

Article

Semantic Partitioning and Machine Learning in Sentiment Analysis

Ebaa Fayyouni *  and Sahar Idwan

Department of Computer Science and Applications, Faculty of Prince Al-Hussein Bin Abdallah II for Information Technology, The Hashemite University, P.O. Box 330127, Zarqa 13133, Jordan; sahar@hu.edu.jo

* Correspondence: enfayyouni@hu.edu.jo; Tel.: +962-79-909-3441

Abstract: This paper investigates sentiment analysis in Arabic tweets that have the presence of Jordanian dialect. A new dataset was collected during the coronavirus disease (COVID-19) pandemic. We demonstrate two models: the Traditional Arabic Language (TAL) model and the Semantic Partitioning Arabic Language (SPAL) model to envisage the polarity of the collected tweets by invoking several, well-known classifiers. The extraction and allocation of numerous Arabic features, such as lexical features, writing style features, grammatical features, and emotional features, have been used to analyze and classify the collected tweets semantically. The partitioning concept was performed on the original dataset by utilizing the hidden semantic meaning between tweets in the SPAL model before invoking various classifiers. The experimentation reveals that the overall performance of the SPAL model competes over and better than the performance of the TAL model due to imposing the genuine idea of semantic partitioning on the collected dataset.

Keywords: semantic partitioning; Jordanian dialect; sentiment analysis; Arabic language; Traditional Arabic Language; Semantic Partitioning Arabic Language



Citation: Fayyouni, E.; Idwan, S. Semantic Partitioning and Machine Learning in Sentiment Analysis. *Data* **2021**, *6*, 67. <https://doi.org/10.3390/data6060067>

Academic Editor: Erik Cambria

Received: 28 May 2021
Accepted: 18 June 2021
Published: 21 June 2021

Publisher's Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright: © 2021 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

Social networks, nowadays, are just like beating hearts—people cannot live without them. Social networks affect various fields, such as health, marketing, politics, businesses, management, etc. Scientists/researchers mine for hidden knowledge amongst the vast amount of content posted via Twitter, Instagram, or Facebook, to facilitate decision making [1]. Twitter, with 330 million users [2], was a fertile source of research in this study. Twitter allows users to share their opinions in short-term messages, with a maximum of 280 characters [3].

Sentiment Analysis (SA) is a vital technique used to gain insight human opinions, emotions, and attitudes regarding particular topics in specific, written languages [4–6]. SA is the most actively researched field in Natural Language Processing (NLP) [7], and it is involved in data mining and text mining studies [8]. Further details about SA applications and challenges can be found in [9,10]. The influence of social media has increased throughout the years, directly impacting the importance of this field [11]. SA helps provide insight into whether society is positively or negatively impacted by an international or national event [12].

SA for English text has been extensively studied and investigated using public datasets [13,14]. On the other hand, SA for foreign language, such as Arabic language, has received very little attention [4,15]. The Arabic language is the sixth official language of the United Nations [16]. Twenty-seven countries use the Arabic language as a primary language; approximately 422 million people worldwide speak it [4]. However, the Arabic language is still at the beginning stage in the NLP field due to insufficient resources and tools [11,17]. This presents a vast challenge for researchers in this field, in regards to its complex structure, history, different cultures, and dialect [17–19]. In general, the Arabic language is categorized into three different types: (1) Modern Standard Arabic (MSA);

(2) Classical Arabic (CA); and (3) Dialectal Arabic (DA). Dialectal Arabic is the most used type on social networks. Arab dialectologists have classified it into five major groups: Peninsular, Mesopotamian, Levantine, Egypto-Sudanic, and Maghrebi [20].

Most studies on SA are based on machine learning approaches [21–23]. The lexicon analysis, machine learning-based analysis, and combined analysis are three different techniques that have been employed in posting tweets for sentiment extraction [21]. Machine translation in English sentiment resources has been extensively used in several studies [24–26]. Unfortunately, this approach is structurally and grammatically inadequate to capture Arabic features [27–29]. Some studies have focused on MSA, but ignored the dialect and Arabic text containing Latin letters that decreased the accuracy in real applications [30]. Moreover, various dialects and free-writing styles increase the challenge and complexity in this research field [31].

Considerable research is aimed at providing a complete survey of state-of-the-art works related to Arabic SA [27,32–34], which employs different classification techniques in machine learning to determine the polarity of different tweets scripted in the Arabic language [4,28,35,36]. Abdulla et al. [37] collected 2000 tweets on various topics, such as politics and arts. These tweets include opinions written in both MSA and the Jordanian dialect. They built an Arabic lexicon from a seed list of 300 words. The new lexicon-based tool was proposed and implemented to automatically determine a tweet's polarity based on the total summation weight for each word in the tweet. The authors in [38] proposed a corpus from a news domain consisting of MSA and dialect, with a total of 3015; this breaks down to 1549 different dialects and 1466 MSN collected from a piece of news. The polarity lexicon used was Arab SentiNet, which contains 3982 extracted adjectives. A systematic examination of more than 20 studies in Arabic sentiment analysis can be found in [39]. Those studies utilized the recurrent neural network due to its talent in textual analysis and significant morphological variations.

Further details on the Arabic language (either MSA, CA, or DA) sentiment analysis with different corpuses are found in [40]. Partitioning is a prime concept used in designing many applications, such as network [41], database [42], security [43,44], web applications [45], etc. To the best of our knowledge, none of the studies have incorporated the partitioning concept in the collected tweets before invoking various machine-learning classifiers and merged the obtained results.

Recently, the entire world suffered from the coronavirus disease 2019 (COVID-19) pandemic. People extensively used social media to convey their ideas, opinions, and feelings regarding precautionary procedures forced by the government. Therefore, many research studies were conducted on utilizing Twitter to discuss the extent of societal acceptance of government procedures during the pandemic by utilizing the power of SA [46–48]. The authors in [46] built a model by using machine learning to predict people's responsiveness to government procedures, taken during the COVID-19 pandemic in five different areas of Saudi Arabia. Their dataset model utilized Arabic language tweets as an input to classify, predict, and compare the degree of citizen awareness in these areas. Other researchers [47] focused on estimating the feelings of people in Gulf countries, examining their actions regarding the pandemic by applying semantic analysis on Arabic language tweets. Three levels of semantic analysis have been listed in their study to measure citizens' feelings. The authors in [48] analyzed the Arabic language tweets posted during the pandemic using machine learning and deep learning models, to measure the correlation between sentiments in these tweets and the actual number of people infected with COVID-19.

Jordan implemented many precautions in March 2020 to reduce the spread of COVID-19, such as complete/partial lockdowns, monitoring and isolating the infected cases, sirens, and a curfew [49]. Consequently, Jordanians and residents extensively used social media to convey their ideas, opinions, and feelings regarding the measures taken by the government. This inspired us to investigate these tweets, in order to induce the reaction of Jordanian nationals. Therefore, the present study addresses the following three research questions:

- Research question 1: What are the performances of the various machine-learning classifiers of the collected tweets that utilize various extracted features?
- Research question 2: What is the impact of applying semantic partitioning on the collected data prior to invoking machine-learning classifiers?
- Research question 3: What are the reactions of Jordanians toward government precautionary measures during the COVID-19 pandemic?

We proposed and implemented two models to address these questions—the Traditional Arabic Language model (TAL) and the Semantic Partitioning Arabic Language (SPAL) model. The inputs included tweets posted by individual Jordanian nationals, via a Tweet Collector Tool (TCT). For both models, we started with tokenization and feature extraction on the processed and collected tweets to generate a CSV file with the same number of extracted features. Afterwards, both models integrated semantic analysis with various machine-learning classifiers to classify the tweets into positive or negative. The main difference between these two models is whether the classifiers are invoked on the entire dataset, as in TAL, or on mutually exclusive (“disjoint”) subsets, as in SPAL. It is worth noting that the partitioning process inside SPAL depends on the hidden semantic meaning in the Jordanian dialect tweets. Therefore, the main objective of this paper was to design and implement trustworthy models to categorize Jordanian opinions, whether they were for or against governmental actions during the COVID-19 pandemic. Many experiments have been conducted to measure the performances of both models using various well-known classifiers, including Support Vector Machine (SVM), Naïve Bayes (NB), J48, Multi-Layer Perceptron (MLP), and Logistic Regression (LR).

This paper is organized as follows: Section 2 illustrates the informal and algorithmic expressions of the newly proposed models. Section 3 discusses the experimental results on the real collected datasets. Finally, the conclusion and future work are presented in Section 4.

2. Proposed Methodology

It is common to invoke machine-learning models to classify collected tweets into positive or negative, but it is tedious, routine work. The novelty of this work lies in assimilating the semantic analysis of Jordanian dialect and enacting the semantic partitioning of the collected tweets to enhance the overall performance of the generated model. This section will present the Traditional Arabic Language (TAL) model and Semantic Partitioning Arabic Language (SPAL) model on Jordanian dialect in the following subsections.

2.1. Traditional Arabic Language (TAL) Model

The TAL model is implemented, as shown in Figure 1. The input of our model is the Jordanian dialect collected tweets, which were processed using computational strategies: textual processing and feature extraction, which will be explained later. The classification stage was applied to classify the number of positive and negative tweets. Finally, the overall performance of the model was measured as listed in Algorithm 1.

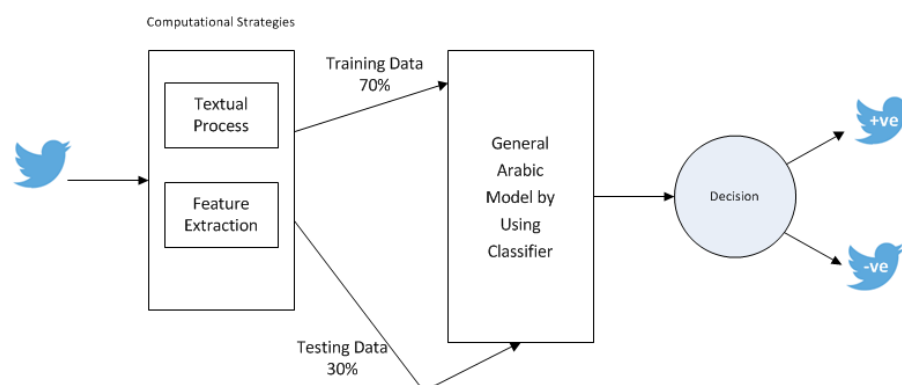


Figure 1. Traditional Arabic Language (TAL).

For the sake of building a sentimental dataset for Jordanian dialect and modern Arabic text, we considered the Arabic content published on Twitter. The dataset was collected from 1 March, 2020 to 21 May, 2020, during the COVID-19 pandemic. We collected the dataset by using a specially developed tool called the Tweet Collector Tool (TCT), in order to search for a specific hashtag (examples of the hashtag: (“#كورونا” “#kuruna”, “#Corona”); (“#الحكومة_كورونا”, “#alhukumatu_kuruna”, “#Corona_government”); (“#الشعب_كورونا”, “#alshaebi_kuruna”, “# Corona people”); (“#ارتفاع_الاسعار”, “#artifaei_alasear”, “#high prices”); (“#تصريحات_وزير_الاعلام”, “#tasrihat_wzir_alaelam”, “#statements by the Minister of Information”); (“#المؤتمر_الصحفي”, “#almutamir_alsahufii”, “# Press Conference”)). The collected dataset size was equal to 2000 randomly selected tweets with equal positive and negative labels divided among three regions (north, middle, and south in Jordan), called “TAL”, as shown in Table 1. This dataset consists of two columns: the collected tweets and the polarity. The polarity for the extracted tweets is assigned manually into positive and negative. This was achieved by three experts in Arabic linguistics. In case of a disagreement, in regards to tweets that were polarized differently, the final decision of the annotation was taken based on the majority chosen.

Algorithm 1. Traditional_Arabic_Language ()

Input: InSet: Set of collected Jordanian Arabic tweets.

Output: POutSet: Number of correctly classified positive tweets.

NOutSet: Number of correctly classified negative tweets.

Accuracy, Precision, Recall, F_Score: Evaluate the performance of the Arabic Language model.

Method:

1. InSet=TextProcessing(InSet).
 2. InSet=FeatureExtraction(InSet).
 3. (TrainingSet, TestingSet)=Validation_Split(InSet).
 4. ClassifierModel=Classifier(ClassifierName, TrainingSet).
 5. ConfusionMatrix=CalculateConfusionMatrix(ClassifierModel(TestingSet)).
 6. (POutSet, NOutSet, Accuracy, Precision, Recall, F_Score)=CalculateAccuracy(ConfusionMatrix).
 7. END Algorithm 1
-

Table 1. Examples of tweets in the “TAL” dataset.

Polarity	Tweet
Positive	(أطباء ومرضيين الاردن هم خط الدفاع الأول في مواجهة مرض كورونا... اللهم أحفظهم جميعا و ردهم الى اهلهم سالمين); (atibaa' w mumaridin alardin hum khatu aldifae al'awal fi muajahat marad kuruna ... allahumm 'ahfazahum jamieana w radahum alaa ahilahum salimin); (Jordan's doctors and nurses are the first line of defense in the face of corona disease ... May Allah protect them all and return them to their families safely)
Positive	(تجنب المصافحة والاتصال المباشر مع الأشخاص وخل السلام نظر); (tajanub almusafahat wal'iitisal almubashir mae al'ashkhas wakhalu alsalam nazar); (Avoid handshakes and direct contact with people and keep the peace of mind)
Positive	(مع هذا الوباء تعلمنا من فيروس كورونا ان يومنا العادي ماكان ابداً يوم عادي ، إنما كان نعمة من نعم الله سبحانه وتعالى); (mae hadha alwaba' tuealamuna min fayrus kurna 'iina yawmana aleadia makan abdaan yawm eadi, 'iinama kan niematan min nieam allah subhanah wataealaa); (With this epidemic, we learned from the Corona virus that our normal day was never an ordinary day, but it was a blessing from the blessings of God Almighty)
Positive	(تحية فخر واعتزاز بوزير صحتنا الانسان الباشا سعد جابر وبكل طبيب وممرض وصيدلاني وفني وكافة الكوادر في القطاع الصحي الذين يسهروا على سلامتنا رغم تعريض انفسهم للعدوى); (tahiat fahr waetizaz biwazir sihatina alainsan albasha saed jaber wabikuli tabib wamumarid wasaydilaniun wafaniyun wakafat alkawadir fi alqitae alsihiyi aladhin yasharu ealaa salamatina raghm taerid ainfisihim lileadwaa); (Greetings of pride and pride to our Minister of Health, the human being, Pasha Saad Jaber, and to all the doctors, nurses, pharmacists, technicians and all cadres in the health sector who ensure our safety despite exposing themselves to infection)
Negative	(ليس من باب التشاؤم، لكن انخفاض الاصابات ليس مؤشر ايجابي او سلبي، المهم هو الجدول الزمني كامل لخطة احتواء المرض); (lays min bab altashawumu, lakina 'iinkhifad alasabat lays muashir ayjabiy aw salbi, almuhimu hu aljadwal alzamani kamil likhutat ahtiwa' almarad); (Not out of pessimism, but the decrease in infections is not a positive or negative indicator, the important thing is the complete schedule of the disease containment plan)
Negative	(يا ريت الحكومة تصير تعطينا رقم الفحوصات اللي قاعدة تعملها يومياً للتوضيح، لأنه للامانة كيرف الأردن (وكثير دول عربية) رياضياً مش منطقي...); (ya ryt alhukumat tasir tuetina raqm alfuhusat allly qaeidatan taemaluha ywmyaan liltawdihi, li'anah lil'amanat kirf al'urdun (wkathir dual earabiatun) ryadyaan mish mantiqi ...); (I wish the government would give us the number of tests that it does daily for clarification, because to be honest, Jordan (and many Arab countries) is mathematically not logical ...)
Negative	(عيب وعار عليكم...يعني للأسف عندي معلومات عن ناس حضروا حفلة العرس لكنهم مختبئين مع كل اسف ...); (eayb waear ealaykum ... yaeni ilasif eindi maelumat ean nas hadaruu haflat aleurs lakinahum mukhtabiyn mae kuli asf ...); (Shame on you ... I mean, unfortunately, I have information about people who attended the wedding party, but they are hiding, with all the regret ...)
Negative	(سينونا من أعراض #كورونا مين بلشت تظهر عنده أعراض الإفلاس...); (sibuna min 'aerad #kuruna min balisht tazhar eindah 'aerad al'iiflas); (Sibona is one of the symptoms of #Corona, who started showing symptoms of bankruptcy)
Negative	(ناس يا عالم يا هووووو... انا قبل أربع سنوات إنتخبت مجلس نواب #وينهم ؟؟؟؟); (ya nas ya ealam ya huwuwu ... ana qabl 'arbae sanawat 'iintakhabat majlis nuaab #wianhim ???); (Hey people, world, hooooo Four years ago, I elected parliament #Where are they???)

Dealing with tweets as a whole unit is confusing and ambiguous due to the poorly structured Arabic text. Therefore, there is an urgent need to extract and analyze the tweets to reduce the number of resources needed for processing, and simultaneously preserve important irrelevant features. It is considered a predominant step in SA. It begins with the initial collected dataset to build derived values, which are optimally used in the learning stage. Consequently, this leads to a better understanding of the problem domain and a more precise interpretation for the original dataset. In our study, the extraction process was achieved by implementing java code that tokenized the received data from the Twitter API. The tokenization was used to separate the text into smaller units called words, phrases, numbers, non-Arabic words, single characters, and punctuation marks. Unfortunately, the obtained data were not homogeneous; therefore, a normalization procedure was conducted to standardize the obtained data, in order to avoid Tashkeel (diacritics), Tatweel (repeated letter), and contextual letter representation. For example; In the Arabic language, most letters have contextual letter representations, such as Alef (ا | إ | ؤ), which have been transformed into a unified representation (l). The obtained extracted features per tweet are saved in a CSV file. There are 30 extracted features cataloged into four categories: Lexicon features, writing style features, grammatical features, and emotional features as shown in Figure 2.

1. **Lexicon features:** it focuses on the word–character structure and emphasizes its effect on the results by computing the number of words and the number of characters per tweet. Moreover, the number of words by length, varying from five characters to ten characters, were counted.
2. **Writing style features:** the writing style is affected by user mood and the user style. Some users used numerical digits when writing certain Arabic letters, while others used special characters and symbols to represent their feelings. Moreover, some users used punctuation, which altered the tweet contents. Therefore, it is paramount to

count the number of numerical digits, special characters, symbols, delimiters, and punctuation per tweet [50].

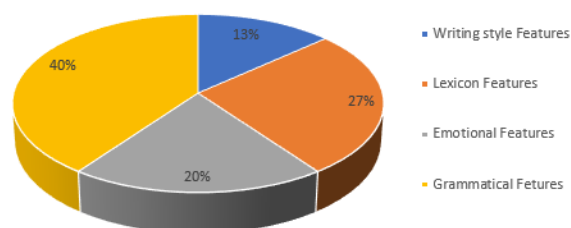


Figure 2. The percentage value of different extracted features used in SA of Jordanian dialect.

3. **Grammatical features:** many researchers utilized grammatical features to understand the language. In our study, we analyzed 11 grammatical rules exercised in the Arabic tweets: Kan and sisters, Enna and sisters, question tools, exception tools, five verbs, five nouns, plural words, imperative clause, Nidaa clause, Eljar letters, and Eatf letters.
4. **Emotional features:** we focused on SA to mine the emotional statuses of the tweet users. Polarity and emotion were identified in words inside the tweets. They are categorized as positive words, negative words, a combination of positive words, and a combination of negative words, as shown in Table 2.

Table 2. Examples of emotional features.

Emotional Features	Example
Positive word feature	الحمد الله; alhamd allah; Praise be to God
	تفاؤل; tafawulu; optimism
	تعاون; taeawuni; cooperation
	يُعاافينا; yaeafina; recovery
	يشفي; yashfi; healing
	أفضل; 'afdal; better
	سعيد; saeid; happy
Negative word feature	فرج; faraj; relief
	خوف; khufa; Fear
	هلع; halaea; panic
	موت; mut; death
	وفيات; wafiat; deaths
	فساد; fasad; corruption
	حرامية; haramiat; thieves
	مش; mash; not
Combination of positive words	ملعون; maleun; cursed
	مغصوس; maghsus; crippled
	لعنة الله عليك; laenat allah ealayk; God damn you
Combination of negative words	ارتفاع الأسعار; airtifae al'asear; Rising prices
	ما نعرف وين; ma binaerif win; We don't know where
	سيب الموضوع; sib almawdue; Leave the topic
Combination of negative words	معك حق; maeak haqun; You are right
	بالله; al'amal biallah; Hope in God
	فسحة أمل; fushat 'amal; a space of hope
	يبعين الله; bieayn allah; God bless

A vitally important part of evaluating various models of the collected tweets is separating the tweets into training and testing datasets. The size of the training and testing

sets are equal to 70% (1400 tweets) and 30% (600 tweets) of the original size of the dataset, respectively.

The classification process is defined to recognize the tweets and separate them into positive and negative categories. There are many predictive machine-learning models of high accuracy and powerful features invoked in several applications. The following models are used in this paper:

Support Vector Machine (SVM): is one of the simplest and most effective neural networks. It is highly recommended to be used in classification problems due to its ability in increasing the predictive correctness by eliminating over-fit to the data [51]. It is based on finding the best hyperplane in multidimensional space to minimize errors [52].

Naïve Bayes (NB): is a probabilistic model based on Bayes theorem, which presumes that each feature generates an independent and equal impact to the target class [53]. It is a fast, accurate, and easy classification model that is used for large dataset size [54].

J48: is a decision tree that partitions the input space of the dataset into mutually exclusive regions, each of which is assigned a label to characterize its data points [55]. Building a decision tree is achieved by following an agreed iterative approach. The algorithm partitions the dataset based on the best informative attribute. The attribute with the maximum gain ratio is selected as the splitting attribute [56]. Generally, decision tree classification models have many advantages, such as being easy to interpret, and obtaining comparable accuracy to other classification models [57].

Logistic Regression (LR): is a statistical predictive model, which uses the observations of one or more independent variables to find the probability of the dependent variable. It is commonly used to solve problems in different applications [58]. It is recommended that it be used with the existence of multi-collinearity and high dimensional datasets [59]. LR is easy to implement and properly clarifies the obtained results [60].

Multi-Layer Perceptron (MLP): is a feed-forward artificial neural network to predict and classify labels in different applications. It generally consists of at least three layers to build relationships between input and output layers in order to compute the required patterns. Each layer consists of a set of neurons with a set of adaptive weights [61].

Performance evaluation metrics are used to evaluate the overall performance of various models. In general, metrics include comparing the polarity of the tweets to the predicated classified tweet class (positive or negative). Accuracy, recall, precision, and F-score are computed from the confusion matrix in order to measure the power of the predictive models. Further details about the definition and equation of each metric can be found in [62].

2.2. Semantic Partitioning Arabic Language (SPAL) Model

This paper concentrates on applying the partitioning concept to the original set of tweets, after applying the computational strategies procedure and before invoking the classification model. The partitioning process is neither randomly nor blindly applied. It depends on the hidden semantic meaning that exists between the Jordanian dialect tweets. We believe that utilizing the semantic meaning between tweets will drastically increase the general accuracy of the model.

The key idea of collecting tweets is to understand Jordanian reactions toward the government measures implemented during the COVID-19 pandemic. It is clearly shown that this study is founded based on the following three domains: coronavirus, government, and people (Jordanian nationals), as shown in Figure 3. The existing semantic correlation among these domains can be viewed as follows: coronavirus and government, coronavirus and people, and, finally, government and people. Each case is derived by taking all possible semantic correlations that exist in the Jordanian tweets according to the predefined domains: $\left(\begin{matrix} D \\ C \end{matrix} \right) = \frac{D!}{C!(D-C)!}$, where D represents the number of studied domains, which is equal to 3, while C represents the number of correlated selected domains, which is equal to

2. Therefore, $\binom{3}{2} = \frac{3!}{2! \cdot 1!} = 3$, three different subsets can be generated from the originally collected datasets that are mutually exclusive, as shown below:

$$S = \bigcup_{i=1}^3 S_i$$

$$S_1 \cap S_2 \cap S_3 = \emptyset$$

where S : presents the original collected tweets, while S_i : presents the subsets of the original set.

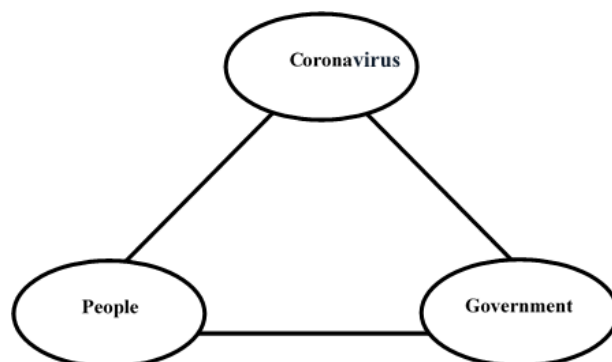


Figure 3. The three domains presented in the collected Jordanian dialect tweets.

Consequently, another copy of the “TAL” dataset was analyzed to partition the collected tweets into three semantic subsets. The collected tweets were manually classified into three mutually exclusive subsets based on their semantic meaning. This process was completed by seven experts in Arabic linguistics. In case of agreement, the tweet was stored in the corresponding classified subset. Otherwise, the final decision of the classification was taken based on the majority chosen subset. Each subset was stored in a separate dataset. The first subset contains the tweets that belongs to governmental responses in Jordan during the COVID-19 pandemic, called the “SPAL_1” dataset, as shown in Table 3. The second subset contains the tweets that belong to people’s reactions during the COVID-19 pandemic, called the “SPAL_2” dataset, as shown in Table 4. The last subset contains the tweets that belong to the interaction between the government and people during the COVID-19 pandemic, called the “SPAL_3” dataset, as shown in Table 5.

Table 3. Examples of tweets in the “SPAL_1” dataset.

Polarity	Tweet
Negative	(ليس من باب التشاؤم، لكن إنخفاض الإصابات ليس مؤشر إيجابي أو سلبي، المهم هو الجدول الزمني كامل لخطة احتواء المرض); (lays min bab altashawumu, lakina ‘iinkhifad alasabat lays muashir ayjaby aw salbi, almuhimu hu aljadwal alzamaniu kamil likhutat ahtiwa’ almarad); (Not out of pessimism, but the decrease in infections is not a positive or negative indicator, the important thing is the complete schedule of the disease containment plan) كامل لخطة احتواء المرض
Positive	(تجنب المصافحة والاتصال المباشر مع الأشخاص وخل السلام نظراً); (tajanub almusafahat wal’iitisa almuwashir mae al’ashkhas wakhalu alsalam nazar); (Avoid handshakes and direct contact with people and keep the peace of mind)
Positive	(أطباء ومرضين الاردن هم خط الدفاع الأول في مواجهة مرض كورونا... اللهم أحفظهم جميعاً و ردهم الى اهلهم سالمين); (atibaa’ w mumaridin alardin hum khatu aldiafa al’awal fi muajahat marad kuruna ... allahumm ‘ahfazahum jamieana w radahum alaa ahilahum salimin); (Jordan’s doctors and nurses are the first line of defense in the face of corona disease ... May Allah protect them all and return them to their families safely)

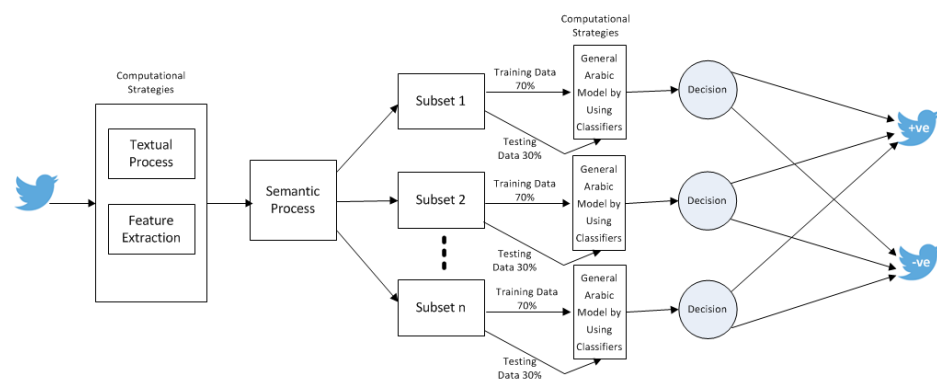
Table 4. Examples of tweets in the “SPAL_2” dataset.

Polarity	Tweet
Negative	(عيب وعار عليكم... يعني للأسف عندي معلومات عن ناس حضروا حفلة العرس لكنهم مختبئين مع كل اسف ...); (eayb waear ealaykum ... yaeni lilasif eindi maelumat ean nas hadaruu hafat aleurs lakinahum mukhtabiyn mae kuli asf ...); (Shame on you ... I mean, unfortunately, I have information about people who attended the wedding party, but they are hiding, with all the regret ...)
Negative	(سيبونا من أعراض #كورونا بلشت تظهر عنده أعراض الإفلاس); (sibuna min ‘aerad #kuruna min balisht tazhar eindah ‘aerad al’iiflas); (Sibona is one of the symptoms of #Corona, who started showing symptoms of bankruptcy)
Positive	(مع هذا الوباء تعلمنا من فيروس كورونا إن يومنا العادي ماكان أبداً يوم عادي ، إنما كان نعمة من نعم الله سبحانه وتعالى); (mae hadha alwaba’ tuealamuna min fayrus kurna ‘iina yawmana aleadia makan abdaan yawm eadi, ‘iinama kan niematan min nieam allah subhanah wataealaa); (With this epidemic, we learned from the Corona virus that our normal day was never an ordinary day, but it was a blessing from the blessings of God Almighty)

Table 5. Examples of tweets in the “SPAL_3” dataset.

Polarity	Tweet
Negative	(يا ريت الحكومة تصير تعطينا رقم الفحوصات اللي قاعدة تعملها يومياً للتوضيح، لأنه للأمانة كيرف الأردن (وكثير دول عربية) رياضياً مش منطقي...); (ya ryt alhukumat tasir tueatina raqm alfuhusat ally qaeidatan taemaluha ywmyaan liltawdihi, li’ anah lil’amanat kirf al’urdun (wkathir dual earabiatun) ryadyaan mish mantiqi ...); (I wish the government would give us the number of tests that it does daily for clarification, because to be honest, Jordan (and many Arab countries) is mathematically not logical ...)
Positive	(تحية فخر واعتزاز بوزير صحتنا الإنسان الباشا سعد جابر وبكل طبيب وممرض وصيدلاني وفني وكافة الكوادر في القطاع الصحي الذين يسهروا على سلامتنا رغم تعريض انفسهم للعوى); (tahiat fakhr waetizaz biwazir sihatina alainsan albasha saed jabir wabikuli tabib wamumarid wasaydilaniun wafaniyun wakafat alkawadir fi alqitae alsihiyi aladhin yasharuu ealaa salamatina raghm taerid ainfsihim lileadwaa); (Greetings of pride and pride to our Minister of Health, the human being, Pasha Saad Jaber, and to all the doctors, nurses, pharmacists, technicians and all cadres in the health sector who ensure our safety despite exposing themselves to infection)
Negative	(ناس يا عالم يا هووووو..... انا قبل أربع سنوات إنتخبت مجلس نواب #وينهم ؟؟؟؟); (ya nas ya ealam ya huwuwu ana qabl ‘arbae sanawat ‘iintakhabat majlis nuaab #wianhim ?????); (Hey people, world, hoooooooooFour years ago, I elected parliament #Where are they????)

The train–test splitting process was independently applied to every subset to estimate the performance of the general model. The training part of each subset was used to build the classification sub-model independently. This entailed independent testing for each sub-model by using its testing part. Each sub-model produced its confusion matrix. All of the generated confusion matrices per sub-model were merged to produce the general confusion matrix for the chosen classification model. Latter, the performance of the whole model was evaluated by measuring F-score, recall, accuracy, and precision on a general confusion matrix. Finally, Algorithm 2 presents the proposed SPAL model, which is visually shown in Figure 4.

**Figure 4.** Semantic partitioning Arabic language (SPAL).

Algorithm 2. SemanticPartitioning_ArabicLanguage ()**Input:** InSet: Set of collected Jordanian Arabic tweets.**Output:** POutSet: Number of correctly classified positive tweets.

NOutSet: Number of correctly classified negative tweets.

Accuracy, Precision, Recall, F_Score: Evaluate the performance of the Semantic Partitioning Arabic Language model.**Method:**

1. InSet=TextProcessing(InSet).
2. InSet=FeatureExtraction(InSet).
3. (SubSet₁, SubSet₂, ..., SubSet_n)= Semantic_Process(InSet).
4. **For each SubSet Do**
 - 4.1. (TrainingSubSet, TestingSubSet)=Validation_Split(SubSet).
 - 4.2. ClassifierModel()=Classifier(ClassifierName,TrainingSubSet).
 - 4.3. SubSetConfusionMatrix=CalculateConfusionMatrix(ClassifierModel(TestingSubSet)).
5. ConfusionMatrix=MergeConfusionMatrix(SubSetConfusionMatrix₁, SubSetConfusionMatrix₂, ..., SubSetConfusionMatrix_n).
6. (POutSubSet,NOutSubSet,Accuracy,Precision,Recall,F_Score)=CalculateAccuracy(ConfusionMatrix).
7. **END Algorithm 2**

The most important step in our newly proposed SPAL model involved merging the produced confusion matrix of each sub-model (i) into a general confusion matrix for the whole classification model, as shown in Figure 5.

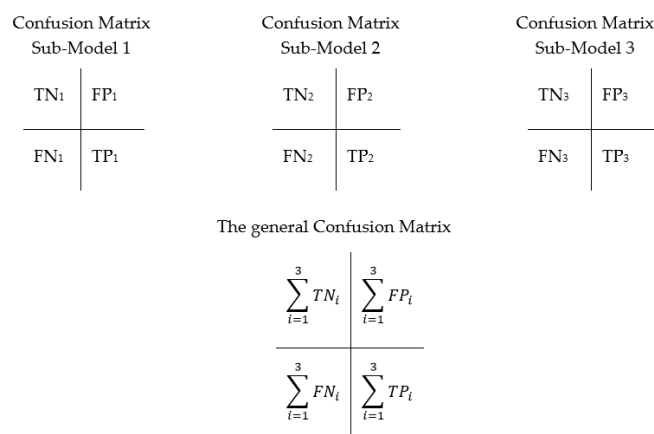


Figure 5. Building the general confusion matrix for the entire classification model of SPAL.

The actual meaning of TN_i, TP_i, FN_i, and FP_i are listed below:

- TN_i: the number of tweets that are negatively classified and their actual polarity is negative for the sub-model i.
- TP_i: the number of tweets that are positively classified and their actual polarity is positive for the sub-model i.
- FN_i: the number of tweets that are negatively classified and their actual polarity is positive for the sub-model i.
- FP_i: the number of tweets that are positively classified and their actual polarity is negative for the sub-model i.

The general confusion matrix is calculated by summing the corresponding values of the confusion matrix for all sub-models, as follows:

$$\sum_{i=1}^3 TN_i = TN_1 + TN_2 + TN_3$$

$$\sum_{i=1}^3 TP_i = TP_1 + TP_2 + TP_3$$

$$\sum_{i=1}^3 FN_i = FN_1 + FN_2 + FN_3$$

$$\sum_{i=1}^3 FP_i = FP_1 + FP_2 + FP_3$$

3. Experimental Results and Discussion

In this paper, all experimental results were conducted by using Weka software. It is open-source, with plenty of machine learning algorithms and visualization tools for data mining. It is one of the chosen languages in data science, particularly NLP [63].

All running experiments were set to use 10-fold cross-validation, and a 0.001 tolerance value. All previously mentioned extracted features equal to 30 were used as test attributes in building each model. Different classifiers were invoked in the proposed models: SVM, NB, J48, MLP, and LR. The sequential minimum optimization algorithm (SMO) was invoked in SVM with a C value equal to 1.0, Epsilon value equal to 1×10^{-12} , and PolyKernel type with E 1.0 and C 250007. In J48, the confidence threshold for pruning was set to 0.25, and the minimum number of instances per leaf was set to 2. The LR was run using the default setting. Finally, the MLP is a shallow deep learning model, constructed using one hidden layer with 12 neurons, and the activation function was sigmoid.

The performance evaluations for the TAL and SPAL models are illustrated in Tables 6 and 7. In Table 6, the SPAL model vies the TAL model concerning the accuracy metric by using different classifiers. The highest percentage of accuracy was obtained by using J48 classifier with values 70.83% and 71.73% in the TAL and SPAL models, respectively. On the other hand, the NB classifier scored the lowest percentage of accuracy, equal to 58.33% in the TAL model and 59.18% in the SPAL model. The J48 classifier is based on the decision tree, while the NB classifier is based on probability [53,55]. Moreover, J48 iteratively partitions the dataset using the best informative attribute (the splitting attribute is the attribute of the maximum gain ratio) [55–57]. While the NB classifier assumes no correlation exists between the different test-attributes, and each attribute has an equal impact on the tweet's polarity [53,54]. The J48 is commonly utilized for some classification tasks, such as emotion recognition from text and Twitter's text categorizations, although it is rarely used for sentiment prediction [64]. The highest accuracy value obtained by invoking the J48 classifier in both models results from a 20% contribution of the emotional features among the other extracted features. These results respond to the first research question addressed in the present study.

Figure 6 shows the number of correctly classified and misclassified tweets for TAL and SPAL models on different classifiers. This is evidence that partitioning the collected data into a number of disjoint subsets will increase the number of correctly classified tweets. The philosophy of partitioning is not applied randomly nor regionally. The partitioning step in the SPAL model properly employs the most important domains that the collected datasets rely on. Therefore, rather than processing all tweets at once using the TAL model, the proposed model SPAL partitions the collected tweets into a number of mutually exclusive subsets by utilizing the hidden semantic meaning among the tweets. The percentage amount of improvement in the accuracy using SPAL model reached up to 3.99% in SVM, 1.46% in NB, 1.27% in J48, 4.22% in MLP, and finally 1.07% in LR compared to the TAL model. It is clearly noticed that the percentage amount of improvement in accuracy metrics using the SVM and MLP classifiers is higher than the percentage amount of improvement using other classifiers. The reason behind this phenomenon is the power of merging the machine learning classifiers with the semantic partitioning in the SPAL model, where machine learning classifiers used a nonlinear stimulation function that helped capture the complex features in the hidden layers [58]. On the other hand, the integration between the

SPAL model, particularly the MLP classifier, scores the highest improvements, with respect to the accuracy metric. This is because the SPAL model uses the semantic partitioning to investigate the hidden correlation between tweets, while the MLP has a strong associative memory and prediction capability after training [65]. This illustrates the power of the SPAL model with respect to the TAL model. These results respond to the second research question addressed in this study.

Table 7 presents the recall, precision, and F-score metrics for TAL and SPAL models. It is well known that the recall and precision metrics present a completely different perspective of mentioned models. These metrics conflict with each other. Therefore, there is a need for a fair index (called F-score) that considers both of them simultaneously. An F-score is considered perfect when it approaches one, while the model is a total failure when it approaches 0. It is clearly noted that the SPAL model scored a higher weighted average F-score value than the TAL model on various classifiers, with a percentage of improvement that reached up to 8.70% in J48, 7.48% in SVM, 4.59% in NB, 6.74 in MLP, and 5.94% in LR. To the best of our knowledge, a good F-score value indicates a minimum number of false positives and a minimum number of false negatives. Consequently, the model correctly identifies real threats, not disturbed by false alarms [9,66].

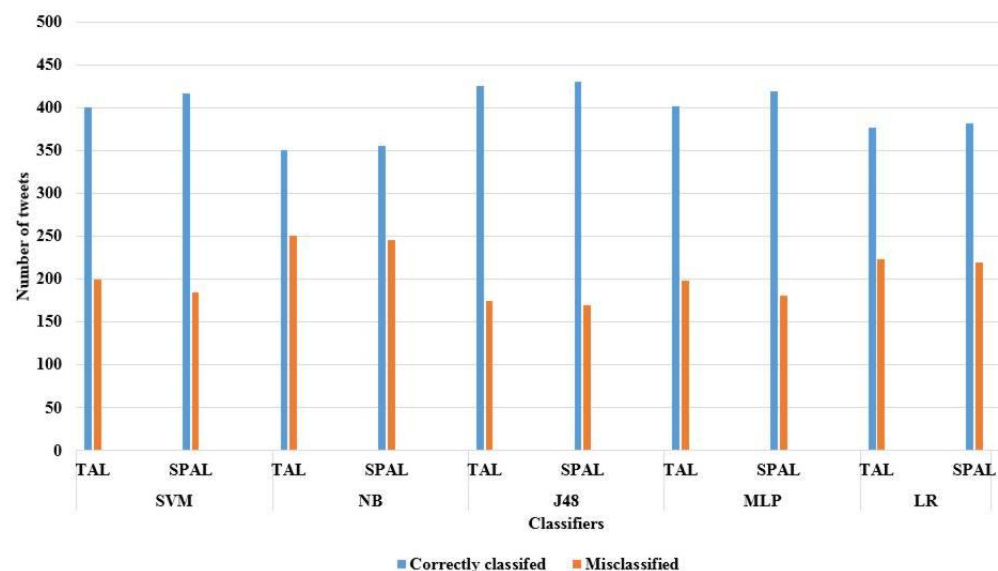


Figure 6. The number of correctly classified and misclassified tweets for TAL and SPAL models.

Table 6. Measuring the accuracy results of the analysis model of over 600 tweets (split 30% of tweets are used for testing).

Classified Method	Analysis Model	Accuracy (%)	Correctly Classified	Misclassified
SVM	TAL	66.67	400	200
	SPAL	69.33	416	184
NB	TAL	58.33	350	250
	SPAL	59.18	355	245
J48	TAL	70.83	425	175
	SPAL	71.73	430	170
MLP	TAL	67.00	402	198
	SPAL	69.83	419	181
LR	TAL	62.83	377	223
	SPAL	63.50	381	219

Table 7. Measuring the performance of the analysis model of over 600 tweets (split 30% of tweets are used for testing).

Classified Method	Analysis Model	Polarity	Precision	Recall	F-Score
SVM	TAL	Positive	0.711	0.844	0.771
		Negative	0.500	0.313	0.385
		Weighted Average	0.640	0.667	0.642
	SPAL	Positive	0.462	0.429	0.444
		Negative	0.778	0.800	0.789
		Weighted Average	0.687	0.694	0.690
NB	TAL	Positive	0.667	0.750	0.706
		Negative	0.333	0.250	0.286
		Weighted Average	0.556	0.583	0.566
	SPAL	Positive	0.286	0.286	0.286
		Negative	0.714	0.714	0.714
		Weighted Average	0.592	0.592	0.592
J48	TAL	Positive	0.705	0.969	0.816
		Negative	0.750	0.188	0.300
		Weighted Average	0.720	0.708	0.644
	SPAL	Positive	0.600	0.400	0.480
		Negative	0.750	0.871	0.806
		Weighted Average	0.701	0.717	0.700
MLP	TAL	Positive	0.742	0.869	0.795
		Negative	0.528	0.354	0.374
		Weighted Average	0.652	0.673	0.658
	SPAL	Positive	0.442	0.471	0.468
		Negative	0.795	0.797	0.792
		Weighted Average	0.696	0.712	0.693
LR	TAL	Positive	0.896	0.670	0.604
		Negative	0.253	0.365	0.374
		Weighted Average	0.623	0.572	0.598
	SPAL	Positive	0.500	0.071	0.625
		Negative	0.723	0.971	0.829
		Weighted Average	0.660	0.714	0.628

The SPAL model superiorly competes with the TAL model in predicting the polarity of the Jordanian dialect tweets based on the extracted lexicon, writing style, grammatical, and emotional features. Utilizing the hidden semantic meaning among the tweets helps enhance the general performance of the used model. The main advantage of the SPAL model is to polarize the dataset in an expressively better performance, especially accuracy and F-score metrics. Another advantage of the proposed model is that such a strategy is applicable in multi-processor machines, particularly shared-memory systems (where there is no need to plan the communication of data between different processors). Moreover, the memory caches will efficiently be used because a subset size is small enough to be stored in cache, and then the partitioning can be achieved without accessing the slower main memory.

Finally, addressing the third research question in this study, which focuses on the reaction of the Jordanian citizens, regarding governmental procedures and precautions during the COVID-19 pandemic, is presented in Figure 7. Jordan is divided into three main regions: north, middle, and south. Statistically, 2000 tweets were divided: 667 tweets in the northern region, 667 tweets in the middle region, and the remaining 666 tweets in the southern region. It is worth noting that the middle region had the highest percentage of negative tweets, 66.27%, in regards to governmental actions and decisions during the COVID-19 pandemic, while the south region registered the highest percentage of positive tweets, 64.71%. This refers to the fact that the middle region encompasses the capital city of Jordan (Amman), which has the most governmental and private companies, as well as institutes that had suffered economically from the lockdown and curfew actions.

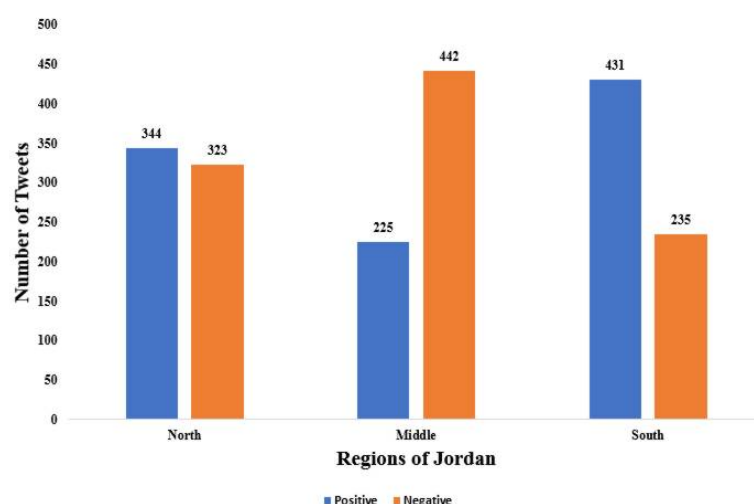


Figure 7. The number of tweet classes (positive and negative) for Jordan based on three regions.

4. Conclusions

In this paper, we considered the SA problem using dialectical Arabic language. We collected 2000 tweets written in Jordan during the COVID-19 pandemic. Herein, we proposed two models to predict the polarity of the collected tweets by invoking SVM, NB, J48, MLP, and LR classifiers. The extraction of different Arabic features and the utilization of the hidden semantic meaning when imposing the partitioning enhance the overall performance of the SPAL model. Our experimental results show an improvement in the accuracy—up to 4.22% in the SPAL model compared to the TAL model when the MLP classifier was invoked. Moreover, an improvement in the weighted F-score, which presents a fair index that optimizes the conflict between recall and precision, the SPAL model against the TAL model reached up to 8.70% in J48 classifier.

The benefits of adapting this proposed model can be itemized into scientific and practical perspectives. The proposed SPAL model polarizes the dataset with better performance, especially accuracy and F-score metrics from the scientific perspective. This is accomplished by discovering the massive amount of knowledge and their hidden relations presented in social networks. Therefore, it is considered a “blossom” tool to enrich the dialectical Arabic SA field. Another advantage of the proposed SPAL model is that such a strategy is applicable in multi-processor machines, particularly shared-memory systems (where there is no need to plan data communication between different processors). Moreover, the memory caches will efficiently be used because the subset size is small enough to be stored in the cache, and then the partitioning can be achieved without accessing the slower main memory. Practically, it is well known that social media is a fertile environment, full of comments, emotions, thoughts, and opinions regarding common events and topics in society [67]. Integrating the SPAL model with social media will help organizations, companies, and governments in mining the hidden metadata, in order to improve the quality of their services to the end-users and to sustain human satisfaction [1]. This could not be achieved without detecting the changes in human opinions.

The limitations of the proposed SPAL model can be described on the following fronts: the partitioning of the collected tweets is achieved manually, not automatically. Moreover, polarity is manually assigned for the collected tweets. Finally, the accuracy of the invoked various classifiers depends on the nature of the data, the number of extracted features, and their types.

We foresee numerous avenues for future work. First, we propose redesigning the proposed SPAL to form parallelism. To the best of our knowledge, focusing on designing a fast parallel Arabic SA model has not yet been explored. Another avenue that can be explored is to study whether the SPAL model can be hierarchically designed to implement semantic partitioning. Moreover, an investigation is required to perform the semantic

partitioning automatically by using artificial intelligence techniques, such as an interactive associative classifier [68] or dependency-based information [69] to deduce the hidden semantic relations in the collected tweets. A final avenue for future work is to investigate the performance of the SPAL model when employing different partitioning techniques, such as random, statistical, or graphical partitioning approaches.

Author Contributions: Conceptualization, E.F.; methodology, E.F. and S.I.; software, E.F. and S.I.; formal analysis, E.F.; investigation, E.F. and S.I.; data curation, S.I.; writing—original draft preparation, E.F. and S.I.; writing—review and editing, E.F. and S.I. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The dataset can be obtained upon request from the author.

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

References

1. Saura, J.R. Using Data Sciences in Digital Marketing: Framework, Methods, and Performance Metrics. *J. Innov. Knowl.* **2020**, *6*, 92–102. [CrossRef]
2. Kastrenakes, J. Twitter’s Final Monthly User Count Shows a Company Still Struggling to Grow. THE VERGE. 23 April 2019. Available online: <https://www.theverge.com/2019/4/23/18511383/twitter-q1-2019-earnings-report-mau> (accessed on 1 April 2021).
3. Boot, A.B.; Sang, E.T.K.; Dijkstra, K.; Zwaan, R.A. How Character Limit Affects Language Usage in Tweets. *Palgrave Commun.* **2019**, *5*, 76. [CrossRef]
4. Boudad, N.; Faizi, R.; Thami, R.O.H.; Chiheb, R. Sentiment Analysis in Arabic: A Review of the Literature. *Ain Shams Eng. J.* **2018**, *9*, 2479–2490. [CrossRef]
5. Liu, B. Sentiment Analysis and Opinion Mining. In *Synthesis Lectures on Human Language Technologies*; Morgan and Claypool Publishers: San Rafael, CA, USA, 2012; Volume 5, pp. 1–167.
6. Groh, G.; Hauffa, J. Social relations via NLP-based sentiment analysis. In Proceedings of the 5th International AAAI Conference on Weblogs and Social Media, Barcelona, Spain, 17–21 July 2011.
7. Gaur, N.; Sharma, N. Sentiment Analysis in Natural Language Processing. *Int. J. Eng. Technol.* **2017**, *3*, 144–148.
8. Al Shamsi, A.A.; Abdallah, S. Text Mining Techniques for sentiment Analysis of Arabic Dialects: Literature Review. *Adv. Sci. Technol. Eng. Syst. J.* **2021**, *6*, 1012–1023. [CrossRef]
9. Tsytsarau, M.; Palpanas, T. Survey on mining subjective data on the web. *Data Min. Knowl. Discov.* **2012**, *24*, 478–514. [CrossRef]
10. Medhat, W.; Hassan, A.; Korashy, H. Sentiment Analysis Algorithms and Applications: A Survey. *Ain Shams Eng. J.* **2014**, *5*, 1093–1113. [CrossRef]
11. Ghallab, A.; Mohsen, A.; Ali, Y. Arabic Sentiment Analysis: A Systematic Literature Review. *Appl. Comput. Intell. Soft Comput.* **2020**, *2020*, 1–21. [CrossRef]
12. Khan, K.; Baharudin, B.; Khan, A.; Ullah, A. Mining Opinion Components from Unstructured Reviews: A Review. *J. King Saud Univ. Comput. Inf. Sci.* **2014**, *26*, 258–275. [CrossRef]
13. Kumar, M.R.P.; Prabhu, J. Role of sentiment classification in sentiment analysis: A survey. *Ann. Libr. Inf. Stud.* **2018**, *65*, 196–209.
14. Alshamsi, A.; Bayari, R.; Salloum, S. Sentiment Analysis in English Texts. *Adv. Sci. Technol. Eng. Syst. J.* **2020**, *5*, 1683–1689. [CrossRef]
15. Halabi, D.; Awajan, A.; Fayyumi, E. Syntactic Annotation in the I3rab Dependency Treebank. *Int. Arab. J. Inf. Technol.* **2021**, *18*, 1.
16. UNESCO. History of the Arabic Language at UNESCO. Available online: <http://www.unesco.org/new/en/unesco/resources/history-of-the-arabic-language-at-unesco/> (accessed on 1 April 2021).
17. Farghaly, A.; Shaalan, K. Arabic natural language processing: Challenges and solutions. *ACM Trans. Asian Lang. Inf. Process.* **2009**, *8*, 1–22. [CrossRef]
18. Rushdi-Saleh, M.; Martin-Valdivia, M.-T.; Ureña-López, L.A.; Perea-Ortega, J.M. OCA: Opinion Corpus for Arabic. *J. Am. Soc. Inf. Sci. Technol.* **2011**, *62*, 2045–2054. [CrossRef]
19. Alotaibi, S.S. Sentiment Analysis in the Arabic Language Using Machine Learning. Ph.D. Thesis, Colorado State University, Fort Collins, CO, USA, 2015.
20. Defradas, M.; Embarki, M. Typology of Modern Arabic Dialects: Features, Methods and Models of Classification. In Proceedings of the Typology of Modern Arabic Dialects: Features, Methods and Models of Classification, Montpellier, France, 14–15 May 2007.
21. Thakkar, H.; Patel, D. Approaches for Sentiment Analysis on Twitter: A State-of-Art study. *arXiv* **2015**, arXiv:1512.01043.

22. Biltawi, M.; Etaiwi, W.; Tedmori, S.; Hudaib, A.; Awajan, A. Sentiment Classification Techniques for Arabic Language: A Survey. In Proceedings of the 7th International Conference on Information and Communication Systems (ICICS), Irbid, Jordan, 5–7 April 2016; pp. 339–346.
23. El-Jawad, M.H.A.; Hodhod, R.; Omar, Y.M.K. Sentiment Analysis of Social Media Networks Using Machine Learning. In Proceedings of the 14th International Computer Engineering Conference (ICENCO), Cairo, Egypt, 29–30 December 2018.
24. Araújo, M.; Diniz, M.L. A comparative study of machine translation for multilingual sentence-level sentiment analysis. *Inf. Sci.* **2020**, *512*, 1078–1102. [\[CrossRef\]](#)
25. Al-Shabi, A.; Adel, A.; Omar, N.; Al-Moslmi, T. Cross-Lingual Sentiment Classification from English to Arabic Using Machine Translation. *Int. J. Adv. Comput. Sci. Appl.* **2017**, *8*, 1. [\[CrossRef\]](#)
26. Barhoumi, A.; Aloulou, C.; Camelin, N.; Estève, Y.; Belguith, L.H. Arabic Sentiment analysis: An empirical study of machine translation's impact. In Proceedings of the Language Processing and Knowledge Management (LPKM), Sfax, Tunisia, 17–18 October 2018.
27. Oueslati, O.; Cambria, E.; Ben HajHmida, M.; Ounelli, H. A Review of Sentiment Analysis Research in Arabic Language. *Futur. Gener. Comput. Syst.* **2020**, *112*, 408–430. [\[CrossRef\]](#)
28. Duwairi, R.; El-Orfali, M. A Study of the Effects of Preprocessing Strategies on Sentiment Analysis for Arabic Text. *J. Inf. Sci.* **2014**, *40*, 501–513. [\[CrossRef\]](#)
29. Balamurali, A.R.; Khapra, M.; Bhattachary, P. Lost in Translation: Viability of Machine Translation for Cross Language Sentiment Analysis. In Proceedings of the 14th International Conference on Computational Linguistics and Intelligent Text Processing, Samos, Greece, 24–30 March 2013.
30. Guellil, I.; Saâdane, H.; Azouaou, F.; Gueni, B.; Nouvel, D. Arabic Natural Language Processing: An Overview. *J. King Saud Univ. Comput. Inf. Sci.* **2021**, *33*, 497–507. [\[CrossRef\]](#)
31. Al-Osaimi, S.; Khan, M.B. Sentiment Analysis Challenges of Informal Arabic Language. *Int. J. Adv. Comput. Sci. Appl.* **2017**, *8*, 1. [\[CrossRef\]](#)
32. Al-Ayyoub, M.; Khamaiseh, A.A.; Jararweh, Y.; Al-Kabi, M.N. A Comprehensive Survey of Arabic Sentiment Analysis. *Inf. Process. Manag.* **2019**, *56*, 320–342. [\[CrossRef\]](#)
33. Badaro, G.; Baly, R.; Hajj, H.; El-Hajj, W.; Shaban, K.B.; Habash, N.; Al-Sallab, A.; Hamdi, A. A survey of opinion mining in arabic: A comprehensive system perspective covering challenges and advances in tools, resources, models, applications, and visualizations. *ACM Trans. Asian Low Resour. Lang. Inf. Process.* **2019**, *18*, 1–52. [\[CrossRef\]](#)
34. Abdallah, E.E.; Abo-Suaileek, S.A.-S. Feature-based Sentiment Analysis for Slang Arabic Text. *Int. J. Adv. Comput. Sci. Appl.* **2019**, *10*, 298–304. [\[CrossRef\]](#)
35. Alomari, K.M.; Elsherif, H.M.; Shaalan, K. Arabic Tweets Sentimental Analysis Using Machine Learning. In Proceedings of the International Conference on Industrial, Engineering and Other Applications of Applied Intelligent Systems, Arras, France, 27–30 June 2017.
36. Al-Harbi, O. Classifying Sentiment of Dialectal Arabic Reviews: A Semi-Supervised Approach. *Int. Arab. J. Inf. Technol.* **2019**, *16*, 995–1002.
37. Abdulla, N.A.; Ahmed, N.A.; Shehab, M.A.; Al-Ayyoub, M. Arabic Sentiment Analysis: Lexicon-based and Corpus-based. In Proceedings of the 2013 IEEE Jordan Conference on Applied Electrical Engineering and Computing Technologies (AEECT), Amman, Jordan, 3–5 December 2013.
38. Abdul-Mageed, M.; Diab, M.; Kübler, S. SAMAR: Subjectivity and Sentiment Analysis for Arabic Social Media. *Comput. Speech Lang.* **2014**, *28*, 20–37. [\[CrossRef\]](#)
39. Alhumoud, S.O.; Al Wazrah, A.A. Arabic Sentiment Analysis Using Recurrent Neural Networks: A Review. *Artif. Intell. Rev.* **2021**, 1–42. [\[CrossRef\]](#)
40. Hussien, I.O.; Dashtipour, K.; Hussain, A. Comparison of Sentiment Analysis Approaches Using Modern Arabic and Sudanese Dialect. In *Lecture Notes in Computer Science*; Springer: Cham, Switzerland, 2018; Volume 10989, pp. 615–624.
41. Huang, S.; Han, T.; Ansari, N. Big-data-driven network partitioning for ultra-dense radio access networks. In Proceedings of the 2017 IEEE International Conference on Communications (ICC), Paris, France, 21–25 May 2017.
42. Sawant, M.; Kinage, K.; Pilankar, P.; Chaudhari, N. Database Partitioning: A Review Paper. *Int. J. Innov. Technol. Explor. Eng.* **2013**, *3*, 82–85.
43. Fayyoubi, E.; ALhiniti, O. Recursive Genetic Micro-Aggregation Technique: Information Loss, Disclosure Risk and Scoring Index. *Data* **2021**, *6*, 53. [\[CrossRef\]](#)
44. Hasan, H.; Chuprat, S. Secured data partitioning in multi cloud environment. In Proceedings of the 2014 4th World Congress on Information and Communication Technologies (WICT 2014), Malacca, Malaysia, 8–11 December 2014.
45. Kaviani, N.; Wohlstadter, E.; Lea, R. Partitioning of Web Applications for Hybrid Cloud Deployment. *J. Internet Serv. Appl.* **2014**, *5*, 14. [\[CrossRef\]](#)
46. Aljameel, S.S.; Alabbad, D.A.; Alzahrani, N.A.; AlQarni, S.M.; AlAmoudi, F.A.; Babili, L.M.; Aljaafary, S.K.; Alshamrani, F.M. A Sentiment Analysis Approach to Predict an Individual's Awareness of the Precautionary Procedures to Prevent COVID-19 Outbreaks in Saudi Arabia. *Int. J. Environ. Res. Public Health* **2020**, *18*, 218. [\[CrossRef\]](#)
47. Albahli, S.; Algsham, A.; Aeraj, S.; Alsaeed, M.; Alrashed, M.; Rauf, H.T.; Arif, M.; Mohammed, M.A. COVID-19 Public Sentiment Insights: A Text Mining Approach to the Gulf Countries. *Comput. Mater. Contin.* **2021**, *67*, 1613–1627. [\[CrossRef\]](#)

48. Alhumoud, S. Arabic Sentiment Analysis using Deep Learning for COVID-19 Twitter Data. *IJCSNS Int. J. Comput. Sci. Netw. Secur.* **2020**, *20*, 132–183.
49. Al-Tammemi, A.B. The Battle against COVID-19 in Jordan: An Early Overview of the Jordanian Experience. *Front. Public Health* **2020**, *8*, 1. [[CrossRef](#)] [[PubMed](#)]
50. Zhao, H.; Zhang, X.; Li, K. A Sentiment Classification Model Using Group Characteristics of Writing Style Features. *Int. J. Pattern Recognit. Artif. Intell.* **2017**, *31*, 1. [[CrossRef](#)]
51. Jakkula, V. *Tutorial on Support Vector Machine (SVM)*; School EECS, Washington State University: Washington, DC, USA, 2011.
52. Awad, M.; Khanna, R. Support Vector Machines for Classification. In *Efficient Learning Machines Theories, Concepts, and Applications for Engineers and System Designers*; Apress Open: New York, NY, USA, 2015.
53. Lowd, D.; Domingos, P. Naive Bayes models for probability estimation. In Proceedings of the 22nd International Conference on Machine Learning, Bonn, Germany, 7–11 August 2005; pp. 529–536.
54. Nguyen, S.T.; Do, P.M.T. Classification optimization for training a large dataset with Naïve Bayes. *J. Comb. Optim.* **2020**, *40*, 141–169. [[CrossRef](#)]
55. Sharma, P. Comparative Analysis of Various Decision Tree Classification Algorithms Using WEKA. *Int. J. Recent Innov. Trends Comput. Commun.* **2015**, *3*, 684–690. [[CrossRef](#)]
56. Patel, N.; Upadhyay, S. Study of Various Decision Tree Pruning Methods with Their Empirical Comparison in WEKA. *Int. J. Comput. Appl.* **2012**, *60*, 20–25. [[CrossRef](#)]
57. Karabulut, E.M.; Özel, S.A.; Ibrikci, T. A Comparative Study on the Effect of Feature Selection on Classification Accuracy. *Procedia Technol.* **2012**, *1*, 323–327. [[CrossRef](#)]
58. Fayyumi, E.; Idwan, S.; AboShindi, H. Machine Learning and Statistical Modelling for Prediction of Novel COVID-19 Patients Case Study: Jordan. *Int. J. Adv. Comput. Sci. Appl.* **2020**, *11*, 122–126. [[CrossRef](#)]
59. Molnar, C. *Interpretable Machine Learning: A Guide for Making Black Box Models Explainable*, 1st ed.; Lulu: Morrisville, NC, USA, 2019.
60. Fang, J. Why Logistic Regression Analyses Are More Reliable than Multiple Regression Analyses. *J. Bus. Econ.* **2013**, *4*, 620–633.
61. Dencelin, L.; Ramkumar, T. Analysis of multilayer perceptron machine learning approach in classifying protein secondary structures. *Biomed. Res. Comput. Life Sci. Smarter Technol. Adv.* **2016**, *1*, S166–S173.
62. Han, J.; Kamber, M. *Data Mining: Concepts and Techniques*; Morgan Kaufmann Publishers: San Francisco, CA, USA, 2001.
63. Frank, E.; Hall, M.A.; Holmes, G.; Kirkby, R.B.; Pfahringer, B.; Witten, I.H.; Trigg, L. Weka-A Machine Learning Workbench for Data Mining. In *Data Mining and Knowledge Discovery Handbook*; Springer: Boston, MA, USA, 2009; pp. 1269–1277.
64. Singh, J.; Singh, G.; Singh, R. Optimization of Sentiment Analysis Using Machine Learning Classifiers. *Human Centric Comput. Inf. Sci.* **2017**, *7*, 32. [[CrossRef](#)]
65. Al-Batah, M.S.; Mrayyen, S.; Alzaqebah, M. Arabic Sentiment Classification Using MLP Network Hybrid with Naive Bayes Algorithm. *J. Comput. Sci.* **2018**, *14*, 1104–1114. [[CrossRef](#)]
66. Sokolova, M.; Japkowicz, N.; Szpakowicz, S. Beyond Accuracy, F-Score and ROC: A Family of Discriminant Measures for Performance Evaluation. In *AI 2006: Advances in Artificial Intelligence*; Springer: Berlin, Germany, 2006; pp. 1015–1021.
67. Furini, M.; Montangero, M. TSentiment: On Gamifying Twitter Sentiment Analysis. In Proceedings of the 2016 IEEE Symposium on Computers and Communication (ISCC), Messina, Italy, 27–30 June 2016; pp. 91–96.
68. Oommen, J.; Fayyumi, E. A novel method for micro-aggregation in secure statistical databases using association and interaction. In Proceedings of the International Conference on Information and Communications Security, Zhengzhou, China, 12–15 December 2007; pp. 126–140.
69. Oommen, J.; Fayyumi, E. On utilizing dependence-based information to enhance micro-aggregation for secure statistical databases. *Pattern Anal. Appl.* **2013**, *16*, 99–116. [[CrossRef](#)]