

Aplicación de la minería de datos en el marketing usando el análisis de sentimientos de los clientes e-commerce

Application of data mining in marketing using e-commerce customer sentiment analysis

Recibido: junio 06 de 2024 | Revisado: junio 14 de 2024 | Aceptado: junio 22 de 2024

IVAN PETRLIK AZABACHE¹
JOSÉ COVEÑAS LALUPU¹
WILFREDO CARRANZA BARRENA¹
LUZ TORRES-TALAVERANO²

RESUMEN

En este estudio se desarrolló y aplicó el análisis de sentimientos respecto a los productos tecnológicos en la red social Twitter/X, asimismo, se determinó las opiniones expresadas por los clientes donde finalmente se identificó el modelo predictivo más conveniente derivado del Machine Learning. Para ello, se recolectaron 7102 tweets relacionados a los productos de Apple y Samsung, empleando la metodología propuesta por Erl, Khattak y Buhler la cual facilitó la implementación de sus fases críticas. Los resultados obtenidos del análisis de sentimientos se evaluaron mediante métricas estándar como Accuracy, Precision, Recall y F1-Score, aplicadas a cuatro modelos de aprendizaje automático: K-Nearest Neighbors (KNN), Logistic Regression (LR), Random Forest (RF) y CatBoost Classifier (CC). De estos, el CatBoost Classifier demostró ser el más efectivo, logrando un 89% en Accuracy, 90% en Precision, 89% en Recall y 88% en F1-Score. Se concluyó que el modelo CatBoost Classifier fue el un modelo óptimo para analizar los sentimientos en Twitter/X, debido a su capacidad de proporcionar insights valiosos sobre la percepción de los productos tecnológicos promocionados permitiendo la eficacia en las campañas de marketing digital.

Palabras clave: Minería de datos, análisis de sentimientos, aprendizaje automático, e-commerce

ABSTRACT

In this study, sentiment analysis was developed and applied to technological products in the Twitter/X social network, also, the opinions expressed by customers were determined and finally the most suitable predictive model derived from Machine Learning was identified. For this purpose, 7102 tweets related to Apple and Samsung products were collected, using the methodology proposed by Erl, Khattak and Buhler which facilitated the implementation of its critical phases. The results obtained from sentiment analysis were evaluated using standard metrics such as Accuracy, Precision, Recall and F1-Score, applied to four machine learning models: K-Nearest Neighbors (KNN), Logistic Regression (LR), Random Forest (RF) and CatBoost Classifier (CC). Of these, the CatBoost Classifier proved to be the most effective, achieving 89% in Accuracy, 90% in Precision, 89% in Recall and 88% in F1-Score. It was concluded that the CatBoost Classifier model was the optimal model for analyzing sentiment on Twitter/X, due to its ability to provide valuable insights into the perception of promoted technology products enabling effectiveness in digital marketing campaigns.

- 1 Universidad Nacional Federico Villarreal, Lima, Perú
- 2 Universidad Nacional Mayor de San Marcos, Lima, Perú

Autor de correspondencia:

Keywords: Data mining, sentiment analysis, machine learning, e-commerce

© Los autores. Este artículo es publicado por la Revista Campus de la Facultad de Ingeniería y Arquitectura de la Universidad de San Martín de Porres. Este artículo se distribuye en los términos de la Licencia Creative Commons Atribución No-Comercial – Compartir-Igual 4.0 Internacional (<https://creativecommons.org/licenses/by/4.0/>), que permite el uso no comercial, distribución y reproducción en cualquier medio siempre que la obra original sea debidamente citada. Para uso comercial contactar a: revistacampus@usmp.pe.

<https://>

Introducción

En los últimos años, se ha observado un aumento significativo en el uso de la Inteligencia Artificial que es una ciencia moderna que pretende crear máquinas que imitan la inteligencia humana en diferentes sectores industriales y de servicios, marcando un hito importante en esta disciplina y sus aplicaciones (Sadeq et al., 2023). Este crecimiento ha dado lugar a desarrollos espectaculares que abarcan cada vez más aspectos de la vida diaria. Entre los sectores transformados por este avance masivo de la Inteligencia Artificial, destaca la evolución en el ámbito del consumo masivo y retail, particularmente en actividades relacionadas con el marketing de productos y publicidad personalizada. Según Medina-Chicaiza y Martínez-Ortega (2020), el marketing se beneficia de las tecnologías de la Inteligencia Artificial, especialmente al utilizar el análisis de mercado y los datos provenientes de la investigación de las necesidades y experiencias del consumidor. La automatización del marketing permite a las empresas mejorar la interacción con su mercado objetivo y aumentar la eficiencia. Es relevante destacar que, según estos autores, la tecnología más utilizada es el aprendizaje automático, que, como sostiene Deisenroth et al. (2020), representa la última innovación en una larga serie de esfuerzos destinados a sintetizar conocimiento y pensamiento lógico en un formato adecuado para la creación de sistemas y máquinas automatizadas. Adicionalmente, Mar et al. (2023) indican que los consumidores frecuentemente evalúan productos y servicios basándose en las opiniones de otros consumidores del mismo producto.

Las nuevas tendencias de marketing para las empresas enfatizan el uso de

las redes sociales como un canal clave para comunicar acciones corporativas, principalmente debido a su enorme potencial para establecer relaciones con los clientes, según Ramona et al. (2018). Esta tendencia se debe a la facilidad con la que se pueden entender y desarrollar estrategias de marketing, así como a la gran disponibilidad de datos. A lo largo de la historia de los negocios, las empresas de productos de consumo masivo han enfrentado constantemente el desafío de comprender las necesidades y opiniones de sus consumidores. La actividad de escuchar se ha convertido en una parte fundamental del proceso de creación de valor de estas compañías. Con el avance de Internet y los medios sociales, el e-marketing ha ganado relevancia, desplazando a las herramientas de marketing tradicionales para un público más alejado del mundo digital, como destacan Marín y Botey (2022). Es importante mencionar que, anteriormente, esta actividad se realizaba de forma manual, utilizando herramientas como encuestas o estudios de opinión. Sin embargo, a pesar de que estas alternativas son consideradas clásicas en los estudios de público y han sido fundamentales para el marketing tradicional, presentan un riesgo elevado de sesgos y pueden perder de vista la verdadera opinión del público. Generalmente, estos métodos entregan resultados con un desfase temporal y a un costo alto, según Mendivelso y Lobos (2019).

Según Taherdoost y Madanchian (2023), diversas empresas invierten considerablemente en investigaciones de mercado para comprender las preferencias y deseos de sus clientes. Además, se predice que la inversión en análisis de datos aumentará un 60% en 2023. La función del marketing siempre ha sido captar la

atención de los clientes y mantener a raya a la competencia. Para ello, las empresas han estado continuamente rediseñando sus estrategias e introduciendo nuevas técnicas para mejorar la atención al cliente. El uso generalizado de las redes sociales ha permitido que un gran número de usuarios de todo el mundo comparta información, como creencias sobre temas y opiniones de productos y servicios, lo que representa una oportunidad importante para las organizaciones de obtener información valiosa sobre el comportamiento del cliente, según Nagarkar et al. (2020). De acuerdo con Herrera-Contreras et al. (2021), Twitter/X se destaca como una de las plataformas más utilizadas globalmente, generando alrededor de 250 millones de tweets diarios, lo que la convierte en el medio idóneo para la aplicación de la minería de opiniones y el análisis de sentimiento.

La presente investigación tiene como objetivo general: desarrollar y aplicar el Análisis de Sentimientos sobre productos tecnológicos en la red social Twitter/X. Los objetivos específicos son: determinar las opiniones expresadas por los clientes e identificar el modelo predictivo más conveniente derivado del Machine Learning.

La investigación aborda la necesidad de obtener información rápida y precisa sobre los resultados de campañas publicitarias recién lanzadas, con el objetivo de evaluar la efectividad de las estrategias de marketing implementadas por empresas que promocionan productos o servicios. Se centra especialmente en explorar las opiniones y reacciones de los usuarios en la red social Twitter/X, empleando tecnologías de Machine Learning y modelos predictivos. Este

enfoque permite analizar las preferencias y críticas de los usuarios hacia los productos o servicios promocionados. La situación planteada brinda la oportunidad de aplicar tecnología avanzada en el campo de la Inteligencia Artificial y herramientas de gestión de contenidos para realizar estudios semánticos y neurolingüísticos, profundizando en el análisis de opiniones y sentimientos de los consumidores.

A continuación, se presentan los antecedentes de la respectiva investigación basado en el análisis de sentimientos de clientes con respecto a diferentes productos, servicios o situaciones, utilizando los modelos de Machine Learning:

En el 2023, se logró comparar diversos modelos supervisados de aprendizaje automático para evaluar cómo los consumidores expresan sus sentimientos. Para ello, se empleó una metodología basada en Inteligencia Artificial que incorporó técnicas de aprendizaje automático, procesamiento del lenguaje natural y análisis de sentimientos. Los datos analizados consistían en comentarios de clientes sobre productos de una tienda de ropa para mujeres. Los resultados revelaron que los modelos lograron una precisión en el rango del 80% al 90%. Específicamente, se encontró que el modelo de regresión logística fue el más preciso, obteniendo una exactitud del 90% (Loukili et al., 2023). De igual forma, Sangeetha y Kumaran (2023) presentaron como solución al análisis de sentimientos aplicado a los consumidores, tomando en cuenta sus opiniones con respecto a los productos o servicios con el propósito de identificar patrones de comportamientos a través de la aplicación de un modelo RNN-LSTM

en la clasificación según su polaridad adecuada. El método propuesto se ha evaluado utilizando los datos recopilados de las reseñas de productos en línea de Amazon.com. El software MATLAB se utilizó para el trabajo propuesto. A partir de los resultados experimentales, el PCCHH-RNNLSTM propuesto demostró mejoras considerables de 95,8% de exactitud, 95,4% de precisión, 95,6% de recuperación y 95,2% de medida F, respectivamente.

En 2022, Sinnasamy y Sjaif (2022) realizaron un experimento que abordó la descripción de productos en Amazon, combinado con el análisis de sentimientos para comprender las preferencias expresadas en las opiniones, con el objetivo de evaluar los servicios, la información y la calidad de los productos. La metodología empleada incluyó un flujo de procesos que comenzó con la recopilación de datos, seguido de la extracción de características, la clasificación de sentimientos y la evaluación de métricas. Los resultados se obtuvieron mediante cuatro métodos de clasificación, utilizando N-Grams para calcular métricas como accuracy, precisión, recall y F1-Score. El método TF-IDF con N-Gram reveló que el modelo SVM alcanzó un accuracy del 82.27%, una precisión del 82%, un recall del 80% y un F1-Score del 72%. Por otra parte, Alroobaea (2022) se enfocó en analizar las opiniones sobre productos de Amazon desde diversos ámbitos empresariales y sociales, aplicando el análisis de sentimientos. El objetivo del estudio fue predecir los sentimientos de los autores utilizando un modelo de red neuronal recurrente (RNN) para clasificar tres conjuntos de datos de reseñas de Amazon (data 1, data 2 y data 3). La metodología incluyó las fases

de adquisición de los datos de Amazon, preprocesamiento, inclusión de palabras y clasificación de las reseñas en positivas o negativas, empleando un modelo RNN propuesto y comparándolo con otros modelos. Los resultados mostraron que el RNN tuvo un mejor accuracy comparado con otros modelos, alcanzando un 85% en data 1, y 70% en data 2 y data 3.

Seguidamente, Zhao et al. (2021) identificaron dificultades en prever con precisión el grado de polaridad de las opiniones de los usuarios en grandes cantidades de compras en línea a través de plataformas de comercio electrónico, atribuidas a cambios en la longitud de la secuencia, el orden del texto y la compleja lógica involucrada. El objetivo de su estudio fue desarrollar un nuevo algoritmo optimizado de aprendizaje automático para el análisis de sentimientos en reseñas de productos en línea. La metodología implementada incluyó la recopilación de datos, el preprocesamiento y la extracción de características. Los resultados destacaron que el algoritmo LSIBA-ENN superó a los algoritmos de primera categoría en la clasificación de sentimientos. Además, se observó que la ENN predominante, utilizando el esquema LTF-MICF propuesto, alcanzó una recuperación del 87.79%, superando los resultados obtenidos por la ENN con otros esquemas como W2V, TF, TF-IDF y TF-DFS, los cuales lograron 83.55, 84.03, 85.48 y 86.04 puntos respectivamente.

Posteriormente, Bansal y Srivastava (2018) evidenciaron un avance en el comercio electrónico al enfatizar la importancia de las opiniones de los consumidores. Estas opiniones fueron analizadas y clasificadas mediante la extracción de datos y el análisis de

sentimientos utilizando técnicas de Machine Learning. El objetivo principal fue transformar las opiniones en representaciones vectoriales usando el modelo Word2Vec para su clasificación. La metodología empleada incluyó la descripción de un conjunto de datos de 400,000 registros de opiniones sobre teléfonos móviles. Además, el proceso abarcó etapas como el preprocesamiento, la identificación de características similares con Word2Vec, la clasificación de sentimientos y la evaluación de métricas. Los resultados mostraron que el modelo Word2Vec, aplicado con los métodos CBOW y Skip-Gram junto con algoritmos de Machine Learning y una validación cruzada de 10 veces, logró clasificar con éxito las reseñas de los clientes, alcanzando una precisión

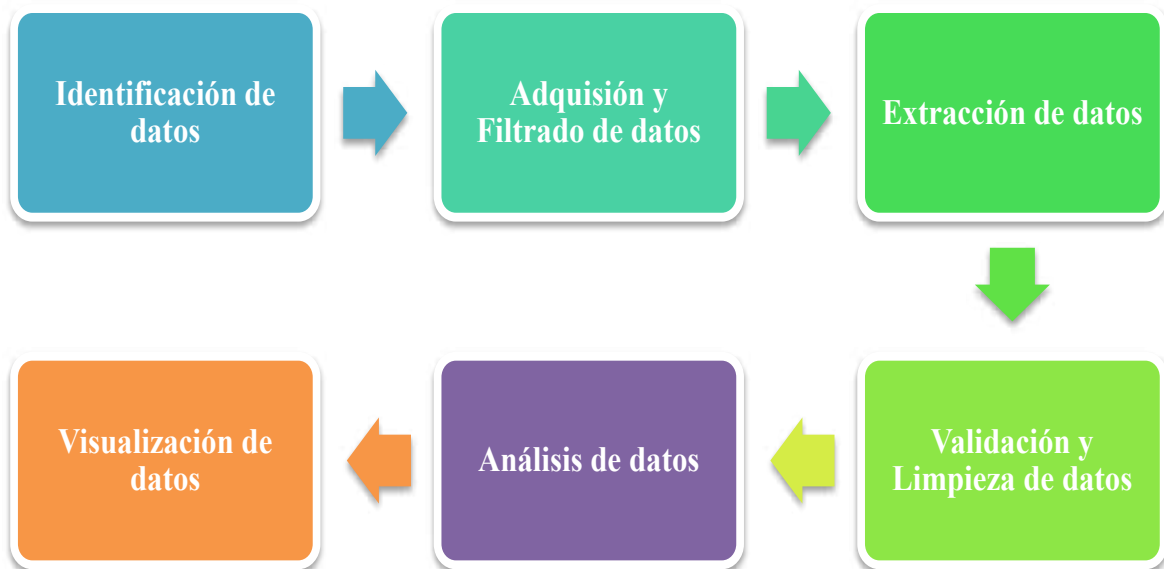
superior al 90.66% con el algoritmo Random Forest y el método CBOW.

Método

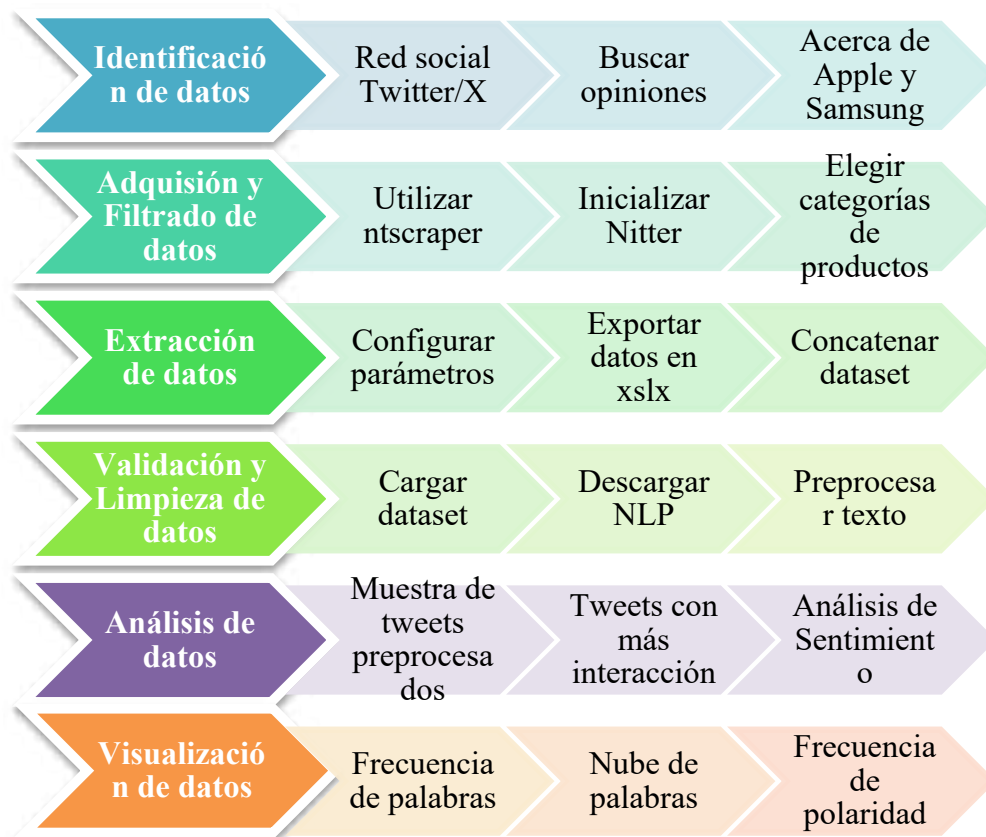
La metodología desarrollada para la minería de datos está basada en la propuesta de Erl et al. (2016), la cual aborda características propias del Big Data como grandes volúmenes de datos, velocidad y variedad de datos. Sin embargo, en el contexto de esta investigación, solo se emplearán seis fases. Esta metodología simplificada es ampliamente utilizada en la minería de datos (Plaza et al., 2022), y se ajusta de manera óptima al ciclo de vida de analítica de datos. La metodología del presente estudio está adecuada según los requerimientos de la investigación, las cuales se presentan en la Figura 1.

Figura 1

Fases de la metodología de investigación



Asimismo, en la Figura 2 se muestra el detalle de cada una de las fases y sus actividades respectivas.

Figura 2*Flujo metodológico de la investigación*

A continuación, se detalla cada una de las fases de la respectiva metodología.

A. Identificación de datos

En esta primera fase, se realizó una exploración del conjunto de datos. Yannis et al. (2018) argumentan que este respectivo flujo de trabajo permite identificar una mayor variedad de fuentes de datos, en las cuales se pueden encontrar patrones y correlaciones ocultas. Además, en este proceso se identificaron datos de productos electrónicos en la red social Twitter/X que, según Kamath (2024) es un servicio de microblogging y redes sociales ofrecido por Twitter Inc, expresados en tweets de opiniones de productos de aparatos tecnológicos de las marcas Apple y Samsung.

B. Adquisición y filtrado de datos

En la fase de Adquisición y Filtrado de Datos, Erl et al. (2016) sostienen que se realiza una meticulosa recolección de información proveniente de todas las fuentes identificadas previamente. Por ende, las fuentes identificadas son las instancias de Nitter que son variantes de Twitter/X que proporcionan una interfaz amigable en el proceso de extracción a través de la biblioteca NtScraper, la cual es una herramienta especializada para este fin. Seguidamente, se realizó un filtrado de datos que un proceso mediante el cual se selecciona una fracción del conjunto total de datos para su visualización o análisis (Capurso, 2021). Esta fracción representa las categorías específicas de productos, de la marca Apple, las

cuales se encuentran sus categorías Mac, iPad, iPhone, Tablet, Smart Watch y Accesorios; y de Samsung; Smartphone, Tablet, Smart Watch y Accesorios, debido a la destacada presencia de tweets en la

red social Twitter/X relacionadas con estas mismas. A continuación, en la Tabla 1, se muestra cada una de las actividades de esta fase.

Tabla 1

Actividades de la fase de adquisición y filtrado de datos

Actividad	Descripción
Instalación <i>Ntscraper</i>	La herramienta <i>ntscraper</i> se incorpora al entorno de Google Colaboratory mediante el comando <i>'pip'</i> , asegurando una instalación directa y eficiente desde el repositorio de PyPI (Python Package Index).
Importación Librerías	Se introducen dos bibliotecas fundamentales: <i>Nitter</i> para la interacción con instancias de <i>Nitter</i> para la obtención de tweets, y <i>pandas</i> , reconocida por su versatilidad, para la manipulación y análisis de datos.
Inicialización Nitter	Se crea una instancia de la clase <i>Nitter</i> , estableciendo parámetros clave. En este contexto, <i>log_level=1</i> configura el nivel de registro para proporcionar información durante la ejecución del código, y <i>skip_intance_check=False</i> asegura que no se omita la verificación de instancias de Nitter.

C. Extracción de datos

La respectiva fase es un proceso de extracción de datos. Está diseñada para recolectar información de manera eficiente a partir de diversas fuentes externas (Han et al., 2022). Su enfoque principal es recolectar datos de diversas fuentes y convertirlos en un formato

compatible con la solución de análisis de datos que se utiliza (Erl et al., 2016). Por consiguiente, en la Tabla 2, se especifica la fase de Extracción de Datos, que consiste en la configuración de parámetros de la librería *ntscraper*, la exportación de datos en formato *xlsx* y la concatenación del dataset.

Tabla 2

Actividades de la fase de extracción de datos

Actividad	Descripción
Configurar parámetros de la librería <i>ntscraper</i>	Se definen cuidadosamente los parámetros de la librería <i>ntscraper</i> , estableciendo límites de 6000 tweets, en idioma inglés y un rango de fechas desde enero de 2018 hasta diciembre de 2023. La información clave extraída abarca el enlace de origen, el texto de tweet, la fecha de publicación, así como el recuento de likes, comentarios y retweets. Este proceso se replica para cada categoría de productos de ambas empresas.
Exportar datos	Se elige el formato <i>xlsx</i> para la exportación del datasets, fundamentado en la presencia de caracteres especiales y emojis en algunos tweets.
Concatenar dataset	Se consolida un dataset combinado que abarca los tweets relacionados con Apple y Samsung. Este conjunto de datos global desempeña un papel fundamental en el proceso de validación y limpieza de datos, garantizando la integridad y calidad del análisis.

Seguidamente en la Tabla 3, se muestra el número de tweets que se extrajeron en base a cada categoría de productos relacionados a Apple y Samsung.

Tabla 3
Cantidad de tweets por cada dataset

Categoría de producto	Tamaño del dataset
Apple Mac	5798
Apple iPad	899
Apple iPhone	53
Smart Watch Apple	4
Tablet Apple	6
Apple Accesories	57
Samsung Smartphone	208
Samsung Tablet	55
Smart Watch Samsung	13
Samsung Accesories	9

Por lo tanto, se deduce que esta fase asegura la disponibilidad de datasets completos y listos para realizar análisis exhaustivos sobre la percepción del usuario en las redes sociales en relación con los productos de Apple y Samsung.

Tabla 4
Actividades de la fase de validación y limpieza de datos

Actividad	Descripción
Cargar dataset global	La importación de datos se realiza desde un archivo Excel mediante la utilización de la biblioteca <i>pandas</i> . Se realiza una exploración inicial del dataset y se identifica la presencia de posibles valores nulos.
Descargar Recursos NLP	Se procede a la instalación de la biblioteca <i>emoji</i> para la gestión eficaz de emojis en el texto. Asimismo, se importan librerías fundamentales para el procesamiento de texto, y se descargan recursos adicionales para el Procesamiento de Lenguaje Natural (NLP) a través de la biblioteca <i>nlk</i> .
Realizar el preprocesamiento de texto	Se define una función específica diseñada para llevar a cabo la limpieza y preprocesamiento en el texto de los tweets. Este proceso incluye desde la eliminación de enlaces y caracteres especiales hasta la conversión de emojis a descripciones de sentimientos. Se aplica la tokenización y eliminación de <i>stopwords</i> en inglés para contribuir en la estructura del texto de manera más coherente.

D. Validación y limpieza de datos

Según Erl et al. (2016), en esta fase se implementan reglas de validación, a menudo complejas, con el objetivo de detectar y eliminar aquellos datos que no se ajustan a los estándares definidos. Además, Naga et al. (2023) afirman que la limpieza de los datos abarca diversas actividades, entre las que se incluyen la gestión de información faltante y la eliminación de datos irrelevantes o atípicos. Específicamente en la investigación, se implementaron tres procedimientos que involucran la carga del dataset a nivel global, la descarga de recursos NLP y, especialmente, el preprocesamiento de texto, este último considerado como el más importante; de acuerdo con Haddi et al. (2013), es esencial el análisis de sentimientos para la clasificación de textos dado que los textos suelen contener ruido, etiquetas, guiones y anuncios, además permite reducir dicho ruido, lo que a su vez mejora el rendimiento y la precisión de la clasificación. Se detalla cada una de las actividades en la Tabla 4 para facilitar un mejor entendimiento.

A continuación, en la Tabla 5, se presenta la descripción detallada del conjunto de datos global 'tweets_ecommerce_products.xlsx' proporcionando una visión concreta de los datos que se utilizan en la investigación.

Tabla 5

Descripción del conjunto de datos global

Descripción	Valor
Rango en años de los datos	2018 – 2023
Nombre del archivo	tweets_ecommerce_products
Formato de archivo	xlsx
Número total de tweets	7102
Campos del dataset	twitter_link, text, date, likes, comments, retweets
Tipo de datos	int64 y object
Cantidad de valores nulos	0
Tamaño de filas y columnas	7102 filas y 6 columnas
Tamaño del archivo	736 KB (754,601 bytes)

E. Análisis de datos

La etapa de análisis de datos implica aplicar diversas técnicas de evaluación, a menudo de manera iterativa, especialmente en análisis exploratorios, hasta hallar patrones o correlaciones claves (Erl et al., 2016). Considerando estos principios conceptuales en esta fase, se

realizaron tres actividades, que involucran la carga del dataset global, la descarga de recursos NLP, y, especialmente, el preprocesamiento de texto, este último considerado como el más importante. Se detalla cada una de las actividades en la Tabla 6 para facilitar un mejor entendimiento.

Tabla 6

Actividades de la fase de análisis de datos

Actividad	Descripción
Explorar de forma aleatoria los tweets preprocesados	Se opta por seleccionar una muestra aleatoria de 10 tweets preprocesados, proporcionando así una visión representativa del conjunto de datos.
Explorar tweets con mayor interacción	Se procede a identificar los tweets que han generado la mayor cantidad de comentarios, likes y retweets en el conjunto de datos. Utilizando la biblioteca <i>tabulate</i> , se presenta esta información en una tabla que incluye el tipo de interacción, la cantidad correspondiente y el texto preprocesado del tweet.
Aplicar el Análisis de Sentimiento	En este procedimiento se emplea la biblioteca <i>TextBlob</i> para realizar el Análisis de Sentimiento en el conjunto de tweets preprocesados. Para ello, se crea la función <i>analyse_sentiment</i> la cual asigna una etiqueta de sentimiento (positivo, negativo o neutro) y un puntaje de sentimiento a cada tweet.

F. Visualización de datos

Esta fase abarca diversas representaciones gráficas para un adecuado análisis de los tweets preprocesados. De acuerdo con la Tabla 7 se especifica la creación de gráficas de tipo barras y nube

de palabras con el objetivo de explorar la frecuencia de palabras en los tweets, identificar patrones específicos de palabras para diferentes tipos de sentimientos y analizar la distribución de los puntajes de sentimiento en el conjunto de datos.

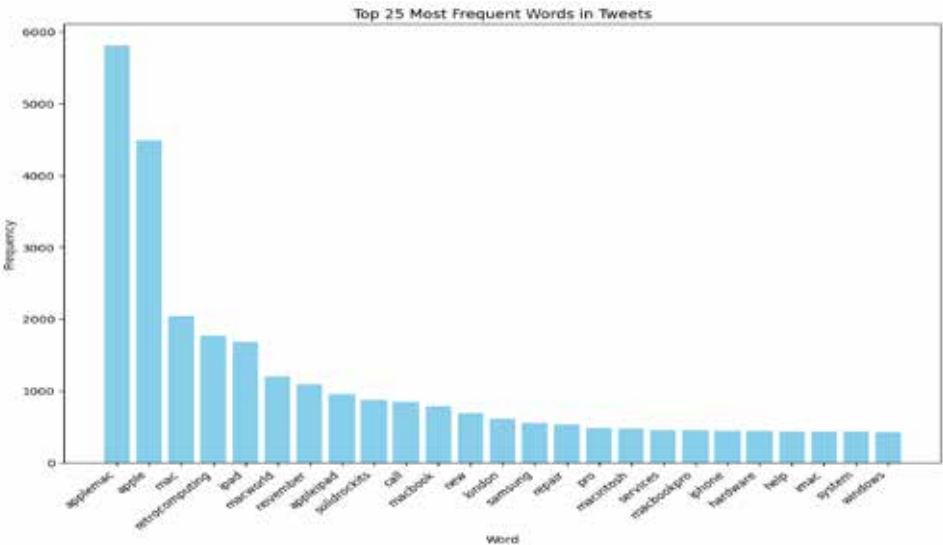
Tabla 7
Descripción de la fase de visualización de datos

Gráfico	Descripción
Gráfico de barras según la Frecuencia de Palabras	Se realiza la tokenización del texto en cada tweet con el fin de calcular la frecuencia de cada palabra y se visualizan las 25 palabras más frecuentes en una gráfica de barras.
Gráfico de barras según la Frecuencia de Palabras para el Tipo de Sentimiento	Se define la función <i>calculate_word_frequency</i> que calcula la frecuencia de palabras para un tipo de sentimiento específico. Para ello, se generan gráficas de barras separadas para palabras frecuentes en tweets positivos, negativos y neutros.
Nube de palabras según la Frecuencia de Palabras	Se utiliza la librería <i>WordCloud</i> para generar una nube de palabras que representa la frecuencia de palabras en todos los tweets preprocesados.
Nubes de palabras según la Frecuencia de palabras para el Tipo de Sentimiento	Se define la función <i>generate_wordcloud</i> , la cual crea nubes de palabras para cada tipo de sentimiento (positivo, negativo, neutral).
Gráfico de barras para identificar la Frecuencia de Polaridad	Se crea un histograma que muestra la frecuencia de los puntajes de sentimiento (Polaridad).

En la Figura 3, se visualiza el gráfico de barras que destaca la Frecuencia de

Palabras en los tweets preprocesados del conjunto de datos.

Figura 3
Gráfico de barras según la Frecuencia de Palabras

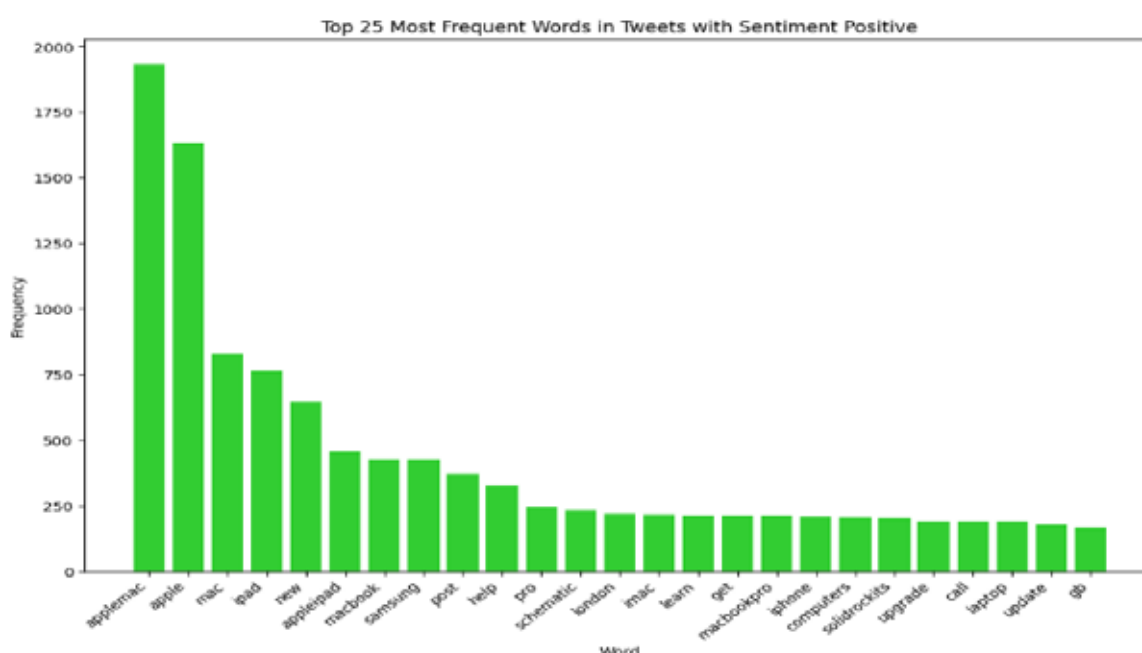


De acuerdo con la Figura 3, se observa un claro dominio de temas relacionados con productos de Apple, evidenciado por términos como *'applemac'*, *'apple'*, *'mac'*, *'ipad'*, *'macbook'* e *'iphone'*. Además, se identifican palabras que indican eventos o tendencias específicas, como *'retrocomputing'*, *'november'* y *'london'*. Por otra parte, la presencia de términos como *'repair'*, *'services'* y *'help'* sugiere

discusiones sobre servicios de reparación o asistencia técnica, proporcionando posibles problemas o consultas de los usuarios. Asimismo, la inclusión de *'samsung'* señala la presencia de productos de la competencia, permitiendo una visión comparativa entre Apple y Samsung. En la Figura 15, se observa el gráfico de barras según la Frecuencia de Palabras para el Tipo de Sentimiento *'positive'*.

Figura 4

Gráfico de barras según la Frecuencia de Palabras para el Tipo de Sentimiento 'positive'



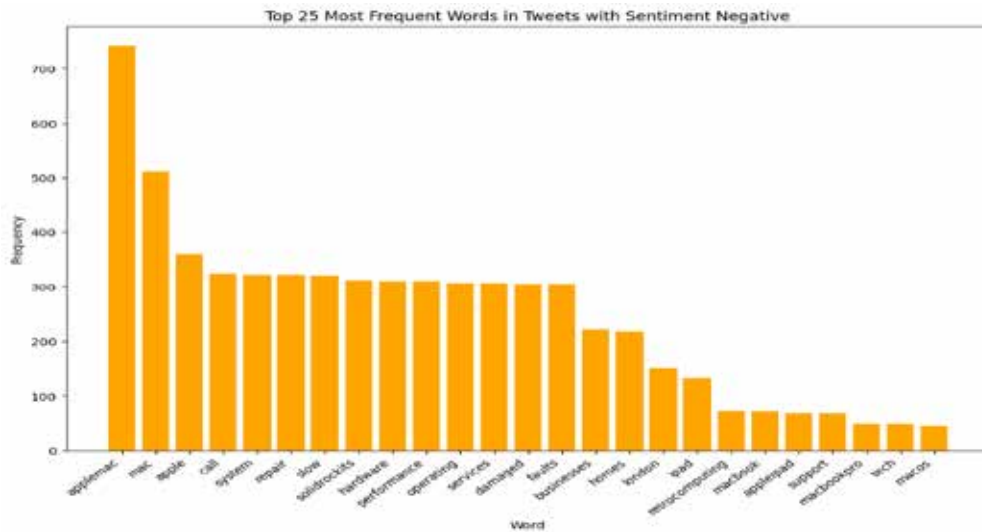
De acuerdo con la Figura 4, las palabras clave *'applemac'* y *'apple'* lideran la lista, indicando una clara asociación positiva con los productos de Apple. Además, términos específicos como *'mac'*, *'ipad'*, y *'macbook'* sugieren menciones directas a productos, indicando que los usuarios expresan opiniones positivas sobre dispositivos específicos. Por otra parte, las palabras como *'post'*, *'help'*, y *'learn'* indican interacciones positivas, donde

los usuarios comparten información útil y buscan aprender más. En contraste, de términos como *'new'*, *'upgrade'*, y *'update'* señalan la discusión de novedades y mejoras, características comunes en tweets positivos sobre productos tecnológicos.

En la Figura 5, se muestra el gráfico de barras según la Frecuencia de Palabras para el Tipo de Sentimiento *'negative'*.

Figura 5

Gráfico de barras según la Frecuencia de Palabras para el Tipo de Sentimiento 'negative'



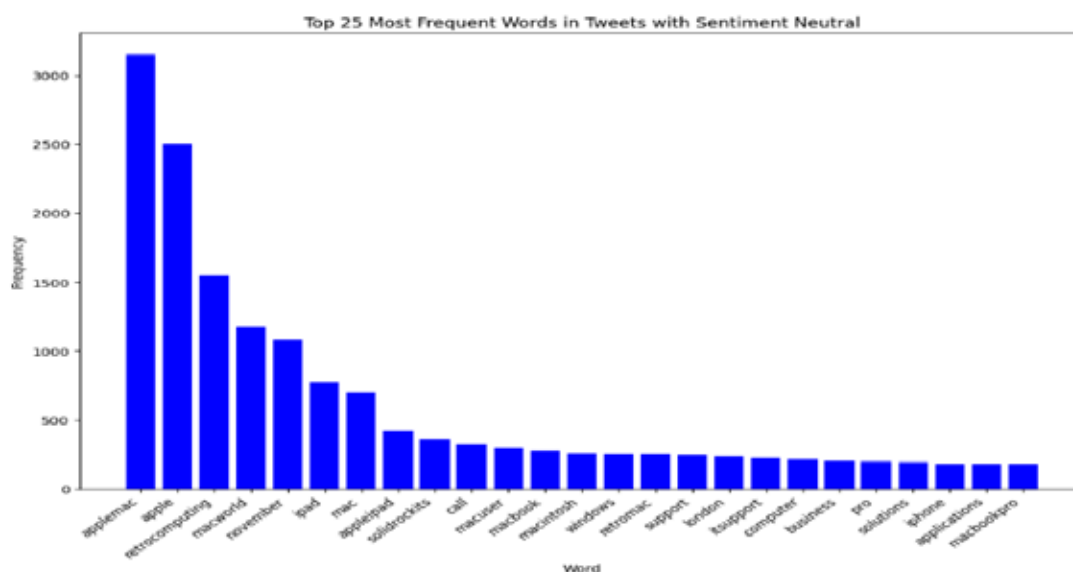
De acuerdo con la Figura 5, se visualiza la presencia prominente de términos como *'applemac,' 'mac,'* y *'apple'* en tweets negativos la cual sugiere problemas o insatisfacciones que están asociados principalmente con productos de la marca Apple. Además, las palabras clave como *'call,' 'system,'* y *'repair'* indican que los clientes pueden estar experimentando dificultades técnicas o problemas operativos, lo que puede requerir atención inmediata de servicio

al cliente o mejoras en los productos. Por otra parte, la repetición de palabras como *'slow,' 'damaged,'* y *'faults'* resalta posibles preocupaciones sobre el rendimiento y la calidad de los productos, ofreciendo una orientación clara sobre las áreas de enfoque para resolver problemas y garantizar la satisfacción del cliente.

En la Figura 6, se muestra el gráfico de barras según la Frecuencia de Palabras para el Tipo de Sentimiento *'neutral'*.

Figura 6

Gráfico de barras según la Frecuencia de Palabras para el Tipo de Sentimiento 'neutral'

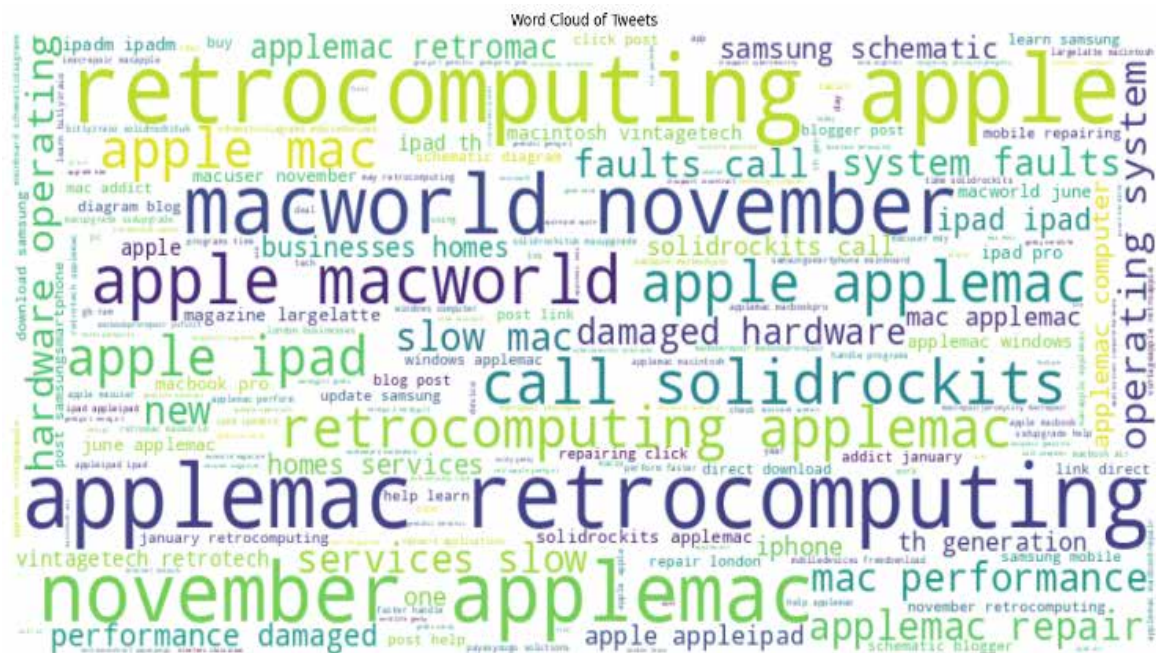


De acuerdo con la Figura 6, en donde la repetición significativa de términos como *'applemac'*, *'apple,'* y *'retrocomputing'* indica que, en este segmento de tweets neutrales, la atención de los usuarios se centra en productos y temas asociados con Apple. Asimismo, la presencia de palabras como *'macworld'*, *'november,'* y *'ipad'* sugiere que los clientes están discutiendo eventos específicos, meses del año y productos de Apple de manera neutral. Por otra parte, los términos como *'solidrockits'*, *'call,'* y *'macbook'* también proporciona información o asistencia relacionada con productos y servicios.

En la Figura 7 se presenta la gráfica de nube de palabras generada a partir de los tweets preprocesados del conjunto de datos. A partir de ello, se afirma que las palabras con mayor notoriedad son ‘*retrocomputing apple*’, ‘*apple retrocomputing*’, ‘*november applemac*’, ‘*macworld november*’ y ‘*apple macworld*’.

Estas expresiones destacadas proporcionan una indicación visual de los temas más relevantes en las conversaciones de los usuarios en relación con los productos de Apple y Samsung en las redes sociales.

Figura 7
Nube de palabras según la Frecuencia de Palabras



En la Figura 8, se muestra la representación gráfica de nube de palabras asociadas al tipo de sentimiento *'positive'*. Por ende, las palabras más destacadas son *'apple'*, *'apple mac'*, *'new'*, *'mac'*, *'apple apple mac'*, *'apple ipad'*, *'samsung schematic'*, *'iphone'* y *'apple mac'*. La presencia de

estas palabras sugiere que los tweets con sentimiento positivo están relacionados con productos y términos específicos como ‘apple’ y ‘mac’, indicando una percepción general positiva respecto a estos elementos en la red social Twitter/X.

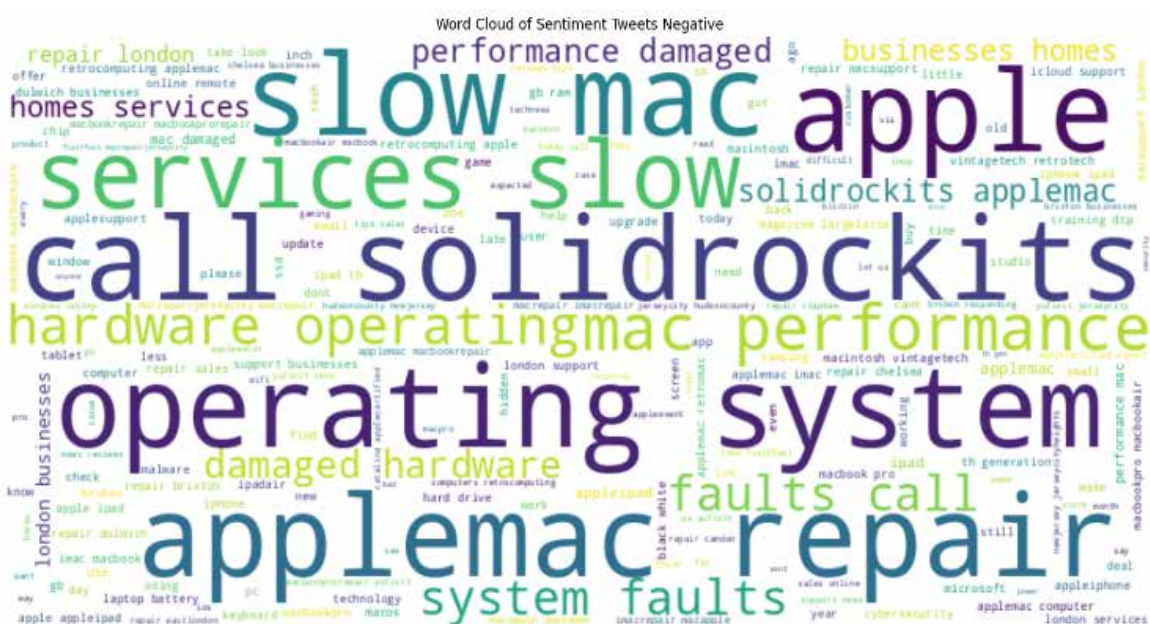
Figura 8
Nube de palabras según la Frecuencia de Palabras para el Tipo de Sentimiento 'positive'



En la Figura 9, se observa la gráfica de nube de palabras asociadas al tipo de sentimiento ‘negative’. En esta visualización, las palabras más destacadas son ‘*slow mac*’, ‘*apple*’, ‘*services slow*’, ‘*call solidrockits*’, ‘*operating system*’ y ‘*apple mac repair*’. Por lo tanto, se deduce que estos tweets están relacionados con experiencias

desfavorables, como problemas de velocidad (*slow mac*), llamadas a servicios (*call solidrockits*) y reparación de productos Apple (*apple mac repair*). Estos aspectos podrían indicar preocupaciones y percepciones negativas expresadas por los usuarios en la red social Twitter/X.

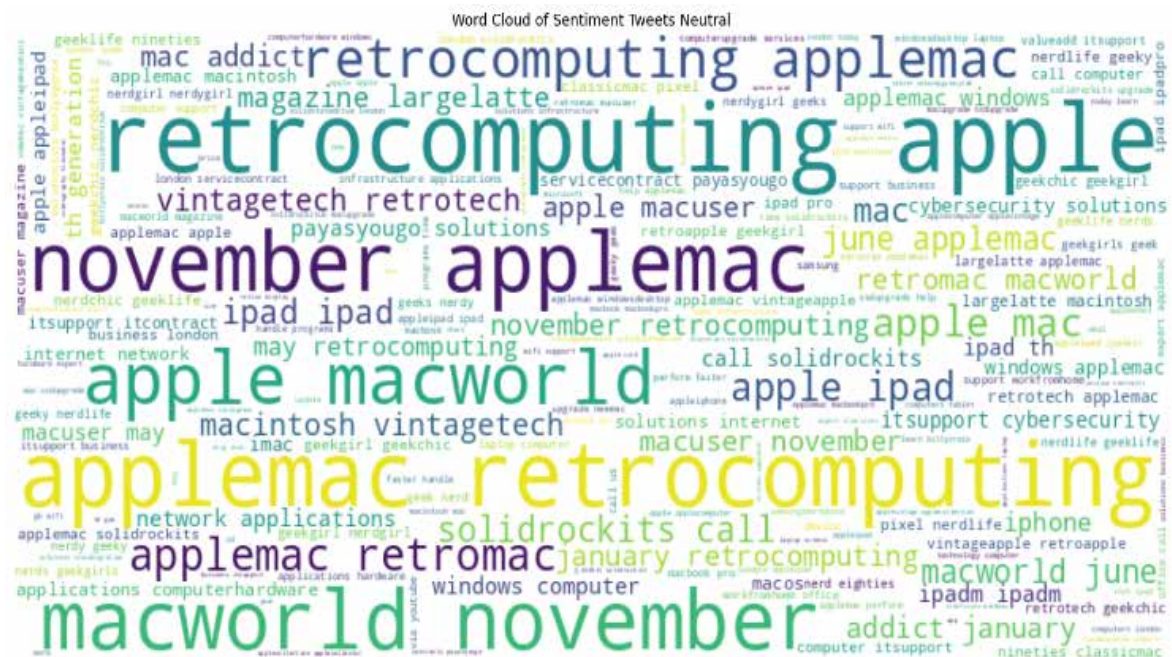
Figura 9
Nube de palabras según la Frecuencia de Palabras para el Tipo de Sentimiento 'negative'



En la Figura 10, se observa la representación gráfica de nube de palabras asociadas al tipo de sentimiento *‘neutral’*. En ella, las palabras más destacadas son *‘retrocomputing apple’*, *‘november applemac’*, *‘applemac retrocomputing’*, *‘macworld november’* y *‘apple macworld’*. En base a ello, estas palabras son temas

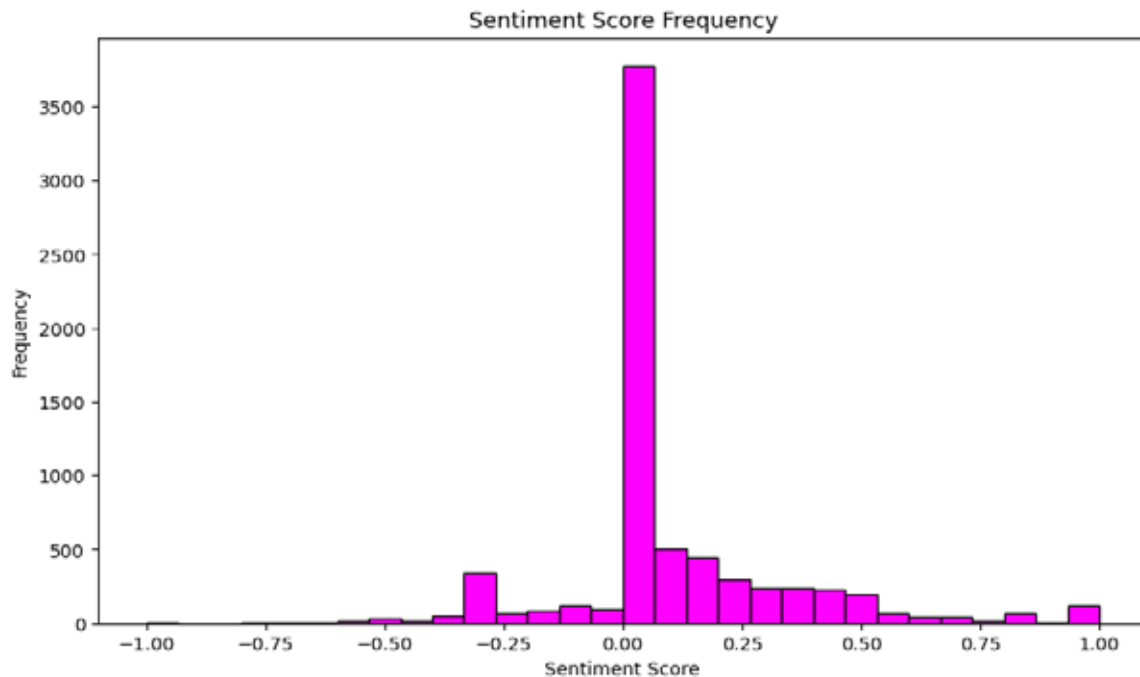
vinculados a retrocomputing, eventos de noviembre y menciones específicas a Apple y Macworld, por lo que estos indican una gama amplia de contenidos neutrales, la cual, ofrece a una perspectiva equilibrada sobre la percepción de los usuarios expresada en la red social Twitter/X.

Figura 10
Nube de palabras según la Frecuencia de Palabras para el Tipo de Sentimiento 'neutral'



En la Figura 11, se observa el gráfico de barras para identificar la frecuencia de polaridad (puntaje del sentimiento), se observa claramente que existe una mayor frecuencia de tweets con puntuaciones de sentimiento en el rango de 0 a 0.50. El histograma muestra una mayor frecuencia en los valores cercanos a 0.0 y 0.25 lo cual indica que la mayoría de los tweets tienden a tener puntajes de sentimientos neutrales

o ligeramente positivos. Por ende, esta distribución sugiere que los tweets experimentan una cantidad sustancial de comentarios con percepciones moderadas o neutrales, proporcionando información valiosa sobre la naturaleza general de los sentimientos expresados por los usuarios en relación con los productos de Apple y Samsung.

Figura 11*Gráfico de barras de la Frecuencia de Polaridad*

Resultados

En esta sección, se examinan los resultados derivados de la aplicación de modelos predictivos en el Análisis de Sentimientos, bajo el enfoque del Machine Learning, utilizando un conjunto de datos extraídos de Twitter/X, referentes a productos tecnológicos de las marcas Apple y Samsung.

A continuación, en la Figura 12 se detalla el desarrollo del enfoque de Machine Learning, a través de un experimento que incluye varias fases críticas. Este proceso comienza con la división del conjunto de datos en partes

de entrenamiento y prueba, seguido de la extracción de características utilizando la técnica TF-IDF. Posteriormente, se procede a la creación y entrenamiento de varios modelos de clasificación. Una vez entrenados, estos modelos se emplean para realizar predicciones sobre el conjunto de prueba, cuyo rendimiento se evalúa mediante métricas clave como Accuracy, Precision, Recall y F1-Score. Además, se utilizan herramientas de evaluación como la Matriz de Confusión y la Curva AUC-ROC para una comprensión más detallada de los resultados. Este procedimiento se lleva a cabo de manera secuencial, siguiendo el orden establecido.

Figura 12

Fases del enfoque de Machine Learning



En las Figuras 13, 14, 15 y 16 se presenta la Matriz de Confusión correspondiente a los modelos k-Nearest Neighbors, Logistic Regression, Random Forest y CatBoost Classifier, respectivamente. Estas visualizaciones ofrecen una representación detallada de los resultados de la clasificación, mostrando la cantidad de predicciones correctas e incorrectas

para cada clase. Cada celda de la matriz indica el número de instancias clasificadas correctamente (diagonal principal) y las clasificaciones erróneas (fuera de la diagonal principal), proporcionando una comprensión más profunda del rendimiento de cada modelo en la tarea de análisis de sentimientos.

Figura 13
Matriz de confusión para el modelo k-Nearest Neighbors

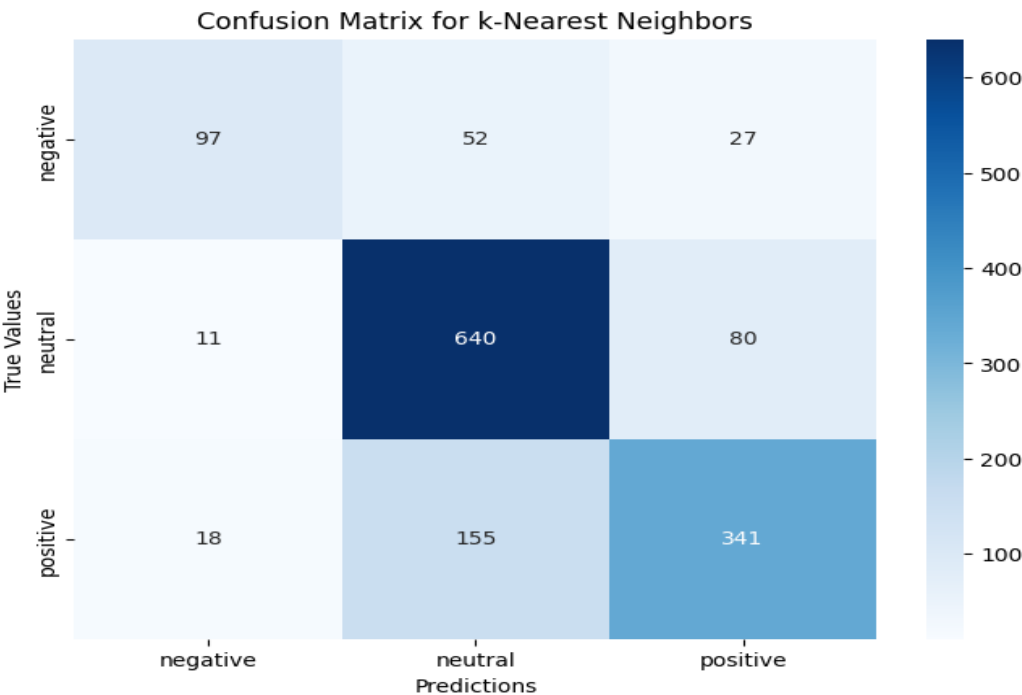


Figura 14
Matriz de confusión para el modelo Logistic Regression

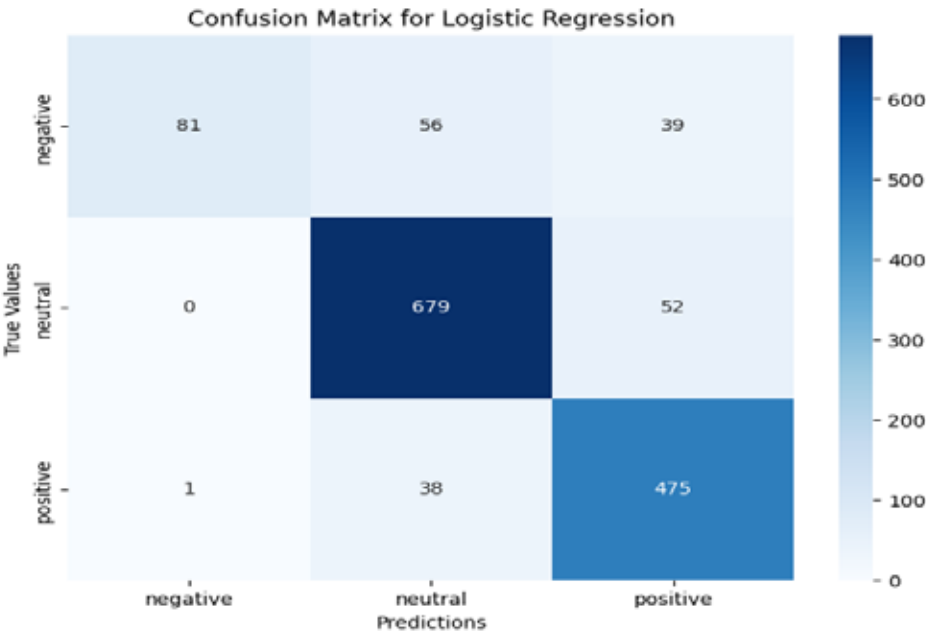


Figura 15

Matriz de Confusión para el modelo Random Forest

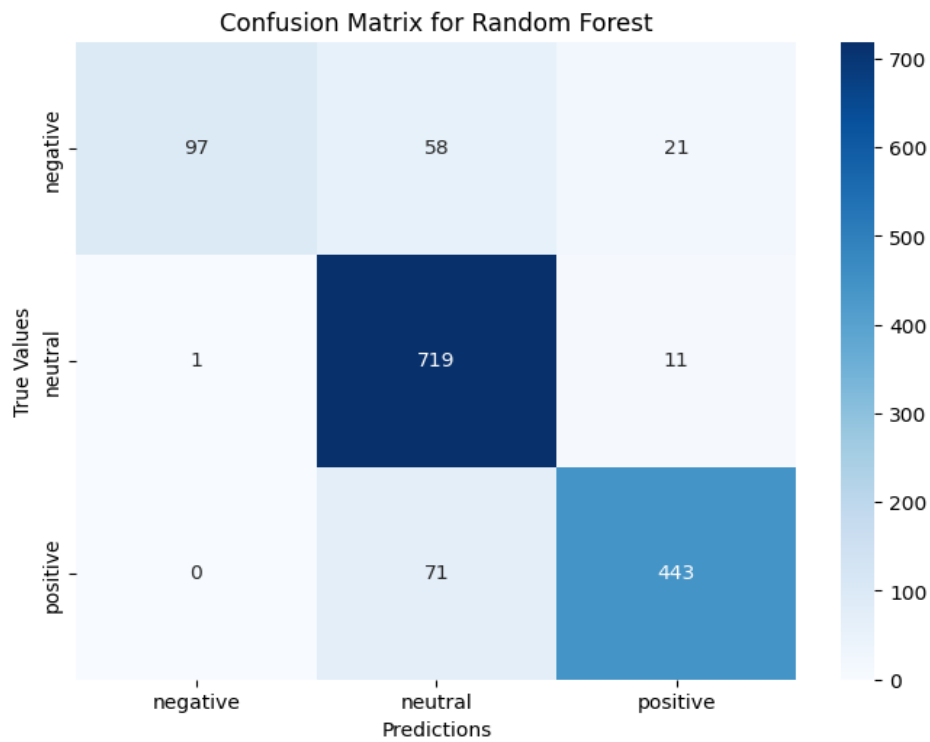
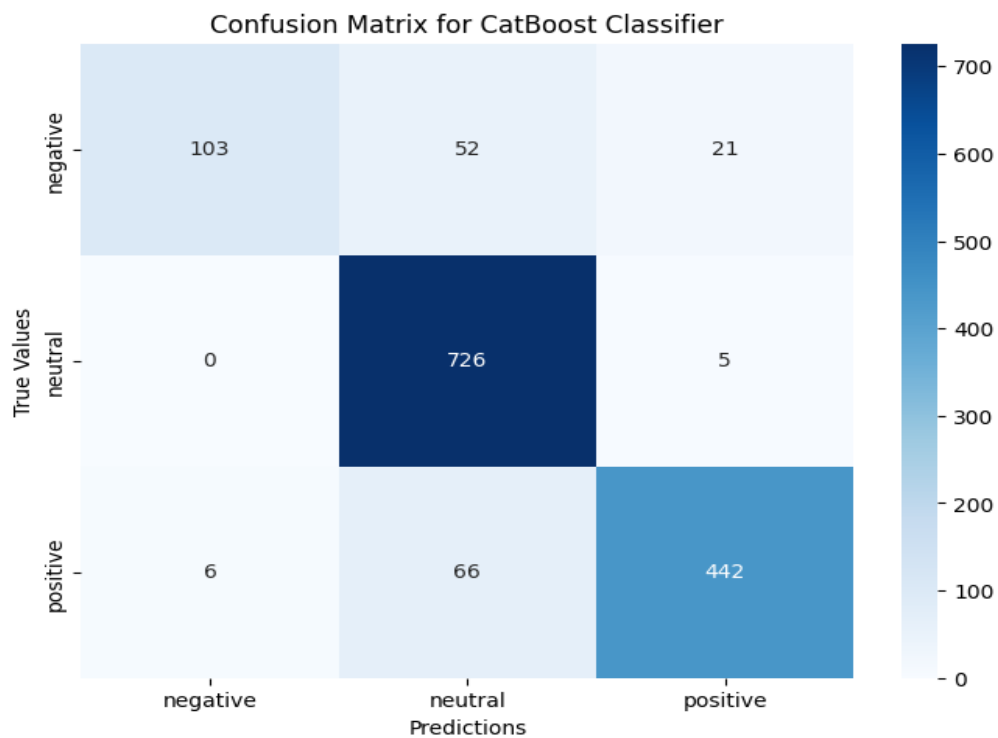


Figura 16

Matriz de confusión para el modelo CatBoost Classifier



En la Tabla 8 se presentan los resultados de las métricas clave para los diferentes modelos de clasificación aplicados al análisis de sentimientos en datos de e-commerce. Estas métricas incluyen Accuracy, Precision, Recall y F1-Score, proporcionando una visión integral del rendimiento de cada modelo.

A partir de ello, se deduce que, para el modelo k-Nearest Neighbors, se observa un Accuracy del 75.86%, indicando una capacidad moderada para clasificar opiniones. La precisión y el recall exhiben valores equilibrados, sugiriendo una capacidad aceptable tanto para identificar sentimientos positivos como negativos, con un F1-Score del 75.33%. En el caso de la Regresión Logística, se destaca su

impresionante Accuracy del 86.91%, indicando una capacidad sobresaliente para realizar predicciones precisas. Las métricas de precisión y recall son equilibradas, resultando en un alto F1-Score del 86.04%. El modelo Random Forest logra un notable Accuracy del 88.60%, demostrando una capacidad robusta para clasificar sentimientos en datos de e-commerce. Las métricas de Precision y Recall son elevadas, resultando en un F1-Score del 88.02%. Finalmente, el CatBoost Classifier lidera con un Accuracy del 89.44%, posicionándose como el modelo más preciso en la clasificación de sentimientos. Exhibe altos valores de Precisión y Recall, con un destacado F1-Score del 88.94%.

Tabla 8

Resultados de las métricas claves

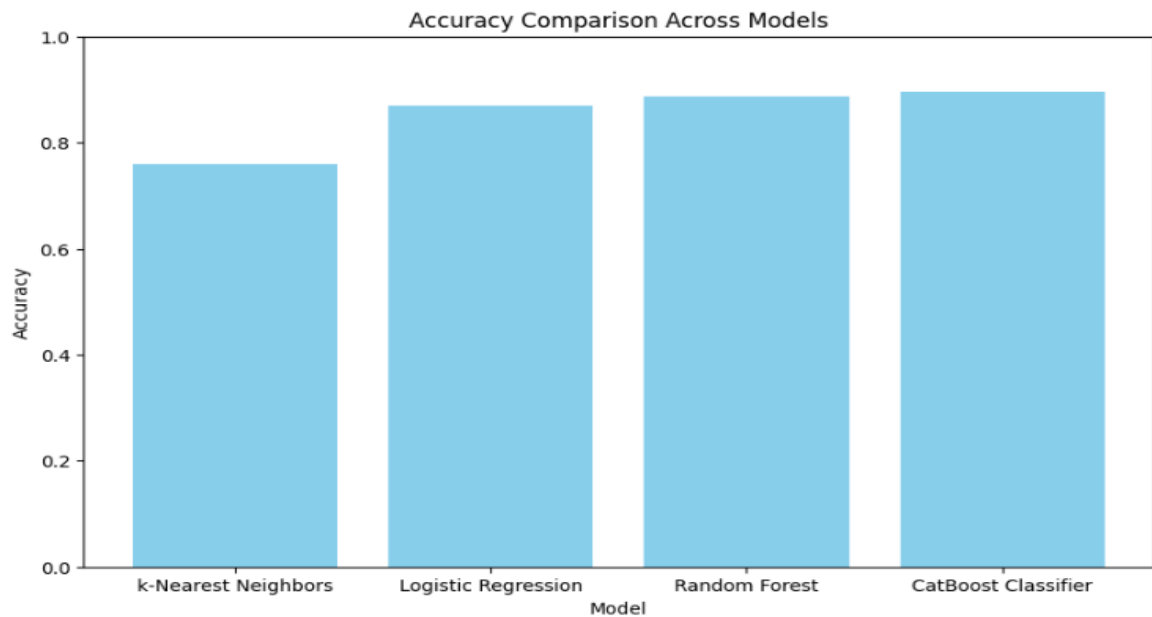
Modelo	Accuracy	Precision	Recall	F1-Score
k-Nearest Neighbors	0.758621	0.759379	0.758621	0.753278
Logistic Regression	0.869106	0.877777	0.869106	0.860436
Random Forest	0.885996	0.896112	0.885996	0.880230
CatBoost Classifier	0.894441	0.901165	0.894441	0.889396

En la Figura 17, se presenta una comparación visual del accuracy de los modelos k-Nearest Neighbors, Logistic Regression, Random Forest y CatBoost Classifier. Estos resultados reflejan la capacidad de cada modelo para realizar predicciones precisas en la tarea de análisis de sentimientos en el contexto del e-commerce. En consecuencia, se destaca que el CatBoost Classifier lidera con un impresionante nivel de accuracy, alcanzando el valor más alto entre los modelos evaluados.

Este resultado subraya la eficacia del CatBoost Classifier en la clasificación de sentimientos en comparación con los otros modelos. Mientras que el k-Nearest Neighbors muestra una capacidad moderada, los modelos Logistic Regression y Random Forest también demuestran un rendimiento destacado, especialmente el Random Forest, que se sitúa como el segundo modelo con la accuracy más elevada.

Figura 17

Comparación de Accuracy de los modelos k-Nearest Neighbors, Logistic Regression, Random Forest y CatBoost Classifier



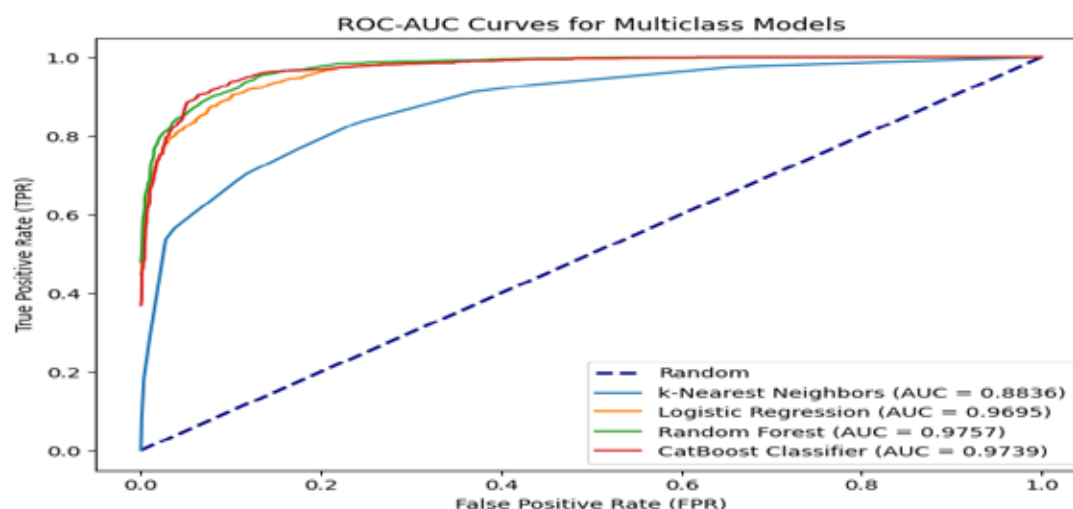
En la Figura 18 se presentan las curvas ROC-AUC de los modelos K-Nearest Neighbors, Logistic Regression, Random Forest y CatBoost Classifier. Estas curvas ofrecen una evaluación más detallada del rendimiento de los modelos en términos de equilibrio entre la tasa de falsos positivos y la tasa de verdaderos positivos. Los resultados destacan especialmente el modelo Random Forest, en la cual, obtiene la curva ROC-AUC más alta entre los modelos evaluados, indicando una capacidad superior para discriminar entre las clases de sentimientos. Por otra parte, los modelos Logistic Regression y CatBoost Classifier también muestran desempeños notables con valores de

ROC-AUC cercanos, lo que sugiere una eficacia consistente en la clasificación de sentimientos.

En contraste, con el modelo k-Nearest Neighbors, aunque presenta un valor de ROC-AUC más bajo en comparación con los otros modelos, aún muestra una capacidad considerable para distinguir entre las clases. Por ende, esta gráfica proporciona una visualización crucial para comprender la capacidad de los modelos en la clasificación de sentimientos y facilita la comparación entre ellos, contribuyendo a una toma de decisiones informada en la elección del modelo más adecuado para la tarea en el contexto del e-commerce.

Figura 18

Curva ROC-AUC de los modelos k-Nearest Neighbors, Logistic Regression, Random Forest y CatBoost Classifier



Discusión

La investigación de Loukili et al. (2023) se centró en realizar un análisis para comparar los resultados (Accuracy, Precision, Recall y F1-Score) de los modelos de K-Nearest Neighbors, Logistic Regression, Random Forest y CatBoost Classifier, basándose en una revisión detallada de la Matriz de Confusión y las Curvas AUC-ROC. Los modelos que obtuvieron los valores más altos fueron Logistic Regression, con un Accuracy del 90%, y Random Forest, con una Precision del 91%, mientras que Logistic Regression también destacó con un 78% en Recall y un 81% en F1-Score. En contraste, la presente investigación, que también incluyó una comparación de modelos como K-Nearest Neighbors, Logistic Regression, Random Forest y CatBoost Classifier considerando una revisión de la Matriz de Confusión y las Curvas de AUC-ROC, encontró que el modelo CatBoost Classifier alcanzó valores altos en todas las métricas: 89% de Accuracy, 90% de Precision, 89% de Recall y 88% de F1-Score.

Si comparamos ambas investigaciones respecto a las métricas, los resultados del Accuracy y Precision no son favorables y una mejora en el Recall del 10.8% y F1-Score del 7.54% para la presente investigación. En la evaluación de las curvas ROC-AUC, ambas investigaciones destacan la capacidad discriminativa de los modelos de análisis de sentimientos en el ámbito de e-commerce. En este estudio, los modelos muestran valores de ROC-AUC que oscilan entre 0.8836 y 0.9757, siendo el Random Forest el modelo con el valor más alto, seguido de cerca por el CatBoost Classifier y el Logistic Regression. En contraste con la investigación previa, presenta valores desde 0.56227 hasta 0.93627, con el Logistic Regression liderando en términos de capacidad discriminativa, seguido por el Random Forest y el CatBoost Classifier. Es importante señalar que, a pesar de las diferencias en los valores exactos, ambos estudios subrayan la efectividad de los modelos en distinguir entre las clases de sentimientos en el contexto del e-commerce. Por lo tanto, las curvas ROC-AUC proporcionan una representación

visual del equilibrio entre la tasa de falsos positivos y la tasa de verdaderos positivos para cada modelo, contribuyendo así a la comprensión integral de su rendimiento en la tarea de análisis de sentimientos.

Asimismo, en la investigación de Bansal y Srivastava (2018), se realizó un análisis de sentimientos utilizando el algoritmo más óptimo, a través de una comparación de modelos que incluía Random Forest, SVM, Logistic Regression y Naive Bayes. Entre ellos, el modelo Random Forest se destacó, alcanzando un Accuracy del 90.2161% con Skip Gram y del 90.6622% con CBOW. En contraste, la presente investigación logró un Accuracy de 89.4441% usando el modelo CatBoost Classifier. Al comparar ambas investigaciones respecto a la métrica de Accuracy, los resultados no fueron favorables para la presente investigación; no obstante, hay una mejora significativa en la curva ROC-AUC, en donde la presente investigación logró un valor de 0.9739 para el modelo CatBoost Classifier y de la investigación previa fue de 0.93 para Random Forest (Skip Gram) y 0.94 para Random Forest (CBOW). Esto indica que, aunque el Accuracy puede ser un indicador importante, no debe ser el único criterio de evaluación. En este caso, la capacidad del modelo para clasificar correctamente los sentimientos, representada por la curva ROC-AUC, también es fundamental.

A continuación, en la investigación de Sinnasamy y Sjaif (2022) en la que se realizó un análisis de sentimientos para comprender las preferencias expresadas en las opiniones, con el propósito de evaluar los servicios, la información y la calidad de los productos, se obtuvieron cuatro modelos de clasificación: Accuracy,

Precisión, Recall y F1-Score. El método TF-IDF con N-Gram reveló un unigrama, que favoreció al modelo Decision Tree, alcanzando un Accuracy de 82.27%. Por otro lado, el modelo SVM obtuvo mejores valores en Precisión (82%), Recall (80%) y F1-Score (72%), contrastando con la presente investigación, que también desarrolló una comparación de modelos en el análisis de sentimientos, empleando el modelo CatBoost Classifier con valores destacados en sus métricas de 89% en Accuracy, 90% en Precisión, 89% en Recall y 88% en F1-Score. Comparando ambas investigaciones en cuanto a métricas, los resultados del CatBoost Classifier muestran una mejora del 6.73% en Accuracy, 8% en Precisión, 9% en Recall y 16% en F1-Score.

Seguidamente, en la investigación de Ramona et al. (2018), se logró recopilar opiniones sobre ofertas localizadas en Twitter/X. Durante este proceso, realizaron un análisis de sentimientos utilizando algoritmos de Machine Learning con 2204 tweets, categorizando el 60.2% como neutros, 32.1% como positivos y 7.7% como negativos. En contraste, la presente investigación analizó 7102 tweets y, mediante modelos similares de Machine Learning, clasificó el 51.4% como neutros, 36.9% como positivos y 11.7% como negativos. Comparando ambas investigaciones en términos de la categorización de opiniones a través del análisis de sentimientos, se observa una mejora del 4.8% en las opiniones positivas en la investigación actual

Por otra parte, en la investigación de Sangeetha y Kumaran (2023), presentaron como solución al análisis de sentimientos aplicado a las opiniones de consumidores, de las cuales se ejecutaron un modelo

de Machine Learning en la clasificación según su polaridad correspondiente, el resultado determinó que el modelo RNN-LSTM demostró mejoras de un 95.8% en el Accuracy, 95.4% de Precisión, 95.6% de Recall y 95.2% de F1-Score, contrastando con la presente investigación que también se aplicó el análisis de sentimientos en diferentes tweets y una comparación de modelos de Machine Learning considerando una pérdida del 6.8% de Accuracy, 5.4% de Precisión, 7.6% de Recall y 7.2% de F1-Score.

Finalmente, en la investigación de Zhao et al. (2021), presentaron la dificultad de prever con precisión el grado de polaridad de las opiniones de los usuarios, para ello, aplicaron el análisis de sentimientos a través de un algoritmo optimizado denominado LSIBA- ENN que logró un mejor rendimiento en la clasificación de 83.55% de Accuracy, 84.03% de Precisión, 85.48% de Recall y 86.04% de F1-Score; en contraste con la presente investigación en la que se utilizó el modelo o algoritmo CatBoost Classifier, logró destacar en todas sus métricas. Si se compara, ambas investigaciones respecto a las métricas, los resultados tuvieron una mejora respecto a la presente investigación del Accuracy del 5.45%, Precisión del 5.97%, Recall del 3.52% y F1-Score del 1.96%.

Conclusiones

Esta investigación concluye que la implementación del análisis de sentimientos para evaluar productos

tecnológicos comentados en la red Twitter/X, ha revelado la capacidad y el potencial de discernir las percepciones de los clientes. Por lo que, esta metodología ha demostrado ser un complemento estratégico valioso en el ámbito del marketing, permitiendo no solo identificar las opiniones y emociones de los usuarios de manera eficaz sino también ofreciendo insights relevantes que enriquecen los estudios de mercado.

Por lo tanto, el Análisis de Sentimientos se posiciona como un proceso necesario para recabar y analizar la información que los clientes opinan sobre la marca en las redes sociales, beneficiando a las empresas tecnológicas que buscan entender mejor y responder a las expectativas de su audiencia.

Asimismo, los recursos del Procesamiento del Lenguaje Natural (PLN) y librerías claves utilizadas facilitaron aplicar técnicas sintácticas y semánticas para comprender la estructura de un texto e identificar su significado. Además, la combinación de dos áreas de la IA, tales como el Procesamiento del Lenguaje Natural (PLN) y Machine Learning permitió realizar el Análisis de Sentimiento de los textos extraídos de la red social Twitter/X.

Finalmente, se logró determinar el modelo predictivo más conveniente a través de las pruebas realizadas con varios modelos, que servirá como patrón para obtener información rápida y precisa sobre los resultados de campañas publicitarias recién lanzadas.

Referencias

- Alroobaea, R. (2022). Sentiment analysis on Amazon product reviews using the recurrent neural network (RNN). *International Journal of Advanced Computer Science and Applications*, 13(4), 37. https://thesai.org/Downloads/Volume13No4/Paper_37-Sentiment_Analysis_on_Amazon_Product_Reviews.pdf
- Bansal, B., & Srivastava, S. (2018). Sentiment classification of online consumer reviews using word vector representations. *Procedia Computer Science*, 132, 1147-1153. <https://doi.org/10.1016/j.procs.2018.05.029>
- Capurso, M. (2021). *Data Science Quick Reference Manual – Methodological Aspects, Data Acquisition, Management and Cleaning: With applications in the Python-based Orange environment*.
- Deisenroth, M., Faisal, A., & Ong, C. (2020). *Mathematics for machine learning*. Cambridge University Press.
- Demircan, M., Seller, A., Abut, F., & Akay, M. (2021). Developing Turkish sentiment analysis models using machine learning and e-commerce data. *International Journal of Cognitive Computing in Engineering*, 2, 202-207. <https://doi.org/10.1016/j.ijcce.2021.11.003>
- Erl, T., Khattak, W., & Buhler, P. (2016). *Big Data Fundamentals: Concepts, Drives and Techniques*. Prentice Hall. <https://dl.icdst.org/pdfs/files3/053e44f660b0c2f405e42ac1f8f1a408.pdf>
- Feizollah, A., Ainin, S., Anuar, N., Abdullah, N., & Hazim, M. (2019). Halal products on Twitter: Data extraction and sentiment analysis using a stack of deep learning algorithms. *IEEE Access*, 7, 83354-83362. <https://ieeexplore.ieee.org/document/8737773>
- Fernandes, S., Sharma, A., & Vidyasagar, A. (2019). Determining customer perceptions using text mining for an FMCG product 'Maggi' in India. <https://www.ijitee.org/portfolio-item/J97800881019/>
- Haddi, E., Liu, X., & Shi, Y. (2013). The role of text pre-processing in sentiment analysis. *Procedia Computer Science*, 17, 26-32. <https://doi.org/10.1016/j.procs.2013.05.005>
- Han, J., Pei, J., & Tong, H. (2022). *Data mining: Concepts and techniques*. Morgan Kaufmann.
- Herrera-Contreras, A., Sánchez-Delacruz, E., Meza, I., & Fuentes-Ramos, M. (2021). SASTuit, software de análisis de sentimiento utilizando aprendizaje automático. Unpublished manuscript. <https://doi.org/10.13140/RG.2.2.30772.37764>
- Iqbal, A., Amin, R., Iqbal, J., Alroobaea, R., Binmahfoudh, A., & Hussain, M. (2022). Sentiment analysis of consumer reviews using deep learning.

- ning. *Sustainability*, 14(17), 10844. <https://www.mdpi.com/2071-1050/14/17/10844>
- Kamath, K. (2024). *Social Media Marketing Essentials*. Vibrant Publishers.
- Loukili, M., Messaoudi, F., & Ghazi, M. (2023). Sentiment analysis of product reviews for E-commerce recommendation based on machine learning. <https://doi.org/10.15849/IJASCA.230320.01>
- Mar, C., Montañez, P., & Simón-Cuevas, A. (2023). Análisis de sentimientos orientado al comercio electrónico de CITMATEL: proyecto de voz del cliente. *Revista Cubana de Transformación Digital*, 4(1), e207-e207. <http://rctd.uic.cu/rctd/article/view/207/109>
- Marín, C., & Botey, M. (2022). Estrategias promocionales de marketing digital en redes sociales: Análisis bibliométrico de estrategias digitales a través de Facebook e Instagram. *revTECHNO*, 12(Monográfico), 1-11. <https://doi.org/10.37467/revtechno.v11.4393>
- Medina-Chicaiza, P., & Martínez-Ortega, A. (2020). Tecnologías en la inteligencia artificial para el marketing: Una revisión de la literatura. *Pro Sci.*, 4(30), 36-47. <https://doi.org/10.29018/issn.2588-1000vol4iss30.2020pp36-47>
- Mendivelso, H., & Lobos, F. (2019). La evolución del marketing: una aproximación integral. *Revista Chilena de Economía y Sociedad*. <http://rches.utem.cl/articulos/la-evolucion-del-marketing-una-aproximacion-integral/>
- Nagarkar, P., Khan, A., Raikar, S., & Zantye, A. (2020). Twitter data mining for targeted marketing. In *Second International Conference on Inventive Research in Computing Applications (ICIRCA)* (pp. 44-50). Coimbatore, India. <http://ieeexplore.ieee.org/document/9183005>
- Naga, R., Pratibha, N., Sk althaf, H., & M, Y. (2023). *Data mining: Concepts and Techniques*. First Edition. AG Publishing House
- Plaza, A., Marcillo, J., Hidalgo, J., Anchundia, O., Pilacuan, L., Parrales, A., & Navas, W. (2022). Text mining for analyzing digital consumer behavior on Twitter. 20th LACCEI International Multi-Conference for Engineering, Education, and Technology. http://laccei.org/LACCEI2022-BocaRaton/full_papers/FP772.pdf
- Ramona, J., Reyes-Menéndez, A., & Palos-Sánchez, P. (2018). Un análisis de sentimiento en Twitter con Machine Learning: Identificando el sentimiento sobre las ofertas de #BlackFriday. *Revista ESPACIOS*, 39(42). <http://www.revistaespacios.com/a18v39n42/18394216.html>
- Sadeq, N., Nassreddine, G., & Younis, J. (2023). Impact of artificial intelligence on E-marketing. *International Journal of Trend in Scientific Research and Development*, 7(1). <https://www.ijtsrd.com/papers/ijtsrd53850.pdf>

- Sangeetha, & Kumaran. (2023). Sentiment analysis of Amazon user reviews using a hybrid approach. *Measurement Sensing*, 27(100790), 100790. <https://doi.org/10.1016/j.measen.2023.100790>
- Sinnasamy, T., & Sjaif, N. (2022). Sentiment analysis using term-based method for customers' reviews in Amazon product. *International Journal of Advanced Computer Science and Applications*, 13(7). <https://doi.org/10.14569/ijacsa.2022.0130780>
- Taherdoost, H., & Madanchian, M. (2023). Artificial intelligence and sentiment analysis: A review in competitive research. *Computers*, 12(2), 37. <https://www.mdpi.com/2073-431X/12/2/37>
- Yannis, A., Charalampos, M., & Thomas, E. (2018). *The World of Open Data: Concepts, Methods, Tools and Experiences*. Springer. <https://doi.org/10.1007/978-3-319-90850-2>
- Zhao, H., Liu, Z., Yao, X., & Yang, Q. (2021). A machine learning-based sentiment analysis of online product reviews with a novel term weighting and feature selection approach. *Information Processing & Management*, 58(5), 102656. <https://doi.org/10.1016/j.ipm.2021.102656>

