags = re.compile(r'\B#(\w+)')
7: text_tuit = text_tuit.lower()
8: text_tuit = mentions.sub("", text_tuit)
9: text_tuit = links.sub("", text_tuit)
10: text_tuit = emoji.sub("", text_tuit)
11: text_tuit = punctuation.sub("", text_tuit)
12: text_tuit = spaces.sub("", text_tuit)
13: matches = hashtags.findall(text)
14: if len(matches):
15:     for match in matches do:
16:         match_extended = infer_spaces(unicode(str(match).lower()))
17:         text = re.sub(r'#+match, match_extended, text)
18: return text.strip()

```

Table 2. Example of preprocessed tweets. The example shows each step for preprocessing including hashtags explanations.

State	Description
Raw input	#Indiferentes ;) a la 🤔 pandemia 😞😞 Miles @ de poblanos @user en las calles sin #SanaDistancia ni medidas sanitarias para evitar el contagio del #COVID-19 ➡ https://t.co/Za2ERxygwY
Pre-processed result	indiferentes a la pandemia miles de poblanos en las calles sin sana distancia ni medidas sanitarias para evitar el contagio del COVID-19

4.3. Fine-Tuning of BERT Architecture

4.3.1. Training Data-Set Selection

In selecting a training dataset for fine-tuning the sentiment analysis model to apply to Spanish tweets, we considered data that had a range of sentiment labels to ensure that the model could learn to distinguish between positive, negative, and neutral sentiments. The Sentiment140 dataset was translated into Spanish, with the dataset containing a total of over 1.6 million tweets with positive, negative, and neutral labels. This dataset had the advantage of being large and diverse, covering a wide range of topics and sentiments. Another option was the Dataset de Twitter en Español—SemEval 2017 Task 4A, which consisted of tweets in Spanish with positive, negative, and neutral labels. This dataset was created for a sentiment analysis task and was specifically designed for training and evaluating sentiment analysis models. Finally, the Emotex dataset, which contained tweets in Spanish annotated with emotional categories, may have been useful for sentiment analysis in the context of emotional recognition. This dataset had the advantage of providing more detailed information about the emotions expressed in the tweets. Table 3 summarizes the key characteristics of these datasets.

Table 3. The dataset descriptions.

Dataset	Size	Sentiment Labels	Topic
Inter-TASS 2020	16 K	Positive, Negative, Neutral	General tweets
SemEval 2017	3.5 K	Positive, Negative, Neutral	General tweets
Twitter Sentiment Dataset	1.6 M	Positive, Negative, Neutral	General tweets

4.3.2. Pre-Trained BERT Model Selection

Various pre-trained models of Bidirectional Encoder Representations from Transformers (BERT) were available for processing text with tweet characteristics, each with their advantages and limitations. BERT-base, one of the most popular models, comprised 110 million parameters and could handle a wide range of natural language processing (NLP) tasks. BERT-large, on the other hand, was capable of processing large volumes of data due to its 340 million parameters. DistilBERT, a lighter version of BERT with approximately half the parameters of BERT-base, was more efficient in terms of memory and processing. RoBERTa, a BERT-based model, utilized a more advanced pre-training algorithm and had slightly better accuracy than BERT in various NLP tasks. Another architecture to consider was ALBERT, which used a factorized embedding parameterization to reduce the number of parameters in the model while maintaining its performance. ALBERT was especially suitable for resource-limited devices due to its reduced memory and computational requirements. Table 4 shows a description of the analyzed models, including their most relevant features.

For this work, we considered using BERT-base as it was a widely used and tested model that could handle a variety of NLP tasks. Additionally, the parameter count of this model was sufficient for effectively processing tweet data. The pre-trained model was initialized, with the parameters released in one study [54], and made available in another [55].

Table 4. Pre-trained BERT-based models.

Model	Param. ⁺	Architecture	T. Data ⁺⁺	Pre-Training Method
BERT-base	110 M	12-layer, 768-hidden, 12-heads	16 GB	Bidirectional transformer
BERT-large	340 M	24-layer, 1024-hidden, 16-heads	16 GB	Bidirectional transformer
RoBERTa-base	125 M	12-layer, 768-hidden, 12-heads	160 GB	BERT without NSP, dynamic masking
RoBERTa-large	355 M	24-layer, 1024-hidden, 16-heads	160 GB	BERT without NSP, dynamic masking
DistilBERT	66 M	6-layer, 768-hidden, 12-heads	16 GB	BERT Distillation
ALBERT-base	12 M	12-layer, 768-hidden, 12-heads	16 GB	BERT with reduced parameters and SOP
ALBERT-base	18 M	24-layer, 1024-hidden, 16-heads	16 GB	BERT with reduced parameters and SOP
BETO-base	110 M	12-layer, 768-hidden, 12-heads	16 GB	Bidirectional transformer, Spanish corpus
RoBERTuito	110 M	12-layer, 768-hidden, 12-heads	16 GB	Bidirectional transformer, Spanish corpus

⁺ Millions of parameters, ⁺⁺ Pre-training data used.

4.3.3. Selecting the Optimal Fine-Tuning Parameters

To select the optimal fine-tuning parameters for a BERT-base architecture, we considered the trade-off between model complexity and generalization. We used configurations based on the original recommendations from the model release. The approach was to use a small learning rate during the initial stages of fine-tuning, followed by a gradual increase in the learning rate as training progresses. Another important consideration was the choice of batch size, which should be large enough to ensure the efficient use of computational resources, but small enough to avoid overfitting. Additionally, the number of epochs used was carefully chosen to prevent overfitting and ensure optimal model performance. Finally, during the fine-tuning process, techniques such as early stopping and dropout were used to prevent overfitting and improve the generalization performance of the model. We performed a grid search on batch sizes of {8, 10, 16, 32} and learning rates of $\{5 \times 10^{-6}, 1 \times 10^{-5}, 31 \times 10^{-5}, 51 \times 10^{-5}\}$ to determine the optimal values for our specific task.

The optimal values are shown in Table 5.

Table 5. The optimal hyperparameter values for the fine-tuning process.

	Parameter	Value
BERT-base	Batch size	16
	Learning rate	1×10^{-5}
	Optimization function	Adam
	Number of epochs	2
	Dropout—Regularization	0.5
	Activation function	RELU

We formed a balanced dataset of 45,000 tweets and utilized a strategy to divide the dataset into 60% for training, 20% for validation, and 20% for the test set. The dataset was carefully designed to ensure an even distribution of the target classes. This partitioning technique was employed to evaluate the performance of the fine-tuned model. The goal was to maintain the generalization of the models by using a subset of the data for validation and testing. The results for the test set, consisting of 9000 tweets, are shown in Table 6.

Table 6. The result of test set performance.

Polarity	Precision	Recall	F1-Score	Pos	Neg	Neu	Acc
positive	0.902	0.920	0.911	2760	30	210	0.89
negative	0.936	0.880	0.907	60	2640	300	
neutral	0.837	0.870	0.853	240	150	2610	

4.4. Sentiment Extraction and Data Analysis

The sentiment analysis phase consisted of evaluating the sentiments using the trained model. Three classes were used regarding the types of sentiments: POSITIVE, NEGATIVE,

and NEUTRAL. The results show that the number of positive tweets was significantly lower than in the other categories. The category with the highest number of tweets was the neutral category, as shown in Figure 5. Some examples are shown in Table 7. Additionally, an emotion was classified in regards to each of the tweets, and the results are shown in Figure 6. It shows that the predominant emotion was anger, followed by the emotion of joy. However, sadness and fear also appeared as relevant sentiments.

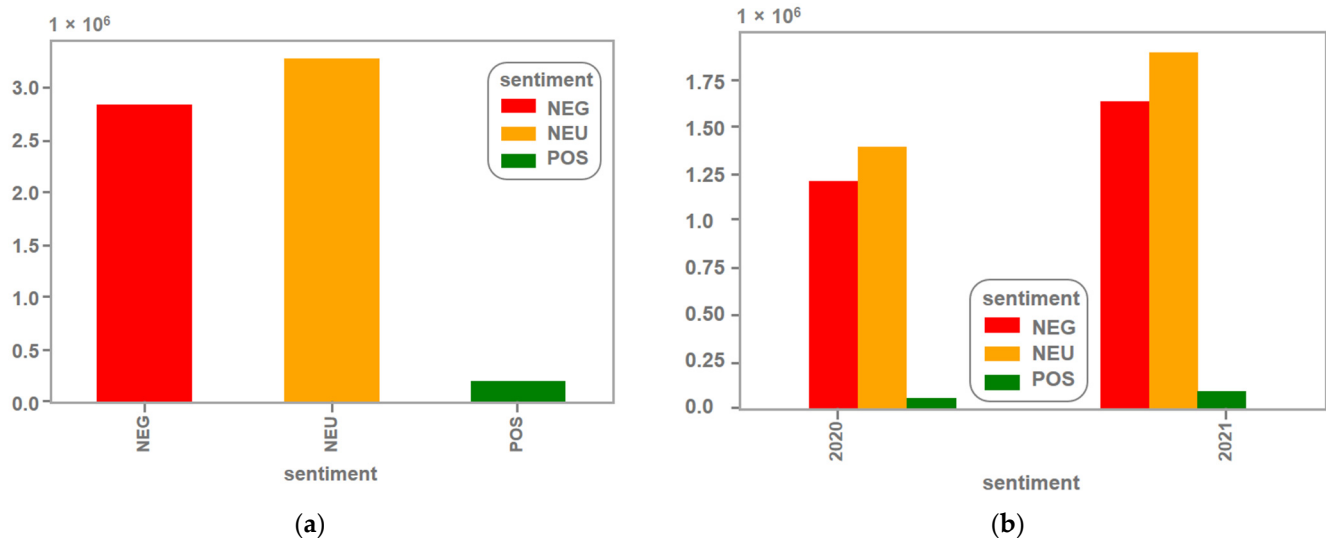


Figure 5. Numbers of tweets, (a) total number of tweets classified by sentiment during the defined time period, (b) numbers of tweets discriminate by year.

Table 7. Results examples of sentiment classification on Spanish tweets.

Sentiment	Example (Raw Input)
POS	Voy a rendir los finales que me quedaron pendientes en marzo por esta pandemia, estoy feliz y mañana arranco a estudiar 😊\nSe nota que me quiero recibir? @user No Jota flashaste, los mejores superheroes son los medicos que trabajaron sin descanso en esta pandemia 🙌 #NosCuidamosEntreTodos #JuntosSalimosDeEsta #QuedateEnCasa https://t.co/Tr6SizA1Yz @user Aldo, como compañía estamos realizando todos nuestros esfuerzos para que nuestros clientes puedan mantener sus cuentas al día en medio de esta pandemia y para eso hemos creado diferentes convenios los que puedes conocer acá (https://t.co/5JtKTKLwE1) Dejen de llorar: total libertad de circulación: CABA\n\nHace rato que joden con la libertad, se mueven como si no estuviéramos en una PANDEMIA\n\nY no cuento las marchas de los Garcas\n\n#QuedateEnCasa https://t.co/NWYu4rtfip
NEG	@user No sabes como me gustaría ver en una realidad paralela a los del otro lado....a ver si lo habrían hecho mejor. Entiende ESTO ES UNA PANDEMIA MUNDIAL...!!!!. Ahora depende de nosotros que no se desbande completamente. Realmente merecemos todo esto que nos ha traído el 2020 🤔😞 como raza humana y más aún viendo todo lo que está pasando en USA us y en plena PANDEMIA para colmo. \n\nNo hemos aprendido 🙄 \n\n#JusticeForGeorgeFloyd
NEU	Cambios en horarios de operación de establecimientos y espacios recreativos para evitar incrementos en contagios por COVID-19, seguimos en naranja: Del Mazo https://t.co/j7CTOH93Xq https://t.co/3CgtFUaBj2 Ministerio de Salud Pública: Cuba reporta 47 nuevos casos de COVID-19 para un acumulado de 3093 en el país https://t.co/HWgInU65hi Via @user Provincia Santa Fe sin nuevos casos de COVID-19 https://t.co/aNilj3ONYy https://t.co/FdDIA8JeaP

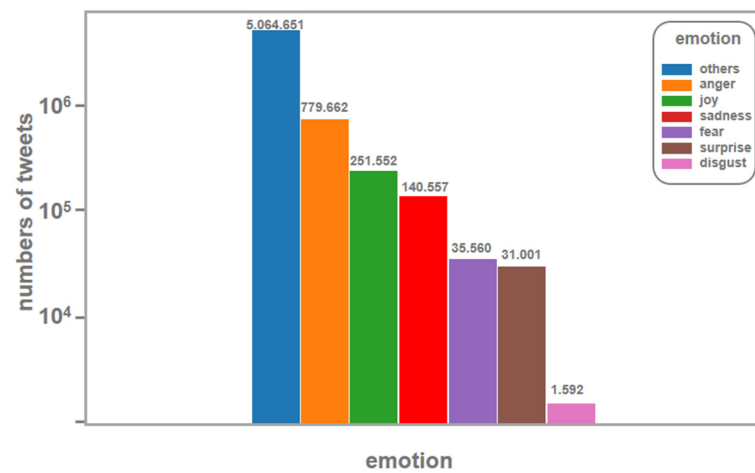


Figure 6. Emotion classified in tweets dataset.

Figure 7 shows the number of tweets in each sentiment class. Positive tweets were few compared to the other classes, while neutral and negative ones were significantly correlated. It is probable that this category is where the classifier presented the most uncertainty. In order to ascertain if any specific users had a significant impact on the daily trend of the data collected, the number of unique users in the group of tweets was also evaluated. Figure 8 shows the daily average of the tweets in blue together with the daily average of unique users in green.



Figure 7. Daily numbers of tweets by class.



Figure 8. Daily numbers of different users (green line) and total number of tweets per day.

The fact that there is a difference of approximately 2000 between these lines suggests that some tweets originated from users who were similar. Yet, there were around 8000 unique users per day, suggesting heterogeneity in the data gathered.

Using Emojis and Pre-Trained Language Model in Spanish

A pre-trained model called roBerTuito [56] was used, having been trained on a specific Spanish corpus for Twitter and incorporating the interpretation of emojis. Therefore, in this section, we will present tweets without removing the emojis and analyze the results.

The roBerTuito model decreased the number of tweets classified as neutral and increased the number of negative and positive tweets (see Figure 9).

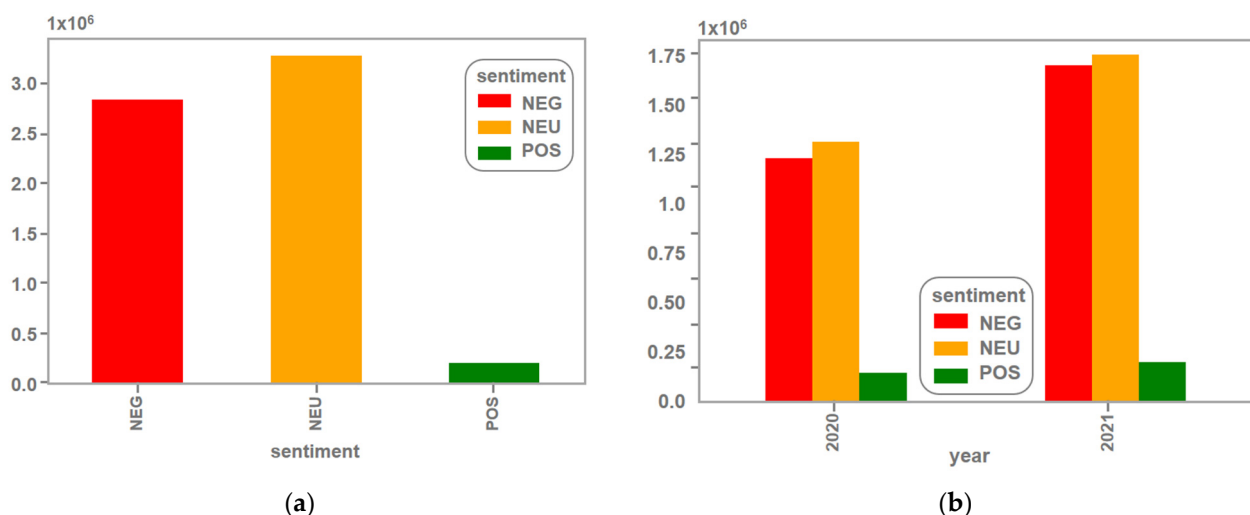


Figure 9. Number of tweets, (a) total number of tweets classified by sentiment during the defined time period, (b) number of tweets discriminate by year using RoBerTuito.

The use of pre-trained models in sentiment analyses poses challenges when attempting to measure their precision, due to the absence of a manually labeled reference dataset. Our collected dataset for sentiment analysis was not manually labeled. Hence, it was difficult to evaluate the precision of the pre-trained models used.

In the fine-tuning process of the models, a specific set of tweets was used for training. According to one researcher [56], the datasets used to train the RoBERTuito model were the same datasets used to fine-tune the BERT-base model, thereby invalidating the statistical comparisons of both models. Due to the impossibility of establishing a comparative statistical framework, qualitative comparisons will be made based on the results obtained in this study (see Table 7). This will allow for a descriptive and subjective evaluation of the performance of the pre-trained models in the task of sentiment analysis.

In Table 8, examples are shown of tweets that were classified as neutral by the Bert-base model and whose sentiment classification varied with the RoBerTuito model. Some examples were accurately classified; however, in regards to others, the sentiment assignment was not as clear. Emojis can impose a burden on the classification that is not easy to quantify due to the high informality and freedom of their use.

Table 8. A comparison of tweet classifications between the two models.

Raw Tweet *	Bert-Base	RoBerTuito
No hay casos de COVID-19 en la Policía Estatal: SSC https://t.co/pfTWw6lCU https://t.co/ppZy	NEU	NEG
🇺🇸 Estados Unidos autoriza el Remdesivir para tratar el COVID-19 y las acciones de Gilead Sciences se disparan 🗨️ #ZirigozaGroup #e-commerce #marketing #ENKIL.es #8020 📞 https://t.co/VRHfUDj	NEU	NEG
#Corrientes 🇵🇷 🚫 “Cerrado por COVID-19”: cercaron una cuadra de un barrio por caso positivo 📞 https://t.co/kutVZbSM	NEU	NEG
Ya empezaron los milagros en las iglesias curando el COVID-19 🙌🙌🙌 https://t.co/HFtsSTG10W	NEU	NEG

Table 8. Cont.

Raw Tweet *	Bert-Base	RoBerTuito
🦠 PANDEMIA DEL COVID-19: ¿Cómo pueden ayudar los periodistas a frenarla? https://t.co/IoIVmwGt	NEU	NEG
Los 🧑🧑🧑🧑 causantes del COVID-19 según los más letrados de este país. https://t.co/sfXAAZxhKk	NEU	NEG
La señora de la alita ya reabrió su negocio, solo delivery, obvio tienda pero es 😊😊😊😊, encima la casa de al lado y 2 casas más abajo de la suya tienen personas con covid. 😞😞😞😞	NEU	NEG
@user 🧑🧑🧑🧑 HB2U en tiempos de COVID 🧐... espero que esté pasando consentido en Familia. Cheers 🍷!!!	NEU	POS
🧑🧑 El día de hoy estaremos realizando la Sanitización de avenidas y calles de nuestro municipio como parte de las acciones para prevenir el COVID-19. #QuédateEnCasa 🏠 #GobiernoConValores #SanPedroPochutla #Oaxaca https://t.co/luyrO2kA	NEU	POS
👉😊 La única forma de combatir al #Covid19 es con el correcto #LavadoDeManos. Si todos seguimos los consejos de sanidad y #SanaDistancia podremos salir más rápido de esta pandemia. #QuédateEnCasa https://t.co/7he8JjNBx	NEU	POS
@COVID-19Time El planeta ha respirado 😊	NEU	POS
Así nomas!! Ya listo listo esperando el banderaso 🇲🇽🇲🇽🇲🇽 #COVID-19 #MiercolesDeGanarSeguidores https://t.co/GiKKzOge	NEU	POS
El uso del tapabocas le ha dado todo el poder a la mirada 👁👁🧐	NEU	POS
#YoMeQuedoEnCasa #COVID-19		
📋👉 Medidas de apoyo al sector cultural y deportivo ante el COVID-19. !!Objetivo: proteger a un sector fundamental en nuestro país para que nadie se quede atrás ante la emergencia sanitaria y social. #ProtegemosLaCultura #ProtegemosElDeporte #EsteVirusLoParamosUnidos https://t.co/EadvNWph	NEU	POS

* Some mentions and web addresses were modified to anonymize the text.

5. Discussion

5.1. Bias and Representativeness of the Data on Twitter

The utilization of only Twitter as a data source in this research may incur a bias due to the platform's characteristics and its usage patterns among users. It is noteworthy that Twitter does not encompass the entire Spanish-speaking population. The profiling of users based on variables such as age, gender, socioeconomic status, and education, among other factors, would be crucial in assessing the data's representativeness on Twitter. Furthermore, certain users may exhibit a more dominant presence on the platform than others, thus impeding the generalizability of the findings to the broader population. Notably, the analysis of Twitter disregards individuals who do not utilize this platform, which may impart a bias on the results and generate incomplete or inaccurate conclusions. Consequently, the outcomes of this study may not be generalizable to the wider Spanish population and may not be a true reflection of reality.

Another interesting aspect regarding the data is the discrimination by country. However, Twitter's own API reports that there is no guarantee that the reported location is accurate. In some cases, if the user has not granted location permissions, this value is often estimated and has low levels of precision.

Unfortunately, such discrimination against social groups accessing the platform is not possible through the API. In cases where it is possible, it may even have significant levels of imprecision, in the sense that some characteristics of users are reported but are also not validated by the social network itself.

Alternative social media and digital platforms may also be indicative of the general population, such as Facebook, Instagram, TikTok, and others. Each one has distinct attributes and attracts unique user groups, thereby resulting in potential variability in the biases based on the platforms employed.

5.2. Considering Emojis in Sentiment Analysis

We consider the fact that emojis play a significant role in modern communication, as they provide additional contextual information and can enrich the interpretations of text. However, we also recognize that the appropriate use of emojis in sentiment analysis is a challenging aspect and does not necessarily guarantee a significant improvement in the accuracy of classification models.

In our research, we have observed that emojis can be useful in certain contexts, as demonstrated by some researchers [57,58] who found that incorporating emojis in NLP models improved sentiment analysis performance. However, other studies have reported less conclusive results. For instance, in [59], the use of emojis and hashtags was found not to significantly improve performance.

One of the most relevant aspects when including emojis is addressing the challenges that arise in sentiment analysis. These challenges include differences in emoji representations across different platforms, variations in emoji usage across cultures and social contexts, and how emojis can affect the polarity of a message. The sentiment of a tweet and the sentiment conveyed by an emoji should be analyzed independently; we should also consider how their combination affects the overall sentiment of the message [60].

In another study [61], the researchers established that emojis, while ubiquitous in modern digital communication, presented differences in their usage across languages and countries. The study analyzed the relationship between emoji usage and linguistic and geographic factors, using millions of tweets to do so. The results showed that the diversity in emoji usage varied across different languages and countries. Another researcher [62] highlighted variations in the interpretations of emojis across different cultures and communities, which could impact the effectiveness of sentiment analysis when emojis are included. In a study carried out in two Spanish cities [63], it was concluded that the overall semantics of the subset of emojis studied were preserved in both of these cities, but some of them were interpreted differently. In another study [64], researchers discussed the ambiguity in the meaning of emojis and the need for a sense network for emojis, indicating that there may be difficulties in interpreting the meaning and sentiment of emojis in sentiment analysis.

In our study, we examined over six million Spanish-language tweets from diverse sources, ranging from Mexico to Argentina, and even including sources in Europe. Given the diversity of cultures represented, we believe that the appropriate use of semantic meaning in terms of emojis is an important topic. However, due to the scope of this study, we made the principal decision to exclude emojis to prevent the analysis from becoming overly complex. Nevertheless, we reported the results of processing the tweets with a model that includes emoji interpretation. The findings on the benefits are inconclusive, and a much deeper analysis is required in this regard. Further research into cultural differences in emoji use and interpretations could provide valuable insights into the effectiveness of incorporating emojis in sentiment analyses for various populations.

6. Conclusions

This study collected a total of 6,306,621 tweets in the Spanish language, using various search terms related to COVID-19. The data was stored using The Hierarchical Data Format version 5 (HDF5), an open-source file format suitable for handling large and complex data.

The collected tweets were pre-processed using the CuDF library for GPU acceleration. The sentiment analysis was performed using a BERT-base architecture trained on a dataset consisting of over 1.6 million tweets with positive, negative, and neutral labels. A balanced dataset of 45,000 tweets was used for fine-tuning the model, which was then evaluated on a test set of 9000 tweets.

The results of this study revealed that the BERT-based model achieved an accuracy of 89%, with the neutral category having the highest number of tweets, followed by the negative category, while positive tweets were relatively few. This study also found that the predominant emotion expressed in the tweets was anger, followed by joy, while sadness

and fear were also significant. The analysis of unique users revealed a high level of heterogeneity in the data collected.

Regarding future work, it is proposed that researchers should replace Emojis with words that correspond to the sentiments that they represent. Replacing Emojis with these words would prevent the loss of representative sentiment information. In the same sense, a sarcasm detector could decrease the false positive rate, considering that this type of text is frequently employed on a social network, such as Twitter.

Author Contributions: Conceptualization, C.H.M. and G.S.-T.; methodology, C.H.M., G.S.-T. and D.S.; software, C.H.M. and G.S.-T.; validation, C.H.M., D.S. and G.S.-T.; investigation, C.H.M.; resources, G.S.-T.; data curation, C.H.M.; writing—original draft preparation, C.H.M. and G.S.-T.; writing—review and editing, C.H.M.; visualization, G.S.-T.; supervision, D.S.; funding acquisition, C.H.M. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by Universidad del Magdalena, grant number VIN 2022117 and the APC was funded by Universidad del Magdalena.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The following supporting information can be downloaded at: <https://github.com/sanchezgt/COVADS> (accessed on 14 March 2023).

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

References

- Smith, B.; Lim, M. How the COVID-19 Pandemic Is Focusing Attention on Loneliness and Social Isolation. *Public Health Res. Pract.* **2020**, *30*, 3022008. [CrossRef] [PubMed]
- Hwang, T.-J.; Rabheru, K.; Peisah, C.; Reichman, W.; Ikeda, M. Loneliness and Social Isolation during the COVID-19 Pandemic. *Int. Psychogeriatr.* **2020**, *32*, 1217–1220. [CrossRef]
- Pokharel, B.P. Twitter Sentiment Analysis During Covid-19 Outbreak in Nepal 2020. *SSRN* **2020**. [CrossRef]
- Henríquez, C.; Guzmán, J.; Salcedo, D. Minería de Opiniones Basado En La Adaptación al Español de ANEW Sobre Opiniones Acerca de Hoteles Opinion. *Proces. Leng. Nat.* **2016**, *41*, 25–32.
- Henríquez Miranda, C.; Guzmán, J. Information Extraction from the Web to Identify Actions of an Automated Planning Domain Mode. *Ingeniare J.* **2015**, *23*, 439–448.
- Pak, A.; Paroubek, P. Twitter as a Corpus for Sentiment Analysis and Opinion Mining. In Proceedings of the Seventh International Conference on Language Resources and Evaluation (LREC'10), Valletta, Malta, 19–21 May 2010; pp. 1320–1326.
- González-Padilla, D.A.; Tortolero-Blanco, L. Social Media Influence in the COVID-19 Pandemic. *Int. Braz. J. Urol.* **2020**, *46*, 120–124. [CrossRef]
- Henríquez, C.; Guzmán, J. A Review of Sentiment Analysis in Spanish. *Tecciencia* **2017**, *12*, 35–48. [CrossRef]
- Hurtado, L.; Pla, F. Análisis de Sentimientos, Detección de Tópicos y Análisis de Sentimientos de Aspectos En Twitter. In Proceedings of the TASS 2014, Girona, Spain, 16–19 September 2014.
- Hung, C.; Chen, S.-J. Word Sense Disambiguation Based Sentiment Lexicons for Sentiment Classification. *Knowl.-Based Syst.* **2016**, *110*, 224–232. [CrossRef]
- Henriquez Miranda, C.; Sanchez, G. Aspect Extraction for Opinion Mining with a Semantic Model. *Eng. Lett.* **2021**, *29*, 61–67.
- Henríquez, C.; Plà, F.; Hurtado, L.F.; Luna, J.A.G. Análisis de Sentimientos a Nivel de Aspecto Usando Ontologías y Aprendizaje Automático. *Proces. Leng. Nat.* **2017**, *59*, 49–56.
- Zhang, L.; Wang, S.; Liu, B. Deep Learning for Sentiment Analysis: A Survey. *WIREs Data Min. Knowl. Discov.* **2018**, *8*, e1253. [CrossRef]
- Devlin, J.; Chang, M.-W.; Lee, K.; Toutanova, K. BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding. *arXiv* **2019**, arXiv:1810.04805.
- Jojoa, M.; Garcia-Zapirain, B.; Gonzalez, M.J.; Perez-Villa, B.; Urizar, E.; Ponce, S.; Tobar-Blandon, M.F. Analysis of the Effects of Lockdown on Staff and Students at Universities in Spain and Colombia Using Natural Language Processing Techniques. *Int. J. Environ. Res. Public Health* **2022**, *19*, 5705. [CrossRef]
- Liu, B. Sentiment Analysis and Opinion Mining. *Sentim. Anal. Opin. Min.* **2012**, *5*, 1–167.
- Vilares, D.; Alonso Perdo, M.Á.; Carlos, G.-R. A Syntactic Approach for Opinion Mining on Spanish Reviews. *Nat. Lang. Eng.* **2013**, *1*, 139–163.
- Plaza-Del-Arco, F.M.; Martín-Valdivia, M.T.; María Jiménez-Zafra, S.; Molina-González, M.D.; Martínez-Cámara, E. COPOS: Corpus Of Patient Opinions in Spanish. Application of Sentiment Analysis Techniques. *Proces. Leng. Nat.* **2016**, *57*, 83–90.

19. Cadilhac, A.; Benamara, F.; Aussenac-Gilles, N. Ontolexical Resources for Feature Based Opinion Mining: A Case-Study. In Proceedings of the 6th Workshop on Ontologies and Lexical Resources, Beijing, China, 22 August 2010; pp. 77–86.
20. Steinberger, J.; Brychcín, T.; Konkol, M. Aspect-Level Sentiment Analysis in Czech. In Proceedings of the Workshop on Computational Approaches to Subjectivity, Sentiment and Social Media Analysis, Baltimore, Maryland, 27 June 2014; pp. 24–30.
21. Pang, B.; Lee, L. Opinion Mining and Sentiment Analysis. *Found. Trends Inf. Retr.* **2008**, *2*, 1–135. [\[CrossRef\]](#)
22. De Freitas, L.A.; Vieira, R. Ontology-Based Feature Level Opinion Mining for Portuguese Reviews. In Proceedings of the Proceedings of the 22nd International Conference on World Wide Web; ACM: New York, NY, USA, 2013; pp. 367–370.
23. Manek, A.S.; Shenoy, P.D.; Mohan, M.C. Aspect Term Extraction for Sentiment Analysis in Large Movie Reviews Using Gini Index Feature Selection Method and SVM Classifier. *World Wide Web* **2016**, *20*, 135–154. [\[CrossRef\]](#)
24. Agüero-Torales, M.M.; Abreu Salas, J.I.; López-Herrera, A.G. Deep Learning and Multilingual Sentiment Analysis on Social Media Data: An Overview. *Appl. Soft Comput.* **2021**, *107*, 107373. [\[CrossRef\]](#)
25. Kaur, H.; Ahsaan, S.U.; Alankar, B.; Chang, V. A Proposed Sentiment Analysis Deep Learning Algorithm for Analyzing COVID-19 Tweets. *Inf. Syst. Front.* **2021**, *23*, 1417–1429. [\[CrossRef\]](#)
26. Jing, N.; Wu, Z.; Wang, H. A Hybrid Model Integrating Deep Learning with Investor Sentiment Analysis for Stock Price Prediction. *Expert Syst. Appl.* **2021**, *178*, 115019. [\[CrossRef\]](#)
27. Litjens, G.; Kooi, T.; Bejnordi, B.E.; Setio, A.A.A.; Ciompi, F.; Ghafoorian, M.; van der Laak, J.A.W.M.; van Ginneken, B.; Sánchez, C.I. A Survey on Deep Learning in Medical Image Analysis. *Med. Image Anal.* **2017**, *42*, 60–88. [\[CrossRef\]](#) [\[PubMed\]](#)
28. Bonifazi, G.; Cauteruccio, F.; Corradini, E.; Marchetti, M.; Sciarretta, L.; Ursino, D.; Virgili, L. A Space-Time Framework for Sentiment Scope Analysis in Social Media. *Big Data Cogn. Comput.* **2022**, *6*, 130. [\[CrossRef\]](#)
29. Bonifazi, G.; Corradini, E.; Ursino, D.; Virgili, L. New Approaches to Extract Information From Posts on COVID-19 Published on Reddit. *Int. J. Info. Tech. Dec. Mak.* **2022**, *21*, 1385–1431. [\[CrossRef\]](#)
30. Manguri, K.H.; Ramadhan, R.N.; Amin, P.R.M. Twitter Sentiment Analysis on Worldwide COVID-19 Outbreaks. *Kurd. J. Appl. Res.* **2020**, *5*, 54–65. [\[CrossRef\]](#)
31. Dubey, A.D. Twitter Sentiment Analysis during COVID-19 Outbreak. *SSRN* **2020**. preprint. [\[CrossRef\]](#)
32. Boon-Itt, S.; Skunkan, Y. Public Perception of the COVID-19 Pandemic on Twitter: Sentiment Analysis and Topic Modeling Study. *JMIR Public Health Surveill.* **2020**, *6*, e21978. [\[CrossRef\]](#)
33. Garcia, K.; Berton, L. Topic Detection and Sentiment Analysis in Twitter Content Related to COVID-19 from Brazil and the USA. *Appl. Soft Comput.* **2021**, *101*, 107057. [\[CrossRef\]](#)
34. Kruspe, A.; Häberle, M.; Kuhn, I.; Zhu, X. Cross-Language Sentiment Analysis of European Twitter Messages during the COVID-19 Pandemic. In Proceedings of the 1st Workshop on NLP for COVID-19 at ACL 2020, Online, 9–10 July 2020.
35. Marcec, R.; Likic, R. Using Twitter for Sentiment Analysis towards AstraZeneca/Oxford, Pfizer/BioNTech and Moderna COVID-19 Vaccines. *Postgrad. Med. J.* **2022**, *98*, 544–550. [\[CrossRef\]](#)
36. Villavicencio, C.; Macrohon, J.J.; Inbaraj, X.A.; Jeng, J.-H.; Hsieh, J.-G. Twitter Sentiment Analysis towards COVID-19 Vaccines in the Philippines Using Naïve Bayes. *Information* **2021**, *12*, 204. [\[CrossRef\]](#)
37. Valle-Cruz, D.; Fernandez, V.; Lopez-Chau, A.; Sandoval Almazan, R. Does Twitter Affect Stock Market Decisions? Financial Sentiment Analysis in Pandemic Seasons: A Comparative Study of H1N1 and COVID-19. *Cogn. Comput.* **2020**. preprint. [\[CrossRef\]](#)
38. Sanders, A.; White, R.; Severson, L.S.; Ma, R.; McQueen, R.; Paulo, H.C.A.; Zhang, Y.; Erickson, J.S.; Bennett, K.P. Unmasking the Conversation on Masks: Natural Language Processing for Topical Sentiment Analysis of COVID-19 Twitter Discourse. *medRxiv* **2020**. [\[CrossRef\]](#)
39. Khan, R.; Shrivastava, P.; Kapoor, A.; Tiwari, A.; Mittal, A. Social Media Analysis with AI: Sentiment Analysis Techniques for the Analysis of Twitter COVID-19 Data. *J. Crit. Rev.* **2020**, *7*, 2020.
40. Rodríguez-Orejuela, A.; Montes-Mora, C.L.; Osorio-Andrade, C.F. Sentimientos hacia la vacunación contra la COVID-19: Panorama colombiano en Twitter. *Palabra Clave* **2022**, *25*, e2514. [\[CrossRef\]](#)
41. Aislamiento Social Obligatorio: Un Análisis de Sentimientos Mediante Machine Learning. Available online: http://www.scielo.org.co/scielo.php?script=sci_arttext&pid=S2215-910X2021000100001 (accessed on 4 August 2022).
42. Ahmad, W.; Wang, B.; Martin, P.; Xu, M.; Xu, H. Enhanced Sentiment Analysis Regarding COVID-19 News from Global Channels. *J. Comput. Soc. Sci.* **2023**, *6*, 19–57. [\[CrossRef\]](#)
43. Kumari, S.; Pushphavathi, T.P. Intelligent Lead-Based Bidirectional Long Short Term Memory for COVID-19 Sentiment Analysis. *Soc. Netw. Anal. Min.* **2022**, *13*, 1. [\[CrossRef\]](#) [\[PubMed\]](#)
44. Cañete, J.; Chaperon, G.; Fuentes, R.; Ho, J.-H.; Kang, H.; Pérez, J. Spanish Pre-Trained BERT Model and Evaluation Data. In Proceedings of the PML4DC at ICLR 2020, Addis Ababa, Ethiopia, 26–30 April 2020.
45. de Arriba Serra, A.; Oriol Hilari, M.; Franch Gutiérrez, J. Applying Sentiment Analysis on Spanish Tweets Using BETO. In Proceedings of the Iberian Languages Evaluation Forum (IberLEF 2021): Co-located with the Conference of the Spanish Society for Natural Language Processing (SEPLN 2021), XXXVII International Conference of the Spanish Society for Natural Language Processing, Málaga, Spain, 24 September 2021; pp. 1–8.
46. Pijal, W.; Armijos, A.; Llumiquinga, J.; Lalvay, S.; Allauca, S.; Cuenca, E. Spanish Pre-Trained CaTrBETO Model for Sentiment Classification in Twitter. In Proceedings of the 2022 Third International Conference on Information Systems and Software Technologies (ICI2ST), Quito, Ecuador, 8–10 November 2022; pp. 93–98.

47. Vernikou, S.; Lyras, A.; Kanavos, A. Multiclass Sentiment Analysis on COVID-19-Related Tweets Using Deep Learning Models. *Neural. Comput. Appl.* **2022**, *34*, 19615–19627. [CrossRef]
48. Jojoa, M.; Eftekhari, P.; Nowrouzi-Kia, B.; Garcia-Zapirain, B. Natural Language Processing Analysis Applied to COVID-19 Open-Text Opinions Using a DistilBERT Model for Sentiment Categorization. *AI Soc.* **2022**, *2022*, 1–8. [CrossRef]
49. Madani, Y.; Erritali, M.; Bouikhalene, B. A New Sentiment Analysis Method to Detect and Analyse Sentiments of COVID-19 Moroccan Tweets Using a Recommender Approach. *Multimed. Tools Appl.* **2023**, 1–20. [CrossRef]
50. Natural Language Processing: State of the Art, Current Trends and Challenges | SpringerLink. Available online: <https://link.springer.com/article/10.1007/s11042-022-13428-4> (accessed on 21 February 2023).
51. Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A.N.; Kaiser, L.; Polosukhin, I. Attention Is All You Need. *arXiv* **2017**, arXiv:1706.03762.
52. Shi, Y.; Wang, J.; Ren, P.; Valizadeh Aslani, T.; Zhang, Y.; Hu, M.; Liang, H. Fine-Tuning BERT for Automatic ADME Semantic Labeling in FDA Drug Labeling to Enhance Product-Specific Guidance Assessment. *J. Biomed. Inform.* **2023**, *138*, 104285. [CrossRef] [PubMed]
53. Kong, J.; Wang, J.; Zhang, X. Hierarchical BERT with an Adaptive Fine-Tuning Strategy for Document Classification. *Knowl.-Based Syst.* **2022**, *238*, 107872. [CrossRef]
54. Raffel, C.; Shazeer, N.; Roberts, A.; Lee, K.; Narang, S.; Matena, M.; Zhou, Y.; Li, W.; Liu, P.J. Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer. *J. Mach. Learn. Res.* **2020**, *21*, 5485–5551.
55. Bert-Base-Uncased. Hugging Face. Available online: <https://huggingface.co/bert-base-uncased> (accessed on 2 December 2022).
56. Pérez, J.M.; Furman, D.A.; Alonso Alemany, L.; Luque, F.M. RoBERTuito: A Pre-Trained Language Model for Social Media Text in Spanish. In Proceedings of the Thirteenth Language Resources and Evaluation Conference; European Language Resources Association, Marseille, France, 20–25 June 2022; pp. 7235–7243.
57. Felbo, B.; Mislove, A.; Søgaard, A.; Rahwan, I.; Lehmann, S. Using Millions of Emoji Occurrences to Learn Any-Domain Representations for Detecting Sentiment, Emotion and Sarcasm. In Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing; Association for Computational Linguistics, Copenhagen, Denmark, 9–11 September 2017; pp. 1615–1625.
58. Novak, P.K.; Smailović, J.; Sluban, B.; Mozetič, I. Sentiment of Emojis. *PLoS ONE* **2015**, *10*, e0144296. [CrossRef]
59. Amrullah, M.S.; Budi, I.; Santoso, A.B.; Putra, P.K. The Effect of Using Emoji and Hashtag in Sentiment Analysis on Twitter Case Study: Indonesian Online Travel Agent. *AIP Conf. Proc.* **2023**, *2654*, 020013. [CrossRef]
60. Ayvaz, S.; Shiha, M.O. The Effects of Emoji in Sentiment Analysis. *Int. J. Comput. Electr. Eng.* **2017**, *9*, 360–369. [CrossRef]
61. Kejriwal, M.; Wang, Q.; Li, H.; Wang, L. An Empirical Study of Emoji Usage on Twitter in Linguistic and National Contexts. *Online Soc. Netw. Media* **2021**, *24*, 100149. [CrossRef]
62. Miller, H.; Thebault-Spieker, J.; Chang, S.; Johnson, I.; Terveen, L.; Hecht, B. “blissfully Happy” or “Ready to Fight”: Varying Interpretations of Emojis. In Proceedings of the 10th International Conference on Web and Social Media, ICWSM 2016, Cologne, Germany, 17–20 May 2016; AAAI Press: Washington, DC, USA, 2016; pp. 259–268.
63. Barbieri, F.; Espinosa-Anke, L.; Saggion, H. Revealing Patterns of Twitter Emoji Usage in Barcelona and Madrid. *Artif. Intell. Res. Dev.* **2016**, *288*, 239–244. [CrossRef]
64. Wijeratne, S.; Balasuriya, L.; Sheth, A.; Doran, D. EmojiNet: Building a Machine Readable Sense Inventory for Emoji. In Proceedings of the Social Informatics, Bellevue, WA, USA, 11–14 November 2016; Spiro, E., Ahn, Y.-Y., Eds.; Springer International Publishing: Cham, Switzerland, 2016; pp. 527–541.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.