

DETECCIÓN DE EMOCIONES EN TEXTOS EN ESPAÑOL USANDO MODELOS DE LENGUAJE BERT

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>

¹Tecnológico Nacional de México/CENIDET, d23ce169@cenidet.tecnm.mx, <https://orcid.org/0000-0003-2941-2482>.

²Tecnológico Nacional de México/CENIDET, noe.cs@cenidet.tecnm.mx, <https://orcid.org/0000-0002-8083-3891>.

³Tecnológico Nacional de México/CENIDET, andrea.ms@cenidet.tecnm.mx, <https://orcid.org/0000-0001-5474-4150>.

⁴Tecnológico Nacional de México/CENIDET, nimrod.gf@cenidet.tecnm.mx, <https://orcid.org/0000-0002-1051-1379>.

⁵Tecnológico Nacional de México/CENIDET, gabriel.gs@cenidet.tecnm.mx, <https://orcid.org/0000-0002-1874-9402>.

Resumen -- La Secretaría de Salud en México recopila información sobre prevalencia de enfermedades de salud mental para la asignación de recursos médicos, destacando la depresión con una presencia del 5.3%. Factores socioeconómicos, la falta de especialistas afectan su diagnóstico y tratamiento. Para mejorar la detección temprana, se propone el uso de inteligencia artificial (IA), específicamente modelos BERT entrenados para detectar emociones. Se aplicó búsqueda de mejores parámetros y validación cruzada con $k=5$ para evitar sobreajuste, las pruebas y validaciones utilizaron métricas como *accuracy*, *precision*, *recall* y *F1-score*, demostrando que el modelo basado en BERT supera a técnicas tradicionales de aprendizaje automático en la detección de emociones.

Palabras Clave: NLP, BERT, emociones, detección.

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 affects 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.

INTRODUCCIÓN

A nivel nacional, la secretaria de salud se encarga de recaudar información sobre la salud mental de las personas con el fin de hacer proyecciones para la asignación de recursos médicos y afrontar estos padecimientos [1], este reporte menciona que el 19.9% de

la población de México puede padecer un trastorno mental. Una de las enfermedades que tiene mayor presencia en este reporte es la depresión con 5.3% [1] que puede manifestarse en diferentes grados de severidad. Existen diferentes fenómenos que pueden afectar a los cuadros de salud mental de las personas relacionados con la depresión [2], desde factores socioeconómicos, educación, el padecimiento de comorbilidades [3] inclusive la falta de especialistas de salud mental [4]. Por ello es muy importante contar con herramientas digitales que ayuden a detectar emociones relacionadas a los padecimientos de depresión utilizando inteligencia artificial como los lenguajes de gran tamaño tales como BERT [2], [5], [6], [7], que estén reentrenados específicamente para buscar este padecimiento en una etapa moderada y las personas puedan ser canalizadas de una mejor manera.

DESARROLLO

Dataset

Para el proceso de detección de emociones se usó el corpus *emotions* con 40MB en archivo y 194,416 observaciones, este archivo fue obtenido de Kaggle. Desafortunadamente está en idioma inglés, y se tuvo que aplicar un proceso de traducción automática a través de un modelo de lenguaje de gran tamaño especializado en proceso de traducción inglés-español llamado Marian semejante al propuesto en el trabajo de [12], y [13] ya que se necesitaba este corpus específicamente en español para el proceso de entrenamiento en el modelo BERT multilinguaje [8], [9], [10], [11], [14], el proceso de traducción se realizó en un servidor Apple M1 con 16G de memoria compartida.

Una vez que se tenía el corpus ya en idioma español se aplicó una separación para las 6 emociones disponibles en el corpus de *emotions*, que cabe mencionar, que este describe el marco de emociones de Paul Ekman: tristeza (0), alegría (1), amor (2), coraje (3), temor (4), sorpresa (5) de 1000 observaciones por clase para el archivo de

entrenamiento, 500 para cada clase para prueba y 100 cada clase para validación.

Entrenamiento del modelo de predicción

Cuando ya se contaba con el corpus en español y separado para las diferentes tareas, se seleccionó el modelo BERT [14] en su versión “bert-base-multilingual-uncased”, con 177,853,440 parámetros y la capacidad de comprender 104 idiomas, entre ellos español.

Teniendo acceso a BERT a través del API de Transformers proporcionada por Hugging Faces. Este modelo se instanció específicamente con el objeto “BertForSequenceClassification” configurado para tareas de clasificación de secuencias de texto.

Con ayuda de herramientas de búsqueda de hiperparámetros con base en métodos bayesianos implementados en la librería Ray Tune se realizó la búsqueda de la mejor combinación de valores, ver tabla 1, estos son: tasa de aprendizaje (*adam lr*), decaimiento de pesos (*weight decay*), tamaño de lote (*batch size*), *dropout*, este proceso se llevó a cabo en Google Colab usando un GPU A100 con 40G de VRAM, 83G de RAM.

Tabla 1. Valores buscados en la librería Ray Tune para el modelo de emotions.

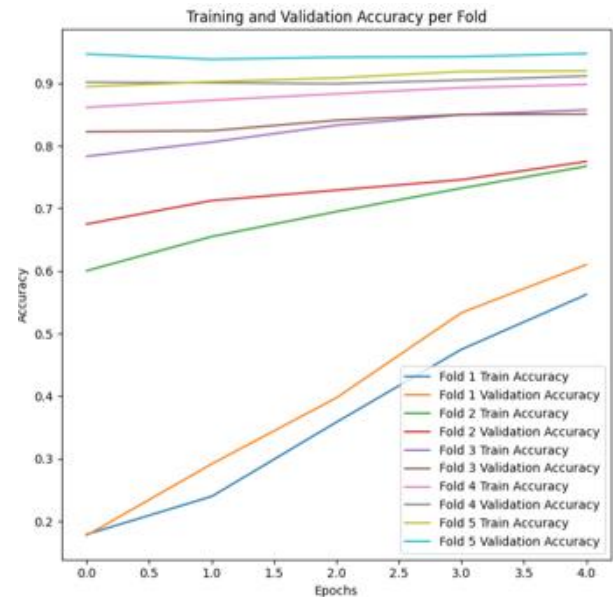
Nombre	Mínimo	Máximo	Sugerido
<i>adam lr</i>	1e-1	1e-7	2e-05
<i>weight decay</i>	0.1	0.5	0.4
<i>batch size</i>	8	32	16
<i>dropout</i>	0.1	0.5	0.3

Fuente. creación propia (2025).

Utilizando los valores de hiperparámetros sugeridos por Ray Tune, también se aplicó la técnica de validación cruzada (*cross-validation*) así tratando de evitar un sobreajuste (*overfitting*) en el modelo de BERT-emotions, con el uso de k=5.

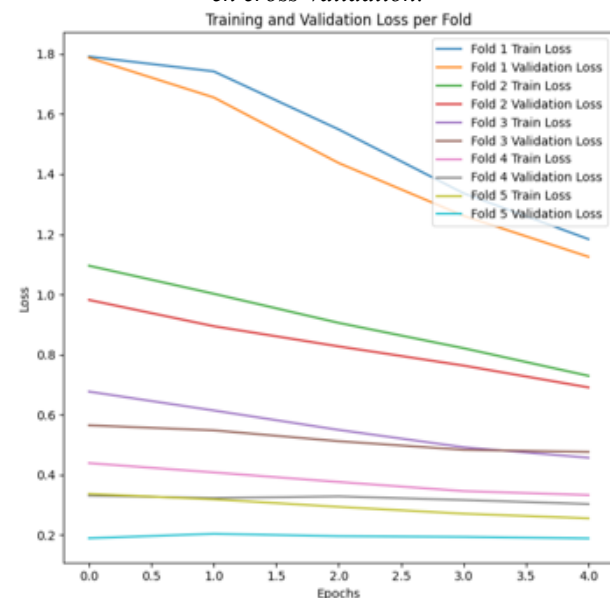
Por cada iteración de k se obtuvo la precisión (*accuracy*) y la pérdida (*loss*), logrando diferentes valores, para el modelo de BERT-emotions, estos se graficaron y se obtuvo lo que se ver figura 1, 2, 3 y 4.

Figura 1. Gráfica de accuracy del modelo de BERT-emotions en cross-validation.



Fuente. creación propia (2025).

Figura 2. Gráfica de loss del modelo de BERT-emotions en cross-validation.



Fuente. creación propia (2025).

Figura 3. Resumen de accuracy del modelo de BERT-emotions en cross-validation.

Train accuracies	Validation accuracies
[0.17916666666666667, 0.23979166666666665, 0.30854166666666665, 0.47479166666666667, 0.5622916666666666]	[0.1775, 0.29166666666666667, 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]

Fuente. creación propia (2025).

Figura 4. Resumen de loss del modelo de BERT-emotions en cross-validation.

Training Loss	Validation Loss
[1.791360681851705, 1.7420288395881653, 1.5490595889091492, 1.336642923351025, 1.1840633503595988]	[1.7876181414252834, 1.6554248207493831, 1.437311511290701, 1.2617667160536115, 1.125373111743676]
[1.0958513792355855, 1.0025018242994945, 0.9058314681053161, 0.8215377628803253, 0.7295879356066386]	[0.9820331178213421, 0.8947605694595137, 0.8275605589151382, 0.7637460631759543, 0.6910714150259369]
[0.6773368521531423, 0.6144728845357895, 0.5498669608434041, 0.4921880375345548, 0.4572153541445732]	[0.5651330528290648, 0.5484175105628214, 0.5122408529645516, 0.48359740094134684, 0.476518167476905]
[0.4394359285632769, 0.4084506352742513, 0.3766838316878, 0.3462496236463377, 0.3332455252110958]	[0.3313105035769312, 0.32341706301820905, 0.32869576819633184, 0.31670198667990535, 0.30331425447213023]
[0.33701842069625854, 0.31895541707674663, 0.29386573274930317, 0.2711464661359787, 0.2558118860423565]	[0.1895500685588837, 0.20440706365594738, 0.1964329890416641, 0.19398186953836366, 0.18894489908492879]

Fuente. creación propia (2025).

Encuesta

Durante los meses de abril y mayo, así como en octubre y noviembre de 2024, se aplicó una encuesta a estudiantes del Tecnológico Nacional de México/Instituto Tecnológico de Zacatepec (TecNM/IT Zacatepec), así como de la Universidad Tecnológica de la Región Norte de Guerrero (UTRNG), para la recolección de las siguientes preguntas:

- 1) Fecha de nacimiento.
- 2) Sexo (masculino, femenino, prefiero no compartir).
- 3) ¿Cuál es la fecha en que sucedió el evento traumático?
- 4) ¿Has podido superar el evento traumático? (si, no).
- 5) Describe con detalle el evento traumático (mínimo 150 palabras).

Con el fin de tener un contexto que describa eventos de una población heterogénea de mayor de 18 años y que el proceso de detección de BERT-emotions pueda predecir con que emoción está relacionada a lo que la persona está mencionando en su evento traumático.

DISCUSIÓN Y ANÁLISIS DE RESULTADOS

Modelo de predicción BERT-emotions

Con respecto a los procesos de prueba y validación del modelo de BERT-emotions se utilizaron las métricas: exactitud (*accuracy*), precisión (*precision*), sensibilidad (*recall*) y puntuación F1 (*F1-score*).

Los resultados de las métricas en prueba se ven en la figura 5, por cada una de las clases.

Figura 5. Resumen de métricas de proceso de prueba de BERT-emotions.

	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

Fuente. creación propia (2025).

Para su proceso de validación, los resultados se ven en la figura 6 por cada una de las clases.

Figura 6. Resumen de métricas de proceso de validación de BERT-emotions.

	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

Fuente. creación propia (2025).

De las métricas anteriores teniendo como base *F1-score* por ser una métrica que está en balance entre las métricas *precision* que representa el número de las observaciones que el modelo clasifica como positivas cuando realmente si lo son y *recall* que representa el número de observaciones reales fueron identificadas de manera correcta con la predicción del modelo.

Así *F1-score* representa la media armónica de estas 2 métricas ver Ec. (1), tabla 2.

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

Con la descripción de *F1-score* el modelo de BERT-emotions, tiene la capacidad de detectar las diferentes emociones en el proceso prueba y validación.

Tabla 2. Emociones relacionadas a su porcentaje según el proceso de prueba y validación.

Emoción predicha	Test	Validation
Porcentaje alto	Tristeza	Alegría, Tristeza

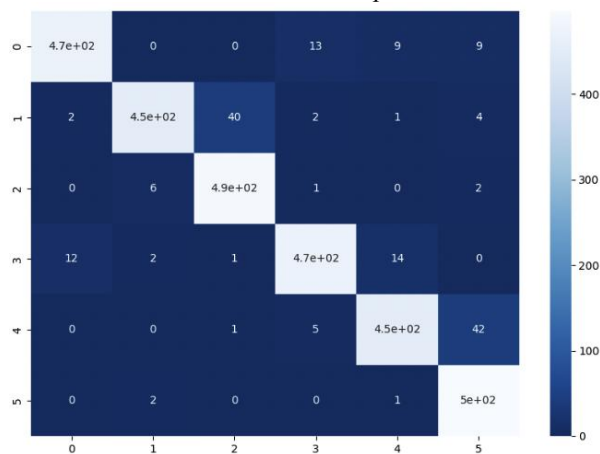
Porcentaje bajo	Temor	Temor
-----------------	-------	-------

Fuente. creación propia (2025).

A su vez, se generaron sus respectivos mapas de calor (*heatmap*) a partir del cálculo de la matriz de confusión (*confusion matrix*) del valor prueba (y_{test}) contra el valor predicho por el modelo de BERT-emotions (y_{pred}) de prueba.

Para el proceso de prueba, ver la figura 7.

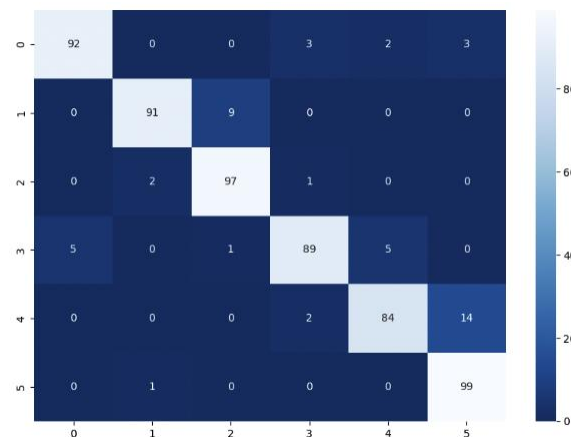
Figura 7. Mapa de calor de la matriz de confusión de BERT-emotions en prueba.



Fuente. creación propia (2025).

La matriz de confusión del valor de validación (y_{val}) contra el valor predicho (y_{pred}) por el modelo de BERT-emotions en el proceso de validación, ver figura 8.

Figura 8. Mapa de calor de la matriz de confusión de BERT-emotions en validación.



Fuente. creación propia (2025).

Modelo BERT-emotions aplicado a la encuesta

En la aplicación de la encuesta se obtuvieron 160 observaciones, de las cuales se eliminaron por preprocesamiento 15 registros resultando solo 145 operativos, se pueden ver en tabla 3.

Tabla 3. Respuestas de la encuesta por sexo.

Sexo	Conteo
Femenino	66
Masculino	76
Prefiero no compartir	3
total	145

Fuente. creación propia (2025).

Después, se analizó el contenido de la descripción del evento traumático con respecto a palabras más frecuentes, para ello, se eliminan las palabras de parada (stopwords) lo que se encontró se puede ver en tabla 4.

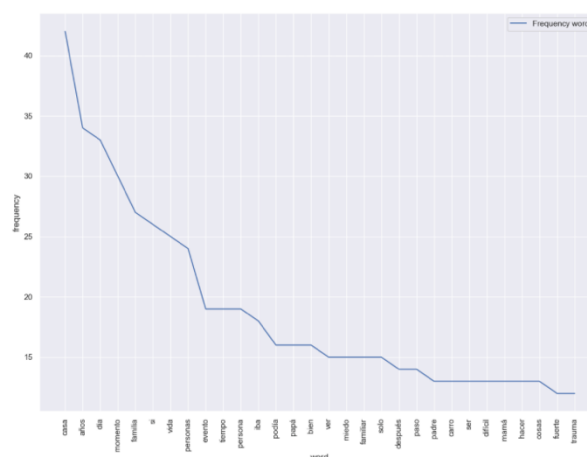
Tabla 4. Las 10 palabras más frecuentes.

Palabra	Conteo
Casa	42
Años	34
Día	33
Momento	30
Familia	27
Si	26
Vida	25
Personas	23
Evento	18
Tiempo	18

Fuente. creación propia (2025).

Usando ese conteo se generó su representación de frecuencias, así como su representación clásica de nube de palabras (*wordcloud*) en donde las palabras más frecuentes tienen una mayor dimensión a comparación a palabras menos frecuentes, ver figura 9 y 10.

Figura 9. Gráfica de frecuencias de palabras.



Fuente. creación propia (2025).

Figura 10. *Gráfica wordcloud.*



Fuente. creación propia (2025).

Se generó un conteo por sexo, y si superó el evento traumático, ver la tabla 5.

Se ejecutó el modelo de BERT-*emotions* sobre la descripción del evento traumático tomando solo donde los participantes que no pudieron superar el evento traumático, y lo que se encontró se puede ver en la tabla 6 y 7.

Tabla 5. *Respuestas de la encuesta por sexo/evento traumático no superado.*

Sexo	Superó evento traumático	Conteo
Femenino	No	27
	Si	39
	<i>total</i>	66
Masculino	No	22
	Si	54
	<i>total</i>	76
Prefiero no compartir	No	1
	Si	2
	<i>total</i>	3

Fuente. creación propia (2025).

Tabla 6. *Respuestas de la encuesta por sexo/evento traumático no superado femenino.*

Sexo	Superó evento traumático	Emoción predicha	Conteo
Femenino	No	Tristeza	10
		Alegría	1
		Amor	4
		Coraje	4
		Temor	8
		Sorpreza	0
		<i>total</i>	27

Fuente. creación propia (2025).

Tabla 7. *Respuestas de la encuesta por sexo/evento traumático no superado masculino.*

Sexo	Superó evento traumático	Emoción predicha	Conteo
Masculino	No	Tristeza	13
		Alegría	3
		Amor	0
		Coraje	1
		Temor	5
		Sorpresa	0
		<i>total</i>	22

Fuente. creación propia (2025).

Las emociones detectadas es la tristeza y temor. En el caso de que si supera el evento traumático ver tabla 8.

Tabla 8. *Respuestas de la encuesta por sexo/evento traumático superado femenino y masculino.*

Sexo	Superó evento traumático	Emoción predicha	Conteo
Femenino	Si	Tristeza	23
		Alegría	1
		Amor	4
		Coraje	4
		Temor	7
		Sorpresa	0
		<i>total</i>	39
Masculino	Si	Tristeza	25
		Alegría	5
		Amor	7
		Coraje	7
		Temor	10
		Sorpresa	0
		<i>total</i>	54

Fuente. creación propia (2025).

Cuando no ha superado el evento traumático para los registros que no compartieron su sexo, ver tabla 9.

Tabla 9. *Respuestas de la encuesta por sexo/evento prefiero no compartir.*

Sexo	Superó evento traumático	Emoción predicha	Conteo
Prefiero no compartir	No	Tristeza	1
		Alegría	0
		Amor	0
		Coraje	0

		Temor	0
		Sorpresa	0
		<i>total</i>	1
	Si	Tristeza	0
		Alegría	0
		Amor	0
		Coraje	0
		Temor	1
		Sorpresa	1
		<i>total</i>	2

Fuente. creación propia (2025).

Teniendo los resultados del modelo de BERT-emotions, se comparó contra ChatGPT-3.5 turbo [15] usando el API de OpenAI, este proceso se llevó a cabo en Google Colab usando un GPU A100 con 40G de VRAM, 83G de RAM.

Para aquellas personas que no pudieron superar el evento traumático y son de sexo femenino la comparación se puede ver en la tabla 10.

Tabla 10. Respuestas de la encuesta por sexo/evento traumático no superado femenino BERT-emotions vs ChatGPT-3.5 turbo.

S e x o	Superó evento	Emoción predicha	BERT- emotions	ChatGPT 3.5-turbo
F	No	Tristeza	10	16
		Alegría	1	0
		Amor	4	1
		Coraje	4	0
		Temor	8	10
		Sorpresa	0	0
		<i>total</i>	27	27

Fuente. creación propia (2025).

Para las personas también de sexo femenino que si pudieron superar el evento traumático su comparación se puede ver en la tabla 11.

Tabla 11. Respuestas de la encuesta por sexo/evento traumático si superado femenino BERT-emotions vs ChatGPT-3.5 turbo.

S e x o	Superó evento	Emoción predicha	BERT- emotions	ChatGPT- 3.5 turbo
F	Si	Tristeza	23	21
		Alegría	1	0
		Amor	4	0

		Coraje	4	0
		Temor	7	16
		Sorpresa	0	2
		<i>total</i>	39	39

Fuente. creación propia (2025).

Así mismo se comparó para personas de sexo masculino con su evento no superado y aquellos que si lo pudieron superar, ver tabla 12, 13.

Tabla 12. Respuestas de la encuesta por sexo/evento traumático no superado masculino BERT-emotions vs ChatGPT-3.5 turbo.

S e x o	Superó evento	Emoción predicha	BERT- emotions	ChatGPT- 3.5 turbo
M	No	Tristeza	13	13
		Alegría	3	0
		Amor	0	0
		Coraje	1	1
		Temor	5	6
		Sorpresa	0	2
		<i>total</i>	22	22

Fuente. creación propia (2025).

Tabla 13. Respuestas de la encuesta por sexo/evento traumático si superado masculino BERT-emotions vs ChatGPT-3.5 turbo.

S e x o	Superó evento	Emoción predicha	BERT- emotions	ChatGPT- 3.5 turbo
M	Si	Tristeza	25	32
		Alegría	5	1
		Amor	7	0
		Coraje	7	2
		Temor	10	17
		Sorpresa	0	2
		<i>total</i>	54	54

Fuente. creación propia (2025).

Por último, la comparación de los participantes que prefirieron no compartir su sexo se puede ver en la tabla 14 y 15.

Tabla 14. Respuestas de la encuesta por sexo/evento traumático no superado prefiero no compartir BERT-emotions vs ChatGPT-3.5 turbo.

S e x o	Superó evento	Emoción predicha	BERT- emotions	ChatGPT- 3.5 turbo
------------------	------------------	---------------------	-------------------	-----------------------

x				
N	No	Tristeza	1	0
		Alegría	0	0
		Amor	0	0
		Coraje	0	0
		Temor	0	1
		Sorpresa	0	0
		total	1	1

Fuente. creación propia (2025).

Tabla 15. Respuestas de la encuesta por sexo/evento traumático si superado prefiero no compartir BERT-emotions vs ChatGPT-3.5 turbo.

S	Superó	Emoción	BERT-	ChatGPT-
e	evento	predicha	emotions	3.5 turbo
x				
o				
N	Si	Tristeza	0	0
		Alegría	0	0
		Amor	0	0
		Coraje	0	0
		Temor	1	2
		Sorpresa	1	0
		total	2	2

Fuente. creación propia (2025).

El proceso de detección de emociones con ChatGPT-3.5 turbo, tiende a perder un poco algunas emociones como alegría, amor, coraje y sorpresa, destacando un rendimiento comparable en tristeza y temor con el modelo entrenado con BERT-emotions para detectar emociones, resaltando que este, si detecta las emociones que su modelo contrincante en la mayoría de los casos omite.

CONCLUSIONES

Los modelos basados en BERT que se han reentrenado para tareas específicas como de clasificación han demostrado un rendimiento superior en esta tarea y en específico de clasificar emociones en comparación con los métodos tradicionales de aprendizaje automático [7], [17].

Esto concuerda con diversos estudios en la literatura, donde modelos de aprendizaje profundo basados en transformers han superado técnicas como Support Vector Machine (SVM), Naïve Bayes y Random Forest y Decision Trees en tareas de clasificación de texto, ya que estos por lo común necesitan mucho trabajo de procesamiento para el tratamiento de las palabras que están en proceso de clasificación, como el tratamiento de lematización que consiste en transformar cada palabra en una representación base o mejor conocida como lema, da

muchas ventajas en ciertos ejercicios, pero su a vez genera una sutil pérdida de contexto que puede ser muy enriquecedor en proceso de clasificación. En comparación con modelos BERT, gracias a su configuración interna y a su mecanismo de tratamiento de palabras propio conocido como wordpiece, no necesitan tanto preprocesamiento previo como el antes mencionado, y lo ocupan tal cual para contextualizarse.

Gracias a el uso de técnicas de búsqueda de hiperparametrización para los mejores parámetros y de validación cruzada se asegura que el modelo de BERT-emotions evite problemas de bajo aprendizaje o en su caso de sobre ajuste y sus predicciones operativas sean consideradas precarias.

Las emociones con mayor prevalencia detectada en la descripción del texto del evento traumático independientemente si se pudo superar o no el evento, la predicción reporta es tristeza y el temor para el modelo de BERT-emotions, pero también tiene cobertura para alegría, amor, coraje y sorpresa. A comparación de ChatGPT-3.5 turbo, pierde estas últimas. Este modelo está entrenado casi para uso general, pero se demuestra con el modelo de BERT-emotions que la especialización en tareas de clasificación con corpus afines mejora el proceso de predicción.

Finalmente, los hallazgos sugieren que la aplicación de modelos de IA en la detección de emociones y problemas de salud mental [15], [16] en español puede ser una herramienta clave para mejorar el acceso a servicios de salud mental.

El trabajo desarrollado por [18] describe el proceso de detección de emociones basado en diferentes lenguajes usando modelos basados en BERT multilinguaje, parte de esas ideas sirvieron como base para nuestra investigación. En este estudio, el modelo es entrenado en una lengua fuente, como por ejemplo en inglés y se evalúa su efectividad de predecir en textos redactados en otros idiomas, como español o árabe, sin que estos contaran con etiquetas anotadas en el lenguaje a predecir. Para esta tarea, emplearon transferencia de conocimiento (transfer learning) como aprendizaje sin ejemplos (zero-shot learning), cabe mencionar para nuestra propuesta, se usó con few-shot learning. También en su estudio, utilizando diferentes modelos BERT como mBERT y XLM-R, que han sido pre-entrenados sobre grandes cantidades de datos multilingües.

Otro estudio [19], propone un enfoque híbrido que combina algoritmos de aprendizaje automático con conocimientos lexicográficos afectivos específicos del español. Para su investigación, primero recolectaron textos obtenidos de tuits en español, agrupándolos en cuatro categorías: ira, miedo, tristeza y alegría. En el tratamiento de características, se extraen vectores tradicionales como bag-of-words o TF-IDF (nuestra

propuesta utilizó *wordpiece* de BERT), junto con indicadores afectivos obtenidos de léxicos emocionales en español. Este conjunto de atributos se utiliza para entrenar clasificadores como regresión logística, SVM, *Naïve Bayes* y perceptrones multicapa.

La integración de estos sistemas en plataformas digitales podría facilitar la detección temprana de casos de depresión, ansiedad, así como posible tendencia suicida permitiendo una intervención oportuna y una mejor canalización de los pacientes hacia profesionales especializados. Para nuestra propuesta, se está desarrollando una plataforma web para integrar el modelo implementado, así como otros modelos que se están entrenando, con el objetivo de brindar ayuda a sectores desprotegidos.

Cabe mencionar que estas herramientas no buscan sustituir a ningún especialista de la salud mental, al contrario, ser un apoyo para ellos y para las personas cuando tenga una crisis.

AGRADECIMIENTOS

Este trabajo ha sido posible gracias al apoyo del antes conocido Consejo Nacional de Humanidades, Ciencias y Tecnologías (CONAHCYT) ahora la Secretaría de Ciencia, Humanidades, Tecnología e Innovación (SECIHTI) a través de la beca otorgada para la realización de mis estudios de doctorado. Agradezco profundamente.

Asimismo, expreso mi gratitud al Centro Nacional de Investigación y Desarrollo Tecnológico (CENIDET) por proporcionarme el conocimiento, los recursos, la infraestructura y el entorno académico necesarios para llevar a cabo este proyecto de investigación. Su apoyo ha sido esencial para el avance y la consolidación de este trabajo. Así como la oportunidad para el desarrollo de mi formación académica.

Mi agradecimiento a los participantes en la encuesta aplicada en el TecNM/IT Zacatepec y la UTRNG por compartir su experiencia y a las personas que apoyaron en difundir y aplicar la misma.

BIBLIOGRAFÍA

- [1] Alcocer Varela J, López-Gatell H. 2o Diagnóstico Operativo de Salud Mental y Adicciones. 2022;
- [2] Mendoza Mojica SA, Márquez Mendoza O, Guadarrama Guadarrama R, Ramos Lira LE. Medición del Trastorno por Estrés Postraumático (TEPT) en universitarios mexicanos. *Salud Ment.* 2013 Jan 1;36(6):493.
- [3] Wagner FA, González-Forteza C, Sánchez-García S, García-Peña C, Gallo JJ. Enfocando la depresión como problema de salud pública en México. 2012;35(1).
- [4] Heinze G, Chapa GDC, Carmona-Huerta J, Departamento de Investigación del Departamento

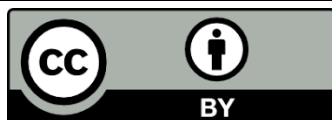
de Psiquiatría y Salud Mental. Facultad de Medicina. Universidad Nacional Autónoma de México., Dirección de Servicios Clínicos. Instituto Nacional de Psiquiatría Ramón de la Fuente Muñiz. Los especialistas en psiquiatría en México: año 2016. sm. 2016 Mar 30;39(2):69–76.

- [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, Carvalho 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.

Tabla de Roles de Contribución

Rol	Autor(es)
Escritura	Luis Muñiz Rascado
Conceptualización	Noé Alejandro Castro Sánchez
Validación	Andrea Magadan Salazar
Supervisión	Nimrod González Franco
Software	Gabriel González Serna



Esta obra está bajo
una licencia internacional
Creative Commons Atribución 4.0.