

Deep Learning Approach for Aspect-Based Sentiment Analysis of Restaurants Reviews in Spanish

Bella-Citlali Martínez-Seis¹, Obdulia Pichardo-Lagunas¹, Sabino Miranda^{2,3},
Israel-Josafat Perez-Cazares¹, Jorge-Armando Rodriguez-González¹

¹ Instituto Politécnico Nacional,
Unidad Profesional Interdisciplinaria en Ingeniería y Tecnologías Avanzadas,
Mexico

² INFOTEC Centro de Investigación e Innovación en Tecnologías de la Información y Comunicación,
Mexico

³ CONACyT Consejo Nacional de Ciencia y Tecnología,
Dirección de Cátedras,
Mexico

{bcmartinez, opichardola}@ipn.mx, sabino.miranda@infotec.mx

Abstract. Online reviews of products and services have become important for customers and enterprises. Recent research focuses on analyzing and managing those kinds of reviews using natural language processing. This paper focuses on aspect-based sentiment analysis for reviews in Spanish. First, the reviews data sets are normalized into different inputs of the neural networks. Then, our approach combines two deep learning models architectures to determine a positive or negative assessment and identify the most important characteristics or aspects of the text. We develop two architectures for aspect detection and three architectures for sentiment analysis. Merging the deep learning models, we tested our approach in restaurant reviews and compared them with state-of-the-art methods.

Keywords. Customer reviews, polarity classification, sentiment analysis.

1 Introduction

Nowadays, it is common to write a review of a product or service. 84% of consumers trust online reviews as much as a personal recommendation [22]; therefore, review sites have become crucial to consumers. Enterprises consider those reviews

as feedback for their products [4]. It allows them to analyze strengths and weaknesses in order to improve the service or product.

Recent research focuses on analyzing and managing those reviews using natural language processing. Aspect-based sentiment analysis not only determines a positive or negative assessment, but also identifies the most important characteristics or aspects of the text [17, 7]. For example, the following review, *A bad service cannot be saved by a good food*, should be classified as positive on the food area, but negative on the service area. It is a major technological challenge [20] because even humans often disagree on the sentiment of a given text, and moreover, on the aspect that the text is talking about.

Enterprises pursue a positive reputation as one of the most powerful marketing assets. This paper focuses on processing, analyzing, and categorizing the large accumulations of information generated from the reviews.

Our approach combines two deep learning models for aspect-based sentiment analysis. The reviews were normalized into five different data

sets inputs for the test of the neural networks. We compare state-of-the-art approaches, and the performance of our aspect classification proposal is promising, but there is still work to do in the sentiment detection.

The rest of the paper is organized as follows. Section 2 is a review of the research involving aspect-based sentiment analysis in Spanish, mainly based on subtask 1 of task 5 within the 2016 edition of SemEval competition [19]. After, Section 3 describes the architecture of our approach, describing the used architectures of the models of deep learning. Following that, Section 4 compares our different architectures, and then, in Section 4.3, we compare our approach with state-of-the-art approaches. Finally, Section 5 concludes the paper and presents future works.

2 Related Work

The interest in sentiment analysis is growing by the need of knowing the polarity of the opinions published on the Internet. Recent research focuses on aspect-based sentiment analysis. For aspect-based sentiment analysis, there are two main tasks: aspect detection and sentiment detection. For aspect detection, we have two possibilities, to recognize the general aspect, for example, "Food" or to identify not only the aspect but its sub-aspect, for instance, "Food, prices."

Earlier approaches for aspect category detection were based on word frequency [13]. Some recent works use Latent Dirichlet Allocation (LDA) where each topic is characterized by a distribution over words [9] [23]. A different approach is in [3], they propose a modular approach focus on Spanish tweets; it is based on a graph-based algorithm for the general aspect classification and a large number of features and polarity lexicons for sentiment detection.

Supervised methods had been used for this task. Some of the most common classifiers are Support Vector Machine (SVM) [2] [15] [1][18], Maximum Entropy (ME) [11, 18], and Conditional Random Field (CRF) [1] [15]. Some of them use more than one. Other hybrids methods had been proposed [10] using rule based methods with optimization.

Considering the ability to learn useful features from low-level data[14], Deep Learning (DL) has become a popular approach for Aspect-Based Sentiment Analysis [8] [10]. It uses multiple layers to progressively extract higher-level features from the raw input. It allows to capture the correlation between non-consecutive words focusing the attention on the specific significant words [24].

2.1 SemEval 2016 Competition

SemEval competition [19] boosts the research on this area. They publish different tasks to be solved. In the 2016 Edition, the subtask 1 of the task 5 was related to aspect-based sentiment analysis. The Spanish language is known for its complexity. The main competitors in Spanish are described below:

Focus on general aspects, [1] achieves a high performance using Support Vector Machine (SVM) with a list of words with a preprocessing stage using the Freeing tagger and dictionaries.

Focus on both tasks, IIT-TUDA group [15] also uses SVM and combines with several tools such as dependency graphs, distributional thesaurus (DT), scores, and a bag of words; they achieve a better result in each task. Similar to IIT-TUDA, the UWB team [11] uses different approaches to optimize the results of the tasks, and they use a Max Entropy Classifiers as their primary classifier.

TGB team [6] uses binary and multi-class linear classifiers. The INSIGHT-1 team [21] uses a Convolutional Neural Networks (CNN) to obtain a similar score. Conditional Random Fields (CRF) improves the performance of the algorithms when they are used in sub-aspects detection [1][15]. Also, having an excellent preprocessing module is of great importance; some works use taggers, parsers, tokenization, filters, dictionaries, among others.

3 System Architecture

The system architecture includes two Convolutional Neural Network (CNN) models with a previous normalization stage. Figure 1 shows that the reviews are preprocessed. The data cleaning process eliminates and replaces emojis, URLs, and special characters.

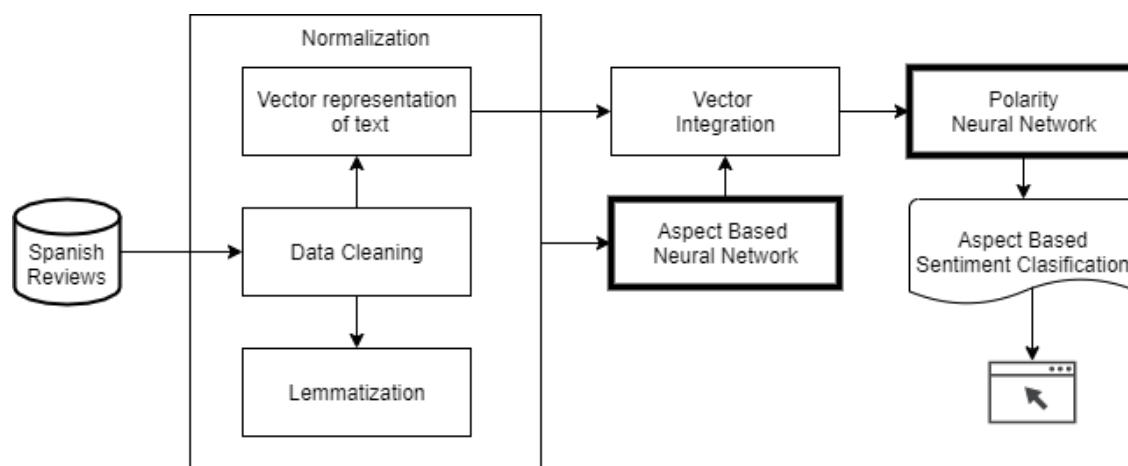


Fig. 1. Simplified block diagram of the system architecture

Token: Buen, Lemma: Buen, Norm: buen , Tag: ADJ__Gender=M
 Token: servicio, Lemma: servicio, Norm: servicio , Tag: N
 Token: ambiente, Lemma: ambiental, Norm: ambiente , Tag: I
 Token: Acogedor, Lemma: Acogedor, Norm: acogedor , Tag: P

Fig. 2. Output of lemmatization process

Once the text is clean, the corresponding word embedding vector is generated using fastText. Also, the cleaned text goes through lemmatization using the spaCy tool, and then it is normalized.

The normalized text is the input of the neural network. The aspect-based neural network model calculates the vector of aspects. The aspect vector combined with the Word Embedding vector are the inputs of the polarity neural network. It produces an aspect-based sentiment classification of restaurant reviews.

It is important to mention that each review has several sentences; then, parallel detection of the classification of aspects and feelings is not recommended. It is because the relation between the polarity and its aspect is lost.

3.1 Normalization

The following processes are essential to prepare the text before processing the reviews into the neural network architecture.

3.1.1 Data Cleaning

First, we remove mentions, URLs, emoticons, and special characters using regular expressions. We keep accented letters and punctuation marks. Because of the unbalanced data sets, the samples of the negative, neutral, and conflict classes were augmented.

3.1.2 Lemmatization

In recent years, the spaCy API [12] has been popular in applications of Natural Language Processing (NLP). We used the model *es_core_wlg*, which is the largest spaCy model for Spanish; it is pre-trained with texts from the web of general purpose.

The first step is to tokenize the text. Then, a tagger process is used to label the tokens of the previous step according to the part-of-speech. Then, the labels of each token are obtained from the parsing process. Finally, the approach detects and labels the entities of the text.

3.1.3 Vector Representation

The reviews were represented as vectors of real numbers using word embeddings. This approach uses the fastText *Multi-lingual word embeddings or word vectors*. This model supports 157 languages, one of them is Spanish. It was previously trained using Common Crawl (a non-profit organization that crawls the web and provides its files and data sets to the public for free) and Wikipedia (free online encyclopedia).

Figure 2 shows an output of this process.

We iterate between 100 and 300 dimensions; and also include the linguistic components obtained from spaCy to improve the performance of the neural networks. The output of the normalization gives five possible inputs for the deep learning phase:

- **A100.** A word embedding vector of 100 dimensions of lemmas from spaCy and fastText,
- **A300.** A word embedding vector of 300 dimensions of lemmas from spaCy and fastText,
- **B100.** A word embedding vector of 100 dimensions of normalized tokens from spaCy and fastText,
- **B300.** A word embedding vector of 300 dimensions of normalized tokens from spaCy and fastText,
- **C300.** A word embedding vector of 300 dimensions extracted with word2Vec using spaCy.

3.2 Aspect-Based Convolutional Neural Network (AB-CNN)

We defined a Convolutional Neural Networks (CNNs) architecture using Tensorflow and Keras for the aspect classification. We defined two architectures, AB1-CNN with a sequential model and AB2-CNN with multi-channel output. Figure 3 shows both architectures, and they are described as follows.

3.2.1 AB1-CNN

The input layer has a One-Dimensional (1D) Convolutional Neural Network with a kernel of size three, and the activation function Rectified Linear Unit (ReLU). The next layer is a 1D Max Pooling; it has an outstanding performance in conjunction with the convolutional layers. To reduce overfitting, a Dropout layer was placed to disable and activate different neurons during the forward-propagation and backward-propagation processes.

For this layer, approximately 20% of neurons remain deactivated. Then, those three layers are repeated. In addition, the architecture has a Flatten Layer and a Dropout of 50%. Finally, two dense layers were added, with a Dropout in the middle, using different activation functions.

3.2.2 AB2-CNN

It is a model with multi-channel output. There are four 1D Convolutional layers with kernels of 3, 3, 5, and 4; all of them with ReLU activation function. After them, a Max Pooling layer and an LSTM (Long-Short Term Memory) are added. The LSTM is an extension of recurrent neural networks, they expand their memory to learn from previous experiences. The output was n dense layers with sigmoid activation function, all of them are connected to the LSTM.

3.3 Model of the Polarity Neural Network (P-CNN)

The polarity was identified by defining Convolutional Neural Networks (CNNs) using TensorFlow and Keras. We defined three architectures: P1-CNN and P3-CNN are sequential models differing in one layer, and P2-CNN has two channels. Figure 4 shows them, and they are described below.

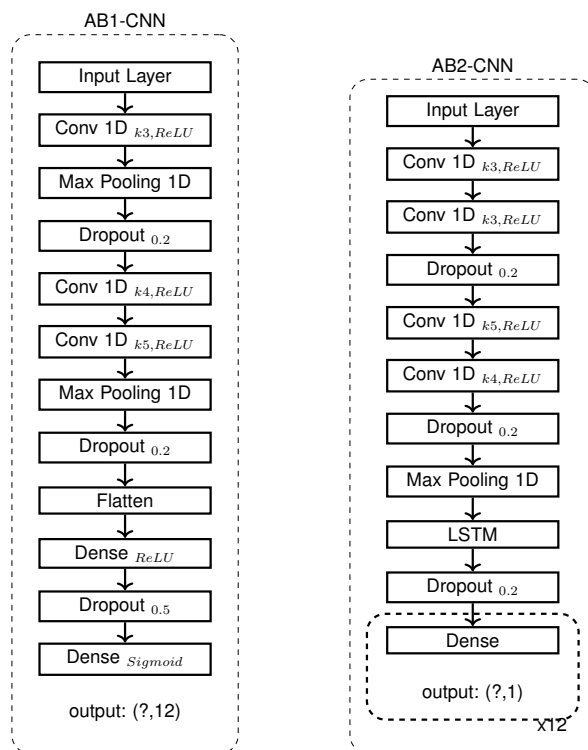


Fig. 3. Two architecture approaches of the Aspect-Based Convolutional Neural Network (AB-CNN)

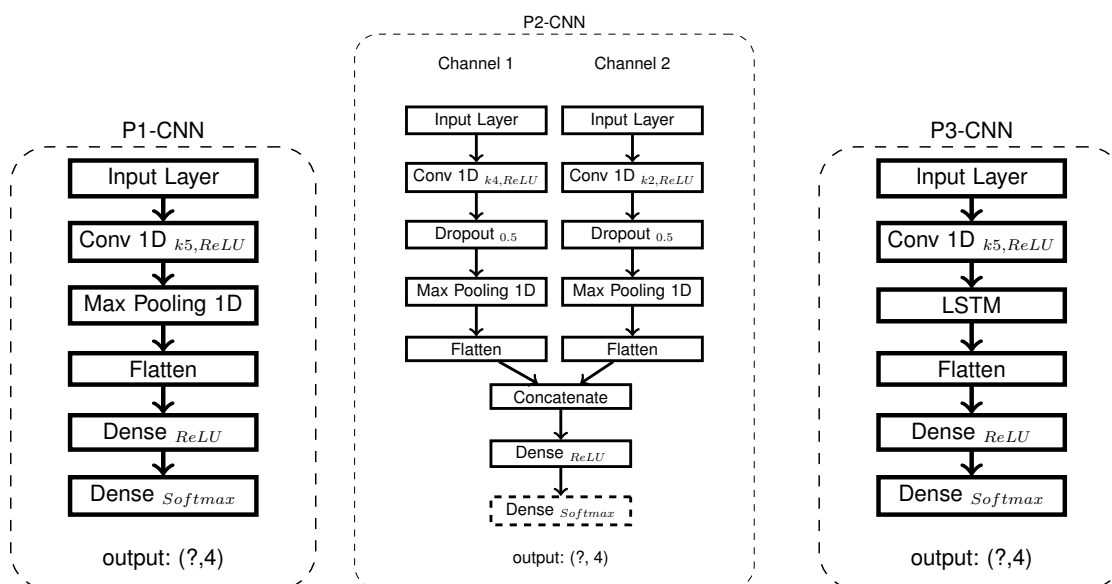


Fig. 4. Three architecture approaches of Polarity Convolutional Neural Network (P-CNN)

3.3.1 P1-CNN

It has a 1D Convolutional Layer with a total of 128 filters and a stride of 5 positions. Then, we used a Max Pooling Layer with a stride of 2 positions to take the most relevant values of the input vector.

Then, there is a Flatten Layer followed by two Dense Layers, the first one with 10 neurons and the activation function of ReLU, and the second one with 4 neurons corresponding to the polarities.

3.3.2 P2-CNN

This architecture has two channels, the first one is related to the vector representation of the text in the normalization phase, and the second one is related to the vector from AB-CNN. Those channels are concatenated, and two Dense Layers are applied.

3.3.3 P3-CNN

The third architecture is similar to the first one. The main difference is in the second layer; instead of the Max Pooling, it has an LSTM layer to expand the memory for learning.

4 Experiments and Validation

The experiments were done in reviews of restaurants given by a data set provided in the 2016 Edition of SemEval competition [19]. The aspect-based sentiment analysis was addressed in subtask 1 of task 5 of the competition. There are three slots in this subtask; this work focuses on slot1 and slot3.

Slot1 refers to the extraction of aspects; those aspects are an Entity **E** and attribute **A** pairs (E#A pair) towards which an opinion is expressed in a text. For example, for the entity *restaurant* and the aspect *price*, the E#A pair is RESTAURANT#PRICE.

Slot3 identifies the sentiment classification. For each pair E#A, will there be a sentiment polarity, such as positive or negative (OTE).

Because of the Spanish language complexity, there were just seven teams in the competition, compared to the 27 teams in the English language. Of those seven, only four of them participated in Slot1 and Slot3, as we present in this research.

4.1 Data Description

SemEval competition provides the training set and the test set. Each review has several sentences, each sentence and each review has several opinions. For each opinion, there is a category and a polarity. The category has an aspect, e.g., *food*, and a subaspect, e.g., *quality* or *price* for the category of *food* (E#A pair).

There are six aspects, namely, *restaurant*, *service*, *ambience*, *food*, *drinks*, and *location*. Considering the subaspects, we tested for 12 different aspects (see Table 1). Each sentence is classified in four possible polarities:

- **Positive.** When the aspect has positive assessment.
- **Negative.** When the aspect has a negative assessment.
- **Neutral.** When the aspect is not positive nor negative.
- **Conflict.** When the aspect has a positive and negative assessment.

4.1.1 Training Data

The provided training data set includes 627 reviews with 2070 sentences. In Table 1, we show the total of sentences for each aspect, and in Table 2 the number of sentences for each polarity.

Our approach uses this training set for the architecture comparison (Section 4.2). We can see that it is an unbalanced data set. It means that there is a difference between the number of examples belonging to each class, e.g., aspects related to *drinks* are minority classes compared to *service*. It affects the model because the learning system may have difficulties to learn the concept related to the minority class. The model should have a predilection to classify a sentence as positive than conflict. In order to solve it, one can look for specific models [16] or balance by eliminating in the majority classes [5] or augmenting the minority classes, as we do.

Table 1. Number of reviews per aspect in the training data set

Aspect	Total
RESTAURANT#GENERAL	602
SERVICE#GENERAL	389
AMBIENCE#GENERAL	219
FOOD#QUALITY	458
FOOD#PRICES	113
RESTAURANT#PRICES	108
FOOD#STYLE_OPTIONS	134
DRINKS#QUALITY	29
DRINKS#STYLE_OPTIONS	19
RESTAURANT#MISCELLANEOUS	13
LOCATION#GENERAL	15
DRINKS#PRICES	10

Table 2. Number of reviews per polarity in the training data set

Polarity	Total
positive	1512
negative	437
neutral	103
conflict	57

4.1.2 Test Data

The test data set of this subtask [19] includes 268 Spanish restaurant reviews with 881 sentences annotated with {E#A, OTE} tuples at the sentence level. Our approach uses this test set for the architecture comparison (Section 4.2) and the experiments (Section 4.3).

4.2 Architecture Comparison

All architectures were compiled using binary cross-entropy loss function, Adam optimizer, and F1-measure for evaluation.

For the AB-CNN architectures we used the outputs of the normalization phase described in Section 3.1. The data sets are *A300*, *A100*, *B300*, *B100*, and *C300*, where numbers *300* and *100* refer to the number of dimensions.

Moreover, data sets with letter *A* correspond to lemmas from spaCy and fastText, data sets with letter *B* correspond to tokens from spaCy and fastText, and data sets with letter *C* use word2Vec representation.

First, we compared the results of the two architectures AB-CNN. We used F1-measure like it is in SemEval challenge [19]. In Table 3, we present two evaluations. **F1 in Training** refers to the output of calling the training method *fit()* which uses a split of the training set to validate, the split is given by *validation_batch_size* with a value of 64. **F1 in Test** corresponds to the evaluation of the model in the Test Data.

Because of the multi-channel output, the AB2-CNN architecture has lower results. Moreover, the unbalance in some classes affects the performance. Then, the selected architecture was AB1-CNN.

For the sentiment classification (slot3) we have three architectures (Section 3.3). In Table 4, we compared with metric accuracy for the training and test data sets. It shows that the P1-CNN architecture has better performance in the training data sets but in the test data set the P2-CNN architecture is better. In both cases, the P3-CNN has the worst results.

In these experiments, we can see that the constructed set *A300* has better results in all the aspect-based experiments. The best data set for each experiment is in italics. Additional to the training and test sets, we did field tests with written reviews and manual classification. It considered the 12 classes of aspects of the AB1-CNN architecture and its combination with P1-CNN and P2-CNN. Table 5 shows the results of the field test with a better result than the performed in the test data set. It allows us to identify that the mistaken classification was related to the minority classes, even if we balanced them.

For the task of sentiment analysis of each aspect, the results of the experiments are shown in Table 6. The results of all the architectures are still similar, so it was decided to use the architecture that obtained the best results with the *A300* training set since this set was the one with the best results for the aspects classifier model, the best with an accuracy of 93.33%.

Table 3. F1-measure in training and test data sets for AB-CNN architectures

Normalized Data Set	F1 in Training		F1 in Test	
	AB1-CNN	AB2-CNN	AB1-CNN	AB2-CNN
A300	0.7995	0.5618	0.6540	0.5570
A100	0.6640	0.5448	0.5017	0.5402
B300	0.6562	0.5497	0.6273	0.5457
B100	0.6530	0.5399	0.5038	0.5357
C300	0.5945	0.5112	0.4456	0.5058

Table 4. Accuracy in training and test data sets for P-CNN architectures

Normalized Data Set	Accuracy in Training			Accuracy in Test		
	P1-CNN	P2-CNN	P3-CNN	P1-CNN	P2-CNN	P3-CNN
A300	0.8640	0.8382	0.7937	0.7789	0.7845	0.6921
A100	0.8504	0.8143	0.8021	0.7553	0.7879	0.7093
B300	0.8494	0.8372	0.8274	0.7811	0.7969	0.7302
B100	0.8593	0.8241	0.8192	0.7520	0.7699	0.7192
C300	0.8427	0.8023	0.8294	0.7710	0.7789	0.7203

Table 5. F1-measure in field tests for AB1-CNN architectures

Normalized Data Set	F1
A300	0.9333
A100	0.9286
B300	0.9226
B100	0.9198
C300	0.9008

Table 7. Results of the aspect extraction

Team	F1
IIT-TUDA	0.5980
INSIGHT-1	0.6137
TGB	0.6355
UWB	0.6196
AB1-CNN	0.6540

Table 6. Accuracy and F1-measure in field tests for P1-CNN and P2-CNN architectures

Normalized Data Set	Accuracy	
	P1-CNN	P2-CNN
A300	0.9333	0.9000
A100	0.9000	0.9333
B300	0.9333	0.9333
B100	0.9333	0.9333
C300	0.9333	0.9120

Table 7 shows the results of those four competitors and our approach AB1-CNN. We can see that in the aspect extraction, we got better results than other competitors.

Table 8 shows the results of the same four competitors and our approach P-CNN with the data set of B300. We did not achieve a better result than most of the competitors in the accuracy metric.

The accuracy measures all the correctly identified polarities, but it is useful when all the classes are equally important.

Although all classes are essential, the training and test data set do not have the same number of samples per class. The positive class represents 71.69% of the complete data set, and the remaining 28.3% is distributed in the other three classes.

4.3 Evaluation

We consider the four teams of the 2016 edition of SemEval that participate in the aspect-based sentiment classification that identifies the aspect and polarity.

Table 8. Results of the sentiment classification

Team	Accuracy
IIT-TUDA	0.8358
INSIGHT-1	0.7957
TGB	0.8209
UWB	0.8134
AB1-CNN	0.7969

5 Conclusions and Future Work

Aspect-based sentiment analysis is a major technological challenge. Our proposal focuses on processing, analyzing, and categorizing reviews and was tested in restaurant reviews.

Our approach combines two deep learning models for aspect-based sentiment analysis. The performance of our aspect classification proposal is promising, but there is still work to do in the sentiment detection.

The preprocessing stage was a core part of the proposal. The text was represented as word vectors of real numbers; it allows us to avoid losing context and relate the polarities to each aspect, even if there were more than one aspect in a sentence. We used lemmas and tokens, but there is also the possibility of using n-grams.

The proposal combines two deep learning models architectures to determine the sub-aspect and the polarity for the classification task. We develop two architectures for the aspect detection and three architectures for the sentiment analysis to get our final architecture merging the preprocessing with the deep learning models.

We got better results than state-of-the-art models in aspect classification but not as good in polarity classification. It was affected by the unbalanced data set and by the CNN, even when we try three different architectures. Other works that also used CNN, such as INSIGHT-1, were also affected in the polarity evaluation.

There is still room for improvement; for instance, including dictionaries in the preprocessing stage to replace core words.

References

1. **Alvarez-López, T., Juncal-Martínez, J., Fernández-Gavilanes, M., Costa-Montenegro, E., González-Castano, F. J. (2016).** Gti at Semeval-2016 task 5: Svm and crf for aspect detection and unsupervised aspect-based sentiment analysis. *Proceedings of the 10th international workshop on semantic evaluation (SemEval-2016)*, pp. 306–311.
2. **Ananiadou, S., Rea, B., Okazaki, N., Procter, R., Thomas, J. (2009).** Supporting systematic reviews using text mining. *Social Science Computer Review*, Vol. 27, No. 4, pp. 509–523.
3. **Araque, O., Corcuera, I., Román, C., Iglesias, C. A., Sanchez-Rada, J. F. (2015).** Aspect based sentiment analysis of Spanish tweets. *TASS@ SEPLN*, pp. 29–34.
4. **Ayuve (2019).** La importancia de obtener buenas reseñas en Google. <https://www.ayuve.net/blog/la-importancia-de-las-resenas-en-google/>. Accessed: March 24, 2021.
5. **Batista, G. E., Carvalho, A. C., Monard, M. C. (2000).** Applying one-sided selection to unbalanced datasets. *Mexican International Conference on Artificial Intelligence*, Springer, pp. 315–325.
6. **Çetin, F. S., Yıldırım, E., Özbey, C., Eryiğit, G. (2016).** TGB at Semeval-2016 task 5: multi-lingual constraint system for aspect based sentiment analysis. *Proceedings of the 10th International Workshop on Semantic Evaluation (SemEval-2016)*, pp. 337–341.
7. **Cruz, I., Gelbukh, A., Sidorov, G. (2014).** Implicit aspect indicator extraction for aspect based opinion mining. *Int. J. Comput. Linguistics Appl.*, Vol. 5, No. 2, pp. 135–152.
8. **Do, H. H., Prasad, P., Maag, A., Alsadoon, A. (2019).** Deep learning for aspect-based sentiment analysis: a comparative review. *Expert Systems with Applications*, Vol. 118, pp. 272–299.
9. **García-Pablos, A., Cuadros, M., Rigau, G. (2018).** W2VLDA: almost unsupervised system for aspect based sentiment analysis. *Expert Systems with Applications*, Vol. 91, pp. 127–137.
10. **Goldberg, Y. (2017).** Neural network methods for natural language processing. *Synthesis lectures on human language technologies*, Vol. 10, No. 1, pp. 1–309.

11. **Hercig, T., Brychcín, T., Svoboda, L., Konkol, M. (2016).** Uwb at Semeval-2016 task 5: Aspect based sentiment analysis. Proceedings of the 10th international workshop on semantic evaluation (SemEval-2016), pp. 342–349.
12. **Honnibal, M., Montani, I. (2017).** spaCy 2: Natural language understanding with Bloom embeddings, convolutional neural networks and incremental parsing.
13. **Hu, M., Liu, B. (2004).** Mining and summarizing customer reviews. Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining, pp. 168–177.
14. **Kelleher, J. D. (2019).** Deep learning. MIT press.
15. **Kumar, A., Kohail, S., Kumar, A., Ekbal, A., Biemann, C. (2016).** lit-tuda at Semeval-2016 task 5: Beyond sentiment lexicon: Combining domain dependency and distributional semantics features for aspect based sentiment analysis. Proceedings of the 10th international workshop on semantic evaluation (SemEval-2016), pp. 1129–1135.
16. **Li, S., Song, W., Qin, H., Hao, A. (2018).** Deep variance network: An iterative, improved CNN framework for unbalanced training datasets. Pattern Recognition, Vol. 81, pp. 294–308.
17. **Miranda, C. H., Buelvas, E. (2019).** AspectSA: Unsupervised system for aspect based sentiment analysis in Spanish. *Prospectiva*, Vol. 17, No. 1, pp. 87–95.
18. **Pannala, N. U., Nawarathna, C. P., Jayakody, J., Rupasinghe, L., Krishnadeva, K. (2016).** Supervised learning based approach to aspect based sentiment analysis. 2016 IEEE International Conference on Computer and Information Technology (CIT), IEEE, pp. 662–666.
19. **Pontiki, M., Galanis, D., Papageorgiou, H., Androutsopoulos, I., Manandhar, S., Al-Smadi, M., Al-Ayyoub, M., Zhao, Y., Qin, B., De Clercq, O., others (2016).** Semeval-2016 task 5: Aspect based sentiment analysis. International workshop on semantic evaluation, pp. 19–30.
20. **Román, J. V., Cámara, E. M., Morera, J. G., Zafra, S. M. J. (2015).** Tass 2014-the challenge of aspect-based sentiment analysis. *Procesamiento del Lenguaje Natural*, Vol. 54, pp. 61–68.
21. **Ruder, S., Ghaffari, P., Breslin, J. G. (2016).** Insight-1 at Semeval-2016 task 5: Deep learning for multilingual aspect-based sentiment analysis. arXiv preprint arXiv:1609.02748.
22. **WebParaRestaurantes (2017).** La importancia de las opiniones online de clientes para tu restaurante. <http://webpararestaurantes.com/resenas-clientes-restaurantes/>. Accessed: August 16, 2019.
23. **Weichselbraun, A., Gindl, S., Fischer, F., Vakulenko, S., Scharl, A. (2017).** Aspect-based extraction and analysis of affective knowledge from social media streams. *IEEE Intelligent Systems*, Vol. 32, No. 3, pp. 80–88.
24. **Zuheros, C., Martínez-Cámara, E., Herrera-Viedma, E., Herrera, F. (2021).** Sentiment analysis based multi-person multi-criteria decision making methodology using natural language processing and deep learning for smarter decision aid. case study of restaurant choice using tripadvisor reviews. *Information Fusion*, Vol. 68, pp. 22–36.

*Article received on 25/06/2021; accepted on 15/11/2021.
Corresponding author is Sabino Miranda.*