

Instituto Tecnológico y de Estudios Superiores de Occidente

Reconocimiento de validez oficial de estudios de nivel superior según acuerdo secretarial
15018, publicado en el Diario Oficial de la Federación del 29 de noviembre de 1976.

Departamento de Matemáticas y Física
Maestría en Ciencia de Datos



Aplicación de Modelos Transformer para la Clasificación y Análisis de Quejas en Atención al Cliente

TRABAJO RECEPCIONAL que para obtener el **GRADO** de
Maestro en Ciencia de Datos

Presenta:
Enrique Ignacio Rojas Villanueva

Director:
Dr. Jaime Emmanuel Alcalá Temores

Tlaquepaque, Jalisco, 12 de mayo de 2024

Aplicación de Modelos Transformer para la Clasificación y Análisis de Quejas en Atención al Cliente

Enrique Ignacio Rojas Villanueva

Resumen

En este trabajo, adoptamos una metodología mixta para explorar el impacto de los modelos de Procesamiento de Lenguaje Natural (Natural Language Processing, NLP, por sus siglas en inglés) basados en *Transformers*, enfocándonos específicamente en su capacidad para analizar e interpretar texto. Nuestra investigación se centra en evaluar la efectividad y eficiencia de estos modelos al procesar y categorizar interacciones textuales específicas. Dicho análisis se realiza sobre una muestra cuidadosamente seleccionada de 3,000 interacciones en forma de *tickets* de atención al cliente, provenientes de diversos canales de comunicación digital.

Este enfoque nos permite no solo comprender cómo los modelos de *Transformers* pueden identificar y clasificar los diferentes tipos de consultas y problemas reportados por los usuarios, sino también evaluar su precisión, la cual se espera alcance al menos el 90 % en la identificación de categorías relevantes. La elección de *tickets* de atención al cliente como objeto de estudio se debe a su riqueza informativa y relevancia para las empresas que buscan optimizar sus estrategias de servicio y soporte al cliente mediante tecnologías de NLP.

Tabla de Contenidos

	Página
1 Introducción	13
1.1. Contexto	13
1.2. Justificación	15
1.3. Problema	16
1.4. Objetivos	17
1.4.1. Objetivo general	17
1.4.2. Objetivos específicos	17
2 Metodología	19
2.1. Descripción de los datos	19
2.2. Análisis exploratorio	20
2.3. Descripción de los modelos	21
2.3.1. Modelos de comparación	22
2.4. Descripción de las métricas	23
2.5. Descripción de los experimentos o simulaciones	25
3 Resultados y discusión.	27
3.1. Resultados	27
3.1.1. Transformer model	27
3.1.2. Regresión Logística	28
3.1.3. Support Vector Machine	28
3.1.4. Random Forest	28
3.1.5. Red Neuronal	29
3.2. Discusión	29
4 Conclusiones y trabajo futuro.	31
4.1. Conclusiones	31
4.2. Trabajo futuro	33
Bibliografía	36

Índice de figuras

	Página
2.1. Distribución de Categorías basada en las Quejas de los Clientes. La gráfica de barras horizontales presenta las diferentes categorías de problemas identificadas a partir de las interacciones de los clientes	20
2.2. Distribución de la Longitud de las Quejas de los Clientes. Esta gráfica representa la distribución de la cantidad de caracteres en las quejas de los clientes. Las barras azules muestran el número de quejas en cada rango de longitud, mientras que la línea suavizada indica la densidad de probabilidad, calculada con kernel density estimation (KDE), de las longitudes de las quejas. La línea roja punteada representa la media de la longitud de las quejas, y la línea azul punteada indica la mediana.	21
2.3. Palabras Más Comunes en Quejas de Consumidores. La imagen muestra una nube de palabras generada a partir de una muestra aleatoria de 5,000 quejas de consumidores. Las palabras que aparecen en tamaños más grandes son aquellas que se mencionan con mayor frecuencia en las quejas.. . . .	21
2.4. Diagrama de funcionamiento para el modelo <i>Transformer</i> , engloba el mecanismo de atención y con ello se remarcan los generadores de texto. En este esquema se especifica la entrada y salida del texto que tiene este tipo de funcionamiento. Tomado de [1].	26
3.1. Distribución del procesamiento del modelo <i>Transformer</i> .	27

Índice de tablas

	Página
2.1. Tabla de Análisis de Datos por Frecuencia	20
3.1. Desempeño del modelo	28
3.2. <i>Accuracy</i> por etiqueta	28
3.3. Desempeño del modelo de Regresión Logística por categoría	28
3.4. Desempeño del modelo SVM por Categoría	29
3.5. Desempeño del modelo RF por Categoría	29
3.6. Desempeño del modelo de Red Neuronal por Categoría	29

AGRADECIMIENTOS

Quiero expresar mi más profundo agradecimiento a mi esposa, cuyo apoyo incondicional ha sido el pilar fundamental para la realización de este proyecto, así como en cada etapa de mi carrera y en los respectivos desafíos que me he propuesto; su paciencia, comprensión y amor con la constante fuente de motivación en este viaje académico.

También agradezco a mis compañeros de curso por su apoyo, colaboración y sobre todo por compartir momentos significativos durante este proceso de formación. Así pues, la sinergia y la camaradería que hemos construido han sido esenciales para superar los retos que se presentaron en el camino.

No puedo dejar de mencionar a mi familia, quienes siempre han sido mi refugio y mi guía. Su fe en mí y su amor inquebrantable me han impulsado a alcanzar metas más altas y a perseverar en los momentos más difíciles.

Especialmente agradezco a Esteban L. Calero, aunque no contribuyó directamente a este documento, su apoyo en mi trabajo diario y sus valiosos consejos en programación han sido esenciales para mi desarrollo profesional y personal.

Mi sincero reconocimiento al Instituto Tecnológico y de Estudios Superiores de Occidente (ITESO) por contribuir en mi formación académica de excelencia y ser el espacio donde he forjado muchas de mis habilidades y conocimientos. Agradezco especialmente a los profesores del posgrado, cuya dedicación, enseñanzas y guía han sido fundamentales en mi trayectoria académica.

1 Introducción

En la intersección de la digitalización y la globalización, se ha reconfigurado el entramado de las relaciones entre cliente-empresa, con oportunidades sin precedentes y desafíos complejos. El nuevo paisaje digital, con su constante interacción y gran cantidad de datos, requiere herramientas sofisticadas para navegar eficazmente por él. A medida que las empresas intentan adaptarse y crecer en este entorno cambiante, se vuelve crucial llevar a cabo un análisis profundo y comprensivo de las interacciones entre el cliente y la empresa para lograr el éxito [2]. En este contexto, los modelos de *Transformers* [3] en el Procesamiento del Lenguaje Natural (*Natural Language Processing*, o NLP, por sus siglas en inglés) representan una innovación revolucionaria, prometiendo descifrar la gran cantidad de datos generados en las plataformas digitales y redes sociales.

Este documento explora la implementación de un modelo de *Transformers* [3] para clasificar las quejas de clientes de una empresa financiera por tema, basándose en el contenido del texto y la relación semántica entre las quejas.

1.1 Contexto

En los últimos años, la digitalización y la globalización han tenido un gran impacto, especialmente en la interacción entre clientes y empresas. Los avances en las telecomunicaciones han modificado notablemente las dinámicas comerciales contemporáneas [4][5]. En la actualidad, las empresas pueden acceder a una cantidad de datos nunca antes vista, lo que marca la tendencia en este ámbito con el florecimiento de las plataformas digitales y las redes sociales [6][7]. Sin embargo, sigue siendo un desafío considerable analizar y comprender estos datos en una escala significativa.

Las empresas pudieron adaptarse a las tendencias emergentes y cambiar sus modelos de negocio en la era digital para satisfacer las demandas cambiantes de los clientes [4]. La eficiencia y la calidad de los servicios ofrecidos a los clientes mejoran mediante esta transformación digital [8]. Al mismo tiempo, la globalización se alinea

con el crecimiento de las empresas, ya que buscan principalmente nuevos mercados y expandir su presencia a nivel mundial [5]. Se han buscado las mejores formas de explorar estos datos que ofrece el mundo globalizado en el campo del Procesamiento del lenguaje.

Es crucial desarrollar nuevas tecnologías y herramientas innovadoras para abordar el reto de analizar y comprender grandes conjuntos de datos de clientes. Los modelos de *Transformers*, que se utilizan en el campo del , son algunas de las herramientas más efectivas para recopilar esta información, lo cual demuestra su eficacia en la extracción de datos valiosos de interacciones entre clientes y empresas [9]. Al enfocarse en el contexto, estos modelos ofrecen una comprensión más profunda de las perspectivas y preocupaciones de los clientes, lo que a su vez permite a las empresas proponer soluciones basadas en datos para mejorar su relación con los clientes [9].

Focalizándose en el contexto, estos modelos ayudan a comprender mejor las perspectivas y preocupaciones de los clientes, lo que habilita a las empresas para ofrecer soluciones basadas en datos con el fin de mejorar la relación cliente-empresa en la era digital. La digitalización y la globalización tienen un impacto notable en clientes y empresas cuando se combinan, esa interacción necesita modelos que puedan manejar la complejidad inherente del lenguaje natural.

En este trabajo se usan modelos *Transformers* para mejorar la comprensión y análisis de las relaciones cliente-empresa. Se espera que este enfoque permita el desarrollo de soluciones basadas en datos cada vez más precisas. El aspecto innovador proviene del codificador con autoatención especial, donde cada capa de autoatención puede tomar una secuencia de vectores y producir una nueva secuencia de vectores.

El proceso innovador de autoatención que define a los modelos *Transformers* es especialmente útil en el contexto de la interacción digital moderna. Mediante la utilización de matrices preentrenadas, estos modelos ajustan dinámicamente el significado de las palabras basándose en su contexto [1]. Esto permite distinguir con precisión entre múltiples significados de una misma palabra, como el término “banco” que puede referirse tanto a una institución financiera como a un objeto físico para sentarse.

Gracias a esta capacidad de discernimiento contextual, los modelos *Transformers* son capaces de ofrecer insights más detallados y precisos sobre las comunicaciones de los clientes. Estos insights se traducen en acciones concretas que las empresas pueden implementar para optimizar sus estrategias de servicio al cliente y marketing, personalizando la experiencia del usuario para satisfacer mejor sus necesidades y expectativas específicas.

1.2 Justificación

En la actualidad, la retención y satisfacción del cliente se destacan como factores críticos para asegurar el éxito y el crecimiento sostenido de cualquier empresa. Numerosos estudios han subrayado la correlación directa entre la satisfacción del cliente y la rentabilidad empresarial. Por ejemplo, según un informe de la *Harvard Business Review*, un aumento del 5% en la retención de clientes puede incrementar las ganancias de una empresa entre un 25% y un 95% [4].

Las empresas que son capaces de identificar y atender de manera proactiva las necesidades y preocupaciones de sus clientes tienden a ganar una ventaja competitiva significativa en el mercado [5]. En un entorno empresarial cada vez más competitivo y globalizado, es vital no solo retener a los clientes actuales, sino también identificar áreas de oportunidad para atraer a nuevos clientes.

Desde una perspectiva económica, trabajar para maximizar la satisfacción del cliente no solo favorece la lealtad hacia la marca, sino que también puede conducir a un incremento en la rentabilidad de la empresa. Un estudio publicado por la *American Marketing Association* demostró que las empresas que se centran en la satisfacción del cliente tienden a disfrutar de una mayor rentabilidad en comparación con aquellas que no lo hacen [6].

Una herramienta poderosa para lograr este objetivo es la adopción de tecnologías avanzadas de procesamiento de datos, como los modelos de atención, que permiten que el modelo se enfoque en diferentes partes de la entrada para cada palabra en la salida, permitiendo así una paralelización masiva que no es posible con las arquitecturas secuenciales típicas como LSTM y GRU [3].

Los *Transformers*, que son un tipo avanzado de modelo de atención, no requieren que los datos se procesen secuencialmente, lo que les permite realizar cálculos en paralelo y, por lo tanto, aprender de las relaciones entre todas las palabras en una secuencia de una vez [3]. El Modelo *Transformers* conduce al éxito porque es capaz de aprender dependencias de largo alcance entre palabras de una oración, facilitando así aplicaciones innovadoras en la personalización de servicios y comunicación con el cliente que pueden directamente impactar en su satisfacción y lealtad.

Dicha tecnología es esencial para muchas tareas de *NLP* (Procesamiento del Lenguaje Natural), sobre todo al permitir que el modelo comprenda el contexto de una palabra en una oración. Este instrumento permite al modelo centrarse en las palabras más relevantes de una oración al decodificar los *tokens* de salida, lo que resulta en mejoras significativas en tareas como traducción automática, análisis de sentimientos y respuesta a preguntas.

Los *Transformers* ofrecen varias ventajas significativas sobre enfoques anteriores en NLP:

- **Paralelización:** Pueden procesar todas las palabras en una sentencia al mismo tiempo, lo que reduce significativamente los tiempos de entrenamiento [3]. Esto se puede realizar debido a que la implementación de *Transformers* hace uso extensivo de *GPUs* (*Graphical Processing Units*). Las *GPUs* están diseñadas para realizar operaciones matemáticas en paralelo, lo que las hace considerablemente más rápidas que las CPUs para tareas que pueden paralelizarse, como es el caso de los cálculos realizados durante el entrenamiento de redes neuronales.
- **Capacidad para capturar dependencias a largo plazo:** Pueden capturar relaciones y dependencias entre palabras o frases que están distantes entre sí en un texto, lo cual es un desafío para las arquitecturas basadas en RNN debido al problema del desvanecimiento del gradiente [3].
- **Mejora del rendimiento en tareas NLP:** Han establecido nuevos estándares de rendimiento en una variedad de tareas NLP.
- **Mayor Ancho de Banda de Memoria:** Las *GPUs* modernas generalmente tienen un mayor ancho de banda de memoria en comparación con las CPUs, lo que permite una transferencia de datos más rápida y, por lo tanto, una aceleración significativa en el tiempo de entrenamiento de modelos grandes y complejos como los *Transformers* [10].

El entrenamiento eficiente de estos modelos avanzados se facilita por el uso de GPUs, que están diseñadas para realizar operaciones matemáticas en paralelo. Esto es debido a que las GPUs tienen muchos más núcleos de procesamiento que las CPUs, permitiendo la ejecución simultánea de múltiples operaciones y proporcionando un mayor ancho de banda de memoria. Esto permite una transferencia de datos más rápida y, por lo tanto, una aceleración significativa en el tiempo de entrenamiento de modelos grandes y complejos como los *Transformers* [10]. Las bibliotecas de *deep learning* como TensorFlow y PyTorch están optimizadas para ejecutarse en GPUs, mejorando así el rendimiento durante el entrenamiento de modelos [11].

1.3 Problema

Problema Práctico: En la era actual, caracterizada por una intensa digitalización y globalización, las empresas generan una gran cantidad de datos a través de sus interacciones con los clientes en diversas

plataformas digitales y redes sociales [4][6][7]. No obstante, uno de los desafíos cruciales que enfrentan muchas empresas es la incapacidad para destilar información valiosa de estas masivas fuentes de datos, específicamente, de las conversaciones que los clientes tienen con los servicios de atención al cliente. El objetivo principal de este estudio es analizar miles de estas interacciones y sintetizarlas para identificar las problemáticas principales que enfrentan los clientes. Así, las empresas pueden dirigir sus esfuerzos eficazmente hacia identificar las causas subyacentes de estas problemáticas y desarrollar soluciones estratégicas que tengan un impacto significativo en la experiencia del cliente, potenciando así su crecimiento sostenible y consolidando una ventaja competitiva en el mercado [5][8][9].

Problema Científico: Las herramientas clásicas de NLP a menudo se basan en enfoques matemáticos que, aunque potentes, suelen carecer de una comprensión profunda y semántica de los textos [6]. Este vacío puede ser abordado con la incorporación de modelos *Transformers*, cuya arquitectura, dotada de capas de atención, facilita una interpretación más rica y contextual de los datos, lo que se traduce en un análisis más detallado de las interacciones entre la empresa y el cliente [7]. La adopción de estos modelos avanzados puede ayudar a superar las limitaciones de los métodos de NLP tradicionales, brindando así nuevas oportunidades para extraer insights valiosos de los datos disponibles.

1.4 Objetivos

1.4.1 Objetivo general

Entrenar y adaptar un modelo de *Transformers* para proporcionar una visión detallada y profunda de las percepciones y preocupaciones de los clientes basada en sus interacciones directas con las empresas.

1.4.2 Objetivos específicos

Desarrollar y Ajustar el Modelo Preentrenado:

- Seleccionar un modelo preentrenado adecuado de Hugging Face para el análisis de las interacciones de los clientes. (los modelos son *open source* <https://huggingface.co/>).
- Ajustar el modelo preentrenado utilizando una muestra etiquetada de la base de datos del cliente para que pueda identificar y categorizar efectivamente las quejas y problemáticas de los clientes.

Analizar y categorizar las Interacciones de los Clientes:

- Utilizar el modelo ajustado para analizar y categorizar las interacciones de los clientes, identificando patrones recurrentes y áreas de oportunidad.
- Consolidar y resumir eficientemente las interacciones con los clientes para destacar las principales áreas de preocupación y oportunidades de mejora.

Evaluar el Desempeño del Modelo:

- Evaluar la efectividad y precisión del modelo en identificar y categorizar las principales áreas de preocupación y percepciones de los clientes.
- Establecer una meta de precisión de al menos el 80 % para las etiquetas que acumulativamente representen el 20 % más crítico de las categorías de problemáticas, en línea con la estrategia basada en el principio de Pareto.

Desarrollar la estrategia Práctica y Aplicable: Un Enfoque Basado en el principio de Pareto:

- Presentar una estrategia que permita a las empresas enfocar sus esfuerzos en las problemáticas que tienen un mayor impacto, en lugar de buscar una solución teóricamente perfecta pero prácticamente inalcanzable.
- Proponer una herramienta realista y aplicable que las empresas pueden implementar para mejorar significativamente la satisfacción del cliente y la eficiencia operativa, fundamentada en el principio de Pareto, que sugiere que el 80 % de los efectos provienen del 20 % de las causas.

2 Metodología

En esta fase de desarrollo del proyecto se aplicaron experimentos y simulaciones computacionales con el propósito de entrenar y validar al modelo seleccionado a fin de asegurar su óptimo rendimiento en la categorización de las quejas de los clientes. Este procedimiento comenzó con la descripción de los datos, el análisis exploratorio para observar características interesantes y ciertas tendencias, y la selección de modelos con base en experimentos.

2.1 Descripción de los datos

El conjunto de datos utilizado en este estudio proviene de una compañía que mantiene un registro detallado de sus interacciones de servicio al cliente. La base de datos consta de 160,000 registros distribuidos en 3 columnas, que incluyen variables como `complaint_id`, `complaint`, `label`. Esta base de datos es particularmente significativa, ya que no solo proporciona una visión numérica, sino también una comprensión profunda de las diversas quejas y retroalimentaciones de los clientes, codificadas en la columna `complaint`. Esta información es invaluable para identificar y abordar los problemas comunes que enfrentan los clientes.

Para visualizar el código, visite el repositorio de GitHub: [Clasificación de quejas](#).

Antes de su análisis, se llevó a cabo un riguroso proceso de limpieza de datos para garantizar la privacidad de los clientes, eliminando cualquier información confidencial como fechas, cuentas y cantidades. Los datos se procesaron para facilitar el análisis semántico, que es fundamental para entender las quejas de los clientes. Las variables clave analizadas incluyen diferentes categorías de quejas y su frecuencia de ocurrencia.

Etiqueta	Frecuencia
Tarjeta de crédito	15564
Reporte de crédito	91171
Cobranza	23145
Hipotecas y Préstamos	18990
Banca Minorista	13535

Tabla 2.1: Tabla de Análisis de Datos por Frecuencia

2.2 Análisis exploratorio

El análisis exploratorio de los datos se hizo con técnicas de procesamiento de lenguaje natural (NLP), según el *script* `first_clean.py`. Al inicio, se purgó profundamente el texto para eliminar caracteres no deseados como fechas o datos confidenciales.

Este preprocesamiento facilitó un análisis de frecuencia efectivo para identificar palabras y frases comunes en las quejas de los clientes. Además, se aplicaron gráficos y tablas para visualizar las tendencias y patrones en los datos, esta elección facilita la toma de decisiones informadas durante las fases posteriores del proyecto. Además, permite visualizar el número de quejas y etiquetas aplicadas hacia los datos sobre los clientes de una organización.

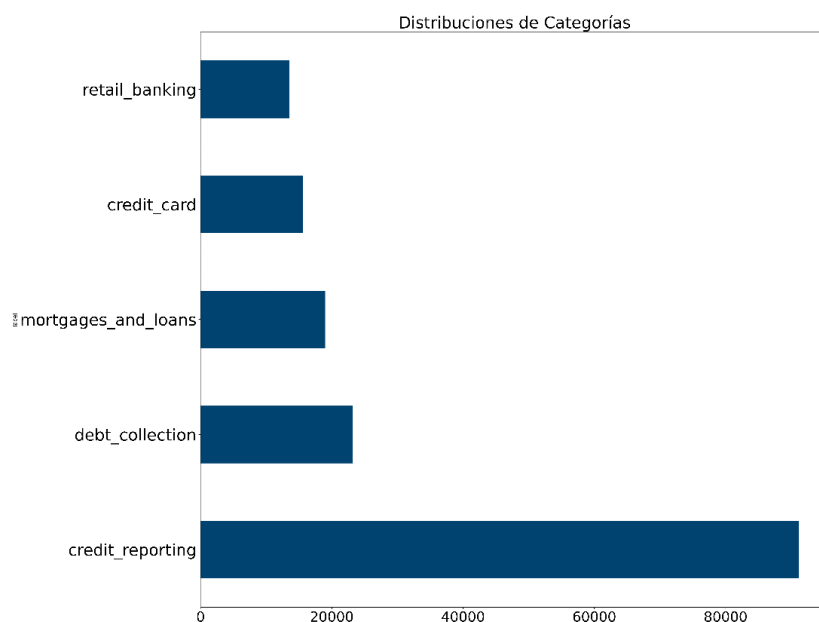


Figura 2.1: Distribución de Categorías basada en las Quejas de los Clientes. La gráfica de barras horizontales presenta las diferentes categorías de problemas identificadas a partir de las interacciones de los clientes

influyen y dependen entre sí.

Como se destacó en la Introducción (sección 1), los *Transformers* cuentan con un mecanismo de atención, permitiendo que el modelo se concentre en segmentos específicos de la entrada para cada palabra en la salida. En esta arquitectura es esencialmente compuesta por múltiples capas de autoatención y redes neuronales *feedforward*, lo cual permite ampliar la comprensión contextual sobre los datos y no avaluarse simplemente de la identificación de palabras por la frecuencia en que aparecen en una frase o párrafo.

2.3.1 Modelos de comparación

Regresión Logística:

- Tipo: Modelo lineal.
- Mecanismo: Calcula la probabilidad de una categoría o evento en función de una o más variables independientes.
- Uso: Comúnmente usado en problemas de clasificación binaria, aunque puede extenderse a clasificación multiclase.
- Ventajas: Fácil interpretación, eficiencia computacional, y buen rendimiento con *datasets* linealmente separables.
- Desventajas: Puede no funcionar bien con relaciones no lineales, no es adecuado para grandes números de características o clases.

SVM (Máquinas de Vectores de Soporte):

- Tipo: Modelo basado en márgenes.
- Mecanismo: Busca el hiperplano que mejor separa las clases en el espacio de características.
- Uso: Ampliamente utilizado para clasificación binaria y multiclase.
- Ventajas: Efectivo en espacios de alta dimensión, versátil a través del uso de diferentes funciones kernel.
- Desventajas: Requiere una buena elección de kernel, no es tan eficiente en conjuntos de datos muy grandes.

Random Forest:

- Tipo: Modelo basado en ensamble.
- Mecanismo: Combina múltiples árboles de decisión para mejorar la precisión y controlar el sobre ajuste.
- Uso: Útil para clasificación y regresión.

- Ventajas: Maneja bien diferentes tipos de datos, es robusto a *outliers*, y proporciona una medida de importancia de las características.
- Desventajas: Menos interpretable que un árbol de decisión individual, puede ser lento para entrenar con grandes datasets.

Red Neuronal:

- Tipo: Modelo basado en aprendizaje profundo.
- Mecanismo: Imita la forma en que los humanos aprenden, ajustando pesos en capas de neuronas interconectadas.
- Uso: Extremadamente versátil, utilizado en clasificación, regresión, generación de imágenes, NLP, entre otros.
- Ventajas: Altamente flexible y capaz de modelar relaciones complejas y no lineales.
- Desventajas: Requiere una gran cantidad de datos para entrenar, es computacionalmente intensivo, y menos interpretable.

2.4 Descripción de las métricas

Para evaluar el rendimiento de los modelos se aplicaron las métricas *precision recall* y el *F1-Score*. Este conjunto de indicadores proporciona una evaluación integral sobre el rendimiento del modelo, lo cual permite evaluar la exactitud de las predicciones, así como la capacidad del modelo para identificar correctamente las categorías reales de las quejas.

Por otro lado, también se utilizó la matriz de confusión, la cual consiste en una herramienta igual de valiosa utilizada en el modelo *Transformers*, proporcionando una representación visual clara de las predicciones del modelo en comparación con las etiquetas reales. Considerando que cada una de estas representan una queja que se expresa mediante el texto, por lo cual se tiene que lograr la identificación precisa de los caracteres con base en los vectores y estas palabras o símbolos se remiten a un significado.

Precisión (*Precision*):

- Definición: La precisión mide la proporción de predicciones correctas entre todas las predicciones positivas realizadas por el modelo.
- Fórmula 1:

$$Precision = \frac{VP}{VP + FP} \quad (2.1)$$

VP: Verdaderos Positivos; FP: Falsos Positivos.

- Interpretación: Una alta precisión indica que el modelo es efectivo en la identificación de una clase en particular con pocas falsas alarmas.

Recall (Sensibilidad o Tasa de Verdaderos Positivos):

- Definición: El *recall* mide la capacidad del modelo para identificar correctamente todas las instancias relevantes (positivas) de una clase.
- Fórmula 2:

$$Recall = \frac{VP}{VP + FN} \quad (2.2)$$

FN: Falsos Negativos.

- Interpretación: Un alto *recall* indica que el modelo es capaz de detectar la mayoría de los casos positivos, pero puede incluir más falsos positivos.

F1-Score:

- Definición: El *F1-Score* es una medida que combina la precisión y el *recall* en una sola métrica, esto hace que se proporcione un balance entre ambos elementos.
- Fórmula 3:

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.3)$$

- Interpretación: Un alto *F1-Score* sugiere un equilibrio adecuado entre precisión y *recall*, lo que es útil en situaciones donde ambos son igualmente importantes.

Matriz de Confusión:

- Definición: Una matriz de confusión es una herramienta que permite la visualización del rendimiento de un algoritmo de clasificación.
- Componentes:
 - Verdaderos Positivos (VP): Instancias correctamente clasificadas como positivas.
 - Falsos Positivos (FP): Instancias incorrectamente clasificadas como positivas.
 - Verdaderos Negativos (VN): Instancias correctamente clasificadas como negativas.
 - Falsos Negativos (FN): Instancias incorrectamente clasificadas como negativas.
- Interpretación: como se puede observar, la matriz de confusión proporciona una visión detallada de cómo el modelo clasifica cada clase y donde puede estar cometiendo errores. Esto implica que el algoritmo se puede precisar desde estos componentes, lo cual indica no solo el margen de error, sino también en qué medida, este puede cambiar.

2.5 Descripción de los experimentos o simulaciones

En esta fase, se realizaron experimentos y simulaciones computacionales para entrenar y validar el modelo seleccionado. El proceso implicó dividir los datos en conjuntos de entrenamiento y validación, configurar argumentos específicos de entrenamiento, y luego entrenar y evaluar el modelo utilizando el conjunto de validación. Estos experimentos fueron cruciales para afinar el modelo y garantizar su capacidad para proporcionar *insights* significativos basados en el análisis de las quejas de los clientes.

Entrenamiento del modelo

Paso 1: Preprocesamiento de Datos Antes de que los datos se alimenten al modelo RoBERTa, se realiza un preprocesamiento esencial que incluye la *tokenización* de texto. En este proceso, el texto se divide en unidades más pequeñas (*tokens*), que pueden ser palabras o sub palabras. Además, se añaden *tokens* especiales como [CLS] al principio de cada secuencia y [SEP] para separar diferentes oraciones en una secuencia [1].

Paso 2: Incorporación de Palabras (Embedding) Una vez que los datos están *tokenizados*, se alimentan al modelo, donde cada *token* se convierte en un vector de alta dimensión utilizando una capa de incrustación (*embedding*). Esta representación vectorial contiene información semántica sobre cada palabra [1].

Paso 3: Codificador (Encoder)

1. Capas de Atención Multi-Cabeza (*Multi-Head Attention Layers*)

Aquí, el modelo RoBERTa utiliza mecanismos de atención para ponderar la importancia de diferentes palabras en una oración durante el entrenamiento. En otras palabras, le permite al modelo ‘prestar atención’ a diferentes partes de la entrada para cada palabra en la oración [12].

2. Redes Feedforward

Después de las capas de atención, la información procesada pasa a través de redes neuronales feedforward, que pueden aprender representaciones complejas de los datos [12].

3. Capas Intermedias

El modelo repite estas operaciones (atención multi-cabeza y redes feedforward) varias veces (las versiones grandes de RoBERTa pueden tener docenas de estas ‘capas’ repetitivas) para aprender representaciones cada vez más abstractas y complejas de los datos de entrada [12].

Paso 4: Salida del Codificador Al final del codificador, obtenemos un conjunto de vectores que representan una comprensión profunda de cada palabra en el contexto de su oración.

Paso 5: Decodificador (Si es aplicable) En el caso de RoBERTa, generalmente se utiliza como un modelo de codificador único, sin una etapa de decodificación. Sin embargo, en tareas que requieren generación de texto, podrías anexar un decodificador para convertir las representaciones aprendidas de vuelta en texto humano-legible [12].

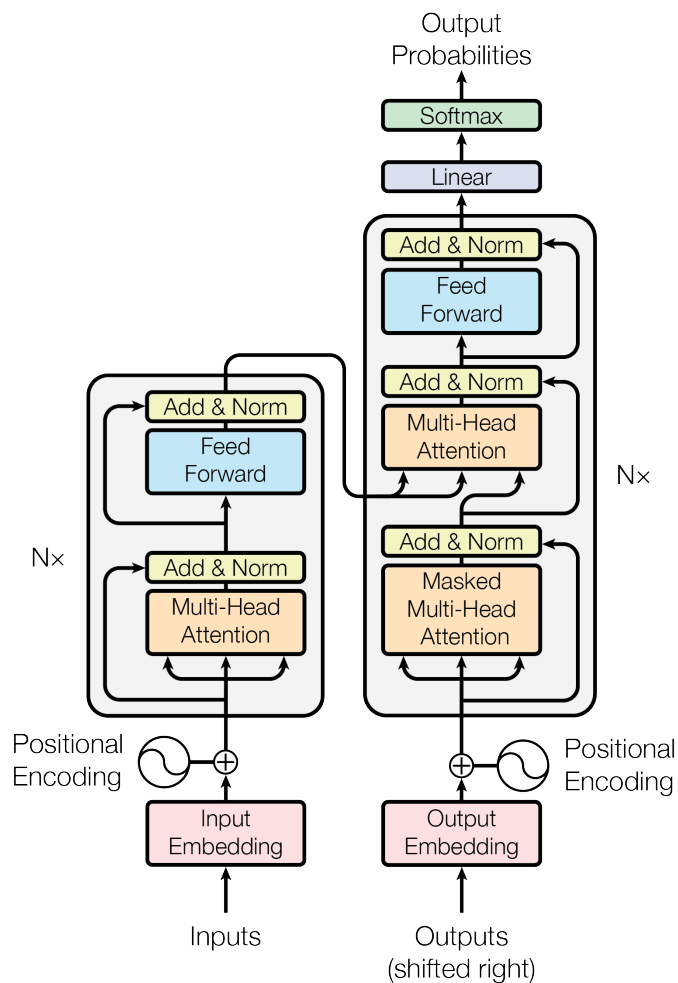


Figura 2.4: Diagrama de funcionamiento para el modelo *Transformer*, engloba el mecanismo de atención y con ello se remarcan los generadores de texto. En este esquema se especifica la entrada y salida del texto que tiene este tipo de funcionamiento. Tomado de [1].

3 Resultados y discusión

En este capítulo se presentan los resultados obtenidos del desarrollo de este trabajo y una discusión sobre el objeto de estudio.

3.1 Resultados

3.1.1 Transformer model

El modelo *Transformer* ha demostrado un desempeño sobresaliente en el procesamiento de lenguaje natural. Como se ilustra en la Figura 3.1, la distribución del procesamiento del modelo muestra una eficiencia notable en diversas categorías.

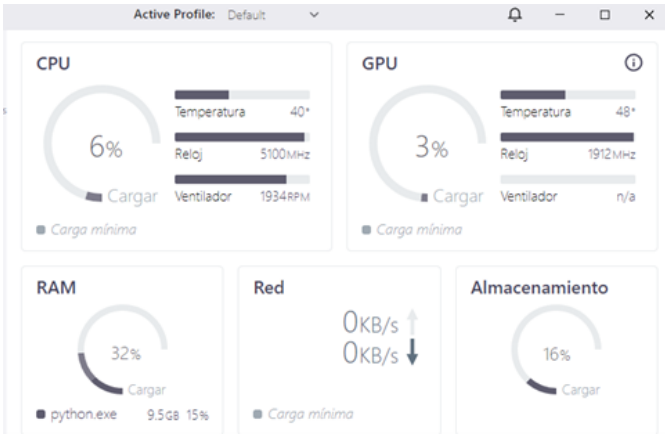


Figura 3.1: Distribución del procesamiento del modelo Transfomer.

La precisión general (*Accuracy*) del modelo *Transformer* es del 90.46 %. A continuación, se presenta una tabla detallada que muestra el desempeño del modelo en distintas categorías, incluyendo precisión (*Precision*), recuperación (*Recall*), puntaje F1 (*F1-Score*) y soporte (*Support*):

Además, la Tabla 3.2 presenta la precisión específica por categoría, lo que permite una comprensión más detallada del rendimiento del modelo en cada segmento.

Categoría	Precision	Recall	F1-Score	Support
Tarjeta de crédito	0.858	0.822	0.840	4801
Reporte de crédito	0.927	0.952	0.940	28197
Cobranza	0.849	0.814	0.831	7080
Hipotecas y Préstamos	0.901	0.841	0.870	5899
Banca Minorista	0.893	0.921	0.907	4023

Tabla 3.1: Desempeño del modelo

Categoría	Accuracy
Tarjeta de crédito	0.822
Reporte de crédito	0.952
Cobranza	0.814
Hipotecas y Préstamos	0.841
Banca Minorista	0.921

Tabla 3.2: Accuracy por etiqueta

3.1.2 Regresión Logística

La Regresión Logística también ofrece un rendimiento notable con una precisión global del 85 %. La Tabla 3.3 desglosa el desempeño de este modelo por categoría, mostrando detalles sobre precisión (*Precision*), recuperación (*Recall*), puntaje F1 (*F1-Score*) y soporte (*Support*).

Categoría	Precision	Recall	F1-Score	Support
Tarjeta de crédito	0.78	0.76	0.77	199
Reporte de crédito	0.88	0.93	0.90	1137
Cobranza	0.80	0.65	0.71	275
Hipotecas y Préstamos	0.80	0.82	0.81	217
Banca Minorista	0.88	0.80	0.84	172
Accuracy	0.85 (2000)			
Macro Avg	0.83	0.79	0.81	2000
Weighted Avg	0.85	0.85	0.85	2000

Tabla 3.3: Desempeño del modelo de Regresión Logística por categoría

3.1.3 Support Vector Machine

El modelo de Máquinas de Vectores de Soporte (SVM) también se evalúa, mostrando un rendimiento similar al de la Regresión Logística con una precisión global del 84 %. A continuación, en la Tabla 3.4 se detalla su rendimiento por categoría.

3.1.4 Random Forest

El modelo de Random Forest presenta un rendimiento inferior comparado con los modelos previamente mencionados, con una precisión global del 75 %. La Tabla 3.5 muestra el desempeño de este modelo en diferentes categorías.

Categoría	Precision	Recall	F1-Score	Support
Tarjeta de crédito	0.76	0.76	0.76	199
Reporte de crédito	0.90	0.91	0.90	1137
Cobranza	0.75	0.68	0.71	275
Hipotecas y Préstamos	0.76	0.81	0.78	217
Banca Minorista	0.84	0.80	0.82	172
Accuracy	0.84 (2000)			
Macro Avg	0.80	0.79	0.80	2000
Weighted Avg	0.84	0.84	0.84	2000

Tabla 3.4: Desempeño del modelo SVM por Categoría

Categoría	Precision	Recall	F1-Score	Support
Tarjeta de crédito	0.81	0.39	0.52	199
Reporte de crédito	0.72	0.99	0.83	1137
Cobranza	0.93	0.32	0.47	275
Hipotecas y Préstamos	0.86	0.54	0.67	217
Banca Minorista	0.82	0.56	0.67	172
Accuracy	0.75 (2000)			
Macro Avg	0.83	0.56	0.63	2000
Weighted Avg	0.78	0.75	0.72	2000

Tabla 3.5: Desempeño del modelo RF por Categoría

3.1.5 Red Neuronal

Finalmente, el rendimiento de la Red Neuronal se analiza mostrando una precisión general del 78 %. A continuación, se detallan las métricas específicas de este modelo por categoría.

Categoría	Precision	Recall	F1-Score	Support
Tarjeta de crédito	0.69	0.69	0.69	199
Reporte de crédito	0.87	0.86	0.87	1137
Cobranza	0.59	0.63	0.61	275
Hipotecas y Préstamos	0.69	0.71	0.70	217
Banca Minorista	0.77	0.68	0.72	172
Accuracy	0.78 (2000)			
Macro Avg	0.72	0.72	0.72	2000
Weighted Avg	0.78	0.78	0.78	2000

Tabla 3.6: Desempeño del modelo de Red Neuronal por Categoría

3.2 Discusión

- Rendimiento del modelo *Transformer*: Este modelo muestra un rendimiento sobresaliente con una precisión global del 90 %. En todas las categorías, el modelo *Transformer* mantiene una alta precisión, recall y puntuación F1, destacando en 'Reporte de crédito' con una precisión de 95 %.
- Comparación con Regresión Logística: La regresión logística, aunque

efectiva, es menos precisa que el modelo *Transformer* en todas las categorías. Su precisión global es del 85 %, notablemente más baja que la del modelo *Transformer*.

- Comparación con SVM (Máquinas de Vectores de Soporte): El modelo SVM muestra un rendimiento similar al de la regresión logística, con una precisión global del 84 %. Aunque competitivo, sigue siendo inferior al modelo *Transformer*.
- Comparación con Random Forest: El modelo de Random Forest muestra un rendimiento inferior en comparación con los otros modelos, con una precisión global del 75 %. Es importante destacar que muestra una baja recall en varias categorías, lo que indica que no es tan eficiente para identificar todas las instancias relevantes.
- Comparación con Red Neuronal: La red neuronal tiene un rendimiento general del 78 %, con un equilibrio razonable entre precisión y recall. Aunque muestra una competencia decente en todas las categorías, sigue siendo inferior al modelo *Transformer*.

4 Conclusiones y trabajo futuro

4.1 Conclusiones

Para terminar, resulta pertinente señalar diversos avances que se pueden lograr mediante esta programación, el modelo *Transformer* aporta múltiples cambios e innovaciones en el lenguaje y en plantear pautas para interactuar con la inteligencia que se despliega en la red. Ahora bien, es difícil recapitular todo lo que se puede hacer con este nuevo modelo, sin embargo, se puede adelantar, remarcando que conlleva beneficios, sobre todo en la publicidad y en aplicaciones sobre servicios financieros.

- Eficacia del modelo *Transformer*: El modelo *Transformer* demuestra ser el más eficaz para esta tarea de clasificación específica, superando a los modelos tradicionales y a la red neuronal en términos de *precision*, *recall* y puntuación F1.
- Equilibrio entre *precision* y *recall*: Aunque otros modelos como la regresión logística y SVM tienen un rendimiento decente, el modelo *Transformer* ofrece el mejor equilibrio entre precisión y recall, crucial para una clasificación efectiva.
- Importancia de la selección del modelo: Esta comparación resalta la importancia de seleccionar el modelo adecuado para tareas específicas. Mientras que los modelos más simples pueden ser suficientes para ciertas aplicaciones, el modelo *Transformer* ofrece ventajas significativas para tareas complejas de clasificación.
- Potencial para ajustes y mejoras: Aunque el modelo *Transformer* es superior, hay espacio para mejorar y ajustar los otros modelos, especialmente en aspectos como el balance de clases y la selección de características.
- Alta precisión global del modelo *Transformer*: Con una precisión global del 90.46 %, el modelo *Transformer* es claramente superior a los otros modelos. Esta alta precisión es crucial en un entorno empresarial, especialmente si la clasificación incorrecta tiene

implicaciones significativas, como en la gestión de quejas de clientes en el sector financiero.

- **Consistencia en diversas categorías:** El modelo *Transformer* no solo tiene una alta precisión global, sino que también mantiene un rendimiento consistente en varias categorías. Esto es esencial para un negocio que maneja múltiples tipos de consultas o quejas, asegurando un nivel uniforme de servicio al cliente.
- **Interpretabilidad y transparencia:** Aunque los modelos *Transformer* son potentes, pueden carecer de la interpretabilidad de modelos más simples como la regresión logística o los árboles de decisión. Esto puede ser un desafío en entornos donde las decisiones basadas en el modelo necesitan ser explicadas o justificadas, como en la conformidad regulatoria.
- **Costo computacional y eficiencia en tiempo real:** Implementar modelos *Transformer* puede ser computacionalmente más costoso que modelos más sencillos. Si el negocio requiere respuestas en tiempo real o tiene limitaciones de recursos computacionales, esto debe ser considerado.
- **Flexibilidad y escalabilidad:** Los modelos *Transformer* son generalmente más flexibles y escalables para manejar grandes volúmenes de datos y tareas más complejas. Esto puede ser una ventaja significativa en un entorno empresarial dinámico y en rápida evolución.
- **Inversión y ROI:** La implementación de un modelo *Transformer* puede requerir una inversión inicial más alta en términos de desarrollo y recursos computacionales. Sin embargo, el retorno de la inversión podría justificarse por la mejora en la precisión y eficiencia en la gestión de quejas de clientes.

Recomendación para la implementación en producción: Dado el alto rendimiento del modelo *Transformer* en términos de precisión, consistencia y equilibrio entre precisión y *recall*, sería la recomendación principal para implementar en un entorno de producción, especialmente si el negocio prioriza la precisión y la eficiencia en la clasificación de quejas de clientes. Sin embargo, es importante tener en cuenta la inversión y los recursos necesarios para su implementación, así como la posible necesidad de explicabilidad en las decisiones del modelo. En resumen, si el negocio puede soportar el costo y necesita la máxima precisión y eficiencia, el modelo *Transformer* sería la opción preferida.

4.2 Trabajo futuro

Para una propuesta sobre líneas de investigación en el futuro se concibe el análisis y la implementación del modelo *Transformer* para la clasificación de quejas de clientes en el sector financiero. En esta línea de investigación se abren varias vías para trabajos futuros, como el diseño y planeación de una empresa tanto en términos de investigación como de aplicaciones prácticas. A continuación, se detallan algunas recomendaciones y descripciones para futuros proyectos, enseguida se enumeran algunos puntos de exposición para otras investigaciones sobre aplicaciones de este modelo.

1. Mejora continua del modelo *Transformer*:
 - Optimización de hiperparámetros: llevar a cabo un análisis más profundo de los hiperparámetros del modelo *Transformer* para mejorar aún más su precisión y eficiencia.
 - Integración de datos actualizados: incorporar datos más recientes y variados para mejorar la generalización del modelo.
2. Exploración de modelos híbridos:
 - Combinar las fortalezas del modelo *Transformer* con modelos más interpretables, como la regresión logística, para mejorar la transparencia de las decisiones del modelo.
3. Análisis de sensibilidad y conformidad regulatoria:
 - Realizar un análisis de sensibilidad para entender mejor cómo diferentes tipos de entradas afectan las predicciones del modelo.
 - Asegurar que el modelo cumpla con las regulaciones y normativas vigentes, especialmente en lo que respecta a la explicabilidad y justicia de las decisiones automatizadas.
4. Desarrollo de interfaces de usuario intuitivas:
 - Diseñar dashboards y herramientas de visualización para que los no expertos puedan interpretar fácilmente los resultados del modelo.
5. Implementación en diferentes contextos:
 - Adaptar y probar el modelo en otros contextos o industrias para evaluar su versatilidad y eficacia en diferentes escenarios.
6. Investigación en reducción de sesgos:
 - Conducir investigaciones adicionales para identificar y mitigar cualquier sesgo inherente en el modelo, garantizando así la equidad en las predicciones.

7. Evaluación del impacto en la satisfacción del cliente:

- Medir cómo la implementación de este modelo afecta la percepción y satisfacción del cliente, ajustando el modelo según sea necesario para mejorar la experiencia del cliente.

8. Desarrollo de capacidades en tiempo real:

- Investigar la posibilidad de adaptar el modelo para aplicaciones en tiempo real, como chatbots o sistemas de respuesta automática.

9. Exploración de nuevas fuentes de datos:

- Integrar y experimentar con diferentes tipos de datos, como datos de redes sociales o transacciones financieras, para enriquecer el análisis.

10. Estudios de caso y documentación de éxitos/desafíos:

- Desarrollar estudios de caso detallados para documentar los éxitos y desafíos enfrentados durante la implementación del modelo, proporcionando aprendizajes clave para futuras implementaciones.

Este trabajo establece una base sólida para futuras investigaciones y desarrollos en el campo de la clasificación automatizada y la inteligencia artificial aplicada al servicio al cliente, con un potencial significativo para mejorar la eficiencia operativa y la experiencia del cliente en el sector financiero y más allá.

Bibliografía

- [1] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention is all you need," *Advances in neural information processing systems*, vol. 30, 2017.
- [2] E. W. Anderson, C. Fornell, and S. K. Mazvancheryl, "Customer satisfaction and shareholder value," *Journal of marketing*, vol. 68, no. 4, pp. 172–185, 2004.
- [3] T. Wolf, L. Debut, V. Sanh, J. Chaumond, C. Delangue, A. Moi, P. Cistac, T. Rault, R. Louf, M. Funtowicz, *et al.*, "Transformers: State-of-the-art natural language processing," in *Proceedings of the 2020 conference on empirical methods in natural language processing: system demonstrations*, pp. 38–45, 2020.
- [4] K. Stalmachova, R. Chinoracky, and M. Strenitzerova, "Changes in business models caused by digital transformation and the covid-19 pandemic and possibilities of their measurement—case study," *Sustainability*, vol. 14, no. 1, p. 127, 2021.
- [5] J. Bughin, B. O'Beirne, and J. Deakin, "The anatomy of successful digital transformation: The role of analytics," *Applied Marketing Analytics*, vol. 5, no. 2, pp. 121–128, 2019.
- [6] J. Reis and N. Melão, "Digital transformation: A meta-review and guidelines for future research," *Heliyon*, 2023.
- [7] J. Amankwah-Amoah, Z. Khan, G. Wood, and G. Knight, "Covid-19 and digitalization: The great acceleration," *Journal of business research*, vol. 136, pp. 602–611, 2021.
- [8] L. Cortellazzo, E. Bruni, and R. Zampieri, "The role of leadership in a digitalized world: A review," *Frontiers in psychology*, vol. 10, p. 1938, 2019.
- [9] L. T. Ha, "Impacts of digital business on global value chain participation in european countries," *AI & SOCIETY*, pp. 1–26, 2022.

- [10] NVIDIA, “Qué es un modelo transformer.” <https://la.blogs.nvidia.com/2022/04/19/que-es-un-modelo-transformer/>, Apr 2022. Último acceso: [30/1/24].
- [11] H. Askary, “Getting started with pytorch.” <https://hassanaskary.medium.com/getting-started-with-pytorch-ad50588a1270>, Feb 2018. Último acceso: [30/1/24].
- [12] T. Wolf, L. Debut, V. Sanh, J. Chaumond, C. Delangue, A. Moi, P. Cistac, T. Rault, R. Louf, M. Funtowicz, J. Davison, S. Shleifer, P. Platen, C. Ma, J. Jelinek, C. Xu, T. Le Scao, S. Gugger, M. Drame, Q. Lhoest, and A. M. Rush, “Hugging face’s transformers: State-of-the-art natural language processing.” <https://arxiv.org/abs/1910.03771>, Oct 2019. Último acceso: [30/1/24].