

APLICACIÓN DE ANÁLISIS DE SENTIMIENTOS EN PQRS PARA LA IDENTIFICACIÓN DE TENDENCIAS EN MODELOS BAYESIANOS EN LAS IA4SG EN BOYACÁ

APPLICATION OF SENTIMENT ANALYSIS IN PQRS FOR THE IDENTIFICATION OF TRENDS IN BAYESIAN MODELS IN IA4SG IN BOYACÁ

Edmundo Arturo Junco Orduz ¹

Carlos Andrés Chaparro Arias ²

Marco Javier Suarez Barón ³

Cómo citar: Junco Orduz, E. A., Chaparro Arias, C. A., & Suarez Barón, M. J. (2026). Aplicación de análisis de sentimientos en PQRS para la identificación de tendencias en modelos bayesianos en las IA4SG en Boyacá. *Conocimiento Global*, 11(S1), 67-91.
<https://doi.org/10.70165/cglobal.v11iS1.670>

Resumen

Este estudio identifica y aplica técnicas de análisis de sentimientos basadas en modelos bayesianos y redes neuronales para clasificar y determinar tendencias en las Peticiones, Quejas, Reclamos y Sugerencias (PQRS) reportadas por entidades de salud del departamento de Boyacá, Colombia, durante 2024-2025, en el marco de la Inteligencia Artificial para el Bien Social (IA4SG). Se entrenaron y evaluaron cuatro modelos computacionales —Naive Bayes, Redes Bayesianas, Procesos Gaussianos y Bayesian Deep Learning— mediante una metodología de enfoque mixto (cuantitativo y cualitativo) y un flujo de Operaciones de Aprendizaje Automático (MLOps) estructurado en cuatro fases: procesamiento de lenguaje natural, entrenamiento de modelos, tratamiento algorítmico y validación cruzada. Los resultados muestran que, en los Procesos Gaussianos, los núcleos RBF y Matern obtuvieron el menor error cuadrático medio (MSE) y error absoluto medio (MAE), con un coeficiente de determinación (R^2) cercano a 0,46; las Redes Bayesianas evidenciaron mayor probabilidad de emociones negativas (tristeza, enfado, miedo) en áreas como vigilancia, consulta externa y radiología, y mayor confianza en vacunación y atención en consultorios; el modelo de Asignación Latente de Dirichlet (LDA) alcanzó su mejor coherencia temática en imagenología y enfermería (0,50-0,52); el clasificador Naive Bayes mantuvo un recall estable entre 0,77 y 0,83 en validación cruzada, y el modelo de Bayesian Deep Learning alcanzó un puntaje F1 de 0,76. Se concluye que estas herramientas optimizan el procesamiento masivo de PQRS, ofrecen a la gestión hospitalaria una alternativa precisa y automatizada para identificar falencias en el servicio, y fortalecen la confiabilidad de las predicciones frente a nuevos datos.

Palabras clave: Análisis de Sentimientos, Asignación Latente de Dirichlet, Red Neuronal Convolucional, PQRS.

Recepción: 29 de mayo de 2026 / Evaluación: 25 de junio de 2026 / Aprobado: 03 de agosto de 2026

¹ Magister en Ingeniería de Sistema y Computación Docente Universidad Pedagógica y Tecnológica de Colombia. Email: edmundo.junco@uptc.edu.co. ORCID: <https://orcid.org/0000-0002-9559-2146>

² Licenciado en Matemáticas y Estadística, Especialista en Estadística Aplicada, Docente Universidad Pedagógica y Tecnológica de Colombia. Email: carlos.chaparroarias@uptc.edu.co ORCID: <https://orcid.org/0009-0007-0778-7810>

³ Doctor en Planeación Estratégica y Dirección de Tecnología Universidad Popular Autónoma del Estado de Puebla. Docente Universidad Pedagógica y Tecnológica de Colombia. Email: marco.suarez@uptc.edu.co. ORCID: <https://orcid.org/0000-0003-1656-4452>

Abstract

This study identifies and applies sentiment analysis techniques based on Bayesian models and neural networks to classify and determine trends in Petitions, Complaints, Claims, and Suggestions (PQRS) reported by health institutions in Boyacá, Colombia, during 2024-2025, within the framework of Artificial Intelligence for Social Good (AI4SG). Four computational models —Naive Bayes, Bayesian Networks, Gaussian Processes, and Bayesian Deep Learning— were trained and evaluated using a mixed-methods approach (quantitative and qualitative) and a Machine Learning Operations (MLOps) workflow structured in four phases: natural language processing, model training, algorithmic treatment, and cross-validation. Results show that, for Gaussian Processes, the RBF and Matern kernels achieved the lowest mean squared error (MSE) and mean absolute error (MAE), with a coefficient of determination (R^2) near 0.46; Bayesian Networks revealed a higher probability of negative emotions (sadness, anger, fear) in areas such as surveillance, outpatient consultation, and radiology, and higher trust in vaccination and outpatient care; the Latent Dirichlet Allocation (LDA) model reached its best topic coherence in imaging and nursing (0.50-0.52); the Naive Bayes classifier maintained a stable recall between 0.77 and 0.83 across cross-validation folds, and the Bayesian Deep Learning model reached an F1-score of 0.76. It is concluded that these tools optimize the massive processing of PQRS, provide hospital management with a precise, automated alternative for identifying service shortcomings, and strengthen the reliability of predictions against new data.

Palabras clave: Sentiment Analysis, Dirichlet Latent Assignment, Convolutional Neural Network, PQRS.

Introducción

En la actualidad, el análisis de las emociones en textos ha tenido una importancia en el ámbito del Procesamiento de Lenguaje Natural (PLN), debido a su capacidad para obtener información acerca de las actitudes, percepciones y estados emocionales que los usuarios expresan en diversas situaciones. La salud es un campo en el que este tipo de investigación es particularmente útil, las experiencias de los pacientes se convierten en una fuente de información que permite evaluar la calidad de la atención, detectar falencias en los procesos y establecer estrategias para mejorar continuamente. Además, las Peticiones, Quejas, Reclamos y Sugerencias (PQRS) se han establecido como una herramienta fundamental de retroalimentación dentro de las entidades de salud, en el cual muestran de manera directa como los usuarios perciben la atención que reciben. Así mismo, en la mayoría de los casos, esta información se utiliza solamente con propósitos administrativos y no para ejecutar investigaciones que faciliten el aprendizaje y la identificación de tendencias (Van Dael et al. 2020).

La importancia del análisis de sentimientos (AS) se está reconociendo en una variedad de dominios diferentes, como productos de consumo, servicios, atención médica, ciencias políticas, ciencias sociales y servicios financieros. Una tarea básica en el análisis de sentimientos es clasificar la polaridad de un texto dado a nivel de documento, oración o característica, es decir, si la opinión expresada en un documento, una oración o una característica de entidad es positiva, negativa o neutral. De esta forma, la clasificación de sentimientos avanzada, "más allá de la polaridad", examina, por ejemplo, estados emocionales como disfrute, ira, disgusto, tristeza, miedo y sorpresa; los textos se componen de palabras significativas que describen las opiniones de las personas. Las opiniones suelen ser expresiones subjetivas que describen los sentimientos y las emociones de las personas hacia entidades, eventos y propiedades el análisis de sentimientos, también conocido como minería de

opiniones, es una tarea que utiliza el procesamiento del lenguaje natural (PLN), el análisis de texto y técnicas computacionales para automatizar la extracción o clasificación de las emociones positivas, negativas y neutras del público general sobre los productos y servicios que han utilizado (Farkhod et al. 2021).

Para Abordar los problemas mencionados, a partir del análisis de sentimientos aplicados en estructuras textuales que contienen información de peticiones, quejas reclamos y sugerencias (PQRS). Para mejorar la calidad de los temas, utilizamos los modelos Naive Bayes (NB), Redes Bayesianas, Procesos Gaussianos, Bayesian Deep Learning y Latent Dirichlet Allocation (LDA).

Análisis de Modelos Bayesianos en PQRS

Naive Bayes (NB)

Naive Bayes es un clasificador probabilístico, lo que significa que para un documento d , de entre todas las clases $c \in C$ El clasificador devuelve la clase $\hat{c}$ que tenga la máxima probabilidad posterior dado el documento. En siguiente ecuación utilizamos la notación de sombrero $\hat{\cdot}$ significa “nuestra estimación de la clase correcta” (Jurafsky and Martin 2026).

$$\hat{c} = \underset{c \in C}{\operatorname{argmax}} P(c|d) \quad (1)$$

Esta idea de inferencia bayesiana se conoce desde el trabajo de Bayes (1763) y fue aplicada por primera vez a la clasificación de textos por Mosteller y Wallace (1964). La intuición de la clasificación bayesiana es utilizar la regla de Bayes para transformar la anterior ecuación en otras probabilidades que tengan algunas propiedades útiles. La regla de Bayes se presenta en la siguiente ecuación; nos da una manera de descomponer cualquier probabilidad condicional $P(x|y)$ en otras tres probabilidades:

$$P(x|y) = \frac{P(y|x)P(x)}{P(y)} \quad (2)$$

Entonces podemos sustituir en la Ecuación 2 en la Ecuación 1 para obtener la Ecuación 3.

$$\hat{c} = \underset{c \in C}{\operatorname{argmax}} P(c|d) = \underset{c \in C}{\operatorname{argmax}} \frac{P(d|c)P(c)}{P(d)} \quad (3)$$

Llamamos a Naive Bayes un modelo generativo porque podemos leer la Ecuación 3 como si se estableciera una especie de suposición implícita sobre cómo se genera un documento: primero se toma una muestra de una clase $P(c)$ y luego las palabras se generan mediante muestreo de $P(d|c)$. (De hecho, podríamos imaginar generar documentos artificiales, o al menos su recuento de palabras, siguiendo este proceso). Hablaremos más sobre esta intuición de los modelos generativos.

Para volver a la clasificación: calculamos la clase más probable $\hat{c}$ dado un documento d , se elige la clase que tiene el mayor producto de dos probabilidades: la probabilidad previa de la clase $P(c)$ y la probabilidad del documento $P(d|c)$:

$$\hat{c} = \underset{c \in C}{\operatorname{argmax}} \underset{\text{likelihood}}{P(d | c)} \underset{\text{prior}}{P(c)} \quad (4)$$

Sin pérdida de generalización, podemos representar un documento d como un conjunto de características $f_1, f_2, \dots, f_n$:

$$\hat{c} = \underset{c \in C}{\operatorname{argmax}} \underset{\text{likelihood}}{P(f_1, f_2, \dots, f_n | c)} \underset{\text{prior}}{P(c)} \quad (5)$$

Desafortunadamente, la Ecuación 5 sigue siendo demasiado difícil de calcular directamente: sin algunas suposiciones simplificadoras, estimar la probabilidad de cada posible combinación de características (por ejemplo, cada posible conjunto de palabras y posiciones) requeriría una enorme cantidad de parámetros y conjuntos de entrenamiento imposiblemente grandes. Por lo tanto, los clasificadores Naive Bayes hacen dos suposiciones simplificadoras.

La primera es la suposición de bolsa de palabras discutida intuitivamente anteriormente: asumimos que la posición no importa y que la palabra. amor “tiene el mismo efecto en la clasificación, ya sea que aparezca como la 1st, 20th, o la última palabra del documento. Por lo tanto, asumimos que las características $f_1, f_2, \dots, f_n$ codifique únicamente la identidad de la palabra y no su posición.

El segundo se llama comúnmente la suposición Naive Bayes: Esta es la suposición de independencia condicional de que las probabilidades $P(f_i|C)$ son independientes dada la clase c y, por lo tanto, pueden multiplicarse “ingenuamente” de la siguiente manera:

$P(f_1, f_2, \dots, f_n|c) = P(f_1|c) \cdot P(f_2|c) \cdot \dots \cdot P(f_n|c)$ La ecuación final para la clase elegida por un clasificador naive Bayes es así: (6)

$$c_{NB} = \operatorname{argmax}_{c \in C} P(c) \prod_{f \in F} P(f | c) \quad (7)$$

Para aplicar el clasificador Naive Bayes al texto, necesitamos considerar las posiciones de las palabras, simplemente recorriendo un índice a través de cada posición de palabra en el documento:

Puestos $\leftarrow$ todas las posiciones de las palabras en el documento de prueba

$$c_{NB} = \operatorname{argmax}_{c \in C} P(c) \prod_{i \in \text{positions}} P(w_i|c) \quad (8)$$

Los cálculos de Naive Bayes, al igual que los cálculos para el modelado del lenguaje, se realizan en espacio logarítmico para evitar el desbordamiento y aumentar la velocidad.

Por lo tanto, la ecuación 8 en general, se expresa como

$$c_{NB} = \operatorname{argmax}_{c \in C} P(c) \prod_{i \in \text{positions}} P(w_i|c) \quad (9)$$

El modelo también incluye una serie de características $X_1, \dots, X_n$ cuyos valores se observan típicamente. La suposición de Naive Bayes es que las características son condicionalmente independientes dada la clase de la instancia. En otras palabras, dentro de cada clase de instancias, las diferentes propiedades pueden determinarse de forma independiente. Más formalmente, tenemos que

$$(X_i \perp X_{-i} | C) \quad \text{para todo } i, \quad (10)$$

donde $X_{-i} = \{X_1, \dots, X_n\} - \{X_i\}$. Este modelo puede representarse utilizando la red bayesiana de la figura 1.

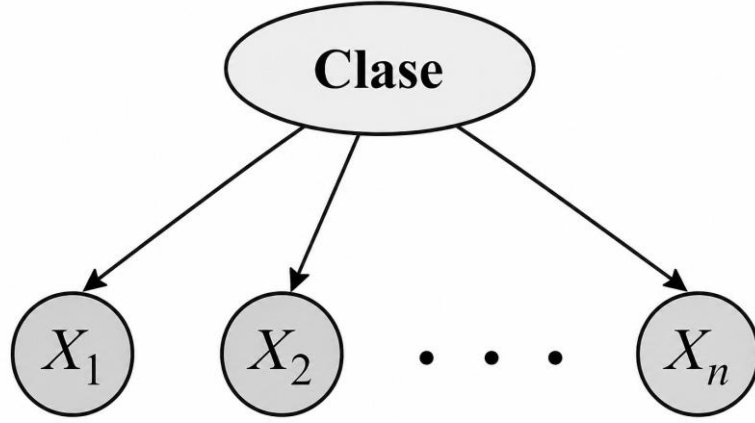


Figura 1: gráfico de red bayesiana para una modelo Naive Bayes Fuente: (Koller y Friedman, 2009).

Redes Bayesianas

Las redes bayesianas son modelos gráficos probabilísticos que pueden representar distribuciones de probabilidad multivariadas factorizándolas mediante distribuciones de probabilidad condicionales locales (CPD). Esta factorización aprovecha la independencia condicional entre las variables aleatorias. Las redes bayesianas son especialmente útiles para representar visualmente distribuciones de probabilidad de alta dimensión con un número reducido de parámetros. Las redes bayesianas pueden ser creadas por expertos humanos o aprendidas a partir de datos (Atienza et al. 2022).

Las redes bayesianas se utilizan habitualmente como modelos estadísticos para datos. En este artículo nos centramos en datos discretos multivariados donde tenemos muchas variables $V = \{V_1, \dots, V_n\}$, y cada variable V_i tiene un número (generalmente muy) finito de valores $(v_{i1}, \dots, v_{in_i})$. Las variables se miden en escala nominal, i.e. El orden de los valores es irrelevante. A menudo nos referimos a los valores por sus índices, de modo que los valores de V_i se toman por $(1, \dots, n_i)$ (Silander and Myllymaki 2012).

Formalmente, una red bayesiana es una tupla $B = (G, \theta)$ donde $G = (V, A)$ es un grafo dirigido acíclico (DAG) con un conjunto de nodos $V = \{1, \dots, n\}$ y un conjunto de arcos $A \subseteq V \times V$, y $\theta = \{P(x_i | \mathbf{x}_{\text{Pa}(i)}), i = 1, \dots, n\}$ es un conjunto de parámetros que con-

tiene una distribución de probabilidad condicional (CPD) para cada variable aleatoria. Una red bayesiana factoriza una distribución de probabilidad conjunta $P(\mathbf{x})$ de un vector de variables aleatorias $\mathbf{X} = (X_1, \dots, X_n)$. El conjunto de nodos V indexa el vector de variables aleatorias, por lo tanto, $\mathbf{X}_V = \mathbf{X}$. La forma en que se factoriza la distribución de probabilidad conjunta depende del conjunto de arcos A de G :

$$P(\mathbf{x}) = \prod_{i=1}^n P(x_i | \mathbf{x}_{\text{Pa}(i)}) \quad (11)$$

Donde $\text{Pa}(i)$ son los índices padres del nodo i en el gráfico G . $P(x_i | \mathbf{x}_{\text{Pa}(i)})$ es el CPD de la variable aleatoria X_i dadas las variables aleatorias $\mathbf{X}_{\text{Pa}(i)}$ contenido en θ .

El número de parámetros de una distribución no factorizada $P(\mathbf{x})$ suele ser más grande que $O(n)$, e.g. es cuadrático, $O(n^2)$, para una distribución gaussiana multivariada, o exponencial $\prod_{i=1}^n |\Omega_i|$, para una tabla de probabilidad condicional (TPC) de una distribución categórica. Dado que normalmente $|\text{Pa}(i)| \ll n$, el conjunto θ contiene n CPDs con un número muy pequeño de parámetros. Por lo tanto, el número de parámetros necesarios para definir una red bayesiana suele ser mucho menor que la representación no factorizada de $P(\mathbf{x})$. Además, el conjunto de

independencias condicionales definidas por la red bayesiana se puede leer directamente del grafo G utilizando el criterio de d-separación.

Procesos Gaussianos

Los procesos gaussianos (PG) son modelos estadísticos flexibles para especificar distribuciones de probabilidad sobre funciones no lineales multidimensionales (Rasmussen y Williams 2006; Neal 1997). Su nombre proviene del hecho de que cualquier conjunto finito de valores de la función se distribuye conjuntamente como una distribución gaussiana multivariada. Los PG se definen mediante una media y una función de covarianza, y se basan en supuestos previos sobre la relación funcional, como la continuidad, la suavidad, la periodicidad y las propiedades de escala. Los PG no solo permiten efectos no lineales, sino que también pueden manejar implícitamente interacciones entre variables de entrada (covariables). Se pueden combinar diferentes tipos de funciones de covarianza para una mayor flexibilidad (Riutort-Mayol et al. 2023).

El paso inicial para aplicar la regresión del proceso gaussiano es definir una media previa $m(\{x\})$ y función de covarianza $k(\{x\}, \{x'\})$, como los GP están completamente especificados por ellos, $\{x\}$ representa el vector de entrada. Para cualquier proceso real $f(\{x\})$ se puede definir:

$$m(\{x\}) = E[f(\{x\})] \quad (12)$$

$$k(\{x\}, \{x'\}) = E[(f(\{x\}) - m(\{x\}))(f(\{x'\}) - m(\{x'\}))] \quad (13)$$

Donde E representa la esperanza. A menudo, por razones prácticas, debido a la notación (simplicidad) y a la falta de conocimiento previo sobre la tendencia general de los datos, la función de media previa se establece en cero. Los procesos gaussianos se pueden definir entonces como

$$f(\{x\}) \sim GP(0, k(\{x\}, \{x'\})) \quad (14)$$

Suponiendo una función de media cero, la función de covarianza podría describirse como

$$\begin{aligned} \text{cov}(f(\{x\}_p), f(\{x\}_q)) &= k(\{x\}_p, \{x\}_q) \\ &= \sigma^2 \exp\left(-\frac{1}{2} \|\{x\}_p - \{x\}_q\|^2\right). \end{aligned} \quad (15)$$

Esta es la función de covarianza exponencial al cuadrado. Es muy importante mencionar una ventaja de la ecuación anterior, ya que la covarianza se expresa como una función únicamente de las entradas. Para la covarianza exponencial al cuadrado, se puede observar que toma valores cercanos a la unidad entre variables cuyas entradas son muy próximas y comienza a disminuir a medida que aumenta la distancia entre las variables en el espacio de entrada (Abdessalem et al. 2017). Suponiendo ahora que se dispone de un conjunto de resultados de entrenamiento $\{f\}$ y un conjunto de resultados de prueba $\{f\}_*$ una tiene la anterior:

$$\begin{bmatrix} \{f\} \\ \{f\}_* \end{bmatrix} \sim N\left(0, \begin{bmatrix} K(X, X) & K(X, X_*) \\ K(X_*, X) & K(X_*, X_*) \end{bmatrix}\right) \quad (16)$$

Donde las letras mayúsculas representan matrices. Se ha utilizado una distribución a priori de media cero para simplificar, y $K(X, X)$ es una matriz cuya, ij -ésimo elemento es igual a $k(x_i, x_j)$. y $K(X, X_*)$ es un vector columna cuyo ij -ésimo elemento es igual a $k(x_i, x_*)$, y $K(X_*, X)$ es la transpuesta de la misma. La matriz de covarianza debe ser simétrica respecto a la diagonal principal.

Dado que la distribución a priori se ha generado mediante las funciones de media y covarianza, para especificar la distribución a posteriori sobre las funciones, es necesario limitar la distribución a priori de manera que incluya solo aquellas funciones que coinciden con los puntos de datos reales. Una forma obvia de hacerlo es generando funciones a partir de la distribución a priori y seleccionando solo aquellas que coinciden con los puntos reales. Por

supuesto, esta no es una forma realista de hacerlo, ya que consumiría mucha potencia computacional. De manera probabilística, la operación se puede realizar fácilmente condicionando la distribución a priori conjunta a las observaciones, lo que dará como resultado (para más detalles, véanse Bishop (1995), Nabney (2002) y Rasmussen y Williams (2006)):

$$\begin{aligned} & \{f\}_* | [X]_*, [X], \{f\} \\ & \sim N \left(\begin{matrix} K([X]_*, [X])K([X], [X])^{-1}\{f\}, K([X]_*, [X]_*) \\ -K([X]_*, [X])K([X], [X])^{-1}K([X], [X]_*) \end{matrix} \right) \end{aligned} \quad (17)$$

Valores de función $\{f\}_*$ se puede generar mediante el muestreo de la distribución posterior conjunta y, al mismo tiempo, evaluando las matrices de media y covarianza de la ecuación anterior.

Las funciones de covarianza utilizadas en este estudio suelen controlarse mediante hiperparámetros para obtener un mejor control sobre los tipos de funciones que se consideran para la inferencia. Uno de los núcleos más empleados para los procesos gaussianos es la función de covarianza exponencial al cuadrado, que puede adoptar la siguiente forma:

$$k_y(x_p, x_q) = \sigma_f^2 \exp \left(-\frac{1}{2l^2} (x_p - x_q)^2 \right) + \sigma_n^2 \delta_{pq} \quad (18)$$

Donde k_y es la covarianza para el conjunto de objetivos ruidosos y (i.e., $y = f(\{x\}) + \epsilon$, donde $\{x\}$ es un vector de entrada y ϵ es el ruido. La escala de longitud l (determina qué tan lejos es necesario moverse en el espacio de entrada para que los valores de la función no se correlacionen), la varianza σ_f^2 de la señal y la varianza del ruido σ_n^2 son parámetros libres que se pueden variar. Estos parámetros libres se llaman hiperparámetros.

La herramienta que se aplica habitualmente para elegir los hiperparámetros 'óptimos' para la regresión GP es la máxima verosimilitud marginal de las predicciones $p(\{y\} | [X], \{\theta\})$ con respecto a los hiperparámetros θ :

$$\begin{aligned} \log p(\{y\} | [X], \{\theta\}) &= -\frac{1}{2} \{y\}^T [K_y]^{-1} \{y\} \\ &\quad -\frac{1}{2} \log | [K_y] | - \frac{n}{2} \log 2\pi \end{aligned} \quad (19)$$

donde $[K_y] = [K_f] + \sigma_n^2 I$ es la matriz de covarianza del conjunto de prueba ruidoso $\{y\}$ y $[K_f]$ es la matriz de covarianza libre de ruido. Para optimizar estos hiperparámetros maximizando la probabilidad logarítmica marginal, las derivadas parciales dan la solución, mediante descenso de gradiente:

$$\begin{aligned} \frac{\partial}{\partial \theta_j} \log p(\{y\} | [X], \{\theta\}) &= \frac{1}{2} \{y\}^T [K]^{-1} \frac{\partial [K]}{\partial \theta_j} [K]^{-1} \{y\} \\ &\quad - \frac{1}{2} \text{tr} \left([K]^{-1} \frac{\partial [K]}{\partial \theta_j} \right) \\ &= \frac{1}{2} \text{tr} \left((\alpha \alpha^T - [K]^{-1}) \frac{\partial [K]}{\partial \theta_j} \right) \end{aligned} \quad (20)$$

donde $\{\alpha\} = [K]^{-1} \{y\}$. Por supuesto, esta solución no es un procedimiento trivial.

Bayesian Deep Learning

Por lo general, un modelo BDL consta de dos componentes: (1) un componente de percepción que es una formulación bayesiana de un cierto tipo de redes neuronales y (2) un componente de tarea específica que describe la relación entre diferentes variables ocultas u observadas utilizando PGM. La regularización es crucial para ambos. Las redes neuronales suelen tener una gran cantidad de parámetros libres que deben regularizarse adecuadamente (Wang and Yeung 2016).

La idea principal detrás del BDL es definir la distribución posterior de los datos para el modelo de redes neuronales para cuantificar la incertidumbre del modelo. Para calcular la

distribución posterior, se utiliza la distribución previa $P(\omega)$ se elige para los parámetros del modelo, que se basa principalmente en la experiencia previa, y la mayoría de las veces se suele utilizar la distribución gaussiana/normal como previo por defecto (Abdullah et al. 2022). La probabilidad $P(D|\omega)$ se calcula utilizando los datos dados antes de calcular la probabilidad posterior (probabilidad condicional). La probabilidad posterior siempre es proporcional a la probabilidad previa multiplicada. El teorema de Bayes se define en

$$P(A | B) = \frac{P(A).P(B|A)}{P(B)} \quad (21)$$

En BDL, la probabilidad de parámetros dados los datos $P(\omega | D)$ se utiliza, y la ecuación bayesiana se puede definir como en

$$P(\omega | D) = \frac{P(\omega).P(D|\omega)}{P(D)} \quad (22)$$

El término probabilidad marginal (evidencia) que es $P(D)$ se utiliza para normalizar la posterior.

Para generar el posterior $P(\omega|D)$ distribución, el producto de la verosimilitud y la distribución a priori se calculan para todos los parámetros, lo que implica encontrar integrales sobre todos los parámetros, lo que se denomina $P(D)$ o evidencia de «verosimilitud marginal». La integral puede calcularse analíticamente si se conoce la forma de la distribución posterior, especialmente cuando se utilizan distribuciones de probabilidad conjugadas para la distribución a priori y la verosimilitud. Sin embargo, en la mayoría de los casos «cuando las distribuciones no son conjugadas», la distribución posterior es intratable, lo que implica que no puede resolverse analíticamente. La distribución posterior se calcula de esta manera mediante métodos de aproximación, a menudo conocidos como muestreo. Con todas sus varianzas, el método de Monte Carlo de cadenas de Markov (MCMC) es el método de muestreo más utilizado para la inferencia bayesiana

Para encontrar la distribución predictiva $P(y^* | x^*, D)$, pesos ω las redes se consideran variables. y^* son resultados previstos, x^* son observaciones, y D son datos de entrenamiento, mientras que $P(\omega|D)$ es la verdadera distribución posterior. es la verdadera distribución posterior. La aproximada posterior $q(\omega)$ se utiliza en lugar de la verdadera distribución posterior para hallar la distribución predictiva, como se muestra en

$$P(y^* | x^*, D) = \int P(y^* | x^*, \omega)P(\omega | D)d\omega \quad (23)$$

$$\approx \int P(y^* | x^*, \omega)q(\omega)d\omega$$

Latent Dirichlet Allocation (LDA)

El modelado de temas se encuentra entre los métodos de clasificación de documentos de texto no supervisados. Consiste en identificar los temas presentes e inferir patrones ocultos en grandes volúmenes de documentos de texto. La asignación latente de Dirichlet (LDA, por sus siglas en inglés) es uno de los métodos de modelado de temas más extendidos y utilizados (Aziz et al., 2022) (Jankova 2023). Para modelar temas, Blei et al. (2003) desarrollaron modelos probabilísticos de temas, como el Análisis Semántico Latente (pLSA) y la Asignación Latente de Dirichlet (LDA). El principio de estos modelos se basa en la suposición de que modelan un documento como una mezcla de temas ocultos (latentes). Los temas individuales se definen mediante la distribución probabilística de las palabras en el vocabulario.

Asignación Latente de Dirichlet

LDA es un modelo estadístico generativo que explica la similitud entre ciertos datos. Primero, se define de antemano el número de temas a buscar y, a continuación, se selecciona el

tema de Feuerriegel y Pröllochs (2021) con base en la mayor probabilidad respecto a un conjunto específico de palabras. Se examinan dos distribuciones de probabilidad para encontrar los temas del modelo LDA:

- $\alpha = P(t|d)$, indicando la probabilidad del tema t en el documento
- $d; \beta = P(w|t)$, indicando la probabilidad de la palabra w en el tema t .

Blei et al. (2003) y Blei (2012) describen este proceso de dos fases con más detalle de la siguiente manera:

1. Para cada documento de texto d en el corpus C , temas θ_d están distribuidos aleatoriamente, donde $\theta_{d,k}$ indica la proporción del número de temas k en el documento d . La variable aleatoria θ_d sigue la distribución de probabilidad de Dirichlet con α dada por la relación:

$$\theta_d \sim \text{Dir}(\alpha), \quad \alpha = (\alpha_1, \alpha_2, \dots, \alpha_N)$$

2. Las palabras se asignan a los temas de tal manera que se selecciona la palabra n del documento d para el tema seleccionado $t_{d,n}$ de θ_d . Posteriormente, las palabras son $W_{d,n}$ seleccionado de un diccionario fijo, que depende del tema seleccionado $t_{d,n}$. La distribución viene dada por la relación $\beta_k \sim \text{Dir}(\eta)$ con previa η .

La probabilidad uniforme viene dada por:

$$P(\theta, \beta, w, t) = \prod_{d=1}^D P(\theta_d | \alpha) \prod_{k=1}^K P(\beta_k | \eta) \prod_{n=1}^N P(t_{d,n} | \theta_d) P(w_{d,n} | t_{d,n})$$

Sin embargo, no es posible maximizar directamente la probabilidad, ya que no se pueden observar los temas, solo los documentos.

$$P(\theta, \beta, t | w, \alpha, \eta) = \frac{P(\theta, \beta, w, t | w, \alpha, \eta)}{P(w | \alpha, \eta)}$$

En un paso posterior, LDA asigna un identificador de nombre de tema a cada tema extraído. Según Feuerriegel y Pröllochs (2021), se examina una lista de las tres a treinta palabras más probables en el tema dado para su correcta identificación.

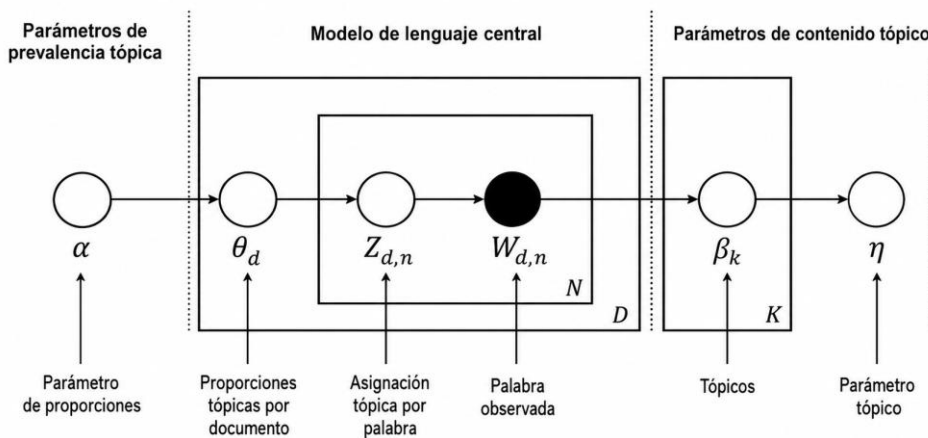


Figura 2: El modelo LDA. Fuente: (Blei, 2012).

Como se muestra en la figura anterior, el modelo gráfico para la asignación latente de Dirichlet. Cada nodo es una variable aleatoria y está etiquetado según su función en el proceso generativo. Los nodos ocultos (proporciones de temas, asignaciones y temas) no están sombreados. Los nodos observados (las palabras de los documentos) están sombreados. Los rectángulos representan la notación de "placa", que indica replicación. La placa N representa la colección de palabras dentro de los documentos; la placa D representa la colección de documentos dentro de la colección (Blei 2012).

Mediante el aprendizaje a partir de muestras sin etiquetar, el análisis de clústeres busca descubrir patrones ocultos, subconjuntos y correlaciones dentro de un conjunto de datos, lo que permite clasificar los objetos de datos sin etiquetar en clases correspondientes, denominadas clústeres de clase. Un método fiable de partición de clústeres debe lograr que el resultado de la partición sea lo más similar posible entre los objetos de datos dentro del mismo clúster de clase (Xiao et al., 2022). Por el contrario, debe ser lo más diferente posible entre los objetos de datos dentro de diferentes clústeres de clase (Zou et al. 2022).

El algoritmo de agrupamiento K-means es el más fundamental y ampliamente utilizado. Su concepto básico consiste en encontrar, iterativamente, un esquema de clasificación para K clústeres que minimice la función de pérdida correspondiente a los resultados del agrupamiento. La función de pérdida se define como la suma de los cuadrados del error de cada muestra con respecto al centro del clúster (Likas et al., 2003).

$$J(c, /m\mu) = \sum_{i=1}^M ||x_i - \mu_{c_i}||^2$$

Donde x_i representa las muestras del clúster al que pertenece, $/\mu_{c_i}$ representa el punto central correspondiente al clúster, y M representa el número total de muestras. Pasos que implica la operación:

1. Seleccionar K puntos al azar de la muestra como centroide inicial
2. Divida las muestras en grupos que correspondan al centroide más cercano calculándola distancia entre cada muestra y cada centroide.
3. Calcula el valor medio de las muestras de cada clúster y actualiza el centro de masa del clúster con dicho valor medio.
4. Los pasos 2 y 3 deben repetirse hasta que se cumpla una de las siguientes condiciones: Se alcanza el número máximo de iteraciones cuando el cambio en la posición del centroide es menor que el valor umbral (por defecto, 0,0001).

Metodología

El propósito estructural de flujo de trabajo en los modelos Bayesianos se orienta en 6 etapas representativas, en donde abarca el análisis y patrones de los datos suministrados por las entidades de hospitales en Boyacá. el proceso de la metodología propone un proceso enfoque de integración en técnicas para el reconocimiento de patrones de cada etapa representada en la Arquitectura del flujo del análisis, diseño, desarrollo y la evaluación del modelo de Bayesianos (Contreras Contreras et al. 2022).

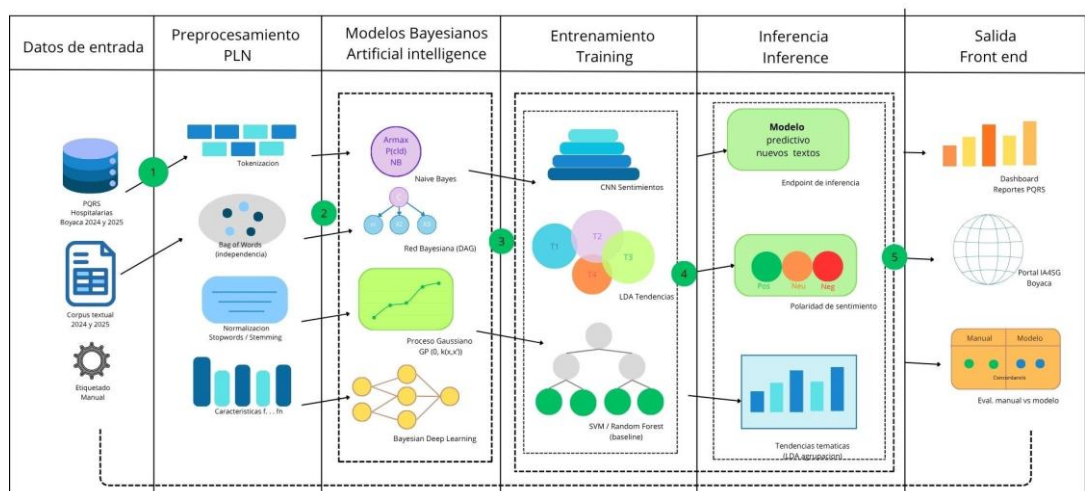


Figura 3: Arquitectura. Fuente: elaboración propia

la Figura 3 representa las seis etapas del procesamiento de datos en los modelos computacionales. Este flujo tiene como objetivo identificar el análisis de sentimientos en la gestión de Peticiones, Quejas, Reclamos y Sugerencias (PQRS), con un enfoque de Inteligencia Artificial para el Bien Social (IA4SG), aplicado a los datos suministrados por prestadoras de salud del departamento de Boyacá durante los años 2024 y 2025.

se identifican seis secciones:

1. (Datos de entrada): Información suministrada por las PQRS de los hospitales de Boyacá'.
2. (Procesamiento de Lenguaje Natural - PLN): Fase en la que se identifican las palabras para realizar la tokenización y la representación vectorial.
3. (Modelos Bayesianos): Identificación y aplicación de algoritmos como Naive Bayes, Redes Bayesianas, Procesos Gaussianos y Bayesian Deep Learning.
4. (Tratamiento algorítmico): Uso de Redes Neuronales Convencionales (CNN) para la clasificación de textos emocionales, Asignación Latente de Dirichlet (LDA) para su agrupación, y Máquinas de Vector de Soporte (SVM) junto con Random Forest para la clasificación y predicción de Machine Learning en los modelos bayesianos.
5. (Inferencia y Tendencia): Determinación de la inferencia para la polaridad de los sentimientos y análisis de la tendencia para la validación cruzada de la agrupación.
6. (Resultados finales): Evaluación definitiva de los modelos bayesianos implementados.

Procesamiento de lenguaje natural (PLN)

En (Bird et al. 2009), el campo del procesamiento de lenguaje natural interpreta en este artículo, los procesos de manipulación computacional en la interpretación del lenguaje humano y las palabras en una comprensión profunda en elementos lingüísticos que son utilizados en el aprendizaje de maquina dentro de la Inteligencia Artificial, como aportes significativos para su transformación estructurada a través de un flujo claro : se inicia con los accesos limpios en un entorno web, de igual manera las tokenizaciones que por medio de cadena de caracteres se realiza los llamados, la normalización contribuye técnicas como stemming o lematización, donde se categoriza los patrones en modelo de aprendizaje supervisado para la extracción de la

información en procesos de segmentación (chunking) y entidades nombradas y ser analizados los datos Lingüísticos en el contexto del análisis de sentimientos de las PQRS.

Modelos Bayesianos

Los modelos bayesianos (Jospin et al. 2022), cumple emplearon cuatro algoritmos de aprendizaje automático tradicionales: Random Forest, Naive Bayes, Árbol de Decisión y el clasificador Perceptrón Multicapa. Los resultados finales mostraron que los algoritmos más adecuados para la clasificación fueron Random Forest y Perceptrón Multicapa.

En el libro de Kevin P. Murphy (Murphy 2023) identifica los modelos lineales y probabilísticos donde, su volumen 1 y 2 describe algunos algoritmos de los modelos bayesianos más utilizados bajo la premisa de la inteligencia artificial, su objetivo es interpretar los modelos Naive Bayes, Redes Bayesianas, Procesos Gaussianos y Bayesian Deep, para predecir datos supervisados en el aprendizaje profundo.

Tratamiento algorítmico

Los (Chahar and Dubey 2026), procesos algorítmicos para procesar los datos en los sentimientos de las PQRS de los hospitales regionales de Boyacá, se tiene en cuenta tres modelos fundamentales como Red Neuronal Convolucional (CNN) para procesar las peticiones, el Análisis Discriminante Lineal (LDA) para clasificar y reducir la bidimensionalidad y Máquina de Vectores de Soporte (SVM) para encontrar fronteras lineales y separación de los datos en un vector. Todo esto está orientado en procesos supervisados.

En (Altarturi et al. 2023), el modelo Latent Dirichlet Allocation (LDA) como se muestra en el 2,5 se toma etiquetas para la coherencia de la parte estadística para identificar los patrones en la clasificación y agrupamiento (Clustering) en el contexto de clasificación de sentimientos de los pacientes al utilizar las PQRS como: anticipación, asco, confianza, enfado, miedo, sorpresa, tristeza y alegría.

Marco metodológico

Los procesos de diseñar, implementar y evaluar dentro de un marco metodológico, para el análisis de emociones en textos de PQRS provenientes de entidades prestadoras de servicios de salud en Boyacá. Su análisis en emociones se centra en anticipación, asco, confianza, enfado, miedo, sorpresa, tristeza, alegría, entre otras que son identificadas para las emociones negativas en donde tiene mayor presencia.

Enfoque de Investigación

La (Younas et al. 2025), investigación acogida en el marco del desarrollo y estratégico del presente artículo son metodologías mixtas, donde los resultados del enfoque en los procesos son (Cuantitativo y Cualitativo), busca la interpretación de la metainferencias en métodos mixtos para la integración, la sinergia y la validación cruzada. por ello los resultados esperados adquiridos de los datos de los hospitales regionales de Boyacá, en su comprensión en fenómenos del comportamiento en la clasificación del análisis de sentimientos en la PQRS, esto conlleva al proceso de integración de estudios Mixed Methods Research (MMR), finalmente se obtiene después de las fases cuantitativa y cualitativas la Metainferencia para su integración es el nivel de sentimientos que son solicitados por los usuarios de los hospitales.

Metodología de Tecnológico de Desarrollo

En los elementos fundamentales de la investigación se utilizaron herramientas tecnológicas que se muestran las diferentes etapas del desarrollo de la presente investigación. Se usó el lenguaje de programación Python, como entorno de desarrollo, fueron usados Google

Colab y Visual Studio Code; adicionalmente, para la parte de contenerización y despliegue, se usó Docker, GCP y Streamlit Cloud.

Metodología de Desarrollo MLOps

La (SACRISTÁN, n.d.), metodología integra los procesos del análisis de emociones en PQRS de entidades prestadoras de servicios de salud en Boyacá. Este flujo de trabajo abarca etapas interconectadas, desde la recolección y preprocesamiento de datos ciclo que especifique sus procesos de despliegue y escalabilidad de la investigación, la cual se establece en tres fases de desarrollo del modelo, la implementación, el monitoreo y el mantenimiento, en el caso de estudio para evaluar calidad y los pronosticar del bienestar social.

Recolección y preparación de datos, Entrenamiento de los modelos, Evaluación de los modelos y Implementación de API, visualización de resultados y despliegue.

Integración de Herramientas Tecnológicas

Para la integración (Demin and Ponomaryov 2022), de los modelos en herramientas tecnológicas se implementó estrategias y librerías de Python entre otras que fundamentan las librerías de manejo para los modelos de Naive Bayes, Redes Bayesianas, Procesos Gaussianos y Bayesian Deep Learning.

Fases de Desarrollo

En las fases que se implementó para obtener los resultados esperados de esta investigación, se manejó las 6 secciones como se evidencio en la figura 3 de la arquitectura, lo cual se debe tener en cuenta un proceso de actividades específicas tales como:

1. Fase 1 Procesamiento de Lenguaje Natural
2. Fase 2 Utilizar modelo Naive Bayes, Redes Bayesianas, Procesos Gaussianos y Bayesian Deep Learning.
3. Fase 3 Tratamiento algorítmico
4. Fase 4 validación cruzada Inferencia y Tendencia

Etiquetado de emociones

La etapa de etiquetado de emociones se realizó sobre el conjunto de datos consolidado con el fin de reconocer la presencia de emociones concretas en los textos. Para este fin, se utilizó el diccionario léxico del NRC Emotion Lexicon, modificado al español [5], que relaciona palabras con 8 emociones: confianza, anticipación, asco, enfado, miedo, sorpresa, tristeza y alegría. El etiquetado de emociones, que comprendió la tokenización de textos y su comparación con el léxico emocional. Se utilizó un procedimiento de clasificación multietiqueta para asignar etiquetas emocionales, lo cual posibilitó que se identificaran diversas emociones en cada registro examinado. Esto se debe a que los relatos de experiencias en servicios de salud pueden expresar diversas emociones en un único registro. Se categorizaron los textos sin una carga emocional significativa en 9 casos como “indeterminado”, los cuales reflejan principalmente contenidos descriptivos o técnicos sin evidencia de emociones evidentes. El entrenamiento de los modelos de inferencia bayesiana se fundamentó en el conjunto de datos etiquetados, que se obtuvo como resultado de este procedimiento.

Modelos de lenguaje grande (LLMs)

En los sistemas de inteligencia artificial se considera modelos que identifican contextos de lenguaje humano para el procesamiento en técnicas de aprendizaje profundo en pro de un lenguaje natural, por ello son utilizadas en contextos de efectos de emociones de los usuarios en las PQRS.

Métodos bayesianos en LLMs

En (Fernández et al. 2024), se aborda el desarrollo de métodos computacionales bayesianos para la creación y automatización de chatbots aplicados a la gestión de PQRS en entidades hospitalarias del departamento de Boyacá, empleando herramientas de código abierto como Dialogflow. De manera similar.

De igual manera (Zapata-Ros, n.d.), se analizan métodos bayesianos basados en estrategias de LLM, los cuales utilizan la distribución de longitudes de palabras y frases, así como la frecuencia y probabilidad de términos (específicamente las emisiones en los parámetros de texto) mediante distribuciones estadísticas en enfoques híbridos bayesianos.

Aprendizaje Profundo en LLMs

El (Patil y Gudivada 2024), Modelo de Lenguaje Grande LLMs, construye técnicas de aprendizaje profundo con el funcionamiento de una arquitectura Transformer para identificar palabras dentro un contexto completo orientado a componentes y poder hacer entrenamientos. Se puede clasificar diferentes mecanismos en contextos multimodales para la integración en enfoques supervisados durante su transferencia dentro del conjunto de datos, esto puede ser aplicado en procesos Transformador Generativo Preentrenado (GPT) permite realizar análisis de sentimientos avanzados sobre el comportamiento y las emociones de los usuarios en entornos hospitalarios de Boyacá. De este modo, la investigación aprovecha la transferencia de datos para generar un impacto positivo orientado al bien común en el sector salud.

Resultados y Discusión

Procesos Gaussianos

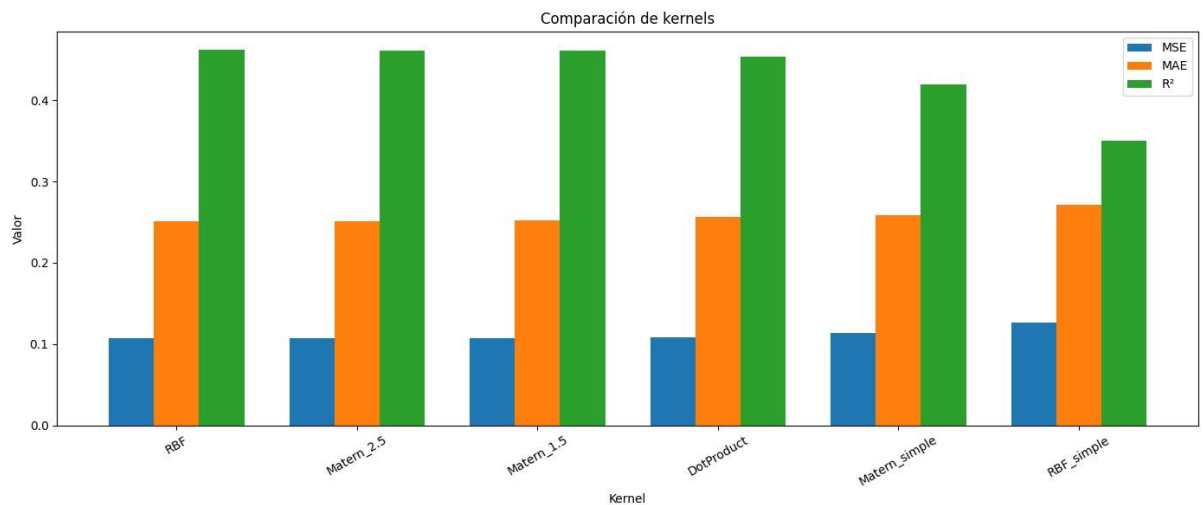


Figura 4: Comparación de Kernels. Fuente: elaboración propia

La gráfica muestra una comparación del rendimiento de diferentes Kernels usados en un modelo de regresión probabilístico. Así mismo, el funcionamiento general de cada Kernel se define como el modelo interpreta la relación de los datos. Interpretación típica:

- menor MSE → menor error cuadrático.
- menor MAE → menor error absoluto.
- mayor R^2 → mejor ajuste predictivo.

Por lo cual, RBF, Matern2.5 y Matern1.5 presentan el menor MSE, de esta forma, RBF Y Matern predicen más cerca los valores reales. De igual forma, el error absoluto promedio se

aproxima a 0.25 unidades, es así que los mejores MAE son: RBF y Matern2.5. Finalmente, R^2 el coeficiente de determinación, explica la variabilidad de los datos teniendo un mejor desempeño RBF, Matern2.5 y Matern1.5 aproximadamente con 0.46. De esta forma, los Kernels RBF y Matern modelan mejor la estructura de los datos, llegando a la conclusión que el modelo logra exponer aproximadamente entre 35% y 46% la variabilidad de los datos.

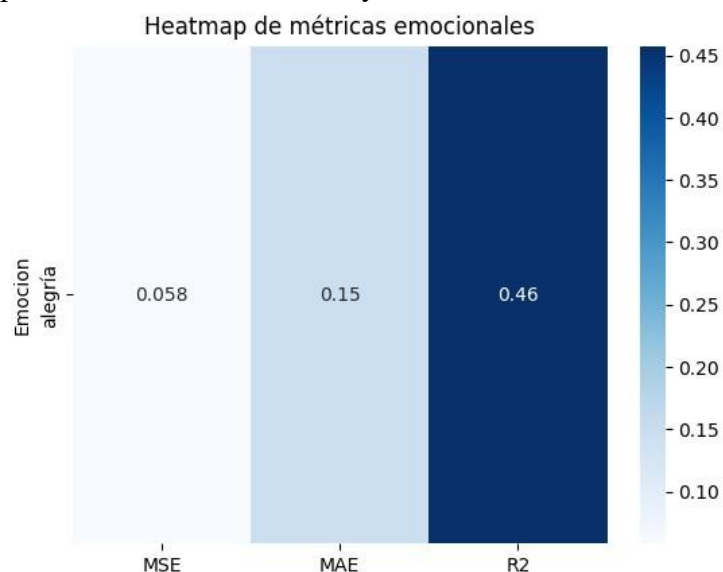


Figura 5: Heatmap métricas emocionales. Fuente: elaboración propia

En la gráfica heatmap (mapa de calor), se resume un modelo para la emoción alegría utilizando tres métricas: MSE, MAE y R^2 . En un primer momento, el MSE (Error cuadrático medio) con un 0.058 es relativamente bajo lo que propone que el modelo no comete errores muy grandes al predecir alegría. De igual forma, el MAE (Error absoluto medio) del modelo se equivoca en aproximadamente 0.15 unidades lo cual se considera un error moderado. Finalmente, R^2 con el 0,46 significa que el modelo explica aproximadamente el 46% de variabilidad de la emoción “Alegría” logrando un desempeño moderado.

Redes Bayesianas

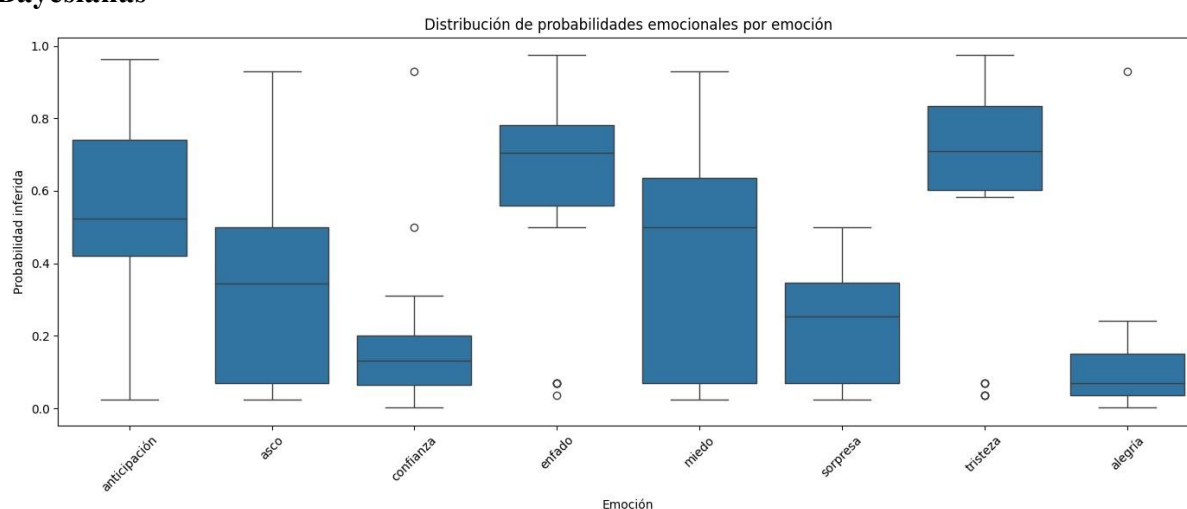


Figura 6: Distribuciones de probabilidades. Fuente: elaboración propia

En el diagrama de cajas y bigotes (boxplot) se puede observar la distribución de las probabilidades inferidas para las diferentes emociones. De esta forma, la emoción por anticipación tiene una mediana cercana a 0.52, presenta dispersión e indica que la emoción

aparece con una intensidad media-alta considerable entre muestras. En la emoción de asco, se presenta una mediana de 0.35 con una distribución amplia desde valores muy bajos y Altos. En la emoción de confianza la mayoría de los valores es baja con una mediana cercana a 0.13 y algunos valores atípicos. En la emoción de enfado, presenta una mediana de 0.70 con una distribución relativamente concentrada lo cual sugiere que el enfado es una emoción consistente en los datos; respecto a la emoción de miedo tiene una mediana cercana a 0.50 con valores bajos y altos con algunas fluctuaciones importante. Respecto a la emoción de sorpresa tiene una mediana de 0.25 variabilidad moderada, respecto con la tristeza, es una de las emociones más altas con una mediana cercana a 0.70, lo cual sugiere que es una emoción dominante en gran parte de las observaciones y finalmente alegría principalmente baja con una mediana cercana al 0.06, presentando un valor atípico de 0.93.

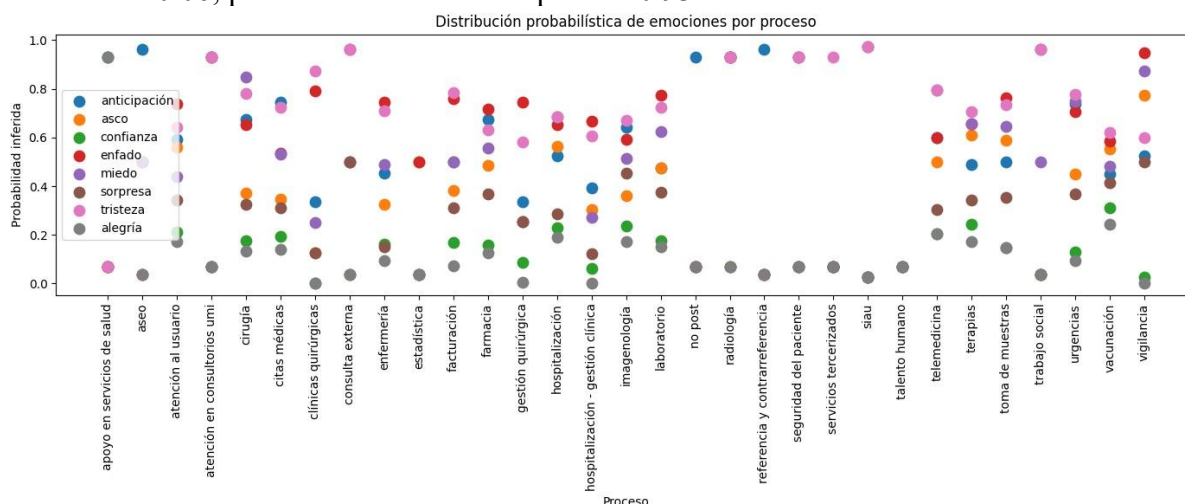


Figura 7: Distribuciones de probabilidades por proceso. Fuente: elaboración propia

En la gráfica se puede observar una distribución probabilística por proceso dentro de las diferentes áreas o servicios, en el eje horizontal los procesos (cirugía, farmacia, urgencias, laboratorio, vacunación, facturación, etc.) y en el eje vertical la probabilidad inferida de cada emoción, con valores entre 0 y 1. De igual forma, las emocionantes predominantes son: tristeza, enfado, anticipación y miedo. De esta forma, se puede mencionar aquellos procesos con mayor carga emocional negativa:

- Vigilancia: presenta valores altos en enfado, miedo, asco y tristeza.
- Servicios tercerizados y seguridad del paciente: tristeza, anticipación y miedo.
- Consulta externa y radiología: tristeza, anticipación y miedo; posiblemente relacionados con espera de resultados, diagnósticos o incertidumbre clínica.

Así mismo, los procesos con mayor confianza o emociones positivas son: Atención en consultorios umi: tiene una combinación alta de confianza y alegría moderada Vacunación: aunque presenta emociones mixtas, mantiene niveles moderados de confianza, algo de alegría y anticipación controlada. Respecto a la alegría; se mantiene baja en casi todos los procesos excepto en casos puntuales como: apoyo en servicios de salud, telemedicina y vacunación. Finalmente, en la gráfica se evidencia que en la mayoría de los procesos analizados están asociados a emociones como: tristeza, miedo y anticipación.

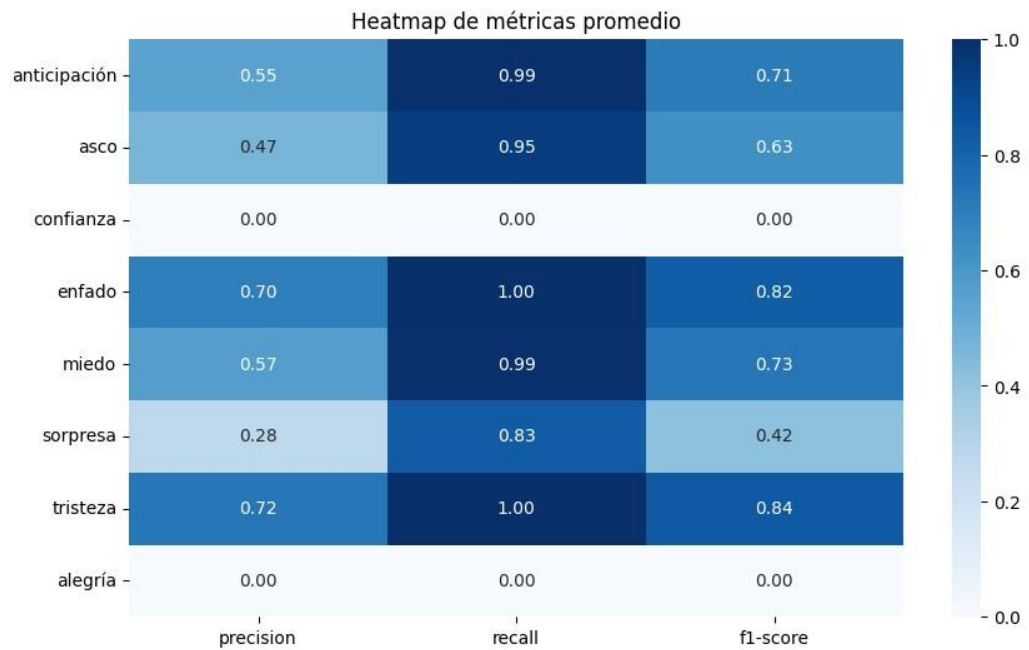


Figura 8: Métricas promedio. Fuente: elaboración propia

Como se puede observar, la gráfica de mapa de calor (heatmap) de métricas promedio del modelo de emociones, los valores van de 0 a 1, donde valores cercanos a 1 indica mejor desempeño y valores cercanos a 0 indican mal desempeño. Así mismo, en la emoción de enfado con una precisión de 0.70, además, el modelo detecta muy bien el enfadado, con un recall de 1.00 identificó los casos reales de enfado; la emoción de tristeza con una precisión de 0.72 en lo cual se puede evidenciar que es una mejor clasificada identificando todos los casos y manteniendo una buena precisión. Finalmente, la emoción del miedo una precisión de 0.57, el modelo detecta casi siempre el miedo, pero se puede confundir con otras emociones.

Métrica LDA procesos

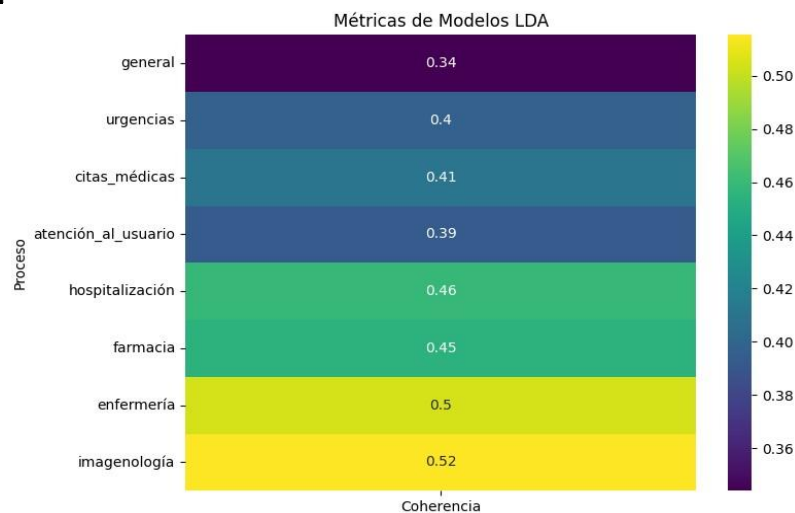


Figura 9: Métricas modelos LDA Fuente: elaboración propia

En la anterior gráfica heatmap, las métricas de coherencia con varios modelos LDA según los procesos en los cuales se mencionan: urgencias, citas médicas, hospitalización, farmacia, etc. Además, colores oscuros (menor coherencia), colores claros/amarillos (mayor coherencia) y los valores numéricos muestran la coherencia de cada modelo. En los modelos

LDA (Latent Dirichlet Allocation), mide que tan relacionados y comprensibles son los temas encontrados; los resultados obtenidos son los siguientes, los modelos con mayor coherencia son:

- imagenología: 0.52
- enfermería: 0.5
- hospitalización: 0.46
- farmacia: 0.45

De esta forma, estos procesos generan temas más consistentes por lo cual el modelo LDA funciona mejor en estas categorías. En este orden de ideas, se encuentran procesos con desempeños medio:

- citas médicas: 0.41
- urgencias: 0.40
- atención al usuario: 0.39

Los procesos encontrados presentan una coherencia moderada, puede existir una mezcla de vocabulario o menor claridad en la temática. Finalmente, el modelo LDA identifica mejor los temas en imagenología y enfermería, el área general presenta alguna dificultad para generar temas coherentes; así mismo, en términos generales los valores entre 0.4 y 0.5 se consideran aceptables en muchos análisis LDA.

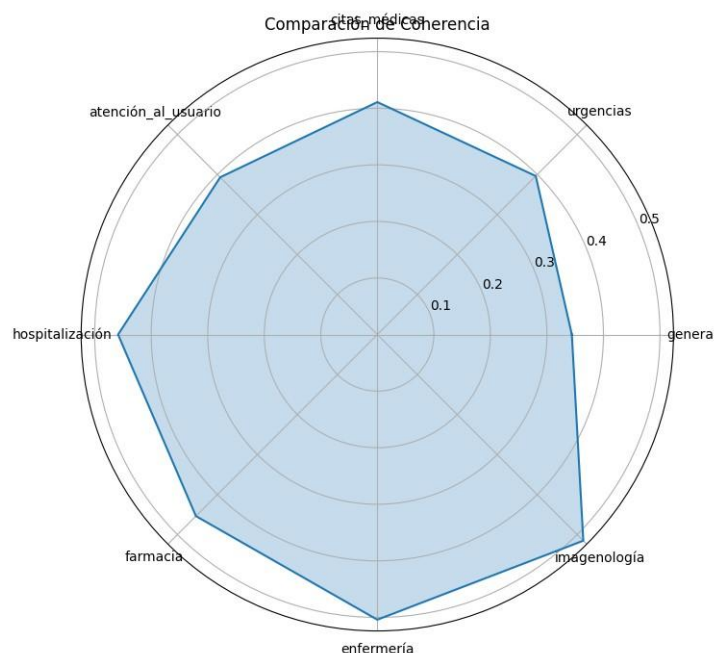


Figura 10: Comparación de coherencia. Fuente: elaboración propia

En la gráfica radar, se comparan diferentes áreas de atención en salud según un nivel de coherencia o desempeño; cada eje representa un área evaluada: urgencias, citas médicas, hospitalización y farmacia, etc. Los valores están aproximados de 0 a 0.5, mientras más lejos este el punto del centro mayor es la coherencia o desempeño en esa área. De esta forma, se puede observar en los resultados que el área de imagenología presenta el valor más alto (0.52), esto indica una mayor coherencia dentro del conjunto analizado; el área de enfermería también presenta un valor alto 0.50 la cual refleja una evidencia sólida y consistente. En conclusión, en la gráfica se puede observar áreas técnicas y asistenciales como imagenología y enfermería presentan mejor coherencia, mientras que el área general de algunos servicios presenta oportunidades de mejora; en conjunto, los resultados muestran un desempeño positivo con las diferentes áreas.

Naive bayes

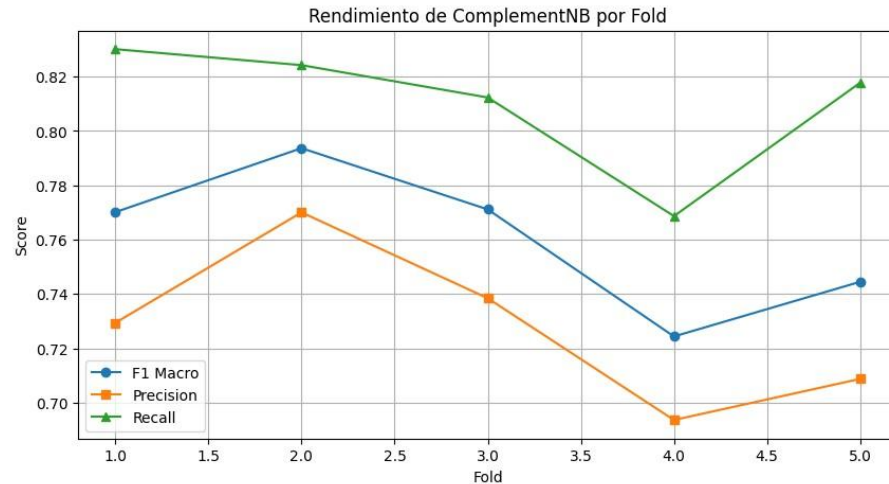


Figura 11: Rendimiento de complementoNB. Fuente: elaboración propia

La gráfica muestra el rendimiento del complemento del modeloNB con estructura de 5 particiones o folds de validación donde se evaluaron tres métricas: F1Macro, Precisión y Recall. En un contexto general, el modelo presenta un rendimiento estable y relativamente alto en los diferentes folds, la métrica con mejores resultados es de Recall, seguida de F1 Macro y luego precisión; además existe una evidente disminución en el folds 4 donde las métricas bajan. A continuación, se va a realizar un análisis en detalle de cada una de las métricas:

- recall: presenta todos los valores altos en los folds, comienza al valor cercano de 0.83 en el fold 1, baja gradualmente en el fold 4 (0.77) y luego recupera aumenta en el fold 5 (0.82).
- fl macro: oscila entre 0.72 y 0.79, su mejor desempeño está en el fold 2 (0.79) y su peor resultado es en el fold 4 (0.72).
- precisión: es la métrica más baja, su punto más alto esta en fold 2 (0.77), y su peor resultado se encuentra en fold 4 (0.69).

En términos generales, el modelo ComplementNB muestra un rendimiento consistente especialmente en el Recall lo que propone que es eficaz detectando la mayoría de las clases positivas.



Figura 12: Tf-Idf. Fuente: elaboración propia.

La gráfica que se presenta es una nube de palabras construida con el método TF-IDF (Term Frequency–Inverse Document Frequency), en la cual se identifica las palabras más importantes dentro de un conjunto de textos, probablemente comentarios o encuestas de

usuarios de un servicio de salud. De esta forma, las palabras con un tamaño grande son aquellas con mayor relevancia: personal, atención, servicio, inconformidad, cita, urgencias, paciente, grosero, demora y amable. En términos generales, los aspectos negativos predominantes se destaca palabras como: inconformidad, grosero, demora, espera, falta, perjudicial, mala; lo cual indica que muchos usuarios expresan insatisfacción relacionado con tiempos largo de espera, mala organización, deficiente trato personal y retrasos en citas y urgencias. Finalmente, la nube de palabras muestra que el tema dominante es la percepción del servicio de salud, términos asociados con inconformidad, demoras y mala atención.

Bayesian Deep Learning

Entrenamiento con 256 neuronas

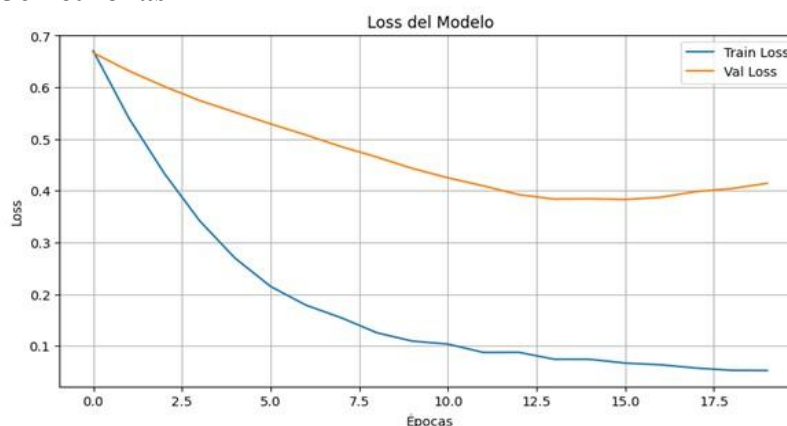


Figura 13: Loss del Modelo. Fuente: elaboración propia

En la anterior gráfica se muestra la evolución de la pérdida (Loss) del modelo durante el entrenamiento. Los Train Loss disminuyen continuamente pasando de 0.67 a 0.05. Por lo cual, el modelo aprende mejor con los datos de entrenamiento. De igual forma, el Val Loss disminuye de 0.67 a 0.38 entre las épocas 0 y 13. Esto significa que el modelo mejora su capacidad de generalización durante las primeras épocas. A partir de la época 13-15 el Val Loss se estabiliza y luego aumenta ligeramente. Este comportamiento es una señal de sobreajuste (overfitting). Finalmente, el mejor desempeño sobre los datos de validación se alcanza entre las épocas 13 y 15 donde el Val Loss es mínimo (0.38). De esta forma, el modelo continúa ajustándose a los datos de entrenamiento, pero pierde capacidad de generalización, lo que indica inicio de sobreajuste.

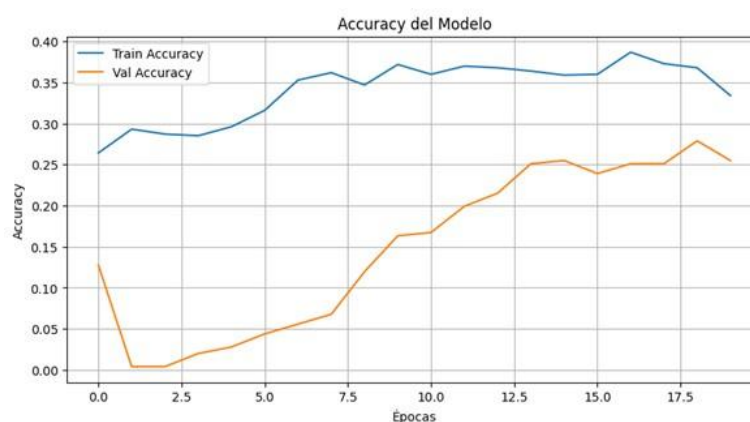


Figura 14: Accuracy del Modelo. Fuente: elaboración propia

La gráfica muestra la evolución de la precisión (accuracy) del modelo durante 20 épocas, comparando el desempeño en los datos de entrenamiento (Train Accuracy) y validación (Val

Accuracy). Durante el comportamiento del entrenamiento, la precisión aumenta de 0.27 (27%) al inicio hasta cerca de 0.38 (38%) en las últimas épocas, así mismo, después de la época 10 la precisión se estabiliza entre 0.36 y 0.39, mostrando una mejora lenta. Durante todo el proceso, la precisión de entrenamiento es superior a la de validación. La diferencia oscila entre 0.08 y 0.15 puntos, lo que sugiere que el modelo aprende mejor los datos de entrenamiento que los datos nuevos. Esto puede indicar un ligero sobreajuste (overfitting), aunque no es excesivamente pronunciado. En términos generales, el modelo muestra un aprendizaje progresivo ya que ambas curvas tienen una tendencia ascendente. Sin embargo, los niveles de precisión son relativamente bajos (menos del 40% en entrenamiento y menos del 30% en validación). Finalmente, la configuración de 256 neuronas presenta un F1 Macro promedio de 0.7260, lo que evidencia un desempeño bueno y equilibrado entre las clases, sugiriendo que el modelo tiene una capacidad adecuada de generalización y clasificación.

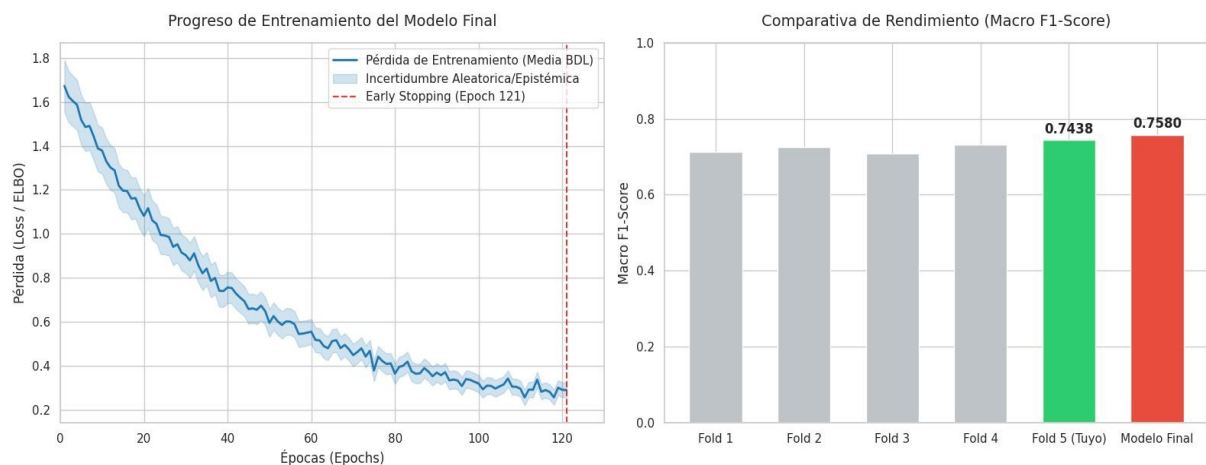


Figura 15: Entrenamiento del Modelo. Fuente: elaboración propia

Como se puede observar en la anterior gráfica, se presentan dos escenarios del modelo; primer escenario, el eje de las X se representa las épocas de entrenamiento (0 a 120) y en el eje de las Y se muestra la pérdida (Loss), una medida del error del modelo. así mismo, la línea azul indica que la pérdida disminuye progresivamente desde 1.75 hasta cerca de 0.28, lo que significa que el modelo va mejorando sus predicciones a medida que avanza el entrenamiento. La zona sombreada de azul corresponde a la incertidumbre del entrenamiento (desviación estándar) mostrando una mayor variabilidad estable; La línea roja punteada marca el punto de Early Stopping (Época 121), indicando que el entrenamiento se detuvo automáticamente donde se observaban mejoras significativas sin provocar un sobreajuste. Por lo cual, el modelo presenta una convergencia adecuada, ya que el error disminuye de forma continua y se estabiliza hacia las últimas épocas. Segundo escenario, Se comparo el desempeño obtenido en diferentes particiones de validación cruzada (folds) utilizando la métrica Macro F1-Score.

- Fold 1: $\approx 0,71$
- Fold 2: $\approx 0,72$
- Fold 3: $\approx 0,71$
- Fold 4: $\approx 0,73$
- Fold 5 (Yu): 0,7438
- Modelo Final: 0,7580

En términos generales, Los cinco folds muestran resultados bastante consistentes (entre 0.71 y 0.74), lo que indica que el modelo es estable; además el Fold 5 alcanza un F1-Score de 0.7438, superior al de los demás folds; así mismo, el modelo final obtiene el mejor resultado

(0.7580), superando ligeramente a todos los folds individuales los resultados de validación cruzada son consistentes.

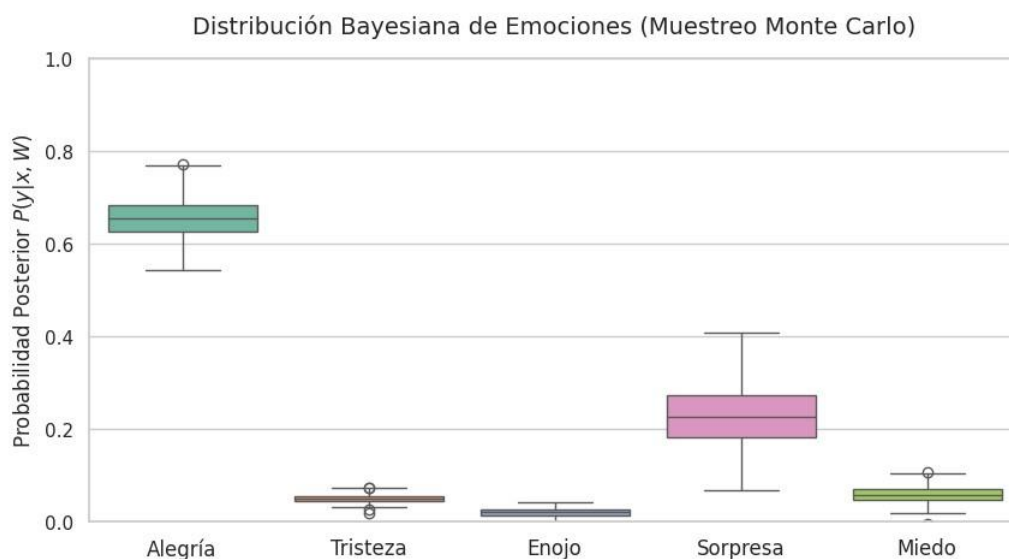


Figura 16: Distribución Bayesiana. Fuente: elaboración propia

Esta gráfica representa un diagrama de cajas de bigote con una distribución de las probabilidades $P(y|x, W)$ obtenidas a partir de un muestreo de Monte-Carlo Bayesiano para cinco emociones: Alegría, Tristeza, Enajo, Sorpresa y Miedo. En este orden de ideas, la alegría presenta la mayor probabilidad con una mediana cercana a 0.65, los valores se concentran aproximada entre 0.62 y 0.68 lo cual demuestra una poca variabilidad dando una alta confianza a esta emoción y finalmente el valor atípico cercano a 0,77 muestra simulaciones con probabilidad aún más alta. Respecto a las demás emociones, sorpresa aparece como una emoción secundaria con una mayor incertidumbre, las emociones tristeza, miedo y enajo presentan emociones muy bajas por lo cual el modelo considera poco probable que correspondan a la emoción observada.

Discusión

Los hallazgos obtenidos confirman que el análisis de sentimientos aplicado a las PQRS, combinado con modelos bayesianos, constituye una vía viable para transformar la retroalimentación de los usuarios en un insumo de gestión activo, superando el uso principalmente administrativo que históricamente se le ha dado a esta información (Van Dael et al. 2020). En primer lugar, la comparación entre los modelos evaluados muestra que no existe un método único óptimo, sino que cada técnica aporta una perspectiva complementaria: los Procesos Gaussianos, mediante los núcleos RBF y Matern, ofrecieron el mejor ajuste predictivo ($R^2 \approx 0,46$) para relaciones no lineales entre variables continuas, el clasificador Naive Bayes se mostró más adecuado para la detección consistente de clases positivas (recall entre 0,77 y 0,83), y el modelo Bayesian Deep Learning, con un F1-Score de 0,76, logró la mejor capacidad de generalización general, a costa de una mayor exigencia computacional y de señales de sobreajuste hacia las últimas épocas de entrenamiento. Esta complementariedad sugiere que la selección del modelo bayesiano debe orientarse por el tipo de decisión que se busca apoyar —predicción puntual, clasificación robusta o cuantificación de incertidumbre— más que por un único criterio de desempeño.

En segundo lugar, la distribución de emociones negativas identificada mediante Redes Bayesianas en procesos como vigilancia, consulta externa y radiología aporta evidencia concreta sobre en qué áreas del sistema hospitalario se concentra la insatisfacción de los usuarios, información más específica que los indicadores administrativos generales que suelen

manejar las entidades de salud. De manera similar, la mayor coherencia temática obtenida mediante LDA en los procesos de imagenología y enfermería (0,50-0,52) permite anticipar que estas áreas son susceptibles de un monitoreo automatizado más confiable que otras, donde la variabilidad del lenguaje de los usuarios limita la interpretabilidad de los temas extraídos. Estos resultados respaldan la utilidad de la Inteligencia Artificial para el Bien Social (IA4SG) como marco de referencia, en la medida en que traducen datos no estructurados en insumos accionables para la mejora continua de los servicios de salud.

En cuanto a la sostenibilidad técnica de la propuesta, la incorporación de una metodología MLOps en cuatro fases evidencia que el valor de estos modelos no reside únicamente en su exactitud, sino en la posibilidad de mantenerlos actualizados y reentrenados frente a nuevos flujos de PQRS, un aspecto habitualmente ausente en estudios centrados solo en el desempeño puntual de los algoritmos.

No obstante, los resultados deben interpretarse considerando varias limitaciones. Primero, el corpus proviene únicamente de entidades de salud de Boyacá durante 2024-2025, lo que introduce un posible sesgo regional y temporal que limita la generalización de los modelos a otros contextos hospitalarios o periodos. Segundo, el etiquetado de emociones se apoyó en una versión del NRC Emotion Lexicon adaptada al español, cuya traducción puede no capturar completamente matices idiomáticos o coloquiales propios de las PQRS analizadas. Tercero, el desempeño moderado observado en varias métricas (R^2 cercano a 0,46; coherencia LDA entre 0,39 y 0,52) indica que una parte relevante de la variabilidad de los datos permanece sin explicar, y los indicios de sobreajuste detectados en el modelo Bayesian Deep Learning sugieren la necesidad de conjuntos de entrenamiento más amplios o de estrategias de regularización adicionales. Estas limitaciones no invalidan los hallazgos, pero delimitan el alcance de las conclusiones y señalan direcciones concretas para validar el enfoque en otras regiones o instituciones de salud.

Conclusiones y Trabajos Futuros

En este estudio, se utilizó modelos Bayesianos en el análisis de sentimientos en Peticiones, Quejas, Reclamos y Sugerencias (PQRS) en las IA4SG en Boyacá. En un primer momento, los resultados obtenidos mediante Procesos Gaussianos, mostraron que los kernels RBF y Matern tienen mejor comportamiento predictivo, alcanzando los menores errores (MSE y MAE) y coeficientes de determinación cercanos a 0,46. Además, el modelo explica una proporción moderada de los datos donde demuestra que este enfoque probabilístico constituye una alternativa adecuada para modelar fenómenos asociados al comportamiento emocional presente en las PQRS.

Las Redes Bayesianas, permitieron modelar adecuadamente los datos y analizar la distribución probabilística de las emociones en los diferentes procesos hospitalarios; así mismo, se identificó que las emociones como tristeza, enfado, miedo y anticipación predominan en áreas críticas como vigilancia, consulta externa, radiología y servicios tercerizados. Finalmente, procesos como vacunación y atención en consultorios presentan mayores niveles de confianza y emociones positivas. El análisis temático mediante Latent Dirichlet Allocation (LDA) permitió identificar grupos de temas con niveles de coherencia aceptables en los procesos de imagenología, enfermería, hospitalización y farmacia. De esta forma, los resultados muestran que el modelado de temas constituye una herramienta para descubrir patrones ocultos dentro de grandes volúmenes de información textual, facilitando la identificación automática de problemas recurrentes en las PQRS. Por su parte, el modelo Naive Bayes presentó un desempeño destacado durante la validación cruzada representando por sus altos valores de recall y un comportamiento consistente entre las diferentes particiones del conjunto de datos. Además, el análisis mediante TF-IDF permitió identificar las principales preocupaciones de los

usuarios que están relacionadas con tiempos de espera, demoras, atención del personal y calidad del servicio, aspectos que representan oportunidades de mejora para las instituciones de salud.

Finalmente, el modelo de Bayesian Deep Learning mostro' un aprendizaje progresivo y un desempeño interesante alcanzando un valor de F1-Score cercano a 0,76 y resultados consistentes en las diferentes particiones de validación. De esta forma, se identificaron indicios de sobreajuste en las últimas 'épocas de entrenamiento, el uso de inferencia bayesiana permitió' incorporar estimaciones de incertidumbre que fortalecen la confiabilidad de las predicciones frente a nuevos datos.

Referencias bibliográficas

- Abdessalem, Anis Ben, Nikolaos Dervilis, David J Wagg, and Keith Worden. (2017). Automatic Kernel Selection for Gaussian Processes Regression with Approximate Bayesian Computation and Sequential Monte Carlo. *Frontiers in Built Environment* 3: 52. doi: 10.3389/fbuil.2017.00052
- Abdullah, Abdullah A, Masoud M Hassan, and Yaseen T Mustafa. (2022). A Review on Bayesian Deep Learning in Healthcare: Applications and Challenges. *IEEE Access* 10: 36538–62.
- Altarturi, Hamza HM, Muntadher Saadoon, and Nor Badrul Anuar. (2023). Web Content Topic Modeling Using LDA and HTML Tags. *PeerJ Computer Science* 9: e1459. DOI 10.7717/peerj-cs.1459
- Atienza, David, Concha Bielza, and Pedro Larrañaga. (2022). PyBNesian: An Extensible Python Package for Bayesian Networks. *Neurocomputing* 504: 204–9. <https://doi.org/10.1016/j.neucom.2022.06.112>
- Barrachina, Estela González, Christian Fortanet van Assendelft de Coningh, Tatiana Hidalgo-Marí, and Cristina González Díaz. (2026). Medición Del Sentimiento Hacia La Marca: Estudio de Caso de La Marca Vicio En Instagram. *Redmarka. Revista de Marketing Aplicado* 30 (1): 126–45. https://github.com/jboscomendoza/lexicos-nrc-afinn/blob/master/lexico_nrc.csv.
- Bird, Steven, Ewan Klein, and Edward Loper. (2009). *Natural Language Processing with Python: Analyzing Text with the Natural Language Toolkit*. " O'Reilly Media, Inc.
- Blei, David M. 2012. Probabilistic Topic Models. *Communications of the ACM* 55 (4): 77–84. [10.1145/2133806](https://doi.org/10.1145/2133806)
- Chahar, Ravita, and Ashutosh Kumar Dubey. (2026). Machine Intelligence for Mental Health Diagnosis: A Systematic Review of Methods, Algorithms, and Key Challenges. *Computers, Materials, & Continua* 86 (1): 1. <https://doi.org/10.32604/cmc.2025.066990>
- Contreras Contreras, Ghiordy Ferney, Byron Medina Delgado, Brayan René Acevedo Jaimes, and Dinael Guevara Ibarra. (2022). Metodología de Desarrollo de técnicas de Agrupamiento de Datos Usando Aprendizaje Automático. *Tecnura* 26 (72): 5–6. <https://doi.org/10.14483/22487638.17246>
- Demin, Alexander, and Denis Ponomaryov. (2022). Machine Learning with Probabilistic Law Discovery: A Concise Introduction. *arXiv Preprint arXiv:2212.11901*. <https://doi.org/10.48550/arXiv.2212.11901>
- Farkhod, Akhmedov, Akmalbek Abdusalomov, Fazliddin Makhmudov, and Young Im Cho. (2021). LDA-Based Topic Modeling Sentiment Analysis Using Topic/Document/Sentence (TDS) Model. *Applied Sciences* 11 (23): 11091. <https://doi.org/10.3390/app112311091>
- Fernández, Laura Rodríguez, Ana Laura Fernández Mena, María Patricia Torres Magaña, Manuel Antonio Rodríguez Magaña, and Manuel Antonio Rodríguez Fernández. (2024).

- Inteligencia Artificial En La Educación: Modelo de Lenguaje de Gran Tamaño (LLM) Como Recurso Educativo. *Revista Ipsilon* 7 (2): 157–64.
- Jankova, Zuzana. (2023). Latent Dirichlet Allocation (LDA) Approximation Analysis of Financial-Related Text Messages. *Economic Computation & Economic Cybernetics Studies & Research* 57 (1).
- Jospin, Laurent Valentin, Hamid Laga, Farid Boussaid, Wray Buntine, and Mohammed Bennamoun. (2022). Hands-on Bayesian Neural Networks—a Tutorial for Deep Learning Users. *IEEE Computational Intelligence Magazine* 17 (2): 29–48. <https://doi.org/10.1109/MCI.2022.3155327>
- Jurafsky, Daniel, and James H. Martin. (2026). *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition, with Language Models*. 3rd ed. <https://web.stanford.edu/~jurafsky/slp3/>.
- Murphy, Kevin P. (2023). *Probabilistic Machine Learning: Advanced Topics*. MIT press.
- Patil, Rajvardhan, and Venkat Gudivada. (2024). A Review of Current Trends, Techniques, and Challenges in Large Language Models (Llms). *Applied Sciences* 14 (5): 2074. <https://doi.org/10.3390/app14052074>
- Riutort-Mayol, Gabriel, Paul-Christian Bürkner, Michael R Andersen, Arno Solin, and Aki Vehtari. (2023). Practical Hilbert Space Approximate Bayesian Gaussian Processes for Probabilistic Programming.” *Statistics and Computing* 33 (1): 17.
- SACRISTÁN, MILLÁN SANTAMARÍA. (n.d). Machine Learning Operations (MLOps): Contexto Actual y Tendencias Futuras. PhD thesis, Universidad de La Rioja.
- Silander, Tomi, and Petri Myllymaki. (2012). A Simple Approach for Finding the Globally Optimal Bayesian Network Structure. *arXiv Preprint arXiv:1206.6875*. <https://scifaro.com/en/abs/a-simple-approach-for-finding-the-globally-optimal-bayesian-network-structure-1206.6875>
- Van Dael, Jackie, Tom W Reader, Alex Gillespie, Ana Luisa Neves, Ara Darzi, and Erik K Mayer. (2020). Learning from Complaints in Healthcare: A Realist Review of Academic Literature, Policy Evidence and Front-Line Insights. *BMJ Quality & Safety* 29 (8): 684–95. doi: 10.1136/bmjqs-2019-009704. Epub 2020 Feb 4. PMID: 32019824; PMCID: PMC7398301.
- Wang, Hao, and Dit-Yan Yeung. (2016). Towards Bayesian Deep Learning: A Framework and Some Existing Methods. *IEEE Transactions on Knowledge and Data Engineering* 28 (12): 3395–408. https://www.researchgate.net/publication/330053760_Constructing_Site-Specific_Multivariate_Probability_Distribution_Model_Using_Bayesian_Machine_Learning
- Younas, Ahtisham, Sergi Fàbregues, Sarah Munce, and John W Creswell. (2025). Framework for Types of Meta-inferences in Mixed Methods Research. *BMC Medical Research Methodology* 25 (1): 18. doi: 10.1186/s12874-025-02475-8. PMID: 39856568; PMCID: PMC11758751.
- Zapata-Ros, Miguel. (n.d). *Métodos Bayesianos Para Detectar Artículos Científicos Creados Con IA*.
- Zou, Yixin, Ding-Bang Luh, and Shizhu Lu. (2022). Public Perceptions of Digital Fashion: An Analysis of Sentiment and Latent Dirichlet Allocation Topic Modeling. *Frontiers in Psychology* 13: 986838. doi: 10.3389/fpsyg.2022.986838