

DETECTION OF EMOTIONS IN SPANISH TEXTS USING BERT LANGUAGE MODELS

DETECTION OF EMOTIONS IN SPANISH TEXT USING BERT-BASED LANGUAGE MODEL

Muñiz Rascado Luis¹, Castro Sánchez Noé Alejandro², Magadan Salazar Andrea³,
González Franco Nimrod⁴, González Serna

Gabriel⁵<https://doi.org/10.61117/ipsumtec.v8i2.377>

¹ National Technological Institute of Mexico/CENIDET, d23ce169@cenidet.tecnm.mx , <https://orcid.org/0000-0003-2941-2482>.

² National Technological Institute of Mexico/CENIDET, noe.cs@cenidet.tecnm.mx , <https://orcid.org/0000-0002-8083-3891>. ³

National Technological Institute of Mexico/CENIDET, andrea.ms@cenidet.tecnm.mx , <https://orcid.org/0000-0001-5474-4150>.

⁽⁴⁾ National Technological Institute of Mexico/CENIDET, nimrod.gf@cenidet.tecnm.mx . <https://orcid.org/0000-0002-1051-1379>. ⁽⁵⁾ National Technological Institute of Mexico/CENIDET, gabriel.gs@cenidet.tecnm.mx , <https://orcid.org/0000-0002-1874-9402>.

Abstract -- The Mexican Ministry of Health collects information on the prevalence of mental health disorders for the allocation of medical resources, highlighting depression with a prevalence of 5.3%. Socioeconomic factors and the lack of specialists affect its diagnosis and treatment. To improve early detection, the use of artificial intelligence (AI) is proposed, specifically BERT models trained to detect emotions. A search for the best parameters and cross-validation with $k=5$ were applied to avoid overfitting. The tests and validations used metrics such as *accuracy*, *precision*, *recall*, and *F1-score*, demonstrating that the BERT-based model outperforms traditional machine learning techniques in emotion detection.

Keywords: NLP, BERT, emotions, detection.

Abstract-- The Mexican Ministry of Health collects data on the prevalence of mental health disorders to allocate medical resources, highlighting depression with a coverage of 5.3%. Socioeconomic factors and the lack of specialists affect its diagnosis and treatment. To improve early detection, the use of artificial intelligence (AI) is proposed, specifically BERT-based models trained to detect emotions. Applied search for best parameters and 5-fold cross-validation approach was applied to prevent overfitting, the tests and validations employed metrics such as accuracy, precision, recall, and F1-score, demonstrating that BERT-based models outperform traditional machine learning techniques in emotion detection.

Key words: NLP, BERT, emotions, detection.

INTRODUCTION

At the national level, the Secretary of Health is responsible for collecting information on people's mental health in order to make projections for the allocation of medical resources and address these conditions [1]. This report mentions that 19.9% of

The Mexican population may suffer from mental disorders. One of the most prevalent illnesses in this report is depression, with 5.3% [1] of the population affected, which can manifest itself in varying degrees of severity. There are various factors that can affect people's mental health in relation to depression [2], ranging from socioeconomic factors, education, and comorbidities [3] to a lack of mental health specialists [4]. Therefore, it is very important to have digital tools that help detect emotions related to depression using artificial intelligence such as large language models such as BERT [2], [5], [6], [7], which are specifically retrained to search for this condition at a moderate stage so that people can be better referred for treatment.

DEVELOPMENT

Dataset

For the emotion detection process, we used the *emotions* corpus with a 40MB file and 194,416 observations. This file was obtained from Kaggle. Unfortunately, it is in English, and an automatic translation process had to be applied using a large language model specialized in English-Spanish translation called Marian, similar to the one proposed in the work of [12], and [13], as this corpus was specifically needed in Spanish for the training process in the multilingual BERT model [8], [9], [10], [11], [14]. The translation process was performed on an Apple M1 server with 16G of shared memory.

Once the corpus was available in Spanish, a separation was applied for the six emotions available in the *emotions* corpus, which, it should be noted, describes Paul Ekman's emotional framework: sadness (0), joy (1), love (2), anger (3), fear (4), surprise (5) of 1000 observations per class for the training file

, 500 for each class for testing, and 100 for each class for validation.

Training the prediction model

Once the Spanish corpus was available and separated for the different tasks, the BERT model [14] model in its "bert-base-multilingual-uncased" version, with 177,853,440 parameters and the ability to understand 104 languages, including Spanish.

Access to BERT was provided through the Transformers API provided by Hugging Faces. This model was specifically instantiated with the "BertForSequenceClassification" object configured for text sequence classification tasks.

With the help of Bayesian-based hyperparameter search tools implemented in the Ray Tune library, the search for the best combination of values was enhanced, see Table 1. These are: learning rate (*adam lr*), *weight decay*, *batch size*, and *dropout*. This process was carried out in Google Colab using an A100 GPU with 40G of VRAM and 83G of RAM.

Table 1. Values sought in the Ray Tune library for the emotions model.

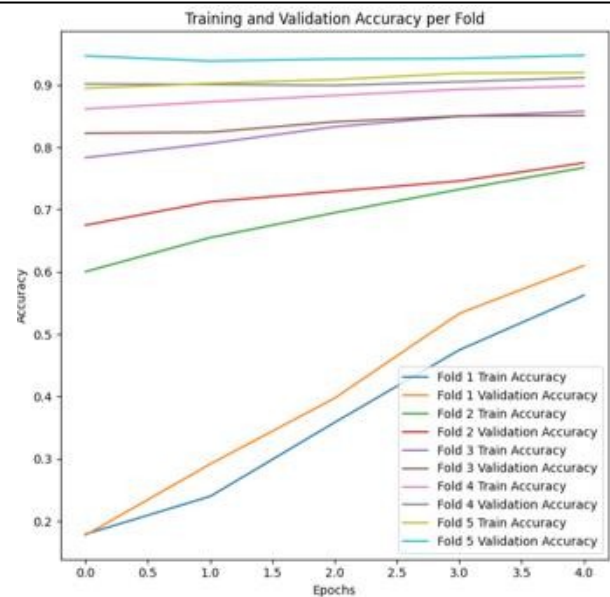
Name	Minimum	Maximum	Suggested
<i>adam lr</i>	1e-1	1e-7	2e-05
<i>weight decay</i>	0.1	0.5	0.4
<i>batch size</i>	8	32	16
<i>dropout</i>	0.1	0.5	0.3

Source: own creation (2025).

Using the hyperparameter values suggested by Ray Tune, the cross-validation technique was also applied in an attempt to avoid overfitting in the *BERT-emotions* model, using $k=5$.

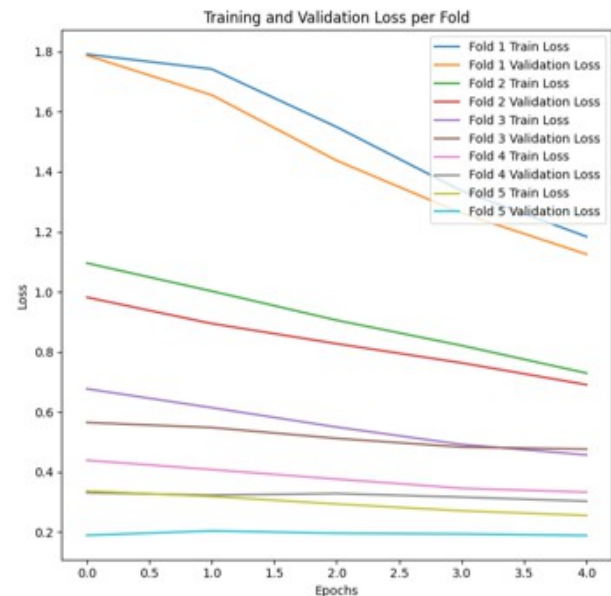
For each iteration of k , the *accuracy* and *loss* were obtained, achieving different values for the *BERT-emotions* model. These were plotted, and the results are shown in Figures 1, 2, 3, and 4.

Figure 1. Accuracy graph of the *BERT-emotions* model in cross-validation.



Source: own creation (2025).

Figure 2. Loss graph of the *BERT-emotions* model in cross-validation.



Source: own creation (2025).

Figure 3. Summary of the accuracy of the *BERT-emotions* model in cross-validation.

Train accuracies	Validation accuracies
[0.17916666666666667, 0.23979166666666665, 0.35854166666666665, 0.47470166666666667, 0.5622916666666666]	[0.1775, 0.2916666666666667, 0.3975, 0.5333333333333333, 0.61]
[0.6004166666666667, 0.655, 0.695, 0.7322916666666667, 0.7672916666666667]	[0.675, 0.7125, 0.7291666666666666, 0.7458333333333333, 0.775]
[0.7833333333333333, 0.8060416666666667, 0.8329166666666666, 0.8497916666666666, 0.8575]	[0.8225, 0.8241666666666667, 0.8408333333333333, 0.85, 0.8508333333333333]
[0.8614583333333333, 0.8729166666666667, 0.883125, 0.8929166666666667, 0.898125]	[0.9016666666666666, 0.9008333333333334, 0.8991666666666667, 0.905, 0.9116666666666666]
[0.8945833333333333, 0.9025, 0.9085416666666667, 0.91875, 0.9197916666666667]	[0.9466666666666667, 0.9383333333333334, 0.9416666666666667, 0.9425, 0.9475]

Source: own creation (2025).

Figure 4. Summary of loss for the BERT-emotions model in cross-validation.

Training Loss	Validation Loss
[1.791360681851705, 1.7420288395881653, 1.5490595889091492, 1.3366429233551025, 1.1840635059595988]	[1.7876181414252834, 1.6554248207493831, 1.437311511290701, 1.2617667160536115, 1.125373111743676]
[1.0958513792355855, 1.0025018242994945, 0.9058314681053161, 0.8215377628803253, 0.7295879356066386]	[0.9820331178213421, 0.8947605994595137, 0.8275605589151382, 0.7637460631759543, 0.6910714150259369]
[0.6773368521531423, 0.6144728645357895, 0.549866960643041, 0.4921880375345548, 0.4572153541445732]	[0.5651330528290648, 0.5484175105628214, 0.5122408529645518, 0.48359740094134684, 0.476518167476905]
[0.4394359285632769, 0.4084506352742513, 0.3766383816878, 0.346249623463377, 0.3332455252110958]	[0.3313105035769312, 0.32341706301820905, 0.32869576819633184, 0.31670198667990535, 0.30331425447213023]
[0.33701842069625854, 0.31895541707674663, 0.29386573274930317, 0.2711464661359787, 0.2558118860423565]	[0.1895500868558837, 0.20440706365594738, 0.1964329890416641, 0.19398186953836366, 0.18894489908492879]

Source: own creation (2025).

Survey

During April and May, as well as October and November 2024, a survey was administered to students at the National Technological Institute of Mexico/Zacatepec Technological Institute (TecNM/IT Zacatepec) and the Technological University of the Northern Region of Guerrero (UTRNG) to collect the following information:

- 1) Date of birth.
- 2) Gender (male, female, prefer not to share).
- 3) When did the traumatic event occur?
- 4) Have you been able to overcome the traumatic event? (yes, no).
- 5) Describe the traumatic event in detail (minimum 150 words).

In order to have a context that describes events in a heterogeneous population of individuals over the age of 18, and so that the *BERT-emotions* detection process can predict which emotion is related to what the person is mentioning in their traumatic event.

DISCUSSION AND ANALYSIS OF RESULTS

BERT-emotions prediction model

The following metrics were used in the testing and validation processes of the *BERT-emotions* model: *accuracy*, *precision*, *recall*, and *F1-score*.

The results of the test metrics are shown in Figure 5 for each of the classes.

Figure 5. Summary of BERT-emotions test process metrics.

	accuracy	precision	recall	f1_scores
0	0.943667	0.971014	0.938	0.954222
1	0.943667	0.978308	0.902	0.938606
2	0.943667	0.921201	0.982	0.950629
3	0.943667	0.957317	0.942	0.949597
4	0.943667	0.947589	0.904	0.925281
5	0.943667	0.897112	0.994	0.943074

Source: own creation (2025).

For the validation process, the results are shown in Figure 6 for each of the classes.

Figure 6. Summary of BERT-emotions validation process metrics.

	accuracy	precision	recall	f1_scores
0	0.92	0.948454	0.92	0.93401
1	0.92	0.968085	0.91	0.938144
2	0.92	0.906542	0.97	0.937198
3	0.92	0.936842	0.89	0.912821
4	0.92	0.923077	0.84	0.879581
5	0.92	0.853448	0.99	0.916667

Source: own creation (2025).

Of the above metrics, we use the *F1-score* as a basis because it is a metric that balances the *precision* metric, which represents the number of observations that the model classifies as positive when they actually are, and the *recall* metric, which represents the number of actual observations that were correctly identified by the model's prediction.

Thus, the *F1-score* represents the harmonic mean of these two metrics, see Eq. (1), Table 2.

$$(F1) = 2 \times \frac{\text{precision} \times \text{recall}}{\text{precision} + \text{recall}} \quad \text{Eq. (1)}$$

With the description of the *F1-score*, the *BERT-emotions* model has the ability to detect different emotions in the testing and validation process.

Table 2. Emotions related to their percentage according to the testing and validation process.

Predicted emotion	Test	Validation
High percentage	Sadness	Joy, Sadness

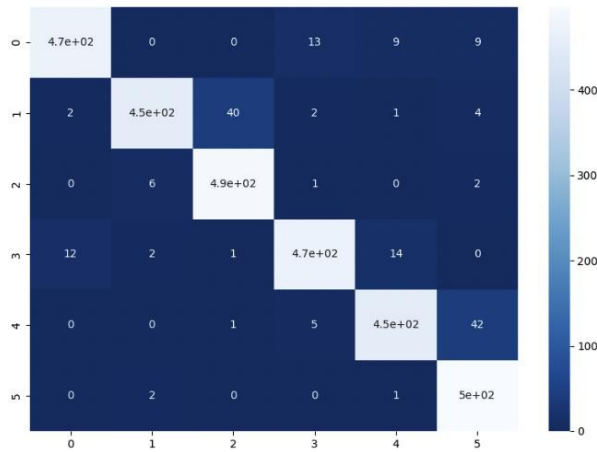
Low percentage	Fear	Fear
----------------	------	------

Source: own creation (2025).

In turn, their respective heatmaps were generated based on the calculation of the *confusion matrix* of the test value (y_{test}) against the value predicted by the *BERT-emotions* test model (y_{pred}).

For the testing process, see Figure 7.

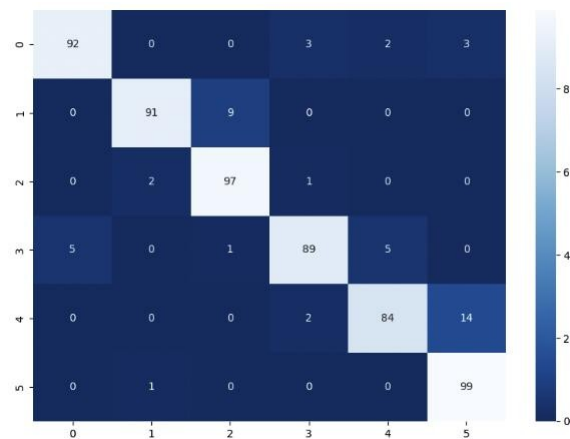
Figure 7. Heatmap of the *BERT-emotions* confusion matrix in testing.



Source: own creation (2025).

The confusion matrix of the validation value (y_{val}) against the value predicted (y_{pred}) by the *BERT-emotions* model in the validation process, see Figure 8.

Figure 8. Heat map of the *BERT-emotions* confusion matrix in validation.



Source: own creation (2025).

BERT-emotions model applied to the survey

In the application of the survey, 160 observations were obtained, of which 15 records were eliminated by preprocessing, resulting in only 145 operational records, as shown in Table 3.

Table 3. Survey responses by gender.

Gender	Count
Women	66
Male	76
I prefer not to share	3
Total	145

Source: own creation (2025).

Next, the content of the description of the traumatic event was analyzed with regard to the most frequent words. To do this, stop words were eliminated. The results can be seen in Table 4.

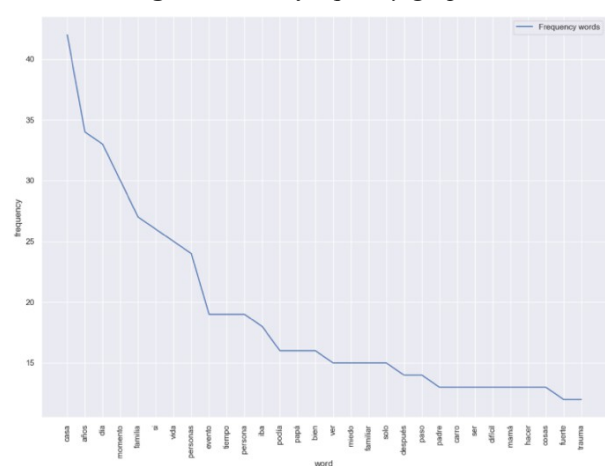
Table 4. The 10 most frequent words.

Word	Count
House	42
Years	34
Day	33
Time	30
Family	27
Yes	26
Life	25
People	23
Event	18
Time	18

Source: own creation (2025).

Using this count, a frequency representation was generated, as well as a classic word cloud representation, where the most frequent words are larger than the less frequent words (see Figures 9 and 10).

Figure 9. Word frequency graph.



Source: own creation (2025).

Figure 10. *Word cloud graph.*



Source: own creation (2025).

A count was generated by gender, and if the traumatic event was overcome, see Table 5.

The *BERT-emotions* model was run on the description of the traumatic event, taking only those participants who were unable to overcome the traumatic event, and the findings can be seen in Tables 6 and 7.

Table 5. Survey responses by gender/traumatic event not overcome.

Gender	Overcame traumatic event	Count
Female	No	27
	Yes	39
	<i>Total</i>	66
Male	No	22
	Yes	54
	<i>Total</i>	76
I prefer not to share	No	1
	Yes	2
	<i>Total</i>	3

Source: own creation (2025).

Table 6. Survey responses by gender/female traumatic event not overcome.

Gender	Overcame traumatic event	Predicted emotion	Count
Female	No	Sadness	10
		Joy	1
		Love	4
		Courage	4
		Fear	8
		Surprise	0
		<i>Total</i>	27

Source: own creation (2025)

Table 7. Survey responses by gender/traumatic event not overcome by males.

Gender	Overcame traumatic event	Predicted emotion	Count
Male	No	Sadness	13
		Joy	3
		Love	0
		Courage	1
		Fear	5
		Surprise	0
		<i>Total</i>	22

Source: own creation (2025).

The emotions detected are sadness and fear. In the case of overcoming the traumatic event, see Table 8.

Table 8. Survey responses by gender/traumatic event overcome, female and male.

Gender	Overcame traumatic event	Predicted emotion	Count
Female	Yes	Sadness	23
		Joy	1
		Love	4
		Courage	4
		Fear	7
		Surprise	0
		<i>Total</i>	39
Male	Yes	Sadness	25
		Happiness	5
		Love	7
		Courage	7
		Fear	10
		Surprise	0
		<i>Total</i>	54

Source: own creation (2025).

When they have not overcome the traumatic event for records that did not share their gender, see Table 9.

Table 9. Survey responses by gender/event I prefer not to share.

Gender	Overcame traumatic event	Predicted emotion	Count
I prefer not to share	No	Sadness	1
		Joy	0
		Love	0
		Courage	0

		Fear	0
		Surprise	0
		Total	1
	Yes	Sadness	0
		Joy	0
		Love	0
		Courage	0
		Fear	1
		Surprise	1
		Total	2

Source: own creation (2025).

Using the results from the *BERT-emotions* model, a comparison was made with ChatGPT-3.5 turbo [15] using the OpenAI API. This process was carried out in Google Colab using an A100 GPU with 40G of VRAM and 83G of RAM.

For those who were unable to overcome the traumatic event and are female, the comparison can be seen in Table 10.

Table 10. Survey responses by gender/unresolved traumatic event female BERT-emotions vs ChatGPT-3.5 turbo.

Sex	Overcame event	Predicted emotion	BERT-emotions	ChatGPT 3.5-turbo
F	No	Sadness	10	16
		Joy	1	0
		Love	4	1
		Courage	4	0
		Fear	8	10
		Surprise	0	0
		Total	27	27

Source: own creation (2025).

For females who were able to overcome the traumatic event, their comparison can be seen in Table 11.

Table 11. Survey responses by gender/traumatic event if overcome female BERT-emotions vs ChatGPT-3.5 turbo.

Sex	Overcame event	Predicted emotion	BERT-emotions	ChatGPT-3.5 turbo
F	Yes	Sadness	23	21
		Joy	1	0
		Love	4	0

		Courage	4	0
		Fear	7	16
		Surprise	0	2
		Total	39	39

Source: own creation (2025).

A comparison was also made between males who had not overcome their traumatic event and those who had, see Tables 12 and 13.

Table 12. Survey responses by gender/unresolved traumatic event male BERT-emotions vs ChatGPT-3.5 turbo.

Sex	Overcame event	Predicted emotion	BERT-emotions	ChatGPT-3.5 turbo
M	No	Sadness	13	13
		Joy	3	0
		Love	0	0
		Courage	1	1
		Fear	5	6
		Surprise	0	2
		Total	22	22

Source: own creation (2025).

Table 13. Survey responses by gender/traumatic event if overcome male BERT-emotions vs ChatGPT-3.5 turbo.

Sex	Event overcome	Predicted emotion	BERT-emotions	ChatGPT-3.5 turbo
M	Yes	Sadness	25	32
		Joy	5	1
		Love	7	0
		Courage	7	2
		Fear	10	17
		Surprise	0	2
		Total	54	54

Source: own creation (2025).

Finally, the comparison of participants who preferred not to share their gender can be seen in Tables 14 and 15.

Table 14. Survey responses by gender/unresolved traumatic event I prefer not to share BERT-emotions vs ChatGPT-3.5 turbo.

Sex	Surpassed event	Predicted excitement	BERT-emotions	ChatGPT-3.5 turbo
-----	-----------------	----------------------	---------------	-------------------

x o r				
N	No	Sadness	1	0
		Joy	0	0
		Love	0	0
		Courage	0	0
		Fear	0	1
		Surprise	0	0
		<i>Total</i>	1	1

Source: own creation (2025).

Table 15. Survey responses by gender/traumatic event if overcome I prefer not to share BERT-emotions vs ChatGPT-3.5 turbo.

Se x o	Event overco me	Predict ed emotion	BERT- emotions	ChatGPT- 3.5 turbo
N	Yes	Sadness	0	0
		Joy	0	0
		Love	0	0
		Courage	0	0
		Fear	1	2
		Surprise	1	0
		<i>Total</i>	2	2

Source: own creation (2025).

The emotion detection process with ChatGPT-3.5 turbo tends to miss some emotions such as joy, love, anger, and surprise, while performing comparably to the model trained with *BERT-emotions* for detecting emotions in terms of sadness and fear, highlighting that the latter detects emotions that its rival model omits in most cases.

CONCLUSIONS

BERT-based models that have been retrained for specific tasks such as classification have demonstrated superior performance in this task and specifically in classifying emotions compared to traditional machine learning methods [7], [17].

This is consistent with various studies in the literature, where deep learning models based on *transformers* have outperformed techniques such as *Support Vector Machine* (SVM), *Naïve Bayes*, *Random Forest*, and *Decision Trees* in text classification tasks, as these usually require a lot of processing work to handle the words being classified, such as lemmatization, which consists of transforming each word into a base representation, better known as a lemma.

many advantages in certain exercises, but at the same time generates a subtle loss of context that can be very enriching in the classification process. Compared to BERT models, thanks to their internal configuration and their own word processing mechanism known as *wordpiece*, they do not need as much pre-processing as the aforementioned, and use it as is for contextualization.

Thanks to the use of hyperparameter search techniques for the best parameters and cross-validation, the *BERT-emotions* model avoids problems of underfitting or, where applicable, overfitting, and its operational predictions are considered reliable.

The emotions with the highest prevalence detected in the description of the traumatic event, regardless of whether or not the event was overcome, are sadness and fear for the *BERT-emotions* model, but it also covers joy, love, courage, and surprise. Compared to ChatGPT-3.5 turbo, it lacks the latter. This model is trained for almost general use, but the *BERT-emotions* model demonstrates that specialization in classification tasks with related corpora improves the prediction process.

Finally, the findings suggest that the application of AI models in the detection of emotions and mental health problems [15], [16] in Spanish can be a key tool for improving access to mental health services.

The work developed by [18] describes the process of emotion detection based on different languages using models based on multilingual BERT. Some of these ideas served as the basis for our research. In this study, the model is trained in a source language, such as English, and its effectiveness in predicting texts written in other languages, such as Spanish or Arabic, is evaluated without these having labels annotated in the language to be predicted. For this task, they used *transfer learning* as *zero-shot learning*. It is worth mentioning that for our proposal, *few-shot learning* was used. Also in their study, they used different BERT models such as mBERT and XLM-R, which have been pre-trained on large amounts of multilingual data.

Another study [19] proposes a hybrid approach that combines machine learning algorithms with affective lexicographic knowledge specific to Spanish. For their research, they first collected texts obtained from tweets in Spanish, grouping them into four categories: anger, fear, sadness, and joy. In feature engineering, traditional vectors such as *bag-of-words* or TF-IDF (our

proposal used BERT's *wordpiece*), along with affective indicators obtained from Spanish emotional lexicons. This set of attributes is used to train classifiers such as logistic regression, SVM, *Naïve Bayes*, and multilayer perceptrons.

The integration of these systems into digital platforms could facilitate the early detection of cases of depression, anxiety, and possible suicidal tendencies, allowing for timely intervention and better referral of patients to specialized professionals. For our proposal, a web platform is being developed to integrate the implemented model, as well as other models that are being trained, with the aim of providing assistance to vulnerable sectors.

It should be noted that these tools are not intended to replace mental health specialists, but rather to support them and individuals in crisis.

ACKNOWLEDGMENTS

This work has been made possible thanks to the support of the former National Council for Humanities, Sciences, and Technologies (CONAHCYT), now the Secretariat of Science, Humanities, Technology, and Innovation (SECIHTI), through the scholarship awarded for my doctoral studies. I am deeply grateful.

I would also like to express my gratitude to the National Center for Research and Technological Development (CENIDET) for providing me with the knowledge, resources, infrastructure, and academic environment necessary to carry out this research project. Their support has been essential for the advancement and consolidation of this work, as well as for the opportunity to further my academic training.

My thanks to the participants in the survey conducted at TecNM/IT Zacatepec and UTRNG for sharing their experience, and to the people who helped disseminate and implement it.

BIBLIOGRAPHY

- [1] Alcocer Varela J, López-Gatell H. 2nd Operational Diagnosis of Mental Health and Addictions. 2022;
- [2] Mendoza Mojica SA, Márquez Mendoza O, Guadarrama Guadarrama R, Ramos Lira LE. Measurement of Post-Traumatic Stress Disorder (PTSD) in Mexican university students. *Salud Ment*. 2013 Jan 1;36(6):493.
- [3] Wagner FA, González-Forteza C, Sánchez-García S, García-Peña C, Gallo JJ. Focusing on depression as a public health problem in Mexico. 2012;35(1).
- [4] Heinze G, Chapa GDC, Carmona-Huerta J, Research Department of the Department

Department of Psychiatry and Mental Health. Faculty of Medicine. National Autonomous University of Mexico, Clinical Services Directorate. Ramón de la Fuente Muñiz National Institute of Psychiatry. *Psychiatrists in Mexico*: 2016. sm. 2016 Mar 30;39(2):69–76.

- [5] Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al. Attention Is All You Need [Internet]. arXiv; 2017 [cited 2023 Mar 18]. Available from: <http://arxiv.org/abs/1706.03762>
- [6] Devlin J, Chang MW, Lee K, Toutanova K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding.
- [7] Garrido-Merchan EC, Gozalo-Brizuela R, Gonzalez-Carvajal S. Comparing BERT Against Traditional Machine Learning Models in Text Classification. *JCCE*. 2023 Apr 21;2(4):352–6.
- [8] Tiedemann J, Thottingal S. OPUS-MT – Building open translation services for the World. European Association for Machine Translation. 2020;Proceedings of the 22nd Annual Conference of the European Association for Machine Translation:479–80.
- [9] Cañete J, Chaperon G, Fuentes R, Ho JH, Kang H, Pérez J. Spanish Pre-trained BERT Model and Evaluation Data [Internet]. arXiv; 2023 [cited 2025 May 5]. Available from: <https://arxiv.org/abs/2308.02976>
- [10] de la Rosa J, Ponferrada EG, Villegas P, Salas PG de P, Romero M, Grandury M. BERTIN: Efficient Pre-Training of a Spanish Language Model using Perplexity Sampling. 2022 [cited 2023 Dec 4]; Available from: <https://arxiv.org/abs/2207.06814>
- [11] Serrano AV, Subies GG, Zamorano HM, Garcia NA, Samy D, Sanchez DB, et al. RigoBERTa: A State-of-the-Art Language Model For Spanish. 2022 [cited 2023 Dec 4]; Available from: <https://arxiv.org/abs/2205.10233>
- [12] Liu Y, Ott M, Goyal N, Du J, Joshi M, Chen D, et al. RoBERTa: A Robustly Optimized BERT Pretraining Approach [Internet]. arXiv; 2019 [cited 2025 May 5]. Available from: <https://arxiv.org/abs/1907.11692>
- [13] Cañete J, Donoso S, Bravo-Marquez F, Carvallo A, Araujo V. ALBETO and DistilBETO: Lightweight Spanish Language Models [Internet]. arXiv; 2022 [cited 2025 May 5]. Available from: <https://arxiv.org/abs/2204.09145>
- [14] Pires T, Schlinger E, Garrette D. How multilingual is Multilingual BERT? [Internet]. arXiv; 2019 [cited 2025 Feb 6]. Available from: <https://arxiv.org/abs/1906.01502>
- [15] Brown TB, Mann B, Ryder N, Subbiah M, Kaplan J, Dhariwal P, et al. Language Models are Few-Shot Learners. 2020 [cited 2023 Dec 9]; Available from: <https://arxiv.org/abs/2005.14165>
- [16] Azad MS, Leon SI, Khan R, Mohammed N, Momen S. SAD: Self-assessment of depression for

- Bangladeshi university students using machine learning and NLP. Array. 2025 Mar;25:100372.
- [17] Ghosal S, Jain A. Depression and Suicide Risk Detection on Social Media using fastText Embedding and XGBoost Classifier. Procedia Computer Science. 2023;218:1631–9.
- [18] Hassan S, Shaar S, Darwish K. Cross-lingual Emotion Detection [Internet]. arXiv; 2021 [cited 2025 Jul 15]. Available from: <https://arxiv.org/abs/2106.06017>
- [19] Plaza-del-Arco FM, Martín-Valdivia MT, Ureña-López LA, Mitkov R. Improved emotion recognition in Spanish social media through incorporation of lexical knowledge. Future Generation Computer Systems. 2020 Sep;110:1000–8.

Contribution Role Table

Role	Author(s)
Writing	Luis Muñiz Rascado
Conceptualization	Noé Alejandro Castro Sánchez
Validation	Andrea Magadan Salazar
Supervision	Nimrod González Franco
Software	Gabriel González Serna



This work is
licensed under an
international
Creative Commons Attribution 4.0.