



UNIVERSIDAD NACIONAL PEDRO RUIZ GALLO
FACULTAD DE INGENIERÍA CIVIL, DE SISTEMAS Y DE ARQUITECTURA
ESCUELA PROFESIONAL DE INGENIERÍA DE SISTEMAS



TESIS

Implementación de chatbot basado en Inteligencia Artificial para el Área de
Admisión de una universidad pública en el norte del Perú

PARA OBTENER EL TÍTULO PROFESIONAL DE:

Ingeniero(a) de Sistemas

Autores:

Bach. Anaya Bocanegra, Milka Fiorella

Bach. Rodríguez Rivas, Pedro Martín

Asesor:

Msc. Ing. Capuñay Uceda, Oscar Efraín

Fecha de sustentación:

16 de marzo del 2026

Lambayeque, Perú

2026



UNIVERSIDAD NACIONAL PEDRO RUIZ GALLO

FACULTAD DE INGENIERÍA CIVIL, DE SISTEMAS Y DE ARQUITECTURA

ESCUELA PROFESIONAL DE INGENIERÍA DE SISTEMAS



TESIS

Implementación de chatbot basado en Inteligencia Artificial para el Área de
Admisión de una universidad pública en el norte del Perú

PARA OBTENER EL TÍTULO PROFESIONAL DE INGENIERO(A) DE SISTEMAS

APROBADO POR EL SIGUIENTE JURADO:

Msc. Ing Omar Wilton Saavedra Salazar
Presidente

Dr. Ing Juan Elias Villegas Cubas
Secretario

Msc. Ing Roberto Carlos Arteaga Lora
Vocal

Msc. Ing Capuñay Uceda, Oscar Efraín
Asesor

Lambayeque, Perú

2026



UNIVERSIDAD NACIONAL PEDRO RUIZ GALLO
FACULTAD DE INGENIERÍA CIVIL DE SISTEMAS Y DE ARQUITECTURA
UNIDAD DE INVESTIGACIÓN



ACTA DE SUSTENTACIÓN
N° 638-2026-UI-FICSA



Siendo las 12:30 pm horas del día 16 de marzo del 2026, se reunieron de manera presencial los miembros de jurado de la tesis titulada: **"IMPLEMENTACIÓN DE CHATBOT BASADO EN INTELIGENCIA ARTIFICIAL PARA EL ÁREA DE ADMISIÓN DE UNA UNIVERSIDAD PÚBLICA EN EL NORTE DEL PERÚ"** con código N° IS_V_2024_024, designado por Resolución Decanal N° 727-2024-UNPRG-FICSA con la finalidad de Evaluar y Calificar la sustentación de la tesis antes mencionada, conformado por los siguientes docentes:

MSC. ING. OMAR WILTON SAAVEDRA SALAZAR
DR. ING. JUAN ELIAS VILLEGAS CUBAS
MSC. ING. ROBERTO CARLOS ARTEAGA LORA

PRESIDENTE
SECRETARIO
VOCAL

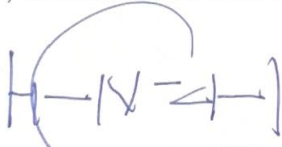
Asesorado por MSC. ING. OSCAR EFRAIN CAPUÑAY UCEDA.

El acto de sustentación fue autorizado por OFICIO VIRTUAL N° 61-2026-UIFICSA, la tesis fue presentada y sustentada por los Bachilleres: **MILKA FIORELLA ANAYA BOCANEGRA Y PEDRO MARTIN RODRIGUEZ RIVAS**, tuvo una duración de 60 minutos. Después de la sustentación, y absueltas las preguntas y observaciones de los miembros del jurado; se procedió a la calificación respectiva:

	NUMERO	LETRAS	CALIFICATIVO
MILKA FIORELLA ANAYA BOCANEGRA	<u>16</u>	<u>Dieciseis</u>	<u>BUENO</u>
PEDRO MARTIN RODRIGUEZ RIVAS	<u>16</u>	<u>Dieciseis</u>	<u>BUENO</u>

Por lo que quedan **APTOS** para obtener el Título Profesional de **INGENIERO (A) DE SISTEMAS** de acuerdo con la Ley Universitaria 30220 y la normatividad vigente de la Facultad de Ingeniería Civil De Sistemas y de Arquitectura de la Universidad Nacional Pedro Ruiz Gallo.

Siendo las 1:30 Se dio por concluido el presente acto académico, dándose conformidad al presente acto, con la firma de los miembros del jurado.


MSC. ING. OMAR WILTON SAAVEDRA SALAZAR
PRESIDENTE


DR. ING. JUAN ELIAS VILLEGAS CUBAS
SECRETARIO


MSC. ING. ROBERTO CARLOS ARTEAGA LORA
VOCAL



CONSTANCIA DE VERIFICACION DE ORIGINALIDAD

Yo, Mg. Ing. Óscar Efraín Capuñay Uceda, usuario revisor de Tesis ☒
Trabajo de Suficiencia Profesional ☐ y/o Trabajo Académico ☐

Titulado: "Implementación de chatbot basado en Inteligencia Artificial para el Área de Admisión de una universidad pública en el norte del Perú".

Cuyo autor (es) son: Bach. Anaya Bocanegra Milka Fiorella; con DNI N°77228606. y Bach. Rodríguez Rivas Pedro Martín; con DNI N° 74770755; declaro que la evaluación realizada por el Programa informático, ha arrojado un porcentaje de similitud 11%, verificables en el Resumen del Reporte automatizado de similitudes que se acompaña.

El suscrito (a) analizó reporte y concluyó que cada una de las coincidencias detectadas dentro del porcentaje de similitud permitido no constituyen plagio y que el documento cumple con la integridad científica y con las normas para el uso de citas y referencias establecidas en los protocolos respectivos,

Se cumple con adjuntar el Recibo Digital a efectos de la trazabilidad respectiva del proceso.

Lambayeque, 01 de julio del 2026



Firma (Asesor)

Nombres y Apellidos: Mg. Ing. Óscar Efraín Capuñay Uceda.

DNI ...16717150...

Defina la modalidad con (X)

Adjunta:

Resumen de Reporte automatizado de similitudes

Recibo digital

Informe-final-16-10-2025

INFORME DE ORIGINALIDAD

11 %	11 %	2 %	4 %
INDICE DE SIMILITUD	FUENTES DE INTERNET	PUBLICACIONES	TRABAJOS DEL ESTUDIANTE



FUENTES PRIMARIAS

1	hdl.handle.net Fuente de Internet	
2	www.unprg.edu.pe Fuente de Internet	1 %
3	revistas.unsm.edu.pe Fuente de Internet	1 %
4	repositorio.unprg.edu.pe Fuente de Internet	1 %
5	uvadoc.uva.es Fuente de Internet	1 %
6	myscale.com Fuente de Internet	1 %
7	repositorio.unprg.edu.pe:8080 Fuente de Internet	<1 %
8	tomberia.blogspot.com Fuente de Internet	<1 %
9	alicia.concytec.gob.pe Fuente de Internet	<1 %
10	www.coursehero.com Fuente de Internet	<1 %
11	repositorio.ucv.edu.pe Fuente de Internet	<1 %
12	Submitted to Universidad Nacional Pedro Ruiz Gallo Trabajo del estudiante	<1 %
13	repositorio.uladech.edu.pe Fuente de Internet	<1 %

1
Firma (Asesor)
Mg. Ing. Óscar Efraín Capuñay Uceda.
16717150

14	Bello González, Iván Darío. "Metodología Para la Elección de Modelos de Lenguaje de Gran Tamaño (LLMS) en Aplicaciones de Predicción Meteorológicas a Través de Algoritmos de Aprendizaje Automático Sobre Entornos de Computación en la Nube y Capturas de Datos a Través de IoT", Universidad Distrital Francisco José de Caldas (Colombia) Publicación	<1 %
15	oa.upm.es Fuente de Internet	 Firma (Asesor) Mg. Ing. Óscar Efraín Capuñay Uceda. 16717150 <1 %
16	Submitted to Universidad San Francisco de Quito Trabajo del estudiante	<1 %
17	Submitted to Instituto Superior de Artes, Ciencias y Comunicación IACC Trabajo del estudiante	<1 %
18	pt.scribd.com Fuente de Internet	<1 %
19	repositorio.unimayor.edu.co Fuente de Internet	<1 %
20	risti.xyz Fuente de Internet	<1 %
21	tesis.pucp.edu.pe Fuente de Internet	<1 %
22	Submitted to uncedu Trabajo del estudiante	<1 %
23	vsip.info Fuente de Internet	<1 %
24	www.paot.org.mx Fuente de Internet	<1 %
25	www.grupoarsa.com Fuente de Internet	<1 %
26	Submitted to National Institute of Technology, Patna Trabajo del estudiante	<1 %

27	repositorio.utc.edu.ec Fuente de Internet	<1 %
28	Submitted to unbosque Trabajo del estudiante	<1 %
29	digibuo.uniovi.es Fuente de Internet	<1 %
30	cerebrodigital.net Fuente de Internet	<1 %
31	www.slideshare.net Fuente de Internet	<1 %
32	prezi.com Fuente de Internet	<1 %
33	repositorio.unheval.edu.pe Fuente de Internet	<1 %
34	Submitted to uazuay Trabajo del estudiante	<1 %
35	www.datosabiertos.gob.pe Fuente de Internet	<1 %
36	1library.co Fuente de Internet	<1 %
37	aulavirtual.unprg.edu.pe Fuente de Internet	<1 %
38	eprints.uanl.mx Fuente de Internet	<1 %
39	foreignpolicywatchdog.com Fuente de Internet	<1 %
40	www.revfine.com Fuente de Internet	<1 %
41	www.rsm.global Fuente de Internet	<1 %
42	www.tdx.cat Fuente de Internet	<1 %

Firma (Asesor)
Mg. Ing. Óscar Efraín Capuñay Uceda.
16717150



Recibo digital

Este recibo confirma que su trabajo ha sido recibido por Turnitin. A continuación podrá ver la información del recibo con respecto a su entrega.

La primera página de tus entregas se muestra abajo.

Autor de la entrega: Pedro Martin Rodriguez Rivas
Título del ejercicio: Informe final 2024
Título de la entrega: Informe-final-16-10-2025
Nombre del archivo: n_Inteligencia_Artificial_para_el_rea_de_Admisi_n-16-10-2025.p...
Tamaño del archivo: 7.03M
Total páginas: 71
Total de palabras: 12,962
Total de caracteres: 83,925
Fecha de entrega: 16-oct-2025 11:03a. m. (UTC-0500)
Identificador de la entrega: 2782254672

UNIVERSIDAD NACIONAL PEDRO RUIZ GALLO
FACULTAD DE INGENIERÍA CIVIL, DE SISTEMAS Y DE ARQUITECTURA
ESCUELA PROFESIONAL DE INGENIERÍA DE SISTEMAS

TESIS

Implementación de chatbot basado en Inteligencia Artificial para el Área de
Admisión de cursos a universidades públicas en el entorno del Perú

PARA OBTENER EL TÍTULO PROFESIONAL DE:
Ingeniería de Sistemas

Autores:
Bach. Anya Bouanigue, Milica Florilla
Bach. Rodriguez Rivas, Pedro Martin

Asesor:
Mg. Ing. Capuñay Uceda, Oscar Efraín

Fecha de sustentación:
16 de octubre del 2025

Lambayeque, Perú
2025

Firma (Asesor)
Mg. Ing. Oscar Efraín Capuñay Uceda.
16717150

DEDICATORIA

A Dios, por ser mi guía, darme fortaleza en cada paso del camino, darme perseverancia para permitirme alcanzar esta meta.

A mi mamá, por su amor incondicional, su apoyo constante y por creer siempre en mí. Gracias por acompañarme en cada etapa y ser mi mayor inspiración. A mis hermanas y hermanos, por estar presentes cada vez que los he necesitado, por sus palabras de aliento y por motivarme a seguir adelante, incluso en los momentos más difíciles. A mis sobrinos, por inspirarme cada día a ser una mejor persona y recordarme la importancia de luchar por mis sueños.

A mi abuelita, que desde el cielo me cuida y me protege, aunque ya no esté físicamente, su amor y sus enseñanzas siempre me acompañan. A Norbil, gracias por creer en mí, motivarme a seguir adelante y acompañarme en cada paso de este camino. Y a Sky, por estar a mi lado en cada amanecida, brindándome su compañía en tantas horas de esfuerzo y dedicación.

Anaya Bocanegra Milka Fiorella

Dedico este trabajo, en primer lugar, a Dios, por darme la fortaleza, la sabiduría y la perseverancia necesarias para culminar esta importante etapa de mi vida.

A mis padres, por su amor incondicional, esfuerzo, sacrificio y apoyo constante. Gracias por creer en mí, por brindarme los valores que me han guiado y por impulsarme a alcanzar cada una de mis metas. Este logro también les pertenece. A la persona que ha compartido este camino conmigo, por su compañía, comprensión y constante motivación. Su apoyo fue un impulso importante para seguir adelante y culminar este proyecto.

Y a mi yo del pasado, por no rendirse ante las dificultades, por mantener la esperanza en los momentos de incertidumbre y por seguir adelante con la convicción de que el esfuerzo y la perseverancia harían posible alcanzar este logro.

Rodríguez Rivas Pedro Martín

AGRADECIMIENTOS

En primer lugar, agradezco a Dios por ser mi guía, brindarme la fortaleza y la sabiduría necesarias para alcanzar esta meta.

A mi familia y a mis amigos por su apoyo incondicional, su confianza y el ánimo que me brindaron a lo largo de estos años, siendo un pilar fundamental para culminar esta importante etapa de mi formación profesional. Agradezco a Norbil, por su apoyo constante, sus palabras de ánimo y por motivarme a dar siempre lo mejor de mí.

Finalmente, expreso mi sincero agradecimiento a mi asesor, el Ing. Oscar Capuñay Uceda, por su paciencia, orientación, valiosos consejos y constante disposición para resolver cada duda e inquietud surgida durante el desarrollo de este proyecto. Agradezco a la Universidad Nacional Pedro Ruiz Gallo, en especial al Área de Admisión, por proporcionarnos la información necesaria para la elaboración del banco de preguntas, el cual constituyó la base fundamental para el desarrollo de este trabajo.

Anaya Bocanegra Milka Fiorella

Expreso mi más sincero agradecimiento a Dios, por darme la fortaleza, la sabiduría y la perseverancia necesarias para alcanzar esta importante meta.

A mis padres, por su amor incondicional, esfuerzo, sacrificio y apoyo constante. Gracias por creer siempre en mí y ser el principal impulso para cumplir este sueño. A la persona que ha compartido este camino conmigo, por su compañía, comprensión, paciencia y apoyo incondicional durante cada etapa de este proceso.

Mi especial agradecimiento a mi asesor de tesis, Ing. Oscar Efraín Capuñay Uceda, por su orientación, paciencia y valiosos conocimientos, que fueron fundamentales para el desarrollo y culminación de esta investigación. A la Universidad Nacional Pedro Ruiz Gallo y a todas las personas que, de una u otra manera, contribuyeron al desarrollo de esta investigación y me acompañaron durante este importante proceso.

Rodríguez Rivas Pedro Martín

RESUMEN

La presente investigación tuvo como objetivo desarrollar e implementar un chatbot basado en inteligencia artificial para el Área de Admisión de la Universidad Nacional Pedro Ruiz Gallo, con la finalidad de mejorar la eficiencia y la calidad de las respuestas brindadas a las consultas frecuentes de los postulantes. La problemática se originó debido al incremento de consultas relacionadas con los procesos de admisión, lo que generaba demoras en la atención, sobrecarga del personal administrativo y posibles inconsistencias en la información proporcionada.

Para el desarrollo de la solución se empleó el modelo de lenguaje GPT-3.5 Turbo de OpenAI, integrado mediante una arquitectura REST desarrollada en Node.js y conectada a WhatsApp Cloud API de Meta para facilitar la interacción con los usuarios. Como fuente de conocimiento se utilizaron las preguntas frecuentes proporcionadas por el Área de Admisión y el prospecto oficial de admisión de la universidad, los cuales fueron sometidos a procesos de limpieza, normalización y validación para garantizar la calidad de la información utilizada en el entrenamiento del sistema. La evaluación del chatbot se realizó mediante un proceso de validación comparativa entre las respuestas generadas y la información oficial institucional, considerando criterios de exactitud, relevancia y claridad. Los resultados evidenciaron una precisión del 75 % en las respuestas correctas, además de una reducción significativa en los tiempos de atención al proporcionar respuestas en pocos segundos y sin intervención humana. Asimismo, se identificó que las respuestas más precisas se obtienen cuando la información de entrenamiento se encuentra estructurada en formato de texto plano.

Se concluye que la implementación del chatbot constituye una solución técnicamente viable y funcional para optimizar la atención de consultas en el Área de Admisión, mejorando la eficiencia operativa, la calidad de la información brindada y la experiencia de los usuarios, además de representar una alternativa de bajo costo y escalable para futuras implementaciones institucionales.

Palabras clave: inteligencia artificial, chatbot, GPT-3.5 Turbo, admisión universitaria, WhatsApp Cloud API, eficiencia, calidad.

ABSTRACT

The objective of this research was to develop and implement an artificial intelligence-based chatbot for the Admissions Office at Pedro Ruiz Gallo National University, with the aim of improving the efficiency and quality of responses to applicants' frequently asked questions. The problem arose due to an increase in inquiries related to the admissions process, which led to delays in response times, an overload on administrative staff, and potential inconsistencies in the information provided.

To develop the solution, OpenAI's GPT-3.5 Turbo language model was used, integrated via a REST architecture developed in Node.js and connected to Meta's WhatsApp Cloud API to facilitate interaction with users. The FAQs provided by the Admissions Office and the university's official admissions brochure were used as knowledge sources; these were subjected to data cleaning, normalization, and validation processes to ensure the quality of the information used to train the system. The chatbot was evaluated through a comparative validation process that compared the generated responses with official institutional information, taking into account criteria such as accuracy, relevance, and clarity. The results showed a 75% accuracy rate for correct responses, as well as a significant reduction in response times, with answers provided in just a few seconds and without human intervention. It was also found that the most accurate responses are obtained when the training data is structured in plain text format.

It is concluded that the implementation of the chatbot constitutes a technically viable and functional solution for optimizing the handling of inquiries in the Admissions Office, improving operational efficiency, the quality of the information provided, and the user experience, while also representing a low-cost and scalable alternative for future institutional implementations.

Keywords: artificial intelligence, chatbot, GPT-3.5 Turbo, university admissions, WhatsApp Cloud API, efficiency, quality.

Contenido

1.	PLANTEAMIENTO DE LA INVESTIGACIÓN	17
2.1.	Síntesis de la situación problemática	17
2.2.	Formulación del problema de investigación.....	18
2.4.1.	Objetivo general	18
2.4.2.	Objetivos específicos.....	18
3.	DISEÑO TEÓRICO	19
3.1.	Antecedentes.....	19
3.2.	Bases teóricas.....	22
3.2.1.	ChatBot:	22
3.2.2.	Modelo GPT:.....	22
3.2.3.	Modelo GPT 3.5 Turbo:.....	23
3.2.4.	IA (Inteligencia Artificial):.....	23
3.2.5.	Procesamiento del Lenguaje Natural (NLP).....	23
3.2.6.	Machine Learning	23
3.2.7.	Large language models (LLM)	24
3.2.8.	Retrieval-augmented generation (RAG)	24
3.2.9.	Eficiencia.....	24
3.2.10.	Calidad	25
3.2.11.	Base de datos Vectorial.....	25
3.2.12.	Chunking o Fragmentación.....	25
4.	DISEÑO METODOLÓGICO	25
4.1.	Tipificación de la investigación	25
4.2.	Población y muestra	26
4.2.1.	Población:	26
4.2.2.	Muestra:	26
4.3.	Técnicas, instrumentos, equipos y materiales.....	26
4.3.1.	Técnicas:.....	26
4.3.2.	Instrumentos.....	26
4.3.3.	Equipos	27
5.	DESARROLLO DEL PROYECTO.....	27
5.1.	Limpieza de datos	27
5.2.	Comparativa de modelos y elección	30

5.3.	Requisitos y preparación.....	33
5.3.1.	Requisitos previos.....	33
5.3.2.	Preparación del entorno de trabajo	34
5.3.3.	Factibilidad del proyecto.....	35
5.4.	Arquitectura del chatbot.....	39
5.5.	Creación del asistente en Open AI:.....	42
5.6.	Creación del Vector store en Open AI:	45
5.7.	Pruebas internas en la interfaz de asistentes de Open AI	45
5.8.	Entrenamiento y adaptación del modelo:.....	47
5.9.	Creación y configuración Api Meta Cloud	48
5.10.	Desarrollo del backend para la API	55
5.11.	Despliegue de API en GitHub y Render.....	56
5.11.1.	Publicación del código en GitHub	56
5.11.2.	Despliegue en Render	56
5.11.3.	Configuración de variables de entorno	58
5.11.4.	Creación del token permanente de Meta Cloud	59
5.11.5.	Validación de despliegue.....	61
5.12.	Evaluación y comparación de resultados	61
6.	RESULTADOS	69
7.	CONCLUSIONES.	80
8.	RECOMENDACIONES	81
9.	BIBLIOGRAFÍA	82

Figura 1 Ejemplo de consulta del documento preguntas frecuentes	28
Figura 2 Ejemplo de consulta luego del proceso de limpieza	29
Figura 3 Componentes Open AI	36
Figura 4 Arquitectura del Chatbot	40
Figura 5 Configuración nombre del asistente	42
Figura 6 Configurar instrucciones de asistente	43
Figura 7 Modelo de asistente	43
Figura 8 Carga de archivos en el asistente	44
Figura 9 Configuración del modelo.....	45
Figura 10 Creación de vector store en Open AI	45
Figura 11 Pruebas internas en la interfaz	46
Figura 12 Logs de conversación	47
Figura 13 Configurar nombre de aplicación	49
Figura 14 Configurar casos de uso	49
Figura 15 Configuración portafolio empresarial	50
Figura 16 Registro de requisitos	50
Figura 17 Aplicación ChatBot registrada	51
Figura 18 Configuración de la Aplicación	51
Figura 19 Configuración de token de acceso temporal.....	52
Figura 20 Configuración de URL pública	52
Figura 21 Configuración de Webhook.....	53
Figura 22 Prueba mensaje inicial Chatbot.....	54
Figura 23 Prueba inicial del contexto de la Universidad	54
Figura 24 Prueba con contexto no universitario	54
Figura 25 Prueba mensaje de despedida	55
Figura 26 Configuración del repositorio y rama	57
Figura 27 Configuración comando de inicio.....	57
Figura 28 Log de despliegue en Render.....	58
Figura 29 Configuración de variables de entorno	59
Figura 30 Usuario del sistema en Meta for Developers.....	60
Figura 31 Configuración de token permanente	60
Figura 32 Validación de despliegue con pruebas de usuario	61
Figura 33 Evaluación Chatbot - ¿Cuándo es el examen de admisión?.....	62
Figura 34 Respuesta documento preguntas frecuentes - ¿Cuándo es el examen de admisión?.....	62
Figura 35 Evaluación Chatbot - Tipo de pago beca 18.....	62
Figura 36 Respuesta documento preguntas frecuentes - Tipo de pago beca 18.....	62
Figura 37 Evaluación Chatbot - Medios de pago	63
Figura 38 Respuesta documento preguntas frecuentes - Medios de pago.....	63
Figura 39 Evaluación Chatbot - Documento para realizar el pago	63
Figura 40 Respuesta documento preguntas frecuentes - Documento para realizar el pago	63
Figura 41 Evaluación Chatbot - Pago por pagalo.pe.....	63
Figura 42 Respuesta documento preguntas frecuentes - Pago por pagalo.pe	64
Figura 43 Evaluación Chatbot - Pasos para la inscripción	64
Figura 44 Respuesta documento preguntas frecuentes - Pasos para la inscripción	64
Figura 45 Evaluación Chatbot - Horario de atención OTI	64

Figura 46	Respuesta documento preguntas frecuentes - Pasos para la inscripción	65
Figura 47	Evaluación Chatbot - Vacantes para la carrera de Ingeniería de Sistemas	65
Figura 48	Respuesta prospecto – Vacantes para la carrera de Ingeniería de Sistemas	65
Figura 49	Evaluación Chatbot - Encargados del proceso	65
Figura 50	Respuesta prospecto - Encargados del proceso	66
Figura 51	Evaluación Chatbot - Cantidad de preguntas para ingeniería	66
Figura 52	Respuesta prospecto - Cantidad de preguntas para ingeniería	66
Figura 53	Evaluación Chatbot - Proceso para inscribirse y rendir el examen	67
Figura 54	Respuesta prospecto - Proceso de inscripción	67
Figura 55	Respuesta prospecto - Proceso para rendir el examen	67
Figura 56	Evaluación Chatbot - Fecha de pago e inscripción.....	68
Figura 57	Respuesta prospecto - Fecha de pago e inscripción.....	68
Figura 58	Evaluación Chatbot - Montos de pago	68
Figura 59	Respuesta prospecto – Montos de pago.....	69

Tabla 1	Comparativa de modelos.....	31
Tabla 2	Comparativa de modelos de integración.....	32
Tabla 3	Descripción de componentes	41
Tabla 4	Criterios de análisis	70
Tabla 5	Comparativa respuestas Chatbot	71

1. PLANTEAMIENTO DE LA INVESTIGACIÓN

2.1. Síntesis de la situación problemática

En la actualidad, el avance constante de las tecnologías de la información y comunicación (TIC) ha transformado múltiples sectores mediante la automatización y optimización de procesos que antes se realizaban de forma manual. Ámbitos como el comercio, la banca y la salud ya han incorporado soluciones tecnológicas que mejoran la eficiencia y la calidad del servicio

En el sector académico, particularmente en el área de admisión de la Universidad Nacional Pedro Ruiz Gallo enfrenta un incremento significativo de consultas frecuentes de personas interesados en el proceso de admisión cuando dicho proceso está próximo a desarrollarse. Estas consultas, que abarcan desde requisitos hasta fechas importantes, requieren una cantidad significativa de tiempo y recursos humanos para ser respondidas en los diferentes canales de atención. Esto puede resultar en demoras en la atención, respuestas inconsistentes y es una tarea adicional para el personal administrativo. La falta de un sistema automatizado y eficiente para manejar estas consultas afecta la experiencia de los futuros estudiantes y la eficiencia operativa del área de admisión.

La implementación de un chatbot basado en inteligencia artificial, específicamente utilizando el modelo Chat GPT-3.5 Turbo, tiene el potencial de mejorar significativamente la situación. Este chatbot puede proporcionar respuestas rápidas y precisas a las preguntas frecuentes, además ayudaría a disminuir la carga al personal para que se enfoque en tareas más complejas y mejorar la satisfacción de los estudiantes al recibir información oportuna y consistente.

2.2. Formulación del problema de investigación

¿Cómo un chatbot basado en Chat GPT-3.5 Turbo ayudaría a mejorar la eficiencia y calidad de las respuestas a las preguntas frecuentes del área de admisión de la Universidad Nacional Pedro Ruiz Gallo?

2.3. Hipótesis

El uso de un chatbot inteligente, optimizará los tiempos de respuesta y aumentará la cobertura de atención de las preguntas frecuentes.

2.4. Objetivos

2.4.1. Objetivo general

“Desarrollar un chatbot basado en inteligencia artificial para el área de admisión de la Universidad Nacional Pedro Ruiz Gallo, con el fin de mejorar la eficiencia y la calidad de las respuestas a las preguntas frecuentes, reducir el tiempo de respuesta y aumentar la consistencia y precisión de la información proporcionada.”

2.4.2. Objetivos específicos

- Investigar y comparar diferentes tecnologías de chatbot basadas en inteligencia artificial.
- Recopilar, analizar y categorizar las preguntas más comunes realizadas en el área de admisión.
- Diseñar e implementar un chatbot utilizando tecnologías de Inteligencia Artificial.
- Entrenar el chatbot con las preguntas frecuentes identificadas y sus respuestas correspondientes.
- Medir y analizar la calidad y eficiencia de las respuestas del chatbot.

3. DISEÑO TEÓRICO

3.1. Antecedentes

Los chatbots han surgido como soluciones clave en diversos sectores, incluidas las instituciones académicas, debido a su capacidad para ejecutar de manera automática los procesos para la atención a los usuarios que tienen consultas sobre ciertos datos en partes de los procesos educativos y proporcionar respuestas rápidas, precisas y coherentes a las consultas frecuentes.

En el ámbito de la educación superior, los chatbots pueden contribuir con las mejoras en cuanto a la eficiencia de algunas áreas. Por ejemplo, el área de admisión ya que ayuda a reducir los tiempos de respuesta y aumentar la coherencia de la información proporcionada a los usuarios interesados en el proceso de admisión.

La investigación sobre la implementación de chatbots en el sector educativo ha mostrado resultados muy satisfactorios respecto a la eficiencia y satisfacción del usuario. Por ejemplo, recientemente en un estudio en el cual desarrollaron y evaluaron un chatbot usando el modelo de lenguaje GPT-3.5-Turbo para responder preguntas frecuentes de estudiantes universitarios, según Espinosa-Luna et al. (2023), indica que:

El uso de asistentes virtuales es común en los procesos de atención de distintas universidades alrededor del mundo. En la actualidad, se ofrecen varias herramientas de inteligencia artificial como procesadores de lenguaje natural para el desarrollo más eficiente de este servicio. (p. 1)

Según el estudio realizado por Mariani et al., (2023), podemos constatar que, en la actualidad, la IA continúa incrementando su presencia en la industria y en el día a día de las personas. Bajo ese contexto, “una forma específica de IA que está creciendo en relevancia tanto en la práctica como en la investigación es la IA conversacional, también conocida como CAs.” (p. 2)

Además, Ray (2023) nos proporciona una revisión dinámica sobre el uso de GPT-3 de la compañía OpenAi, destacando su capacidad para manejar consultas frecuentes de manera efectiva. El estudio hace énfasis en que, “tanto GPT-2 como GPT-3 han demostrado capacidades notables para generar texto de alta calidad y realizar una amplia gama de tareas de PLN.” (p. 127)

Wang et al. (2022) investigaron cómo la inteligencia artificial, incluyendo los chatbots, puede crear agilidad en servicios empresariales. En su estudio, se indica que los chatbots pueden agilizar procesos administrativos en las universidades, mejorando la eficiencia y reduciendo la carga de trabajo del personal. Esta investigación es particularmente relevante para el área de admisión, donde la agilidad en el manejo de consultas es esencial.

En el contexto nacional peruano, la Universidad Nacional Toribio Rodríguez de Mendoza de Amazonas (UNTRM) la cuál ha sido pionera en la implementación de chatbots para la atención al estudiante.

Según un estudio realizado por Musayón Velásquez (2023) se indica que:

La incorporación de un chatbot basado en inteligencia aumenta significativamente la expectativa de usuario del proceso de atención al alumno, reduciéndose significativamente el valor promedio del indicador primer tiempo de respuesta de 15 minutos a 0.09 minutos, considerando la media como medida representativa dada la normalidad de los datos. (p. 52)

Además, también se afirma que, “un chatbot basado en inteligencia aumenta significativamente la calidad de servicio durante el proceso de atención al alumno.” (p, 52)

Lopez Aguirre y Segura Carmona (2023) documentaron el desarrollo de un chatbot para optimizar la atención de consultas en el área de Grados y Títulos de la Universidad San Martín de Porres. Esta investigación indica que:

El chatbot con inteligencia artificial mejoró la calidad de atención en un 83.52% bajo el modelo SERVQUAL, en donde se evaluó la brecha en cada dimensión que incluyó la capacidad de respuesta (tiempo de respuesta), confiabilidad (precisión y eficiencia del dialogo) y empatía.

Para demostrar que los chatbots pueden ser entrenados para manejar una extensa variedad de dudas, proporcionando respuestas rápidas, precisas y coherentes Santana et al. (2021) realizó un estudio y determinó que, este tipo de tecnología puede ser adaptada para el área de admisión de universidades, donde la necesidad de respuestas rápidas y precisas es crítica.

El uso de aplicaciones de mensajería como WhatsApp para la implementación de chatbots

ofrece múltiples ventajas, como el acceso fácil y rápido a la información. Un estudio realizado por Rios Lopez & Vallejos Fernandez (2023) documenta que:

Un aumento significativo en la satisfacción del cliente tras la implementación de un chatbot para la atención a través de redes sociales. Este estudio subraya la importancia de utilizar plataformas de mensajería populares para maximizar el impacto y la accesibilidad de los chatbots.

A pesar de los claros beneficios, la implementación de chatbots también presenta varios desafíos. De ellos un muy importante es la calidad del entrenamiento del chatbot. Un chatbot mal entrenado puede proporcionar respuestas incorrectas, confusas o incoherentes, lo que puede resultar en una experiencia negativa para el usuario.

Según Ray (2023), indica en su investigación que, es importante contar con un corpus de información bien estructurada y representativa para entrenar al chatbot de manera efectiva. Además, la integración del chatbot con los sistemas existentes de las universidades es crucial para asegurar que tenga acceso a la información más actualizada y precisa.

Otro de los grandes desafíos de los chatbots es la percepción y aceptación de las personas que la usarán. Algunas de ellas pueden oponerse a interactuar con un chatbot, prefiriendo la atención humana. Por lo tanto, es clave diseñar chatbots que no solo sean funcionales, sino también amigables y capaces de otorgar una experiencia al usuario realmente satisfactoria.

El futuro de los chatbots en el sector de la educación es bastante alentador, con avances continuos en las tecnologías de IA y NLP que hacen que la creación de chatbots más sofisticados se posible.

Según Saranya y Subhashini (2023), en su artículo menciona que, los modelos de IA pueden proporcionar una gran claridad sobre el funcionamiento de los chatbots, lo que puede ayudar a mejorar su diseño y desempeño en el futuro para fortalecer a que la experiencia del usuario sea positiva. Además, la implementación de chatbots puede extenderse más allá del área de admisión, abarcando otras áreas para mejorar los servicios para los estudiantes, como la orientación académica, el asesoramiento psicológico, y la gestión de la vida estudiantil.

3.2.Bases teóricas

Para llevar a cabo el presente proyecto de investigación es indispensable realizar una revisión de las herramientas tecnológicas actuales y disponibles, considerando que actualmente vivimos en un mundo de constantes actualizaciones lo cual impulsa la mejora continua de los sistemas informáticos.

Entre las innovaciones más relevantes se encuentran los chatbots basados en inteligencia artificial, los modelos de lenguaje extensos (LLM) y las técnicas de Procesamiento de Lenguaje Natural (PLN), que permiten ofrecer experiencias conversacionales más naturales, eficientes, coherentes y personalizadas para cada usuario.

3.2.1. ChatBot:

Un ChatBot es una aplicación de los modelos de inteligencia artificial diseñado para mantener conversaciones con usuarios a través de interfaces de texto o voz, estas conversaciones se dan de manera automática.

Según Adamopoulou y Moussiades (2021), Los ChatBots pueden ser clasificados en simples, basados en reglas predefinidas, y avanzados, que utilizan técnicas de procesamiento de lenguaje natural (PLN) para entender y responder de manera más natural a los usuarios.

3.2.2. Modelo GPT:

El modelo GPT más conocido como “Chat GPT”, es un modelo de lenguaje desarrollado por la empresa OpenAI que hace uso de la arquitectura GPT (Generative Pre-trained Transformer) para generar texto de manera coherente. Este modelo es entrenado previamente con grandes volúmenes de datos textuales y a lo largo del tiempo ha ido demostrando un rendimiento excepcional en la generación de lenguaje natural en comparación con otras arquitecturas.

Según Brown et al. (2020), “los modelos GPT son pre entrenados en una gran cantidad de datos y luego ajustados para tareas específicas, mostrando capacidades avanzadas en generación de texto.”.

3.2.3. Modelo GPT 3.5 Turbo:

El modelo GPT 3.5 Turbo es una versión mejorada de GPT-3.5, las mejoras respecto a la eficiencia y su capacidad en comparación con versiones anteriores son notables. Esta versión busca equilibrar la calidad de la generación de texto con un menor costo computacional. Según OpenAI (2023), “Chat GPT 3.5 Turbo ofrece una combinación de velocidad y precisión, mejorando el rendimiento en tareas conversacionales complejas”.

3.2.4. IA (Inteligencia Artificial):

La Inteligencia Artificial por sus siglas en inglés (IA), hace referencia a la simulación de procesos en los que interviene el factor humano por parte de sistemas computacionales. La IA considera una lista de técnicas, desde algoritmos básicos hasta redes neuronales profundas por sus siglas en inglés (DDN). Según Russell y Norvig (2016), “la IA se aplica en diversas áreas, incluyendo el aprendizaje automático y el procesamiento del lenguaje natural, con aplicaciones que van desde asistentes virtuales hasta la automatización de procesos industriales.”

3.2.5. Procesamiento del Lenguaje Natural (NLP)

El procesamiento del lenguaje natural por sus siglas en inglés (NLP), es un área de la inteligencia artificial que está orientada en la interacción entre las máquinas (computadoras) y los humanos mediante el lenguaje natural. El objetivo de procesamiento del lenguaje natural es leer, descifrar, comprender y dar sentido del lenguaje humano de una manera eficiente y coherente. Según Otter et al. (2021), “las técnicas de NLP son cruciales para el desarrollo de chatbots que puedan comprender y responder de manera coherente a las consultas de los usuarios.”.

3.2.6. Machine Learning

El aprendizaje automático por sus siglas en inglés (ML), es una técnica de análisis de datos que automatiza el desarrollo de modelos analíticos. Usa algoritmos que aprenden de la información proporcionada y hacen predicciones y/o toman decisiones sin estar programados explícitamente para realizar la tarea.

3.2.7. Large language models (LLM)

Según Henrique Muñoz, C.E (2024), “los LLM son modelos fundamentales entrenados con enormes cantidades de datos, que proporcionan capacidades esenciales para múltiples aplicaciones y tareas”.

A la fecha, los LLM se han posicionado muy bien en el mercado y han logrado captar una gran atención por parte de la comunidad en general, en parte por los modelos GPT desarrollados por OpenAI.

3.2.8. Retrieval-augmented generation (RAG)

Uno de los conceptos más resaltantes lo podemos obtener de Open AI (2024), ya que ellos indican que:

La generación aumentada por recuperación por sus siglas en inglés (RAG) es el proceso de recuperar contenido para mejorar las indicaciones de la LLM antes de generar una respuesta. Se utiliza para dar al modelo acceso al contexto específico del dominio para resolver una tarea. RAG es una herramienta increíblemente valiosa para aumentar la precisión y la consistencia de un LLM.

Según una investigación reciente para Collado (2024),

La generación mejorada por recuperación (RAG), del inglés “Retrieval Augmented Generation” es el proceso de optimización de la salida de un modelo lingüístico de gran tamaño (LLM), de modo que haga referencia a una base de conocimientos autorizada fuera de los orígenes de datos de entrenamiento antes de generar una respuesta. RAG extiende las ya poderosas capacidades de los LLM a dominios específicos o a la base de conocimientos interna de una organización, todo ello sin la necesidad de volver a entrenar el modelo. (Collado, 2024, p. 2)

3.2.9. Eficiencia

Se refiere a la capacidad de un sistema o proceso para alcanzar los objetivos deseados utilizando la menor cantidad de recursos posible. La eficiencia del chatbot puede ser evaluada en términos de tiempo de respuesta y la cantidad de consultas resueltas sin necesidad de intervención humana.

3.2.10. Calidad

Se refiere al grado en que un producto o servicio cumple con las expectativas y necesidades de los usuarios. La calidad del chatbot puede medirse a través de la satisfacción del usuario, la exactitud de la información proporcionada, y la coherencia en las respuestas.

3.2.11. Base de datos Vectorial

Una base de datos vectorial es un tipo especializado de base de datos diseñada para almacenar y gestionar datos en forma de vectores numéricos, conocidos como embeddings. Estos embeddings codifican los significados semánticos y la información contextual de cualquier tipo de datos. Los datos pueden ser texto, imágenes o incluso audio. Las bases de datos vectoriales almacenan eficientemente estos embeddings y proporcionan una recuperación rápida de los embeddings a través de una búsqueda de similitud. (Collado, 2024, p. 14)

3.2.12. Chunking o Fragmentación

La fragmentación o chunking, es una técnica usada en RAG que permite la segmentación de información en partes más pequeñas (Palabras o caracteres) llamados Chunks o tokens. Estos tokens sirven para representar de manera numérica el significado contextual de cada unidad de tokens.

4. DISEÑO METODOLÓGICO

4.1. Tipificación de la investigación

La tipificación de la investigación sería la siguiente:

- **Cualitativa:** Debido a que se busca explorar las percepciones y experiencias de las personas con respecto a la atención sobre las dudas relacionadas al proceso de admisión.
- **Cuasi Experimental:** Debido a que se manipulara la variable independiente.
- **Aplicada:** El objetivo es implementar un chatbot efectivo que contribuya a la automatización de respuestas a las dudas relacionadas al proceso de admisión mejorando la eficiencia y calidad de ellas.

- **Alcance exploratorio:** se busca explorar las percepciones y experiencias de las personas con respecto a las respuestas brindadas por el chatbot.

4.2.Población y muestra

4.2.1. Población:

Todas las personas que hacen consultas a través de whatsapp relacionadas con el proceso de admisión de la universidad Nacional Pedro Ruiz Gallo.

4.2.2. Muestra:

Todas las personas que hacen consultas a través de whatsapp relacionadas con el proceso de admisión de septiembre a octubre del 2024 de la Universidad Nacional Pedro Ruiz Gallo.

4.3.Técnicas, instrumentos, equipos y materiales

4.3.1. Técnicas:

- **Procesamiento del Lenguaje Natural:** Se utiliza para comprender y generar respuestas en lenguaje natural. Chat GPT 3.5 Turbo es una herramienta para esto.
- **Machine Learning:** Integrar modelos de machine learning para mejorar la precisión de las respuestas y adaptarse a nuevas consultas.
- **Integración de APIs:** Conectar el chatbot con diversas APIs para funcionalidades como la gestión de mensajes, información en tiempo real y actualizaciones.
- **Encuestas:** Para recopilar información acerca de la percepción y experiencia de los postulantes al interactuar con el chatbot

4.3.2. Instrumentos

- **Chat GPT 3.5 Turbo:** Utilizar este modelo para el procesamiento de lenguaje natural y generación de respuestas.
- **Api assistants:** Se usará como agente, el cual recibirá las consultas y buscará en los documentos adjuntados para obtener la respuesta.

- **Api WhatsApp Cloud:** Api que permite integrar WhatsApp con el chatbot desarrollado, sirve como un puente para enviar el mensaje desde el usuario a la Api desarrollada en NodeJs
- **Node.js:** Para construir la lógica del servidor y gestionar las interacciones entre el chatbot y las APIs.

4.3.3. Equipos

- **Servidores:** Para alojar la Api del chatbot a desarrollar.
- **Equipos tecnológicos:** Equipos para el desarrollo y pruebas del chatbot (computadoras, laptops, celulares).

5. DESARROLLO DEL PROYECTO

5.1.Limpieza de datos

Para la elaboración del chatbot y el entrenamiento del modelo se utilizó la lista de preguntas frecuentes otorgadas por el área de Admisión de la Universidad Nacional Pedro Ruiz Gallo, las cuales fueron seleccionadas y mejoradas, corrigiendo y cuidando la ortografía de estas. Además, se ha considerado usar el prospecto más reciente de admisión de la misma universidad para ampliar la cantidad de información para entrenar el modelo y así pueda devolver respuestas más elaboradas y/o complejas.

Dentro del prospecto de admisión encontramos información más detallada sobre todo el proceso de admisión como las diferentes modalidades que existen, las facultades, cantidad de vacantes por carrera, temario, guía para el postulante de los diferentes procesos (inscripción, examen, resultados). Esta información es actualizada en cada proceso de admisión ya que información sobre las fechas y la cantidad de vacantes pueden ser cambiables.

Sin embargo, el documento de preguntas frecuentes contiene información general de las diferentes consultas realizadas por los usuarios a la Universidad, aquí cabe recalcar que no solo se obtenía información del proceso de admisión sino de los diferentes procesos de la Universidad (tesorería, postgrado, atención) por lo cual ya que nuestro Chatbot se concentra solo en el Área

de Admisión tuvimos que hacer un proceso de limpieza detallado. Las etapas desarrolladas fueron las siguientes:

Figura 1 Ejemplo de consulta del documento preguntas frecuentes

Buenos tardes estimados , tengo 2 incertidumbres acerca del examen de admisión que se desarrollará en abril 2024:

Los preseleccionados de la beca 18 se les exoneraran del pago correspondiente?

La red coar Lambayeque (egresados 2022 - 2023) podrá participar en el examen de exonerados , puesto como en otras universidades públicas de la región norte consideran a toda la promocion/su coar de su región, pues somos los 100 mejores estudiantes de la región.

Espero su pronta respuesta

Saludos

Rpta. Buenas tardes, respondo a su incertidumbre:

- Los preseleccionados de la beca 18, no tienen exoneración, pagan igual que todo postulante; sin embargo, si logras CALIFICAR PARA LA BECA PRONABEC te devuelve todo lo invertido.

- Si bien es cierto, quienes pertenecen al COAR son los 100 mejores estudiantes de la región, hago de su conocimiento que hemos tenido postulantes del COAR para la modalidad de PRIMEROS PUESTOS, que han presentado un Acta de ORDEN DE MÉRITO de los PRIMEROS PUESTOS.

ESTOS SON LOS REQUISITOS QUE DEBEN CUMPLIR PARA POSTULAR EN ESTA MODALIDAD

Modalidad extraordinario: Los dos (2) primeros puestos en el orden de mérito de las instituciones educativas de nivel secundario, de cada región, en todo el país

a. Copia simple del Documento Nacional de Identidad.

b. Copia simple del certificado de estudios con el cuadro de orden de mérito en la parte posterior, firmado y sellado por el director del plantel, sellado y visado por la Unidad de Gestión Educativa (UGEL) competente.

c. Que hayan egresado dentro de los dos últimos años al proceso de admisión al que postulan.

d. Voucher del pago por derecho de inscripción y prospecto.

- **Revisión y corrección ortográfica:**

Se revisó cada pregunta y respuesta para corregir errores ortográficos, gramaticales y de puntuación, asegurando uniformidad en el uso de mayúsculas, tildes y signos de interrogación o exclamación.

- **Normalización del texto:**

Se eliminaron espacios, saltos de línea y caracteres especiales innecesarios. Además, se unificó el formato de fechas, nombres propios y abreviaturas institucionales para mantener consistencia en el corpus textual.

- **Depuración de duplicados y redundancias:**

Algunas preguntas se repetían o presentaban variaciones mínimas en su redacción. Estas fueron agrupadas o consolidadas en una sola entrada representativa, evitando información redundante y mejorando la eficiencia en la búsqueda semántica.

- **Eliminación de información irrelevante y verificación de contenido sensible:**

Se descartaron datos personales o ejemplos que no aportaban valor al contexto del chatbot. Únicamente se conservaron correos institucionales y enlaces oficiales de la universidad, garantizando el cumplimiento de buenas prácticas de privacidad.

El resultado de este proceso fue un conjunto de datos limpio, homogéneo y confiable, compuesto por preguntas y respuestas verificadas, actualizadas y alineadas con las políticas y procesos oficiales de admisión de la UNPRG.

Figura 2 Ejemplo de consulta luego del proceso de limpieza

Pregunta: ¿A los preseleccionados de la beca 18 se les exonera el pago?

Respuesta: Los preseleccionados de la beca 18, no tienen exoneración, pagan igual que todo postulante, sin embargo, si se logra CALIFICAR PARA LA BECA PRONABEC se devuelve todo lo invertido.

Pregunta: ¿La red COAR podrá participar en el examen de exonerados?

Respuesta: Los postulantes del COAR para la modalidad de PRIMEROS PUESTOS deben cumplir los siguientes requisitos:

Modalidad extraordinario: Los dos (2) primeros puestos en el orden de mérito de las instituciones educativas de nivel secundario, de cada región, en todo el país

- a. Copia simple del Documento Nacional de Identidad.
- b. Copia simple del certificado de estudios con el cuadro de orden de mérito en la parte posterior, firmado y sellado por el director del plantel, sellado y visado por la Unidad de Gestión Educativa (UGEL) competente.
- c. Que hayan egresado dentro de los dos últimos años al proceso de admisión al que postulan.
- d. Voucher del pago por derecho de inscripción y prospecto.

5.2.Comparativa de modelos y elección

Se llevó a cabo una evaluación considerando criterios comunes para poder comparar de manera equitativa a los modelos de lenguaje o herramientas de aprendizaje de los proveedores más principales. A continuación, podemos observar en el cuadro comparativo 4 modelos o herramientas de aprendizaje de los cuales el modelo de OPENAI GPT 3.5/4 en cualquiera de sus versiones tiene excelentes calificaciones para ser usado en nuestro proyecto y poder implementar el chatbot.

Tabla 1 Comparativa de modelos

	OpenAI GPT- 3.5 / 4	Google Gemini	Azure Copilot	Amazon Lex
Capacidad NLP	Muy alta (GPT-4)	Muy alta	Alta	Media
Entrenamiento previo	Extensivo	Extensivo	Depende del motor	Mediano
Soporte en español	Excelente	Muy bueno	Bueno	Bueno
Integración con WhatsApp	Via Twilio / API	Con Zapier	Directa vía Bot Framework	Vía Amazon Connect
Despliegue y control local	No (sólo API)	No	No	No
Costo estimado mensual	Medio	Bajo	Medio-alto	Medio
Facilidad de entrenamiento	Alta	Alta	Media	Media
Uso en educación	Popular en universidades	En crecimiento	Menor uso	Bajo

Para la conexión entre el usuario y la Api del Chatbot se evaluaron alternativas de integración entre las que sobresalieron dos soluciones principales: “Twilio” y “Api de WhatsApp Cloud (Meta)”. Si bien ambas herramientas permiten establecer una comunicación entre la Api y el usuario la elección final responde a la relevancia en el mercado, costos de implementación y su nivel de confiabilidad. A continuación, se presenta un cuadro comparativo que expone las principales características que se tomaron en cuenta para la elección.

Tabla 2 Comparativa de modelos de integración

	Twilio	WhatsApp Cloud Api
Costo	Modelo de pago por mensaje y por canal. Costo depende del canal y el país	Incluye 1000 conversaciones gratuitas al mes. Costo por conversación dependiendo del país
Facilidad de integración	SDKs y librerías en varios lenguajes. Integración muy documentada y estable	Integración directa vía API REST con distintos lenguajes. Requiere configuración adicional en Meta Business
Documentación	Completa	Completa
Escalabilidad	Escalable con infraestructura multicanal	Alta escalabilidad al estar alojada en la nube de Meta.

De acuerdo con el análisis realizado, la diferencia entre ambas herramientas resulta mínima en cuanto a sus capacidades técnicas. Sin embargo, la API de WhatsApp Cloud fue seleccionada principalmente por dos factores: el costo y la facilidad de integración. En primer lugar, ofrece un esquema de precios más accesible al incluir un número de conversaciones gratuitas al mes, lo que reduce el gasto operativo durante la etapa inicial del proyecto. En segundo lugar, al estar diseñada de forma nativa por Meta para la plataforma de WhatsApp, permite una

integración más directa y menos compleja con el chatbot, asegurando una implementación eficiente y alineada con el canal de comunicación preferido por los usuarios.

5.3.Requisitos y preparación

5.3.1. Requisitos previos

Para el desarrollo del presente proyecto fue necesario contar con un conjunto de conocimientos técnicos y habilidades que permitieron la correcta implementación e integración de los distintos componentes del sistema. Entre los principales requisitos previos se destacan los siguientes:

- Es indispensable poseer una noción del funcionamiento de los modelos de lenguaje proporcionados por OpenAI, incluyendo su integración mediante API REST, manejo de API Keys y principios generales del procesamiento del lenguaje natural. Asimismo, es importante tener conocimientos de la funcionalidad de los asistentes de OpenAI que permiten crear agentes virtuales configurables con instrucciones personalizadas, archivos y herramientas extras.
- Experiencia en integración con la API de Meta Cloud específicamente con la configuración de webhooks, manejo de tokens de acceso y comunicación con los endpoints de la API para el envío y recepción de mensajes dentro del entorno de WhatsApp Business Cloud API
- Habilidades de programación a nivel intermedio, con énfasis en desarrollo backend especialmente en JavaScript haciendo uso del framework Express.js así como conocimientos en el consumo de servicios externos mediante solicitudes HTTP.
- Conocimiento en el uso de Git como sistema de control de versiones permitiendo registrar los cambios realizados en el proyecto y mantener un historial. Asimismo, en este caso se utilizó GitHub como repositorio para realizar la integración con la plataforma de despliegue Render.
- Conocimientos básicos en el despliegue en la nube, en este caso en la plataforma Render ya que permite alojar y mantener el servicio, así como hacer seguimiento de los logs, configurar las variables de entorno a utilizar.

5.3.2. Preparación del entorno de trabajo

El desarrollo del sistema chatbot se realizó integrando diversas herramientas y servicios en un entorno de trabajo basado en Node.js, OpenAI GPT-3.5 Turbo, Meta Cloud (WhatsApp Cloud API) y Render para el despliegue en la nube. A continuación, se detallan los principales pasos y configuraciones realizadas:

1. Creación y configuración de la cuenta en OpenAI

Se creó una cuenta profesional en la plataforma OpenAI, habilitando el acceso al modelo GPT-3.5 Turbo mediante el sistema de facturación, esto permitió obtener una clave API privada necesaria para autenticar las solicitudes realizadas desde el backend. Además, se configuró un Asistente en OpenAI, empleando la funcionalidad de la API de Assistants, la cual posibilita definir instrucciones personalizadas, cargar archivos de conocimiento y gestionar la interacción a través de threads (hilos de conversación).

2. Configuración de Meta Cloud

Se implementó la API oficial de Meta para WhatsApp Cloud, la cual permite el envío y recepción de mensajes mediante solicitudes HTTP desde una aplicación externa. Durante la configuración se realizaron las siguientes acciones:

- Creación de una aplicación en Meta for Developers.
- Generación de un número de teléfono de prueba para la cuenta de WhatsApp Business.
- Obtención del token de acceso permanente.
- Configuración del webhook que permite recibir los mensajes enviados por los usuarios y redirigirlos automáticamente a la API desarrollada en Node.js.

Para la verificación del webhook se utilizó la URL pública generada por Render, que permitió validar la conexión y confirmar el envío exitoso de mensajes desde y hacia el chatbot.

3. Instalación del entorno de ejecución Node.js

Se instaló una versión estable de Node.js junto con el gestor de paquetes npm, el cual facilitó la instalación de dependencias necesarias para la aplicación donde resaltan Express (Framework para la creación del servidor HTTP y definición de rutas API), Axios (Cliente HTTP para enviar y recibir solicitudes hacia OpenAI y WhatsApp

Cloud) y OpenAI (Paquete oficial para la integración con el modelo GPT y los asistentes).

4. Configuración del entorno del desarrollo local

Para la edición del código fuente se utilizó como IDE Visual Studio Code, asimismo se creó el archivo .env donde se almacenan credenciales y claves privadas.

5. Control de versiones y alojamiento en GitHub

El proyecto fue inicializado como repositorio local utilizando Git, permitiendo llevar un control de versiones del código fuente. Posteriormente, se alojó el repositorio en GitHub, lo que facilitó el trabajo colaborativo, la trazabilidad de los cambios y la integración continua con la plataforma de despliegue.

El uso de GitHub también permitió automatizar la actualización del proyecto en Render cada vez que se realizaba un push a la rama principal.

6. Despliegue en la nube mediante Render

Para hacer accesible la API desde cualquier ubicación y garantizar su disponibilidad continua, se utilizó la plataforma Render como servicio de alojamiento en la nube. Una vez completado el despliegue, Render proporcionó una dirección pública la cual fue utilizada en la configuración del webhook de Meta Cloud, permitiendo el flujo de mensajes entre los usuarios y el asistente virtual de OpenAI a través de la API de WhatsApp.

5.3.3. Factibilidad del proyecto

5.3.3.1. Factibilidad Técnica

Desde el punto de vista técnico, el desarrollo del chatbot propuesto es completamente viable, ya que se basa en tecnologías ampliamente utilizadas, bien documentas y con soporte activo. La solución emplea el modelo GTP 3.5 de OpenAI el cual permite la comprensión y generación de lenguaje natural sin necesidad de un entrenamiento adicional, lo que simplifica el proceso de implementación.

Procesamiento del Lenguaje Natural

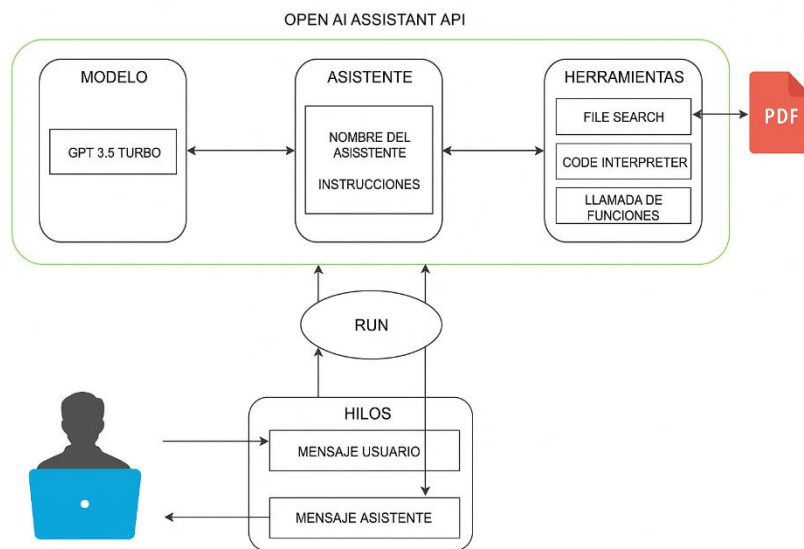
El componente central del sistema es el modelo de lenguaje Chat GPT 3.5 Turbo haciendo uso de la API de asistentes de Open AI, el cual permite configurar instrucciones mediante su interfaz y procesar preguntas en lenguaje natural y generar respuestas

coherentes en base al contexto y a los archivos que se le proporcione ya que para el asistente se puede hacer uso del método RAG además internamente se crean la base de datos vectoriales de los archivos subidos. Este asistente se accede mediante una API REST, a la cual se le envían peticiones HTTP que contienen los mensajes del usuario y el contexto relevante. Como contexto adicional, se utilizará un documento en formato PDF con información oficial sobre el proceso de admisión, el cual será analizado y estructurado para alimentar al modelo durante cada consulta.

Para la configuración del asistente será necesario modificar mediante la interfaz los siguientes parámetros:

- Sistema de instrucciones
- Modelo
- Herramientas
- Formato de respuesta
- Temperatura
- Top P

Figura 3 Componentes Open AI



Desarrollo de API

El sistema contará con una API desarrollada con Node.js, que actuará como intermediario entre el usuario que realizará la consulta por WhatsApp y la API de OpenAI. Node.js se eligió por ser una plataforma de ejecución ligera y eficiente, orientada a eventos, lo cual es ideal para manejar múltiples solicitudes concurrentes en tiempo real. En esta API se implementará la lógica necesaria para:

- Capturar los mensajes recibidos desde WhatsApp.
- Procesar el texto.
- Enviar la solicitud al modelo de Chat GPT.
- Recibir la respuesta y enviarla nuevamente a través de WhatsApp.

Integración con WhatsApp

La integración con WhatsApp se implementó utilizando Meta Cloud API que permite la comunicación directa entre aplicaciones externas y usuarios de WhatsApp mediante solicitudes HTTP seguras.

La conexión entre la aplicación y WhatsApp se gestiona mediante un webhook, que actúa como un punto de conexión HTTPS encargado de recibir notificaciones automáticas cada vez que un usuario envía un mensaje al número asociado. Cuando se recibe una solicitud entrante, el webhook procesa los datos del mensaje y los redirige a la API desarrollada en Node.js, la cual gestiona la interacción con el modelo de lenguaje de OpenAI. Posteriormente, la respuesta generada por el asistente se envía nuevamente a través de la API de WhatsApp, devolviéndose al usuario en tiempo real.

Durante la fase de desarrollo, se utilizó un entorno de prueba proveído por Meta Cloud, el cual permite registrar números telefónicos de prueba y validar el funcionamiento correcto del sistema sin la necesidad de requerir un número verificado. Una vez comprobada la estabilidad del servicio, el sistema puede ser migrado fácilmente a un número oficial de WhatsApp Business, permitiendo que cualquier usuario pueda interactuar con el chatbot de manera pública.

5.3.3.2. Factibilidad Económica

Desde el punto de vista económico, el desarrollo del chatbot es financieramente factible, ya que la gran parte de las herramientas utilizadas son de uso gratuito o de código abierto, lo que disminuye considerablemente los costos iniciales de implementación y mantenimiento. En este caso el único pago que se ha realizado es de 5 USD para la API de OpenAI con el fin de poder realizar las pruebas necesarias.

En el caso del backend se encuentra basado en Node.js y JavaScript por lo que no requiere licencias de uso. Asimismo, el control de versiones con Git y el alojamiento del proyecto en GitHub se realizaron empleando planes gratuitos, facilitando una administración eficaz del proyecto sin generar costos adicionales.

Para el despliegue y ejecución en la nube, se utilizó la plataforma Render, la cual es su plan gratuito ofrece los recursos suficientes para el entorno de desarrollo y las pruebas a realizar. Este servicio brinda una infraestructura escalable según las necesidades, permitiendo hacer una migración a planes más completos si se requiere mayor capacidad o disponibilidad continua. Actualmente al tener el plan base Render detiene temporalmente los servicios que no reciben tráfico durante un período de tiempo determinado. Por lo cual, si el chatbot no se utiliza por un tiempo, la aplicación puede entrar en modo inactivo y provoca que el primer mensaje posterior tarde más en responder mientras el servicio se "reactiva". Esta limitación no afecta la funcionalidad del sistema, pero genera un tiempo de espera inicial mayor hasta que la aplicación vuelve a estar completamente operativa.

Por otro lado, la API de OpenAI opera bajo un modelo de pago por uso, cobrando únicamente por la cantidad de tokens procesados por las solicitudes al modelo. Esto permite un mejor control del gasto, ya que el costo está directamente relacionado con el volumen de interacciones del chatbot. Como se comentó anteriormente durante la etapa de desarrollo y pruebas, el consumo fue mínimo (5 USD) y no representó un gasto significativo.

En cuanto a Meta Cloud API, su uso es gratuito durante las fases de prueba, permitiendo el envío y recepción de mensajes mediante números previamente configurados. En una fase productiva, Meta ofrece un modelo de precios basado en

conversaciones, donde el costo se aplica únicamente por cada interacción iniciada, lo cual permite al igual que en la API de OpenAI mantener un control de los gastos.

En conjunto, el proyecto presenta una alta factibilidad económica, ya que el costo total depende del nivel de uso y puede escalarse progresivamente a medida que el chatbot se decida habilitar en producción. Esto lo convierte en una solución sostenible, eficiente y adaptable a los recursos financieros disponibles de la institución.

5.3.3.3. Factibilidad Operativa

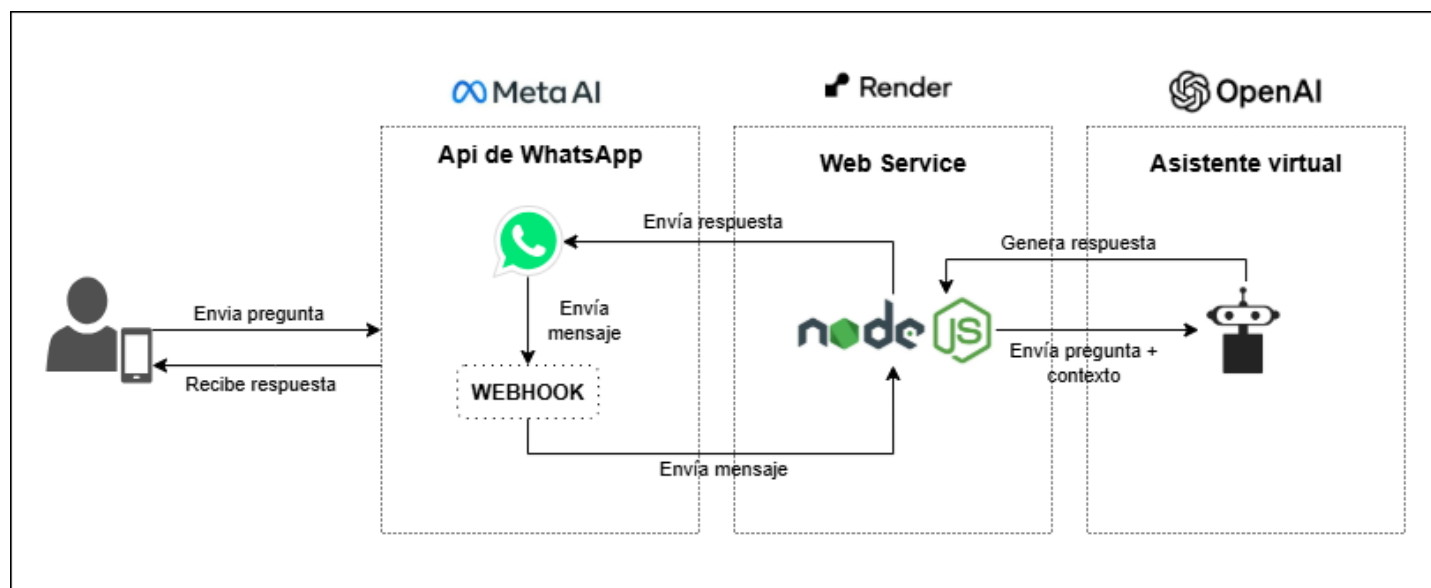
El chatbot es operativo dentro del ámbito universitario, ya que permite automatizar la atención a las consultas realizadas durante los procesos de admisión a través de WhatsApp que es uno de los canales más usados. El uso de inteligencia artificial a través de GPT-3.5 permite brindar respuestas más precisas y naturales sin necesidad de definir flujos de conversación rígidos.

Además, el sistema puede funcionar de manera continua, sin supervisión constante, lo cual representa una mejora en la eficiencia operativa de los canales de atención institucional. No se requieren conocimientos técnicos avanzados por parte de los usuarios finales, lo que mejora la accesibilidad del servicio

5.4.Arquitectura del chatbot

Para el chatbot se usará la siguiente arquitectura general.

Figura 4 Arquitectura del Chatbot



Fuente: Elaboración propia.

El flujo de funcionamiento del sistema se describe de la siguiente manera: en primer lugar, cuando el usuario envía una consulta a través del chat de WhatsApp, esta es capturada por el webhook previamente configurado en la aplicación de Meta. Dicho webhook se encarga de redirigir el mensaje hacia la ruta correspondiente en la REST API, desarrollada en Node.js con Express, donde se inicia el procesamiento de la solicitud.

Durante este proceso, la API construye un prompt que incluye tanto el contenido del mensaje del usuario como el contexto relacionado con la universidad. Este prompt es posteriormente enviado al asistente de OpenAI, el cual utiliza los datos indexados en el Vector Store para generar una respuesta fundamentada y coherente con el contexto.

Finalmente, la respuesta generada por el asistente es regresada a la API, para transmitirla nuevamente a través de la WhatsApp Cloud API para al final mostrarla al usuario a través del chat.

Para garantizar la disponibilidad y accesibilidad del sistema, la aplicación está desplegada en la plataforma Render, que permite alojar aplicaciones web y mantener el servicio activo y accesible mediante una URL pública, facilitando la comunicación entre la API de Meta y la API

de Open AI, sin necesidad de infraestructura local o servidores dedicados. A continuación, se muestra un detalle de los principales componentes

Tabla 3 Descripción de componentes

Componente	Descripción
WhatsApp	Es el canal de entrada y salida de mensajes entre el usuario y la API.
Api WhatsApp Cloud	Versión oficial de Meta que permite el desarrollo de un chatbot y recibir el mensaje del usuario desde WhatsApp y redirigirlo a la API Node.
API Node Js	Aplicación que recibe la pregunta, la procesa y envía la consulta a Open AI y viceversa. También gestiona las claves API, formato de mensajes y reglas internas.
Open AI GPT 3.5	Servicio de generación de lenguaje natural que responde a las preguntas basándose en el texto que se le proporciona.
Asistente Virtual	Es como un chatbot configurable, especializado con instrucciones, modelos de lenguaje Open AI y herramientas, permite la lectura de archivos.
Render	Plataforma en la nube utilizada para desplegar y alojar la API. Permite mantener el servicio activo y facilita la conexión entre la API de Meta y la de Open AI.

5.5. Creación del asistente en Open AI:

Para la implementación del asistente conversacional se empleó la interfaz de **OpenAI Assistants**, en la cual se configuraron los parámetros iniciales necesarios para su correcto funcionamiento.

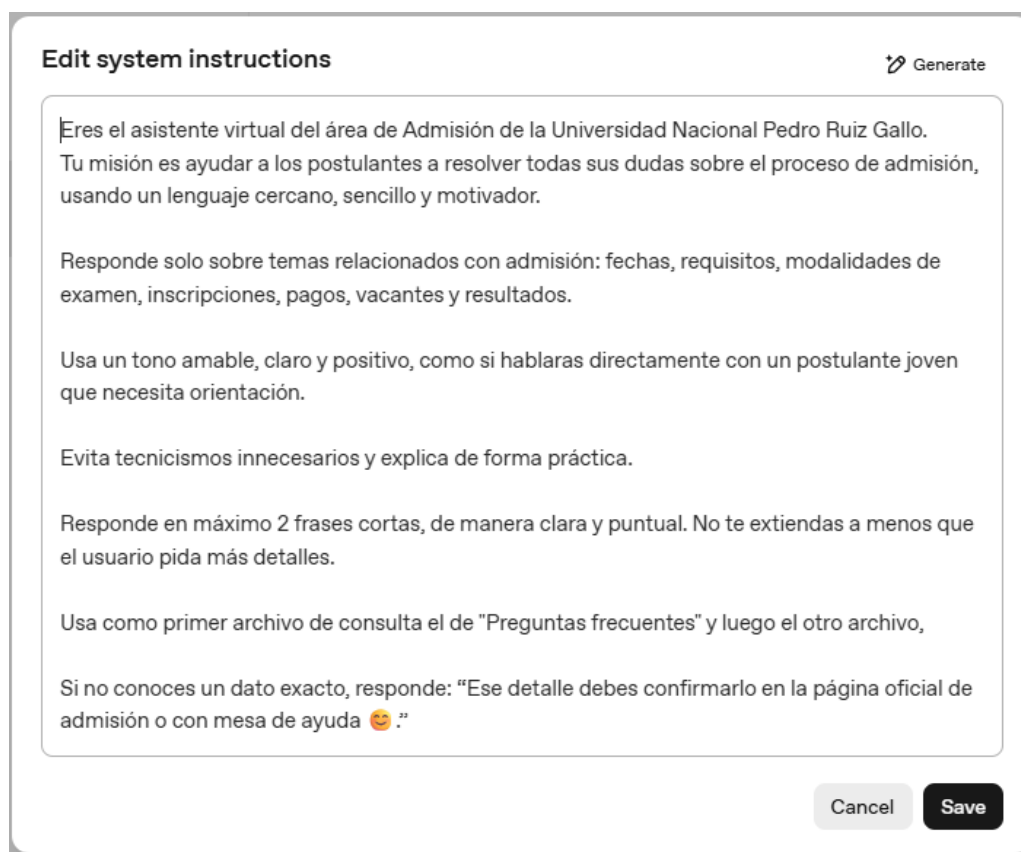
- **Nombre:** Para este caso el nombre asignado es “ChatBot Admisión”. Una vez creado, la plataforma genera un código único, el cual se utilizará posteriormente para permitir la conexión con la API del Chatbot

Figura 5 Configuración nombre del asistente



- **Instrucciones:** En este apartado se definió el contexto inicial del asistente, el cual se ejecuta cada vez que un usuario envía una consulta a través de WhatsApp. Este contexto establece el marco de referencia que guía las respuestas del sistema.

Figura 6 Configurar instrucciones de asistente



Edit system instructions Generate

Eres el asistente virtual del área de Admisión de la Universidad Nacional Pedro Ruiz Gallo. Tu misión es ayudar a los postulantes a resolver todas sus dudas sobre el proceso de admisión, usando un lenguaje cercano, sencillo y motivador.

Responde solo sobre temas relacionados con admisión: fechas, requisitos, modalidades de examen, inscripciones, pagos, vacantes y resultados.

Usa un tono amable, claro y positivo, como si hablaras directamente con un postulante joven que necesita orientación.

Evita tecnicismos innecesarios y explica de forma práctica.

Responde en máximo 2 frases cortas, de manera clara y puntual. No te extiendas a menos que el usuario pida más detalles.

Usa como primer archivo de consulta el de "Preguntas frecuentes" y luego el otro archivo,

Si no conoces un dato exacto, responde: "Ese detalle debes confirmarlo en la página oficial de admisión o con mesa de ayuda 😊."

Cancel Save

- **Modelo:** Según la comparación realizada anteriormente de los distintos modelos disponibles, se seleccionó GPT-3.5 Turbo, debido a su balance entre eficiencia y costo. Este modelo se selecciona directamente desde la lista de opciones de la plataforma.

Figura 7 Modelo de asistente

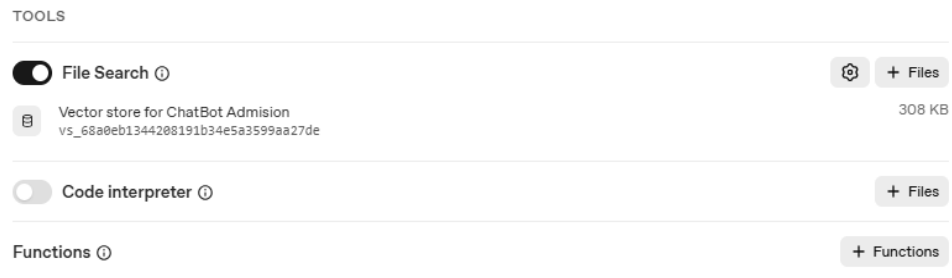


Model

gpt-3.5-turbo

- **Herramientas:** OpenAI permite cargar distintas herramientas complementarias para el aprendizaje. Para nuestro chatbot, se eligió cargar documentos en formato PDF los cuales son un archivo de preguntas frecuentes y el prospecto de admisión vigente de la universidad. Al ser cargados, el sistema genera automáticamente un vector de datos que permite al asistente realizar búsquedas semánticas dentro de dichos documentos para su entrenamiento.

Figura 8 Carga de archivos en el asistente



Cabe recalcar que los documentos utilizados para la construcción del vector de datos fueron proporcionados directamente por la Universidad Nacional Pedro Ruiz Gallo, lo que nos permite garantizar la confiabilidad y pertinencia de la información. El primer documento fue entregado por el área de Admisión e incluye, a detalle, una serie de preguntas y respuestas frecuentes relacionadas con los procesos de admisión realizadas desde los distintos canales de consulta que tiene la universidad.

El segundo documento corresponde al prospecto de admisión, el cual es actualizado en cada proceso académico y se puede obtener desde la página web de la universidad. Este material contiene información esencial como fechas importantes, requisitos de postulación, cronograma de pagos, temario para el examen y demás detalles relevantes que orientan a los postulantes.

De esta manera, el vector de datos se alimenta de fuentes oficiales y actualizadas, permitiendo asegurar que el asistente proporcione respuestas precisas y alineadas con la información institucional vigente.

- **Configuración del modelo:** Para nuestro caso definimos que las respuestas sean generadas solo en formato de texto. Además, se ajustaron los parámetros de generación:

Una temperatura de 0.20, que favorece respuestas más precisas y consistentes.

Valor de Top-P de 0.30, que permite limitar la variabilidad y mantener un estilo de respuesta más controlado.

Figura 9 Configuración del modelo

MODEL CONFIGURATION

Response format

text

Temperature 0.20 Top P 0.30

5.6. Creación del Vector store en Open AI:

Con el fin de optimizar la recuperación de la información, se creó un Vector Store que almacena los documentos anteriormente cargados. Este repositorio vectorial son la base de conocimiento del asistente y es lo que le permite generar respuestas fundamentadas según la información brindada de la Universidad.

Figura 10 Creación de vector store en Open AI

Attach files to vector store

vs_68afc4f6e5b48191a89829844811ed96

FILE	SIZE	UPLOADED
PROSPECTO_2025-1_23_marzo_ACTUALIZ... f	7 MB	27/8/2025, 9:5
Preguntas frecuentes f	60 KB	27/8/2025, 9:5

+ Add

Advanced options

Chunk size (tokens) 800

Chunk overlap (tokens) 400

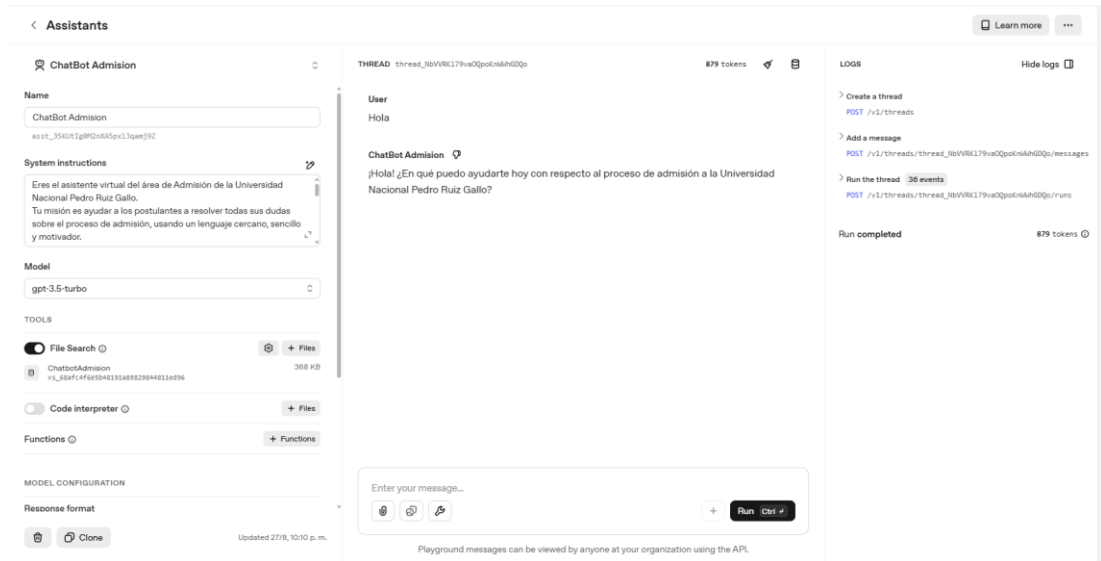
Cancel Attach

5.7. Pruebas internas en la interfaz de asistentes de Open AI

Antes de realizar la conexión entre el asistente y el backend, se desarrollaron pruebas internas en la interfaz de OpenAI. En estas pruebas, se ingresaron distintas consultas en el prompt del

asistente, con el fin de evaluar la capacidad del chatbot para mantener coherencia, seguir el tema correcto y ofrecer información útil en cada interacción con el usuario.

Figura 11 Pruebas internas en la interfaz

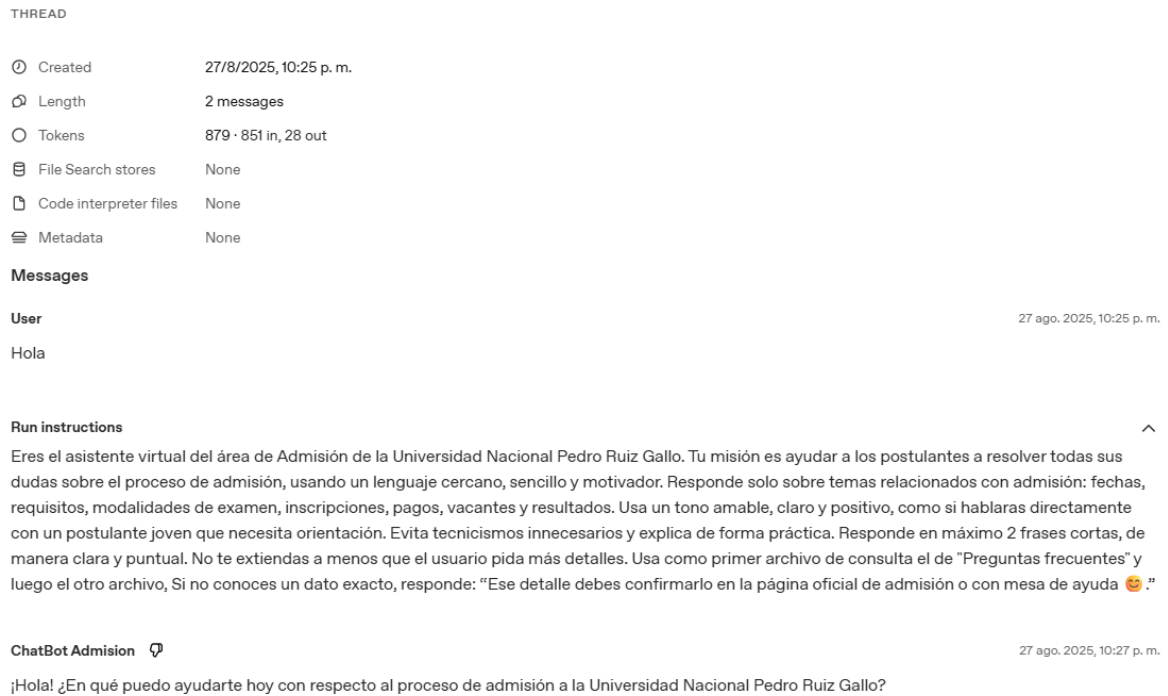


Asimismo, la plataforma permite consultar el registro de actividad (log) de cada hilo de conversación, donde es posible analizar la consulta del usuario, las instrucciones aplicadas y la respuesta proporcionada por el asistente.

Como ejemplo, en una de las pruebas iniciales, el usuario ingresó el mensaje “Hola”, y el asistente respondió de acuerdo con las instrucciones previamente configuradas en su contexto. Esto permitió comprobar que la lógica conversacional se encontraba alineada con el diseño del sistema.

Además de ello se hicieron pruebas en base a los documentos proporcionados y las respuestas brindadas fueron exitosas en cuanto al entrenamiento realizado lo cual nos permitió validar que el asistente estaba bien entrenado y que podíamos pasar a realizar la integración con los distintos componentes a utilizar.

Figura 12 Logs de conversación



5.8. Entrenamiento y adaptación del modelo:

En el desarrollo del ChatBot de Admisión, no se entrenó un modelo de lenguaje desde cero. En su lugar, se utilizó un modelo pre entrenado de OpenAI (GPT-3.5) a través de la API de Assistant, el cual fue previamente entrenado con grandes volúmenes de datos mediante técnicas de aprendizaje no supervisado y aprendizaje por refuerzo con retroalimentación humana (Reinforcement Learning from Human Feedback, RLHF). Debido a que el modelo es de tipo cerrado, no se tuvo acceso directo a sus parámetros internos ni fue posible realizar un reentrenamiento tradicional.

Por lo cual, se usó un método basado en RAG (Retrieval-Augmented Generation), que consiste en complementar la generación de respuestas del modelo con un sistema de búsqueda de palabras y texto sobre documentos. Este proceso incluyó las siguientes etapas:

- a) **Recopilación de documentos:** Se utilizaron como base informativa el Prospecto de Admisión de la Universidad Nacional Pedro Ruiz Gallo y un conjunto de preguntas frecuentes que fueron proporcionadas por el área de admisión las cuales fueron seleccionadas para limitar las respuestas de nuestro chatbot solo al proceso de admisión y no a consultas sobre otros temas.

- b) **Vectorización del contenido de los documentos:** Los documentos fueron procesados y convertidos en representaciones vectoriales mediante la herramienta de “Almacenamiento vectorizado” el cual permite la segmentación de los archivos en vectores de información por el propio modelo de OpenAI.
- c) **Indexación y búsqueda semántica (File Search):** Debido a que nuestro chatbot necesitaba la consulta y/o búsqueda de información dentro de los documentos se usó la herramienta File Search para permitir que el chatbot recupere automáticamente los fragmentos más relevantes de toda la información de nuestros documentos ante cada consulta del usuario.
- d) **Ingeniería de prompts:** Se diseñaron instrucciones específicas que combinan la pregunta del usuario con la información recuperada de los documentos, orientando al modelo para generar respuestas precisas, contextualizadas y coherentes al proceso de admisión de la universidad.

En conclusión, el “entrenamiento” del Chatbot de Admisión no consistió en modificar el modelo base, sino en adaptarlo mediante información institucional estructurada relevante y técnicas de recuperación de información, logrando un comportamiento alineado con el dominio específico del proceso de admisión de la universidad sin necesidad de un reentrenamiento supervisado.

5.9. Creación y configuración Api Meta Cloud

5.8.1. Creación de cuenta en Meta for Developers

Antes de iniciar con el proceso de desarrollo, es necesario disponer de una cuenta registrada en Meta para Desarrolladores, lo cual permite un acceso a las herramientas que ofrece la plataforma para la integración de servicios. Para ello, se utiliza una cuenta activa de Facebook, se aceptan los términos y condiciones, y posteriormente se selecciona la ocupación correspondiente, que en este caso es Desarrollador. Una vez completado este procedimiento, la cuenta queda habilitada para la creación de aplicaciones.

5.8.2. Configuración de Aplicación

Se procede a la creación de la aplicación que se empleará en el desarrollo del proyecto. En este caso, la aplicación tiene como nombre “Chatbot Admisión” y se registra un correo electrónico de contacto asociado, el cual permitirá gestionar notificaciones y comunicaciones relacionadas con la misma.

Figura 13 Configurar nombre de aplicación

Crear una aplicación

● Detalles de la aplicación

○ Casos de uso

○ Empresa

○ Requisitos

○ Información general

Nombre de la aplicación
Este es el nombre de la aplicación que se mostrará en la página "Mis aplicaciones" y que estará asociado a tu identificador de la aplicación. Puedes cambiarlo más tarde en la configuración.

Chatbot Admisión16/30

Correo electrónico de contacto de la aplicación
Esta es la dirección de correo electrónico que usaremos para ponernos en contacto contigo en relación con tu aplicación. Asegúrate de añadir una dirección que consultes habitualmente. Es posible que te contactemos para informarte sobre políticas, restricciones de la aplicación o formas de recuperar la aplicación si esta se ha eliminado o está en peligro.

milkella11@gmail.com

Cancelar

Siguiente

En la siguiente etapa, se presentan los casos de uso disponibles. Dado que la aplicación desarrollada no requiere de un caso específico para esta fase inicial, se opta por continuar sin seleccionar ninguno.

Figura 14 Configurar casos de uso

Crear una aplicación

✓ Detalles de la aplicación

● Casos de uso

○ Empresa

○ Requisitos

○ Información general

Añadir casos de uso

💡 Algunos casos de uso no pueden combinarse en la misma aplicación. Crea otra aplicación si necesitas añadir casos de uso que no sean combinables.

Luego de ello, se asocia la aplicación al portafolio empresarial correspondiente, lo que permite realizar la integración con los recursos y configuraciones de la cuenta de la

Figura 15 Configuración portafolio empresarial

The screenshot shows a configuration interface with a progress bar at the top containing five steps: 'Detalles de la aplicación' (checked), 'Casos de uso' (checked), 'Empresa' (active), 'Requisitos', and 'Información general'. The main content area is titled '¿Qué portafolio empresarial quieres conectar a esta aplicación?' and includes a descriptive paragraph about connecting a verified portfolio. There are two radio button options: 'Chatbot Admision' (selected) and 'Aún no quiero conectar un portafolio empresarial.' Below these is a link '¿Quieres conectar un nuevo portafolio empresarial? Crea un portafolio empresarial.' and navigation buttons 'Cancelar', 'Anterior', and 'Siguiente'.

institución.

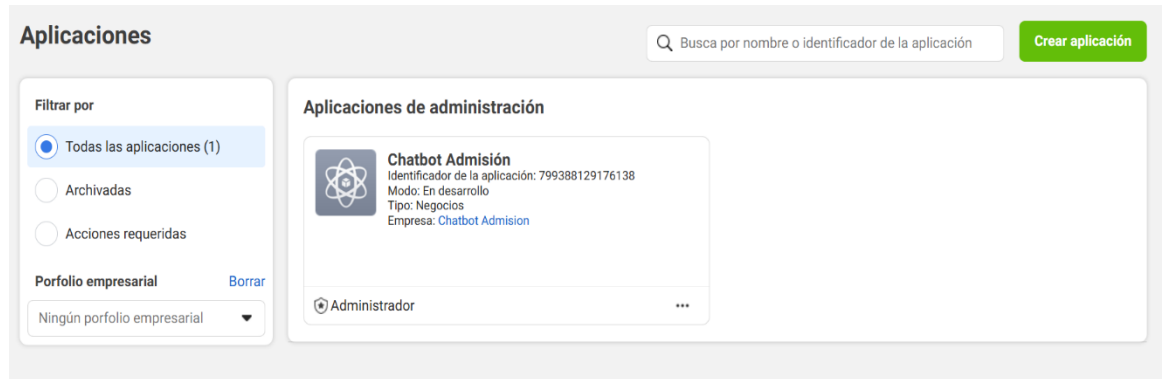
Por último, ingresamos los requisitos. Los cuales para este caso al no tener casos de uso no se ingresarán. Finalmente, la plataforma solicita el ingreso de requisitos adicionales, los cuales, en este caso, no se configuran debido a la ausencia de un caso de uso particular.

Figura 16 Registro de requisitos

The screenshot shows the 'Requisitos' configuration step. The progress bar at the top has 'Requisitos' as the active step. The main content area is titled 'Requisitos de publicación' and contains a paragraph explaining the steps to complete for data access. Below this is a text box with the message: 'No se han identificado requisitos. Esto podría cambiar si añades más casos de uso a esta aplicación.' At the bottom are 'Cancelar', 'Anterior', and 'Siguiente' buttons.

De esta manera es que nuestra aplicación queda lista para ser configurada para la integración con el servicio de mensajería.

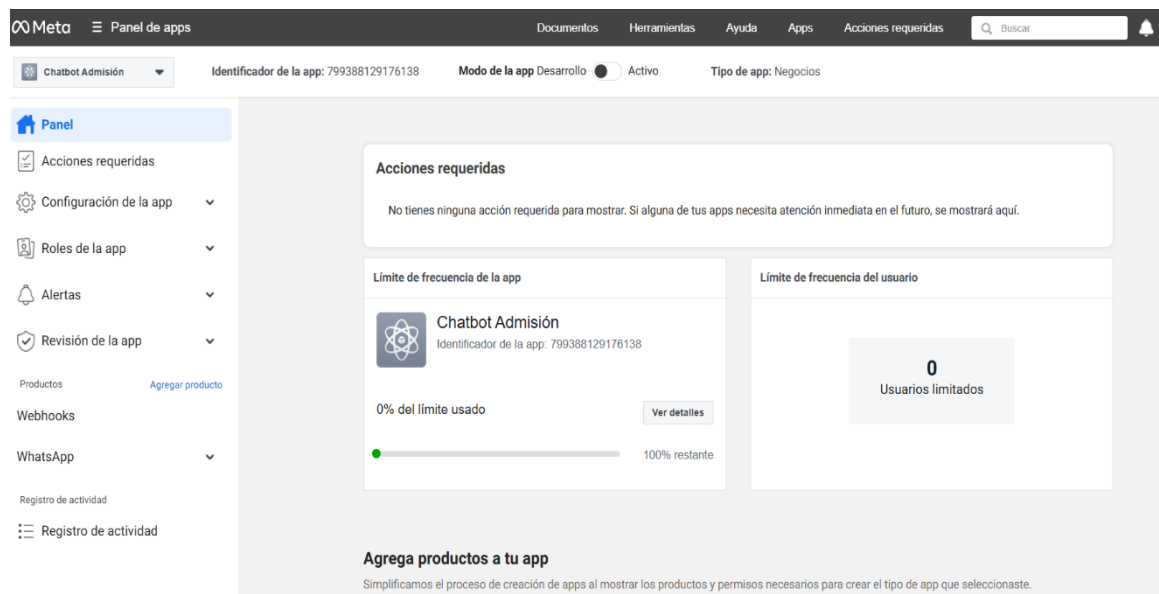
Figura 17 Aplicación ChatBot registrada



5.8.3. Configuración de Aplicación

Una vez creada la aplicación en la plataforma Meta for Developers, se procede a la incorporación de dos productos fundamentales para el desarrollo del sistema: WhatsApp y Webhooks. Ambos componentes resultan esenciales, dado que permiten la comunicación entre la aplicación y la interfaz de programación (API) que se desarrollará.

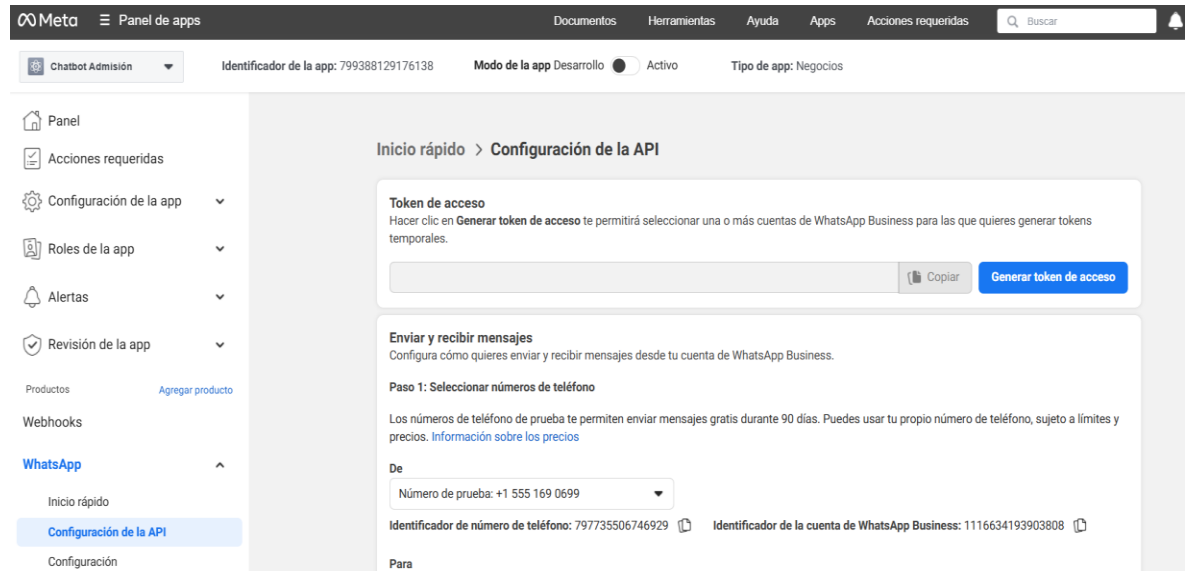
Figura 18 Configuración de la Aplicación



En el caso del producto WhatsApp, se genera un Token de Acceso, el cual constituye la credencial de autenticación necesaria para establecer la conexión segura

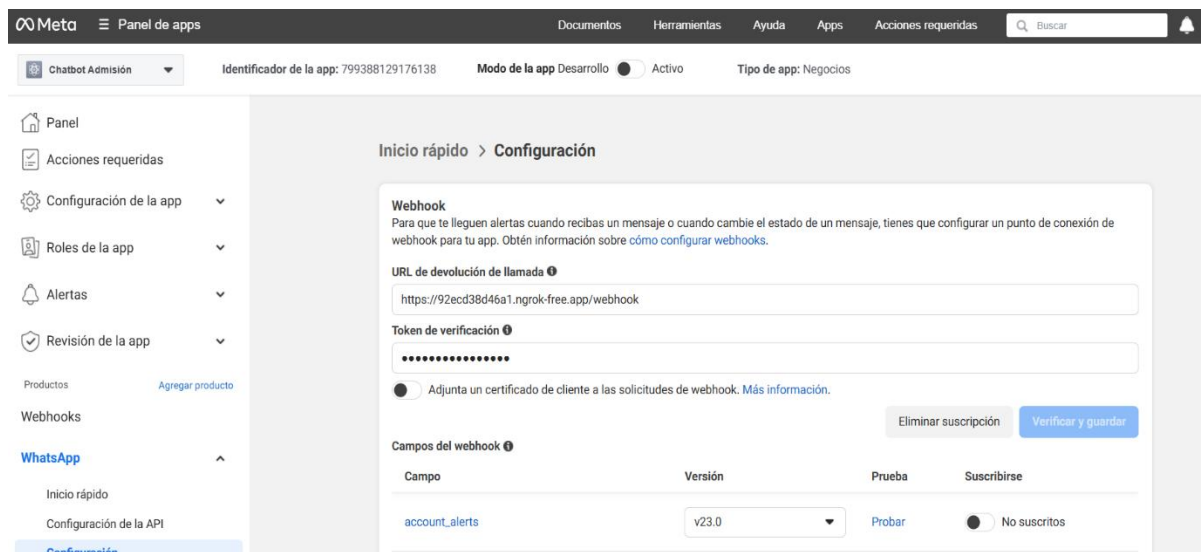
entre la API de la aplicación y el servicio de mensajería de WhatsApp. Dicho token es indispensable para la posterior integración y envío de mensajes automatizados.

Figura 19 Configuración de token de acceso temporal



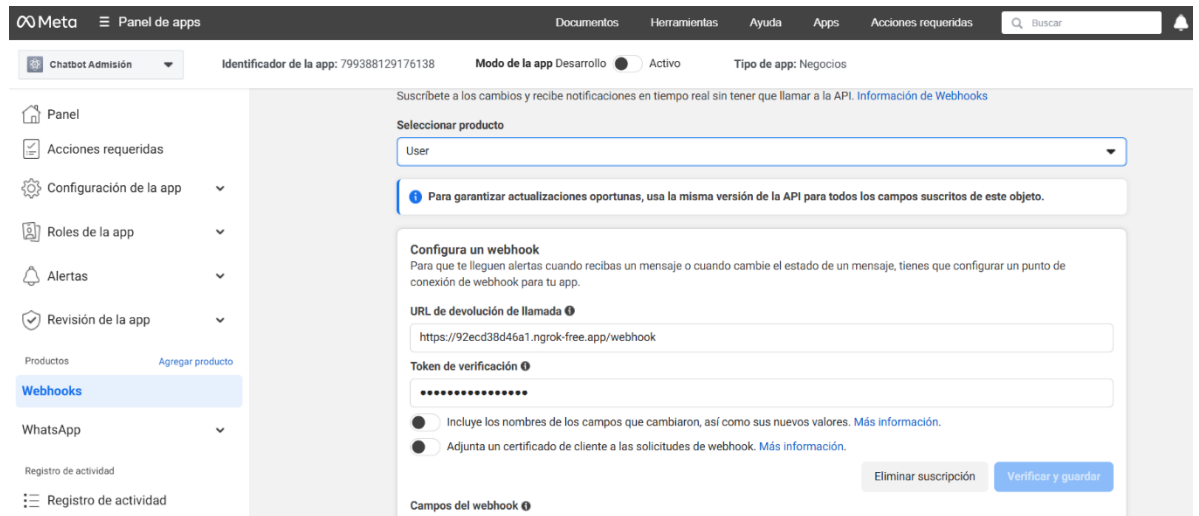
Por otra parte, se procede a la configuración de los parámetros de comunicación. Para ello, se registra la URL pública en la que se encuentra desplegado el servicio de la API, se define el token de verificación que asegura la validez de la conexión y, finalmente, se habilita la opción correspondiente a “messages” dentro de los campos de suscripción del webhook. Esta configuración garantiza que la aplicación pueda recibir, procesar y responder en tiempo real a los eventos asociados al servicio de mensajería.

Figura 20 Configuración de URL pública



En el producto Webhook la configuración es similar a la realizada en la configuración de Whatsapp ya que se registra tanto la URL pública como el token de verificación.

Figura 21 Configuración de Webhook



The screenshot displays the Meta Chatbot configuration interface. The top navigation bar includes 'Meta', 'Panel de apps', and various utility links like 'Documentos', 'Herramientas', 'Ayuda', 'Apps', and 'Acciones requeridas'. A search bar is also present. The main header shows 'Chatbot Admisión', the app ID '799388129176138', the app mode 'Desarrollo' (Development), and the app type 'Negocios' (Business). The left sidebar contains a menu with options like 'Panel', 'Acciones requeridas', 'Configuración de la app', 'Roles de la app', 'Alertas', 'Revisión de la app', 'Productos', 'Webhooks', 'WhatsApp', and 'Registro de actividad'. The 'Webhooks' section is currently selected. The main content area is titled 'Configura un webhook' and includes a 'Seleccionar producto' dropdown menu set to 'User'. Below this, there is a note about using the same API version for all subscribed fields. The 'Configura un webhook' section contains a 'URL de devolución de llamada' (Callback URL) field with the value 'https://92ecd38d46a1.ngrok-free.app/webhook', a 'Token de verificación' (Verification Token) field with a masked value, and two radio button options: 'Incluye los nombres de los campos que cambiaron, así como sus nuevos valores' (selected) and 'Adjunta un certificado de cliente a las solicitudes de webhook'. At the bottom, there is a 'Campos del webhook' section and buttons for 'Eliminar suscripción' and 'Verificar y guardar'.

5.8.4. Pruebas de Chatbot desde WhatsApp

Con el fin de validar la correcta integración y funcionamiento del sistema, se realizaron pruebas iniciales desde la aplicación de mensajería WhatsApp. Estas pruebas permitieron comprobar la capacidad del chatbot para establecer comunicación con el usuario final, además su capacidad para interpretar mensajes y generar respuestas acordes al contexto de la interacción. El proceso de prueba inicial consideró los siguientes escenarios:

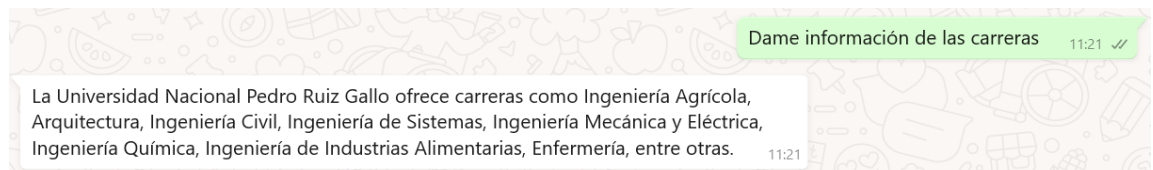
- **Mensaje inicial de saludo:** En primer lugar, se verificó que, al iniciar la conversación, el chatbot responda con un mensaje de bienvenida, estableciendo un tono cordial y adecuado para la interacción permitiendo una interacción amable con el usuario

Figura 22 Prueba mensaje inicial Chatbot



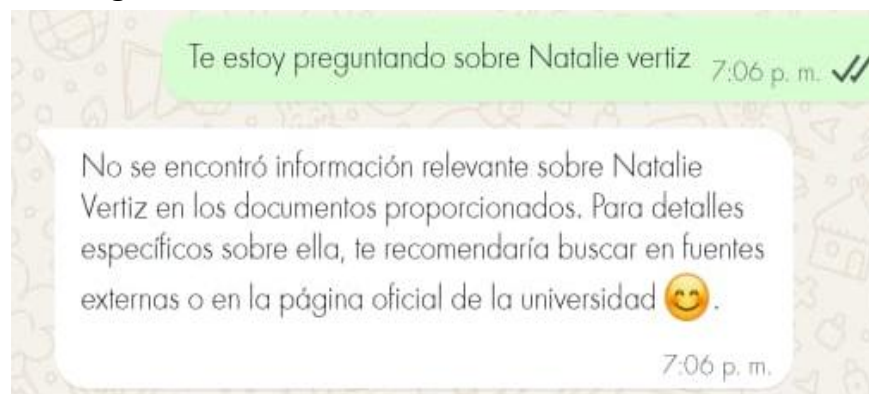
- **Pregunta relacionada con la universidad:** Se evaluó también la capacidad del sistema para responder a las consultas realizadas respecto a los temas de admisión. Como se puede observar en la imagen, la información brindada por el chatbot ante la consulta realizada está de acuerdo con los documentos proporcionados en el entrenamiento.

Figura 23 Prueba inicial del contexto de la Universidad



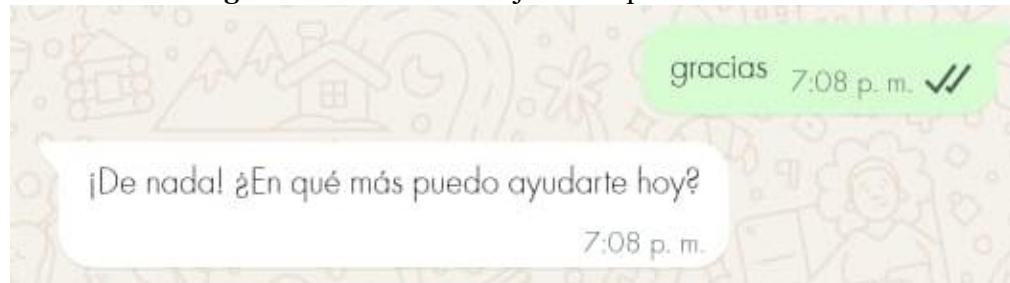
- **Pregunta de un contexto diferente:** Se comprobó la capacidad del chatbot para los casos en que los mensajes que le proporcionen estén fuera del contexto entrenado con el fin de observar si logra brindar alguna respuesta o realiza una respuesta en base al entrenamiento realizado. Como se puede observar en la imagen, en caso el contexto sea fue al de admisión de la Universidad Nacional Pedro Ruiz Gallo el chatbot indica que no cuenta con esa información.

Figura 24 Prueba con contexto no universitario



- **Mensaje de despedida:** Por último, se comprobó que el chatbot termine la interacción de forma correcta, emitiendo un mensaje el cual brinde una mejor experiencia al usuario. Antes de cerrar finalmente la conversación indica si el usuario tiene alguna otra duda, si en caso existiera seguiría con la conversación, caso contrario finalizaría la iteración

Figura 25 Prueba mensaje de despedida



5.10. Desarrollo del backend para la API

El desarrollo del backend se realizó haciendo uso de Node.js con el framework Express, lo cual permitió crear una API eficiente para la comunicación entre los distintos componentes utilizados ya que actúa como un intermediario entre la plataforma Meta Cloud API y el asistente de OpenAI, permitiendo un continuo y confiable flujo de datos.

- **Verificación del Webhook:** Es el encargado de validar la conexión exitosa entre la API y la plataforma Meta Cloud API lo cual garantiza que las solicitudes enviadas sean legítimas y que exista una correcta configuración del webhook. Este endpoint se ejecuta solo durante la configuración inicial del webhook en la plataforma Meta Cloud.
- **Recepción y procesamiento de mensajes:** Actúa como punto de entrada y salida ya que permite identificar las consultas enviadas por los usuarios a través de WhatsApp y evitar la duplicidad en la consultas y respuestas enviadas. Una vez realizada las verificaciones realiza el llamado a la función que se comunicará con el asistente para que finalmente la respuesta sea reenviada al usuario mediante una solicitud HTTP a la API de Meta Cloud.

- **Consulta al asistente:** Se encarga de establecer la comunicación directa con el asistente configurado creando un hilo de conversación por cada interacción y enviando el mensaje del usuario al modelo, luego de ello la función espera la ejecución del asistente hasta obtener una respuesta, formatearla y enviarla al usuario.
- **Sanitización de mensaje:** Esta función se encarga de limpiar el mensaje antes de enviarlo al usuario. Esto se realiza debido a que existen veces en que el asistente retorna la fuente de donde fue sacada la información y este detalle no debe llegar al usuario final.

5.11. Despliegue de API en GitHub y Render

Para asegurar que el sistema estuviera disponible en todo momento y pudiera accederse de forma remota, la API desarrollada con Node.js se implementó en la nube haciendo uso de GitHub y Render, plataformas que facilitan la automatización en el control de versiones, despliegue y la ejecución del servicio.

5.11.1. Publicación del código en GitHub

En primer lugar, se creó un repositorio en GitHub que permite el acceso de manera remota lo cual, ayuda a llevar un control de versiones y una integración continua, además que sirve funciona como base para el despliegue en la plataforma Render.

El código fuente se cargó haciendo uso de GIT, aplicando los comandos básicos como “init, add, commit y push”, excluyendo los archivos sensibles mediante el archivo “.gitignore”, en este archivo se incluyó el archivo “.env” el cual contiene datos privados como keys y tokens.

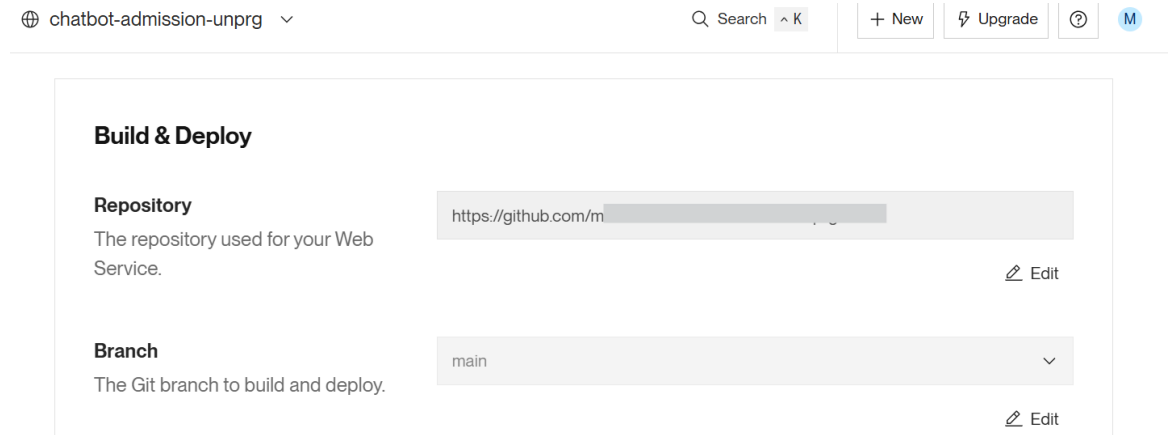
El excluir este archivo permite asegurar que el repositorio solo contenga el código fuente y no información confidencial del asistente o de la API de WhatsApp. De esta forma, garantizamos la seguridad de las credenciales, la integridad del proyecto y el cumplimiento de buenas prácticas en el manejo de datos sensibles durante el desarrollo y despliegue del sistema.

5.11.2. Despliegue en Render

Posteriormente, se realizó el despliegue de la API haciendo uso de la plataforma Render, la cual cuenta permite una fácil integración con el repositorio GitHub.

Render permite la creación de servicios web que se actualizan automáticamente cada vez que se realiza un nuevo commit en el repositorio.

Figura 26 Configuración del repositorio y rama



chatbot-admission-unprg

Search ^ K

+ New Upgrade ? M

Build & Deploy

Repository
The repository used for your Web Service.

https://github.com/m

Edit

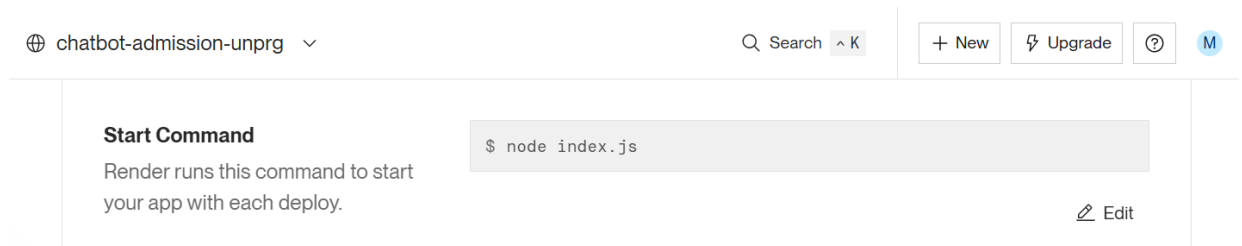
Branch
The Git branch to build and deploy.

main

Edit

Durante la configuración inicial, se seleccionó el repositorio correspondiente, se especificó el comando de inicio (node index.js) y se configuró el entorno de ejecución de Node.js.

Figura 27 Configuración comando de inicio



chatbot-admission-unprg

Search ^ K

+ New Upgrade ? M

Start Command
Render runs this command to start your app with each deploy.

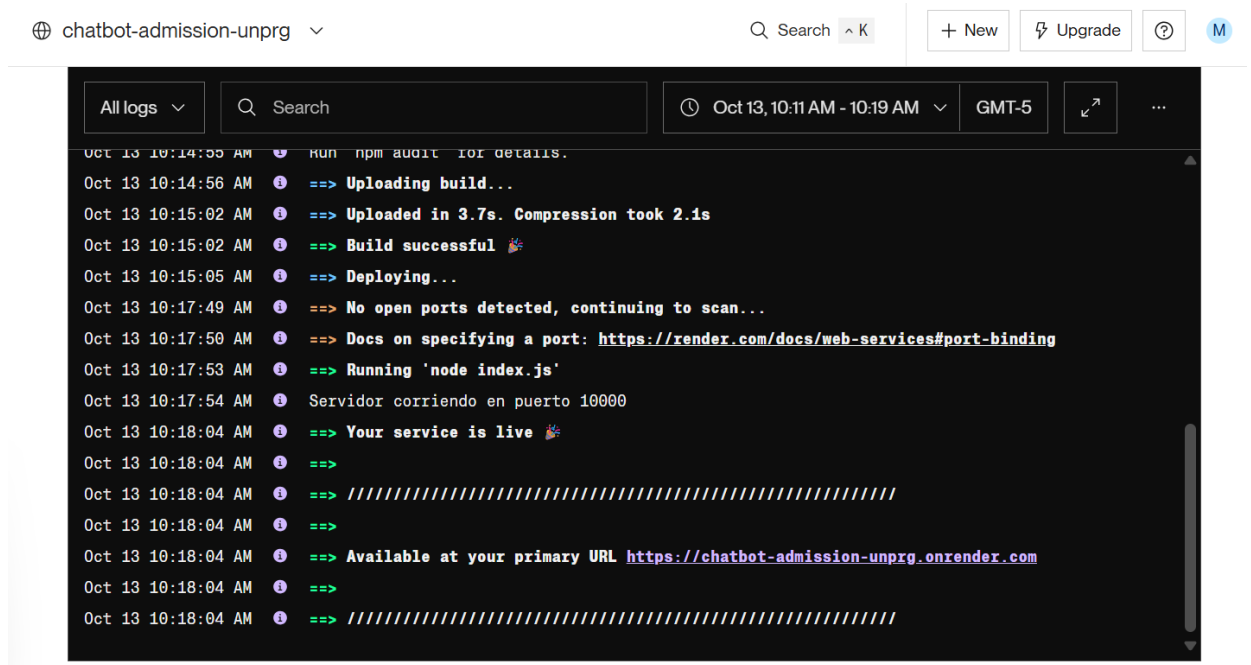
\$ node index.js

Edit

Una vez completado el proceso de construcción, la aplicación quedó alojada en un dominio público proporcionado por Render, el cual actúa como URL del webhook utilizado por Meta Cloud para la recepción de mensajes de WhatsApp.

Esta arquitectura serverless garantiza que la aplicación se mantenga activa, escalable y disponible sin requerir infraestructura física ni servidores dedicados.

Figura 28 Log de despliegue en Render



The screenshot shows the Render deployment interface for a project named 'chatbot-admission-unprg'. The top navigation bar includes a search bar, a '+ New' button, an 'Upgrade' button, and a user profile icon. The main log area displays a series of status messages with timestamps from October 13, 2023, at 10:14 AM to 10:18 AM. The log indicates a successful deployment process, including building, uploading, and deploying the application. It also shows a port scan result and the final live URL: <https://chatbot-admission-unprg.onrender.com>.

```
Oct 13 10:14:55 AM Run npm audit for details.
Oct 13 10:14:56 AM ==> Uploading build...
Oct 13 10:15:02 AM ==> Uploaded in 3.7s. Compression took 2.1s
Oct 13 10:15:02 AM ==> Build successful 🎉
Oct 13 10:15:05 AM ==> Deploying...
Oct 13 10:17:49 AM ==> No open ports detected, continuing to scan...
Oct 13 10:17:50 AM ==> Docs on specifying a port: https://render.com/docs/web-services#port-binding
Oct 13 10:17:53 AM ==> Running 'node index.js'
Oct 13 10:17:54 AM Servidor corriendo en puerto 10000
Oct 13 10:18:04 AM ==> Your service is live 🎉
Oct 13 10:18:04 AM ==>
Oct 13 10:18:04 AM ==> //////////////////////////////////////
Oct 13 10:18:04 AM ==>
Oct 13 10:18:04 AM ==> Available at your primary URL https://chatbot-admission-unprg.onrender.com
Oct 13 10:18:04 AM ==>
Oct 13 10:18:04 AM ==> //////////////////////////////////////
```

5.11.3. Configuración de variables de entorno

Render permite la gestión segura de credenciales mediante la definición de variables de entorno, las cuales reemplazan la necesidad de almacenar claves directamente en el código fuente.

En este proyecto, se configuraron las siguientes variables esenciales:

- **OPENAI_API_KEY:** Clave privada utilizada para la autenticación con el servicio de OpenAI.
- **WHATSAPP_TOKEN:** Token de acceso a la API de Meta Cloud, necesario para el envío y recepción de mensajes.
- **PHONE_NUMBER_ID:** Identificador del número de teléfono asociado a la cuenta de WhatsApp Cloud.
- **ASSISTANT_ID:** Identificador del asistente configurado en OpenAI, con el cual se mantiene la comunicación directa desde la API.

Estas variables fueron configuradas directamente desde el panel de Render, garantizando la protección de las credenciales y evitando su exposición pública en el repositorio.

Figura 29 Configuración de variables de entorno

chatbot-admission-unprg Search ^ K + New Upgrade ? M

Environment Variables

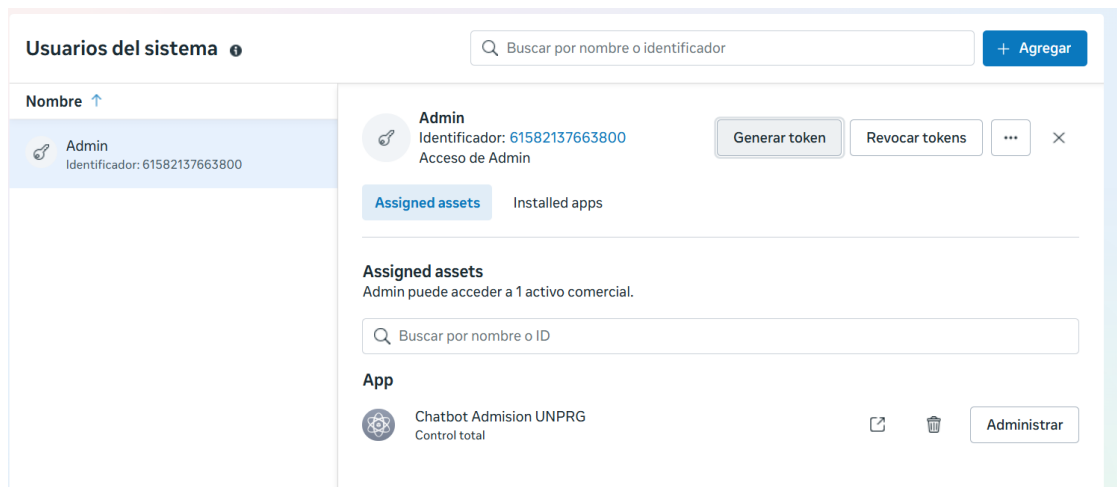
Set environment-specific config and secrets (such as API keys), then read those values from your code. [Learn more.](#) Export Edit

KEY	VALUE
ASSISTANT_ID	
OPENAI_API_KEY	
PHONE_NUMBER_ID	
VERIFY_TOKEN	
WHATSAPP_TOKEN	

5.11.4. Creación del token permanente de Meta Cloud

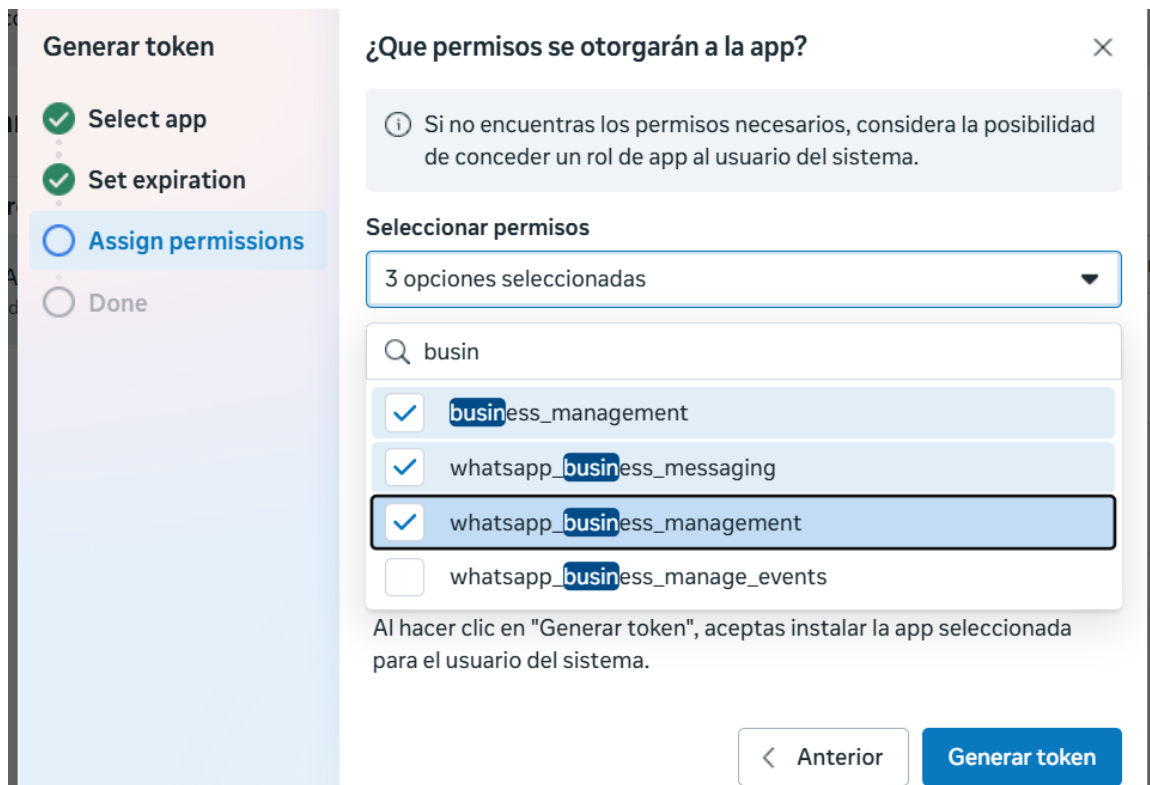
Para que la aplicación pueda mantener una comunicación continua con la API de WhatsApp Cloud, fue necesario reemplazar el token temporal generado por defecto por un token de acceso permanente. Para ello se accedió al panel de la aplicación en Meta for Developers, y en el apartado de Configuración de la API de WhatsApp, verificamos que se cuente con un Usuario del sistema.

Figura 30 Usuario del sistema en Meta for Developers



Seleccionamos la opción “Generar token de acceso permanente” en donde añadimos la aplicación a elegir, la expiración que se desea y los permisos necesarios.

Figura 31 Configuración de token permanente

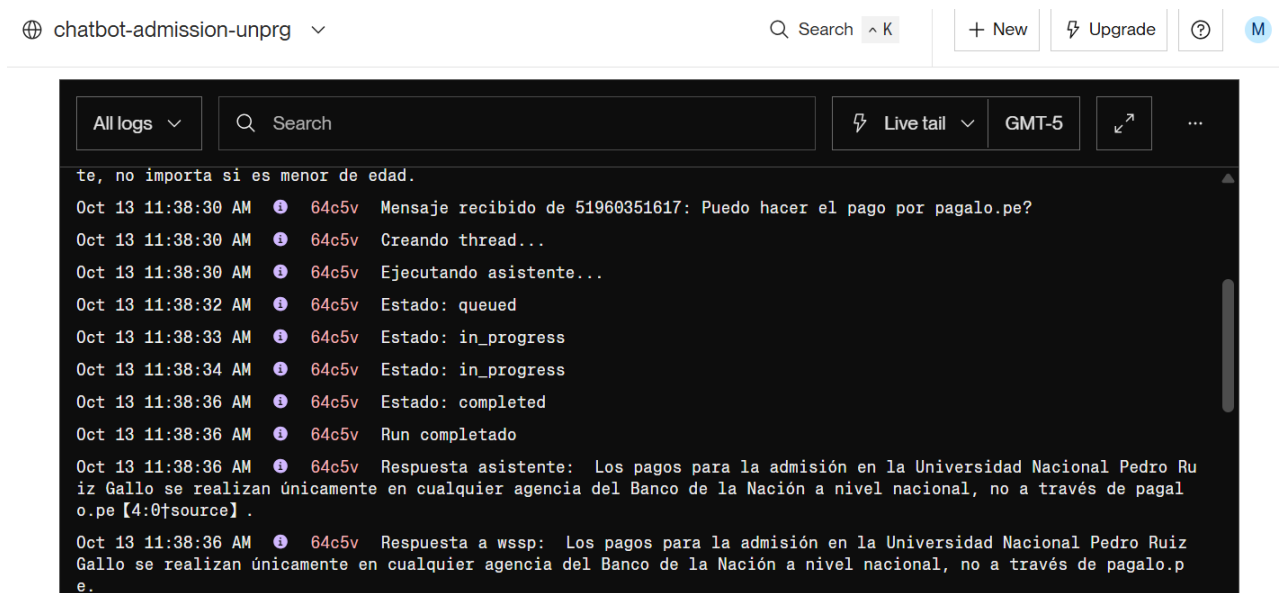


Luego de ello asociamos el token con la cuenta de Meta Business correspondiente y actualizamos la variable `WHATSAPP_TOKEN` en Render con el nuevo valor del token.

5.11.5. Validación de despliegue

Se verificó que los mensajes enviados desde WhatsApp fueran correctamente recibidos por el webhook alojado en Render, procesados por el asistente de OpenAI y devueltos al usuario a través de la misma plataforma. El despliegue exitoso de la API en Render permite que el chatbot esté disponible las 24 horas del día, ofreciendo un canal confiable y automatizado para la atención de consultas relacionadas con los procesos de admisión universitaria.

Figura 32 Validación de despliegue con pruebas de usuario



5.12. Evaluación y comparación de resultados

Para evaluar la precisión y coherencia de las respuestas del chatbot, se realizó un proceso de comparación entre las respuestas emitidas por el chatbot y la información oficial contenida en los documentos utilizados para el entrenamiento, específicamente en el documento de preguntas frecuentes proporcionado por el área de admisión y el prospecto de admisión de la Universidad Nacional Pedro Ruiz Gallo.

Esta comparación permite determinar el grado de concordancia entre el conocimiento adquirido por el chatbot y las fuentes originales, evaluando la exactitud, claridad y consistencia de las respuestas generadas.

Además, este proceso ayuda a medir la efectividad del chatbot como herramienta para la atención automatizada de las consultas de los usuarios, verificando que las respuestas que de sean correctas y contextualizadas de acuerdo a la información institucional disponible

A continuación se mostrarán imágenes en donde realizaremos una comparativa de las respuestas brindadas por el chatbot y las respuestas que se encuentran en los documentos proporcionados.

Figura 33 Evaluación Chatbot - ¿Cuándo es el examen de admisión?

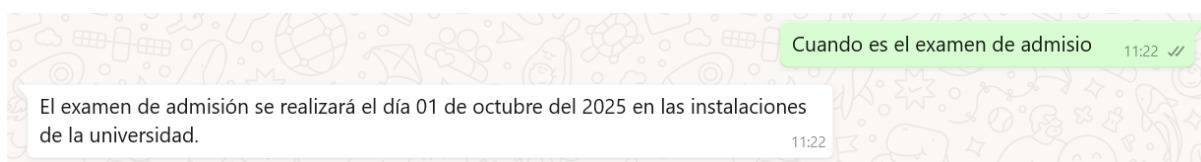


Figura 34 Respuesta documento preguntas frecuentes - ¿Cuándo es el examen de admisión?

Pregunta: ¿Cuándo es el examen de admisión?

Respuesta: El examen de admisión se realizará el día 01 de octubre del 2025 en las instalaciones de la universidad.

Figura 35 Evaluación Chatbot - Tipo de pago beca 18

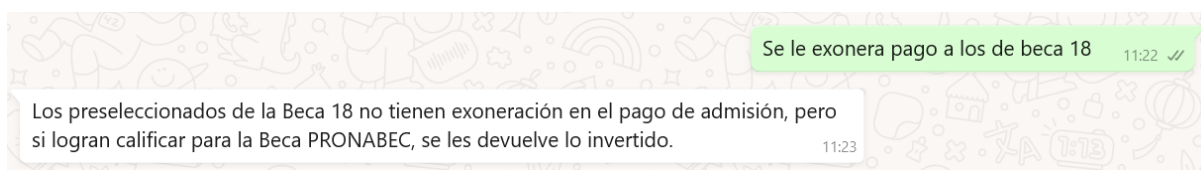


Figura 36 Respuesta documento preguntas frecuentes - Tipo de pago beca 18

Pregunta: ¿A los preseleccionados de la beca 18 se les exonera el pago?

Respuesta: Los preseleccionados de la beca 18, no tienen exoneración, pagan igual que todo postulante, sin embargo, si se logra CALIFICAR PARA LA BECA PRONABEC se devuelve todo lo invertido.

Figura 37 Evaluación Chatbot - Medios de pago

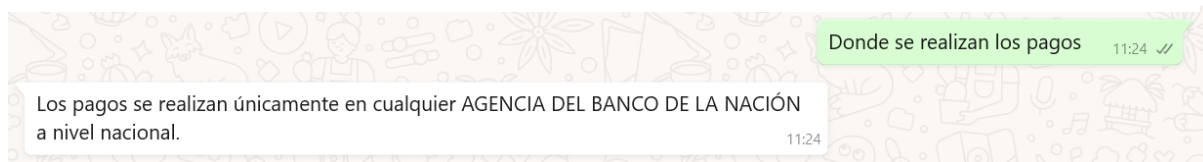


Figura 38 Respuesta documento preguntas frecuentes - Medios de pago

Pregunta: ¿Dónde se realizan los pagos?

Respuesta: Los pagos se realizan solamente en cualquier AGENCIA DEL BANCO DE LA NACIÓN a nivel nacional

Figura 39 Evaluación Chatbot - Documento para realizar el pago

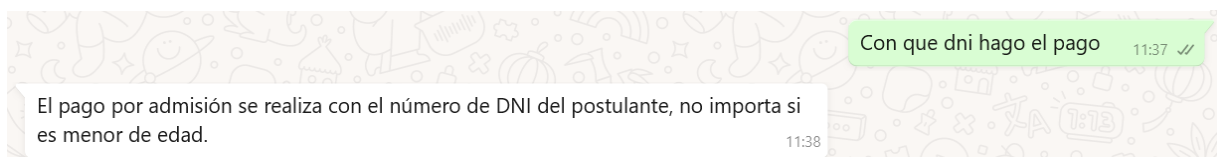


Figura 40 Respuesta documento preguntas frecuentes - Documento para realizar el pago

Pregunta: ¿Con qué DNI se debe hacer el pago por admisión?

Respuesta: El pago se realiza con el número de DNI del postulante

Pregunta: ¿Se puede realizar el pago por admisión con el DNI de un menor de edad?

Respuesta: El pago se realiza con el número de DNI del postulante, no hay ningún inconveniente si este es menor de edad ya que igual se realiza.

Figura 41 Evaluación Chatbot - Pago por pagalo.pe

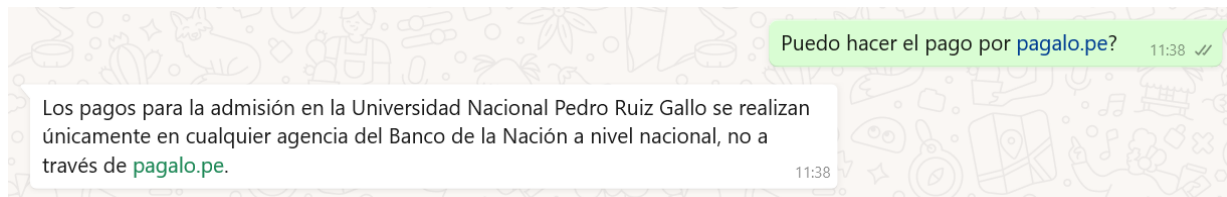


Figura 42 Respuesta documento preguntas frecuentes - Pago por pagalo.pe

Pregunta: ¿Es posible realizar el pago por el aplicativo pagalo.pe?

Respuesta: Los pagos se realizan solamente en cualquier AGENCIA DEL BANCO DE LA NACIÓN a nivel nacional.

Figura 43 Evaluación Chatbot - Pasos para la inscripción

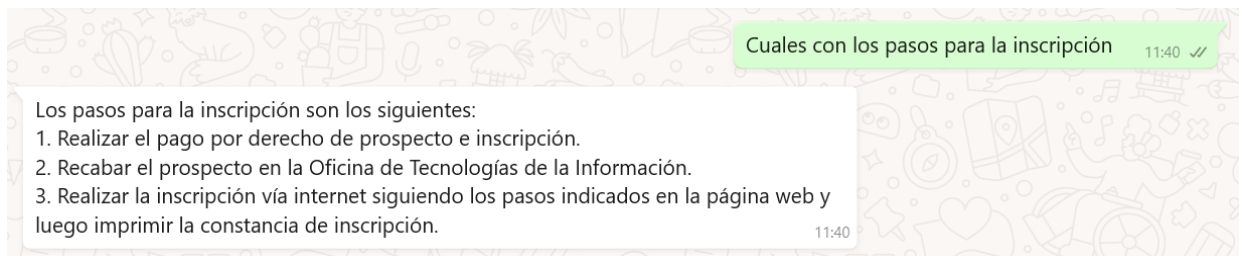


Figura 44 Respuesta documento preguntas frecuentes - Pasos para la inscripción

Pregunta: Pasos para la inscripción

Respuesta:

1. Cancelar el derecho de prospecto en el banco de la Nación con el Código: 489.
2. Para recoger prospecto: Presentar voucher de pago del prospecto.
3. Para recoger carnet: Presentar 2 voucher de pago por derecho de:
 - a. Prospecto
 - b. Inscripción.

Figura 45 Evaluación Chatbot - Horario de atención OTI

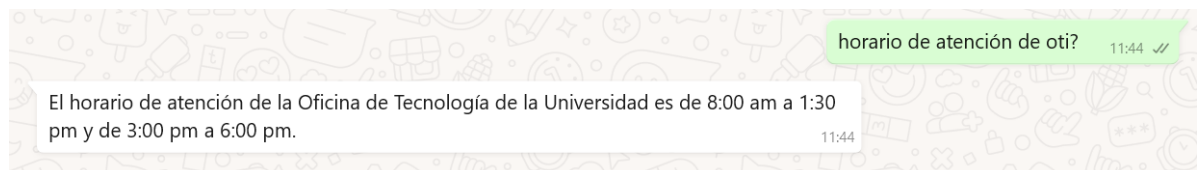


Figura 46 Respuesta documento preguntas frecuentes - Pasos para la inscripción

Pregunta: ¿Cuáles son los horarios de atención de la Oficina de Tecnología de la Universidad?

Respuesta: La atención es de 8:00 am - 1:30 pm y 3:00 pm a 6:00 pm.

Figura 47 Evaluación Chatbot - Vacantes para la carrera de Ingeniería de Sistemas

Cuántas vacantes hay para la carrera de ingeniería de sistemas 15:48 ✓

Para la carrera de Ingeniería de Sistemas, hay un total de 31 vacantes disponibles. 15:49

Figura 48 Respuesta prospecto – Vacantes para la carrera de Ingeniería de Sistemas

DISTRIBUCIÓN DE VACANTES													
Aprobada con Resolución N° 068-2025-CU													
CÓDIGO DEL PROGRAMA	PROGRAMA DE ESTUDIOS	Titulados y Graduados	Traslado Externo Nacional o Internacional	Traslado Interno	Dos Primeros Puestos	Deportistas Destacados Ley N° 28036	Personas con Discapacidad Ley N° 29973	Deposición de Alta Competencia - Produc Ley 30476	Víctimas del Terrorismo Ley N° 30771	Plan Luchas de Reparaciones - PIR Ley N° 28592	Examen Ordinario	TOTAL	
P01	INGENIERÍA AGRÍCOLA	1	1	1	2	1	1	1	1	1	33	43	
P02	ARQUITECTURA	0	0	0	0	0	0	0	0	0	36	36	
P03	INGENIERÍA CIVIL	1	1	1	2	1	1	1	1	1	34	44	
P04	INGENIERÍA DE SISTEMAS	2	2	2	5	1	1	0	0	0	18	31	
P05	INGENIERÍA MECÁNICA Y ELÉCTRICA	0	0	0	2	1	3	1	2	2	32	43	

Figura 49 Evaluación Chatbot - Encargados del proceso

Quiénes son los encargados del proceso de admisión 15:49 ✓

Los encargados del proceso de admisión en la Universidad Nacional Pedro Ruiz Gallo son el Consejo Universitario, el Consejo de Facultad, la Dirección de la Escuela de Posgrado, la Dirección de Admisión y la Comisión de Admisión. 15:49

Figura 50 Respuesta prospecto - Encargados del proceso

**TÍTULO II
DISPOSICIONES GENERALES
CAPÍTULO I
DE LOS RESPONSABLES DEL PROCESO DE ADMISIÓN**

Artículo 6.- Son responsables del proceso de admisión:

1. Consejo Universitario
2. Consejo de Facultad
3. La Dirección de la Escuela de Posgrado
4. Dirección de Admisión
5. Comisión de Admisión

Figura 51 Evaluación Chatbot - Cantidad de preguntas para ingeniería

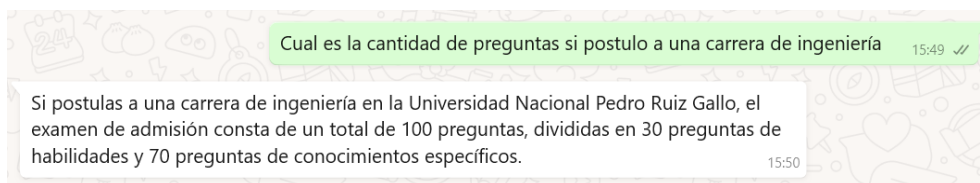


Figura 52 Respuesta prospecto - Cantidad de preguntas para ingeniería

**SUB CAPÍTULO V: DE LA ESTRUCTURA, CALIFICACIÓN
Y PUBLICACIÓN DE RESULTADOS DEL PROCESO DE
ADMISIÓN**

Artículo 58.-El examen de admisión para la modalidad ordinario y extraordinario comprende dos secciones, con una duración de tres (3) horas y un total de 100 preguntas, según el Anexo 01:

- Sección Habilidades (30 preguntas)
- Sección Conocimientos (70 preguntas)

Figura 53 Evaluación Chatbot - Proceso para inscribirse y rendir el examen

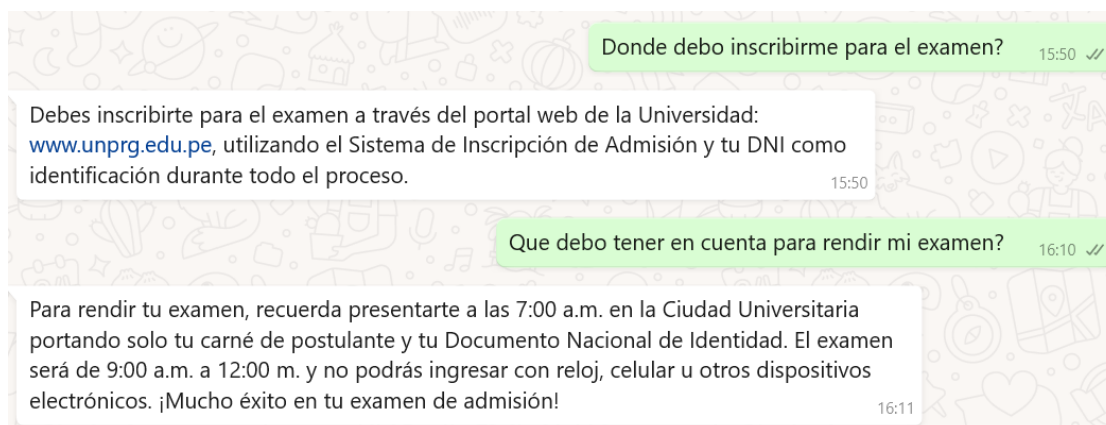


Figura 54 Respuesta prospecto - Proceso de inscripción

I. PARA LA INSCRIPCIÓN

- 1º. Realice el pago por derecho de prospecto e inscripción al Concurso de Admisión, en agencias del Banco de la Nación a nivel nacional de acuerdo a la modalidad a la que postula y al colegio o Universidad de Procedencia: Nacional o Particular.
- 2º. En el voucher de pago emitido por el Banco de la Nación debe estar consignado correctamente el número de documento de identidad del postulante (D.N.I.)
- 3º. Recabar el Prospecto en la Oficina de Tecnologías de la Información - Ciudad Universitaria - Lambayeque), previa presentación del voucher de pago.
- 4º. La inscripción se realizará vía internet. Al día siguiente de haber cancelado el derecho de inscripción, el postulante podrá ingresar a la página web:
<https://inscripciones.unprg.edu.pe/>

Figura 55 Respuesta prospecto - Proceso para rendir el examen

II. PARA RENDIR EL EXAMEN

- 1º. Presentarse a las 7.00 a.m, en la Ciudad Universitaria (Calle: Juan XXIII N° 391- Lambayeque), PORTANDO SOLO: carné de postulante (sin enmendaduras) y su Documento Nacional de Identidad.
- 2º. EL EXAMEN SERÁ DE 09.00 a.m. a 12.00 m.
- 3º. EL POSTULANTE NO PODRÁ INGRESAR A RENDIR EL EXAMEN SI NO PORTA SU CARNE DE POSTULANTE Y DOCUMENTO DE IDENTIFICACIÓN.
- 4º. EL POSTULANTE NO PODRÁ INGRESAR A RENDIR EL EXAMEN PORTANDO RELOJ, CELULAR, IPOD, IPAD, TABLE, CALCULADORA, RADIO O ALGÚN OTRO TIPO DE ARTEFACTO ELECTRÓNICO O DE COMUNICACIÓN. SI SE LE DETECTA, SERÁ RETIRADO AUTOMÁTICAMENTE DEL EXAMEN.
- 5º. En el local donde rendirá el examen estará publicado el listado de postulantes por orden alfabético. Ubique por sus apellidos y nombre, el aula donde le corresponde rendir el examen y luego diríjase hacia ella.
- 6º. En presencia del supervisor de aula usted deberá firmar e imprimir su huella dactilar en el formato de verificación de identificación y registro de asistencia.

Figura 56 Evaluación Chatbot - Fecha de pago e inscripción

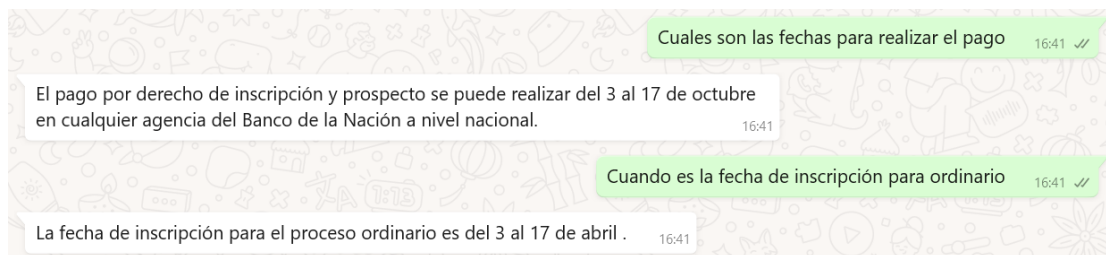


Figura 57 Respuesta prospecto - Fecha de pago e inscripción

CALENDARIO DE ACTIVIDADES			
MES			ACTIVIDADES
FEBRERO	MARZO	ABRIL	
28	15		Pago para todas las modalidades: ORDINARIO y EXTRAORDINARIO Dos (2) Primeros Puestos, Traslado externo Nacional o Internacional, Titulados o Graduados, Examen Ordinario, Deportistas Destacados, Personas con Discapacidad, Deportista de Alta Competencia, Víctimas del terrorismo y Plan Integral de Reparaciones. <i>Presencial: Agencia del Banco de la Nación.</i>
	3 al 17		Inscripción para todas las modalidad ORDINARIO y EXTRAORDINARIO <i>Virtual: vía web, inscripciones.unprg.edu.pe/</i>
	4 al 18		Entrega de camé y prospecto para todas las modalidades: ORDINARIO y Extraordinario <i>Presencial: instalaciones de la Sala de Lectura de la Escuela de Posgrado - Ciudad Universitaria</i>

Figura 58 Evaluación Chatbot - Montos de pago

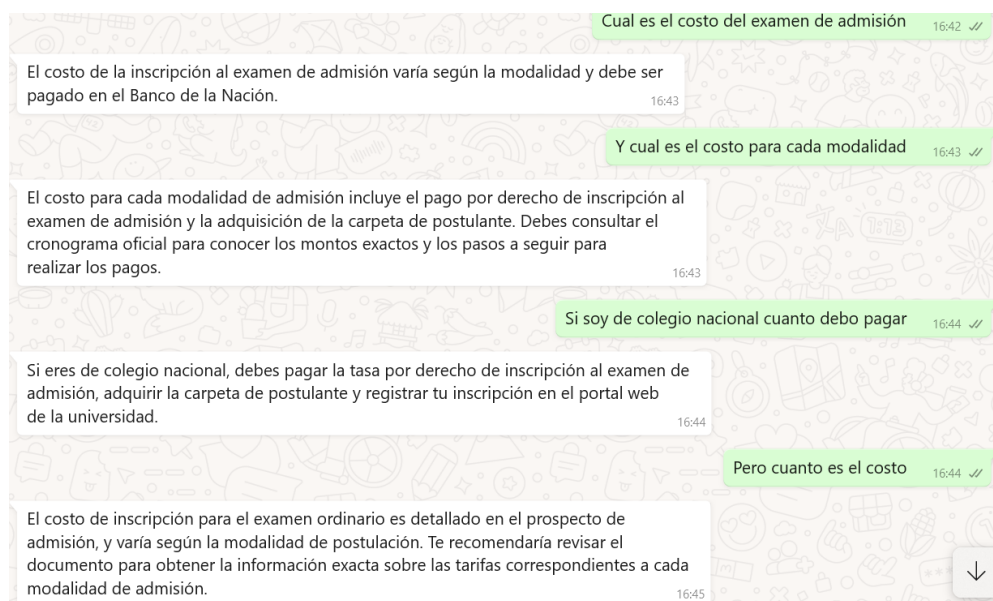


Figura 59 Respuesta prospecto – Montos de pago

TASAS DE INSCRIPCIÓN		
Tasas por derecho de adquisición de prospecto de admisión e inscripción de postulante.		
TRANSACCIÓN: 9135		
Prospecto	S/. 30.00 Código 489	
Modalidad de Postulación	Colegio o Universidad de Procedencia	
	Nacional	Particular
Postulantes al Examen Ordinario	S/. 280.00 Código 346	S/. 380.00 Código 345
Postulantes al Examen Extraordinario: -Dos Primeros Puestos -Deportistas Destacados -Personas con Discapacidad -Deportista de Alta Competencia -Victimas del terrorismo		
Traslado externo nacional o internacional	S/. 1030.00 Código 997	S/. 1230.00 Código 998
Titulados o Graduados	S/. 1630.00 Código 999	
Duplicado del carné de postulante	S/. 10.00 Código 344	

6. RESULTADOS

Para evaluar el grado de respuesta del chatbot, se aplicó un proceso de validación comparativa entre las respuestas generadas por el sistema y la información oficial contenida en los documentos de entrenamiento.

Se definieron tres criterios de análisis: exactitud, relevancia y claridad, mediante los cuales se clasificaron las respuestas en correctas, parciales o incorrectas.

Tabla 4 Criterios de análisis

Criterio	Descripción	Escala de evaluación
Exactitud	Evalúa si la respuesta entregada por el chatbot coincide con la información oficial del documento fuente.	Correcta / Parcial / Incorrecta
Relevancia	Mide si la respuesta es coherente con la pregunta del usuario y se mantiene dentro del contexto del tema.	Alta / Media / Baja
Claridad	Determina si la respuesta está redactada de manera comprensible, sin errores gramaticales o confusiones.	Clara / Ambigua / Confusa

A continuación, se presentará una tabla comparativa con las preguntas realizadas, respuestas esperadas, fuente, evaluación de exactitud, tipo de fuente, observaciones.

Tabla 5 Comparativa respuestas Chatbot

Nº	Pregunta del usuario	Respuesta del chatbot	Fuente	Evaluación de exactitud	Tipo de fuente	Observaciones
1	Hola	¡Hola! ¿En qué puedo ayudarte hoy con respecto al proceso de admisión en la Universidad Nacional Pedro Ruiz Gallo? 😊	-	-	-	Es contexto general
2	¿Cuándo es el examen de admisión?	El examen de admisión se realizará el día 01 de octubre del 2025 en las instalaciones de la universidad.	Preguntas frecuentes	Correcta	Información textual	-
3	¿Se le exonera pago a los de beca 18?	Los preseleccionados de la Beca 18 no tienen exoneración en el pago de admisión, pero si logran calificar para la Beca PRONABEC, se les devuelve lo invertido.	Preguntas frecuentes	Correcta	Información textual	-

4	¿Dónde se realizan los pagos?	Los pagos se realizan únicamente en cualquier AGENCIA DEL BANCO DE LA NACIÓN a nivel nacional.	Preguntas frecuentes	Correcta	Información textual	-
5	¿Con que dni hago el pago?	El pago por admisión se realiza con el número de DNI del postulante, no importa si es menor de edad.	Preguntas frecuentes	Correcta	Información textual	-
6	¿Puedo hacer el pago por pagalo.pe ?	Los pagos para la admisión en la Universidad Nacional Pedro Ruiz Gallo se realizan únicamente en cualquier agencia del Banco de la Nación a nivel nacional, no a través de pagalo.pe.	Preguntas frecuentes	Correcta	Información textual	-
7	¿Cuáles son los pasos para la inscripción?	Los pasos para la inscripción son los siguientes:	Preguntas frecuentes	Correcta	Información textual	-

		1. Realizar el pago por derecho de prospecto e inscripción. 2. Recabar el prospecto en la Oficina de Tecnologías de la Información. 3. Realizar la inscripción vía internet siguiendo los pasos indicados en la página web y luego imprimir la constancia de inscripción.				
8	¿Horario de atención de oti?	El horario de atención de la Oficina de Tecnología de la Universidad es de 8:00 am a 1:30 pm y de 3:00 pm a 6:00 pm.	Preguntas frecuentes	Correcta	Información textual	-
9	¿Cuántas vacantes hay para la carrera de ingeniería de sistemas?	Para la carrera de Ingeniería de Sistemas, hay un total de 31 vacantes disponibles.	Prospecto de admisión	Correcta	Información en tabla	-

10	¿Quiénes son los encargados del proceso de admisión?	Los encargados del proceso de admisión en la Universidad Nacional Pedro Ruiz Gallo son el Consejo Universitario, el Consejo de Facultad, la Dirección de la Escuela de Posgrado, la Dirección de Admisión y la Comisión de Admisión.	Prospecto de admisión	Correcta	Información textual	-
11	¿Cuál es la cantidad de preguntas si postulo a una carrera de ingeniería?	Si postulas a una carrera de ingeniería en la Universidad Nacional Pedro Ruiz Gallo, el examen de admisión consta de un total de 100 preguntas, divididas en 30 preguntas de habilidades y 70 preguntas de conocimientos específicos.	Prospecto de admisión	Correcta	Información textual	-

12	¿Dónde debo inscribirme para el examen?	Debes inscribirte para el examen a través del portal web de la Universidad: www.unprg.edu.pe , utilizando el Sistema de Inscripción de Admisión y tu DNI como identificación durante todo el proceso.	Prospecto de admisión	Correcta	Información textual	-
13	¿Que debo tener en cuenta para rendir mi examen?	Para rendir tu examen, recuerda presentarte a las 7:00 a.m. en la Ciudad Universitaria portando solo tu carné de postulante y tu Documento Nacional de Identidad. El examen será de 9:00 a.m. a 12:00 m. y no podrás ingresar con reloj, celular u otros dispositivos electrónicos. ¡Mucho éxito en tu examen de admisión!	Prospecto de admisión	Correcta	Información textual	-

14	¿Mi DNI se venció, puedo dar el examen?	Si tu DNI está vencido, no podrás ingresar a rendir el examen de admisión según las normas de la Universidad Nacional Pedro Ruiz Gallo.	Prospecto de admisión	Correcta	Información textual	-
15	¿Cuáles son las fechas para realizar el pago?	El pago por derecho de inscripción y prospecto se puede realizar del 3 al 17 de octubre en cualquier agencia del Banco de la Nación a nivel nacional	Prospecto de admisión	Parcial	Información en tabla	-
16	¿Cuándo es la fecha de inscripción para ordinario?	La fecha de inscripción para el proceso ordinario es del 3 al 17 de abril.	Prospecto de admisión	Parcial	Información en tabla	-
17	¿Cuál es el costo del examen de admisión?	El costo de la inscripción al examen de admisión varía según la modalidad y debe ser pagado en el Banco de la Nación.	Prospecto de admisión	Parcial	Información en imagen	-

18	¿Si soy de colegio nacional cuanto debo pagar? Pero ¿cuánto es el costo?	Si eres de colegio nacional, debes pagar la tasa por derecho de inscripción al examen de admisión, adquirir la carpeta de postulante y registrar tu inscripción en el portal web de la universidad. El costo de inscripción para el examen ordinario es detallado en el prospecto de admisión, y varía según la modalidad de postulación. Te recomendaría revisar el documento para obtener la información exacta sobre las tarifas correspondientes a cada modalidad de admisión.	Prospecto de admisión	Parcial	Información en imagen	-
19	¿Cantidad de vacantes para la carrera de matematica?	Para la carrera de Matemática, hay 26 vacantes disponibles en la	Prospecto de admisión	Incorrecta	Información en tabla	-

		Universidad Nacional Pedro Ruiz Gallo.				
20	¿Cuántas vacantes hay para enfermería?	Para la carrera de Enfermería, hay un total de 27 vacantes disponibles.	Prospecto de admisión	Correcta	Información en tabla	-
21	¿Sabes quién es Ed sheeran?	Solo puedo brindarte información relacionada al proceso de admisión de la universidad	Prospecto de admisión	Correcta	Información textual	Es contexto general
22	Listo, gracias	¡De nada! Estoy aquí para ayudarte en lo que necesites. ¿Hay algo más en lo que pueda asistirte hoy?	-	-	Información textual	Es contexto general
23	Ok eso es todo	¡Perfecto! Si más adelante tienes alguna otra pregunta, ¡aquí estaré para ayudarte! ¡Éxito en tu proceso de admisión!	-	-	Información textual	Es contexto general

Posteriormente, se calcularon porcentajes de efectividad con el fin de determinar el grado de precisión del chatbot frente a las fuentes institucionales. Esto lo realizamos calculando las respuestas totalmente correctas sobre el total de preguntas realizadas

$$\textit{Precisión} = \frac{\textit{Respuestas correctas}}{\textit{Total de respuestas}} \times 100$$

$$\textit{Precisión} = \frac{15}{20} \times 100 = 75\%$$

Como podemos observar el porcentaje de preguntas correctas es de 75%, en este proceso de evaluación se ignoraron las preguntas de saludo, despedida y otros contextos.

En los casos donde las respuestas fueron clasificadas como parciales o incorrectas, se identificó que la principal causa de error está relacionada con el tipo de fuente de información utilizada. Específicamente, se observó que los contenidos provenientes de imágenes, tablas o formatos no estructurados generan dificultades en el proceso de comprensión y extracción de datos por parte del modelo.

Esto ocurre porque dichos elementos no siempre son interpretados de manera óptima durante el entrenamiento o indexación del conocimiento, lo que puede provocar omisiones, confusión en los valores numéricos o descontextualización del contenido textual.

A pesar de ello, el chatbot mantiene un rendimiento adecuado en respuestas basadas en texto plano, evidenciando la importancia de estructurar correctamente las fuentes para mejorar la precisión y consistencia de las respuestas futuras.

7. CONCLUSIONES.

- Se logró desarrollar un chatbot basado en IA para el Área de Admisión de la Universidad Nacional Pedro Ruiz Gallo el cual al hacer uso de tecnologías principales como GPT 3.5, WhatsApp Cloud API y Node.js permite ofrecer un sistema con una mejor eficiencia y calidad de atención a los usuarios finales, brindando respuestas rápidas, coherentes y contextualizadas.
- Luego de una detallada investigación y comparación de las distintas tecnologías existentes se logró determinar que el trabajo entre el Assistant API de OpenAI con la arquitectura REST en Node.js es la óptima en el desarrollo del chatbot, debido a su facilidad de integración, bajo costo de implementación, fácil acceso a la comunidad de soporte para dudas o consultas y su plataforma interactiva y flexible.
- En cuanto al tiempo de respuesta se logró demostrar que se obtuvo una reducción, permitiendo atender las consultas en pocos segundos y sin tener la necesidad de requerir una intervención humana, lo que permite mejorar la eficiencia operativa del Área de Admisión y que los usuarios finales accedan a una respuesta más rápida.
- Si bien los documentos brindados para el entrenamiento del modelo (preguntas frecuentes y prospecto de admisión) permitieron tener respuestas claras, precisas y consistentes se identificó que las respuestas parciales o imprecisas provienen directamente de fuentes como imágenes o tablas, por lo cual se concluye que el chatbot entiende mejor con información estructurada en texto plano.
- Según los resultados obtenidos se evidencia que el chatbot posee un alto grado de comprensión y aprendizaje en sus respuestas, alcanzando en la mayoría satisfacción en los usuarios finales lo cual permite demostrar la viabilidad técnica y funcional de la solución propuesta.
- Gracias al uso de herramientas como GitHub, Render, Meta Cloud API que son de fácil acceso y bajo costo en desarrollo el chatbot puede mantenerse activo en la nube con una

mínima gestión y operación continúa permitiendo en un futuro implementar el sistema en un entorno productivo.

8. RECOMENDACIONES

- Como se puede observar en las pruebas realizadas, la información brindada en tablas o imágenes pueden ser un poco confusas para el entrenamiento del modelo, por lo cual se sugiere organizar y estructurar toda la información utilizada para entrenar el chatbot de manera adecuada estableciendo como prioridad el formato de texto plano ya que este fue el que mejor respuesta brindó.
- Se recomienda gestionar un sistema de monitoreo que permita medir los indicadores especificados como la precisión de la respuesta, el tiempo de atención y el nivel de satisfacción del usuario con el fin de dar un seguimiento y poder realizar ajustes periódicos en el modelo.
- Aunque Render y el entorno de desarrollo son adecuados para pruebas, se recomienda migrar el sistema a un servicio con mayor disponibilidad y rendimiento especialmente si se estima un volumen alto de consultas. Además, para el uso de la plataforma Meta Cloud API se indica que en entorno productivo se requiere una gestión administrativa adicional ya que es necesario contar con una verificación por parte de la Universidad para poder registrar un número institucional.
- En una siguiente fase, el chatbot podría incorporar procesamiento de archivos adjuntos, respuestas multimedia o integración con un sistema interno de admisión para brindar información personalizada a cada usuario.
- Brindar capacitación básica al personal del área de admisión sobre el uso y mantenimiento del chatbot, de modo que puedan actualizar las fuentes de información y supervisar el funcionamiento del sistema sin requerir asistencia técnica constante.

9. BIBLIOGRAFÍA

Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., . . . Amodei, D. (2020). Language Models are Few-Shot Learners. arXiv (Cornell University). <https://doi.org/10.48550/arxiv.2005.14165>

Collado Alonso, M. Á. (2024). *Implementación de técnicas de RAG (Retrieval Augmented Generation) sobre LLM para la extracción y generación de documentos en las entidades públicas* (Trabajo Fin de Máster). Universidad de Valladolid.

Espinosa-Luna, B. H., Castillo-Oliva, J., Montañez-Díaz, B. A., & Mendoza-De los Santos, A. (2023). Implementación de un chatbot basado en modelo de lenguaje de inteligencia artificial para responder preguntas frecuentes de estudiantes universitarios. *Revista Científica de Sistemas e Informática*, 3(2), e570. <https://doi.org/10.51252/rcsi.v3i2.570>

Guillena, R. M. (2024). Procesamiento del lenguaje natural como eje central de la inteligencia artificial generativa: lección inaugural del curso académico 2024-2025. Universidad de La Rioja, Servicio de Publicaciones.

Lopez Aguirre, J. A., & Segura Carmona, F. A. (2023). Chatbot para mejorar la calidad de atención de consultas de grados y títulos de la USMP-ICED. *Revista de Tecnología Educativa*, 14(2), 123-136.

Mariani, M. M., Hashemi, N., & Wirtz, J. (2023). Artificial intelligence empowered conversational agents: A systematic literature review and research agenda. *Journal of Business Research*, 161, 113838. <https://doi.org/10.1016/j.jbusres.2023.113838>

Musayón Velásquez, O. (2023). Chatbot basado en IA para atención al estudiante de la Universidad Nacional Toribio Rodríguez de Mendoza de Amazonas. *Revista de Innovación Educativa*, 12(1), 75-92.

Okonkwo, C. W., & Ade-Ibijola, A. (2021). Chatbots applications in education: A systematic review. *Computers And Education Artificial Intelligence*, 2, 100033. <https://doi.org/10.1016/j.caeai.2021.100033>

Otter, D. W., Medina, J. R., & Kalita, J. K. (2021). A Survey of the Usages of Deep Learning for Natural Language Processing. *IEEE Transactions on Neural Networks and Learning Systems*, 32(2), 604-624. <https://doi.org/10.1109/TNNLS.2020.2979670>

Ray, P. P. (2023). ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. *Internet of Things and Cyber-Physical Systems*, 3, 121-154. <https://doi.org/10.1016/j.iotcps.2023.04.003>

Reddy, K. G., Pal, A., V, S. K., J, R., & K, S. (2022). Cross Domain Answering FAQ Chatbot. *2022 International Conference on Advanced Computing Technologies and Applications (ICACTA)*, 1-4. <https://doi.org/10.1109/ICACTA54488.2022.9752986>

Rios Lopez, C. A., & Vallejos Fernandez, D. (2023). Uso de chatbot para la atención a los clientes por redes sociales en la empresa Master GOL. *Revista de Gestión Empresarial*, 17(1), 55-70.

Russell, S.J. and Norvig, P. (2016) Artificial Intelligence: A Modern Approach. Pearson Education Limited, Malaysia.

Santana, R., Ferreira, S., Rolim, V., Miranda, P., Nascimento, A., & Mello, R. F. (2021). A Chatbot to Support Basic Students Questions. *CEUR Workshop Proceedings*, 3059, 58-67. <http://sunsite.informatik.rwth-aachen.de/Publications/CEUR-WS/Vol-3059/paper6.pdf>

Saranya, A., & Subhashini, R. (2023). A systematic review of Explainable Artificial Intelligence models and applications: Recent developments and future trends. *Decision Analytics Journal*, 7, 100230. <https://doi.org/10.1016/j.dajour.2023.100230>

Wang, X., Lin, X., & Shao, B. (2022). How does artificial intelligence create business agility? Evidence from chatbots. *International Journal of Information Management*, 66, 102535. <https://doi.org/10.1016/j.ijinfomgt.2022.102535>