



MÁSTER EN INGENIERÍA DE TELECOMUNICACIÓN

TRABAJO FIN DE MÁSTER

Aplicación de la inteligencia artificial al diseño de
sistemas de voz en tiempo real

Autor: Carlos Martín de Argila Lorente

Director: David Contreras Bárcena

Madrid
Junio 2026

Anexo I

Declaración de originalidad

Declaro bajo mi responsabilidad que el Proyecto presentado con el título **Aplicación de la inteligencia artificial al diseño de sistemas de voz en tiempo real** en la ETS de Ingeniería – ICAI de la Universidad Pontificia Comillas en el curso académico 2025/2026 es de mi autoría y no ha sido presentado con anterioridad a otros efectos. El Proyecto no es plagio de otro, ni total ni parcialmente y la información que ha sido tomada de otros documentos está debidamente referenciada.

Uso de Inteligencia Artificial¹

Declaro bajo mi responsabilidad que (indicar la opción correcta):

- ☐ No he utilizado Inteligencia Artificial en la elaboración del presente documento.
- ☒ He utilizado Inteligencia Artificial en la elaboración del presente documento y/o del Anexo B siempre en las condiciones permitidas por la Universidad Pontificia Comillas, es decir, aplicando el Nivel 2 de la [Escala de Evaluación de Perkins et al. \(2024\)](#): “*La IA puede utilizarse para actividades previas a la tarea, como la lluvia de ideas, la descripción y la investigación inicial. Este nivel se centra en el uso de la IA para la planificación, las síntesis y la generación de ideas, pero las evaluaciones deben hacer hincapié en la capacidad de desarrollar y refinar estas ideas de forma independiente*”. En concreto, la Inteligencia Artificial ha sido empleada para:

1. **Investigación inicial:** Búsqueda rápida de documentación, artículos y referencias sobre arquitecturas de voz y modelos *Open Source*, para acelerar la fase de investigación del estado del arte.
2. **Lluvia de ideas y estructuración:** Apoyo para ordenar y relacionar conceptos técnicos complejos antes de escribirlos en la memoria.
3. **Revisión del documento:** Uso de la IA a modo de corrector para revisar faltas de ortografía, de gramática y de puntuación en la memoria antes de la entrega final.
4. **Ayuda con el formato en LaTeX:** Ayuda con dudas técnicas y para solucionar errores de compilación así como con ciertos comandos específicos (tablas, figuras, ecuaciones, etc.).
5. **Ayuda con la traducción al inglés:** Ayuda con la traducción al inglés del abstract y del resumen de la memoria.

¹Esta declaración se refiere al uso de la Inteligencia Artificial generativa para realizar los documentos del Proyecto (Anexo B y Memoria). No aplica a Proyectos donde, por su naturaleza, deban emplear inteligencia artificial como parte de los mismos (aplicación de técnicas de aprendizaje automático, redes neuronales, análisis de datos. . .)


Firmado (alumno): Carlos Martín de Argila Lorente
Fecha: Junio 2026

Autorización para la entrega del Proyecto

El Director del Proyecto	El co-Director del Proyecto (si aplica)
	
Firmado: David Contreras Bárcena	Firmado:
Fecha: Junio 2026	Fecha:



MÁSTER EN INGENIERÍA DE TELECOMUNICACIÓN

TRABAJO FIN DE MÁSTER

Aplicación de la inteligencia artificial al diseño de
sistemas de voz en tiempo real

Autor: Carlos Martín de Argila Lorente

Director: David Contreras Bárcena

Madrid
Junio 2026

Agradecimientos

Quiero dar las gracias a mis padres y a mi hermana por el apoyo constante durante todo este camino. Han estado siempre ahí, especialmente en los momentos de más carga, y eso ha sido clave para llegar hasta aquí.

También quiero agradecer a mis amigos por acompañarme durante esta etapa. Entre clases, entregas y horas de trabajo, han conseguido que todo fuera bastante más llevadero.

Por último, quiero dar las gracias a mi tutor del TFM, David. Gracias por tu orientación y por compartir tu experiencia durante el desarrollo de este proyecto. También te quería dar las gracias por toda la paciencia que has tenido conmigo, porque sé que te he dado la lata en más de una ocasión y también gracias por conseguir que tuviera acceso al servidor DGX de ICAI, sin el cual el proyecto no habría sido posible.

Muchas gracias a todos.

Aplicación de la inteligencia artificial al diseño de sistemas de voz en tiempo real

Autor: Carlos Martín de Argila Lorente

Director: David Contreras Bárcena

Entidad colaboradora: ICAI - Universidad Pontificia Comillas

Resumen

En los últimos años, la gran revolución de los asistentes de voz ha sido la aparición de modelos unificados que son muy rápidos y naturales, pero que funcionan como cajas negras, ya que generan el audio directamente sin poder entender fácilmente la lógica de una respuesta. En entornos empresariales esto supone un problema ya que no se pueden controlar fácilmente. Además, como la mayoría de avances *Open Source* se centran en LLMs de texto, para acceder a estos modelos es necesario recurrir a APIs comerciales privadas, asumiendo su coste y riesgos de privacidad. Para solucionarlo, en este TFM se ha desarrollado una arquitectura de voz en tiempo real utilizando únicamente modelos *Open Source* con un enfoque modular y en cascada (VAD, ASR, LLM y TTS), lo que permite mayor flexibilidad y control.

El sistema es también un agente capaz de llamar a herramientas (*tools*) para interactuar con bases de datos. Se validó simulando la automatización telefónica de un restaurante aunque, cambiando las herramientas y el *system prompt*, se puede cambiar el escenario. Para conseguir una interacción fluida, se ha implementado un *pipeline* en *streaming* y una máquina de estados (FSM) que gestiona los turnos de palabra y reacciona de forma casi inmediata si el usuario interrumpe a la IA. Para poder elegir el LLM más adecuado, se evaluaron decenas de modelos de código abierto, eligiendo Gemma 4 31B por tener la mejor relación entre inteligencia en español y el tiempo hasta el primer token (TTFT).

Finalmente, se desarrolló un frontend para interactuar con el sistema, permitiendo hablar en tiempo real y visualizar los estados internos. La evaluación automática del *pipeline* sobre 300 llamadas telefónicas demostró que el asistente completaba las tareas con éxito en el 83,7 % de los casos, con una latencia mediana de 812.4 ms.¹

Palabras clave: Inteligencia Artificial, Asistentes de voz en tiempo real, Modelos *Open Source*, Arquitectura en cascada, LLM, Baja latencia.

¹Para ver el funcionamiento real del sistema, se ha preparado una pequeña demo del sistema atendiendo una llamada de teléfono a la que se puede acceder en: <https://youtu.be/ggq2c-R6Pjg>

Introducción

Actualmente, la Inteligencia Artificial conversacional está evolucionando desde los asistentes de texto hacia modelos de voz nativos. Aunque estos sistemas unificados (*Speech-to-Speech*) son muy rápidos, en un entorno empresarial presentan problemas importantes: operan como cajas negras, pueden inventarse datos (alucinaciones) y su uso a través de APIs comerciales puede ser muy caro a gran escala. Por eso, este TFM busca crear una alternativa utilizando únicamente modelos *Open Source*, para así poder mantener el control de los datos, asegurar la privacidad y que además sea viable económicamente.

Definición del proyecto

El objetivo principal de este trabajo es diseñar, implementar y validar una arquitectura modular en cascada para un asistente de voz en tiempo real orientado a entornos empresariales. Además de mantener una conversación fluida en español, el sistema tiene que tener capacidades agénticas, lo que significa que el LLM debe ser capaz de llamar a herramientas (*tools*) para consultar o escribir en una base de datos. Para medir su rendimiento de la forma más realista posible, se ha simulado un caso de uso concreto que es la automatización del teléfono de un restaurante. Se ha elegido porque es un escenario del día a día en el que hay que recoger datos precisos como pueden ser fechas, horas, nombres y números de teléfono.

Descripción del sistema

De forma resumida, el sistema funciona recibiendo el audio del cliente a través de WebSockets. Primero pasa por un detector de voz (VAD) y el fragmento se transcribe usando un modelo de reconocimiento de voz (ASR) (Parakeet-TDT-0.6B). Ese texto llega al LLM. El modelo procesa la información, ejecuta las herramientas si hace falta y genera la respuesta. Finalmente, el modelo de síntesis de voz (TTS) (Qwen3-TTS) convierte el texto en voz y lo envía al cliente.

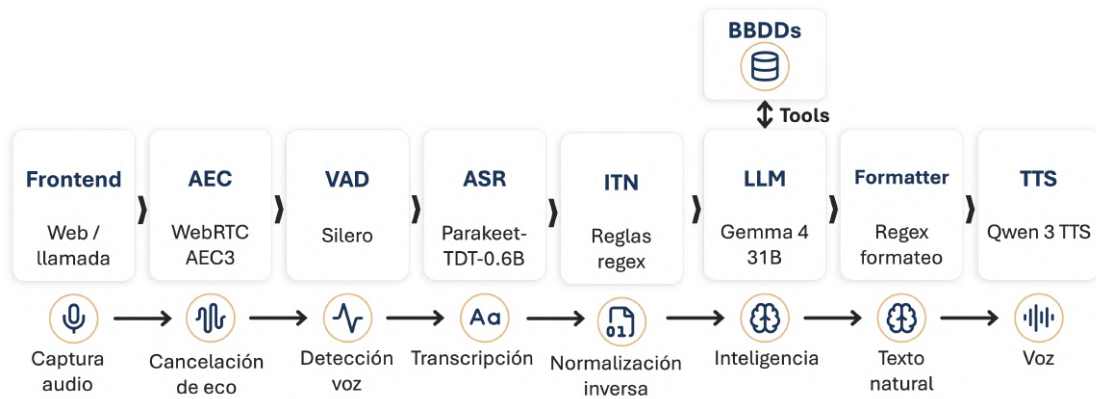


Figura 1: Arquitectura simplificada del sistema modular en cascada.

Todos los bloques se comunican de forma asíncrona y en *streaming*. Se ha diseñado una Máquina de Estados Finitos (FSM) que se encarga de controlar en todo momento los bloques del sistema y las posibles interrupciones del usuario mientras la IA está hablando.

Para poder interactuar con el sistema, se ha desarrollado un *frontend*. En la Figura 2, se puede ver la interfaz web que permite al usuario llamar al agente de IA y también donde se puede ver en tiempo real la transcripción de la conversación y el estado actual de la FSM entre otras cosas.

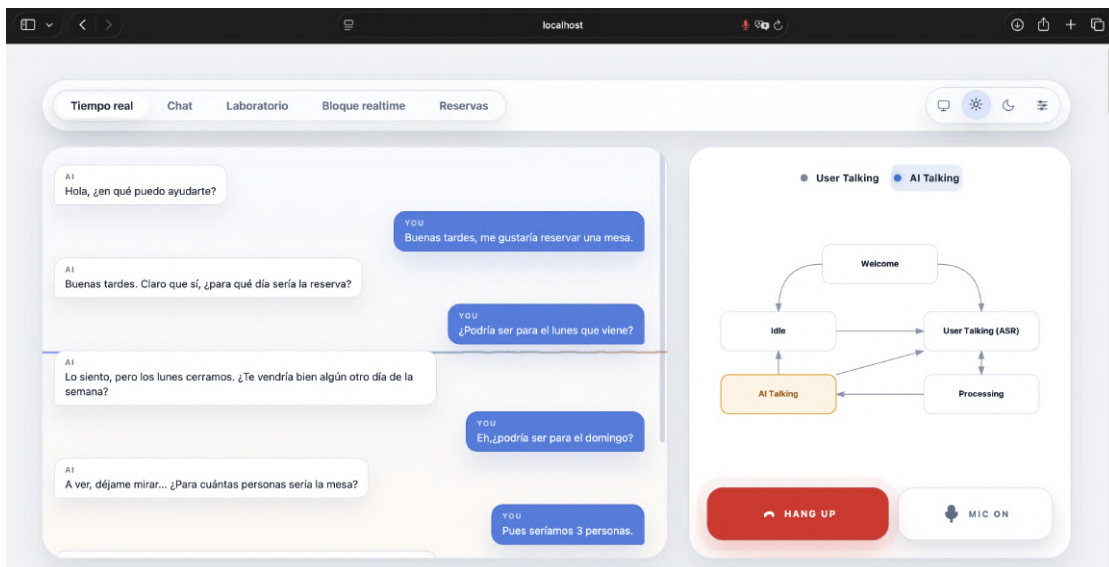


Figura 2: Frontend para interactuar con el sistema desarrollado.

Resultados

La evaluación del sistema completo (*End-to-End*) se realizó mediante un simulador automático que realizó 300 llamadas de teléfono sobre el *pipeline*. Como se ve en la Figura 3, los resultados se midieron en tres niveles. El agente logró completar la tarea correctamente en el 83,7 % de los casos, mantuvo una buena calidad en la conversación en el 88,0 % de las llamadas, y consiguió una validación completa (es decir, hacer bien la tarea y hablar de forma natural a la vez) en el 80,7 % de los escenarios.

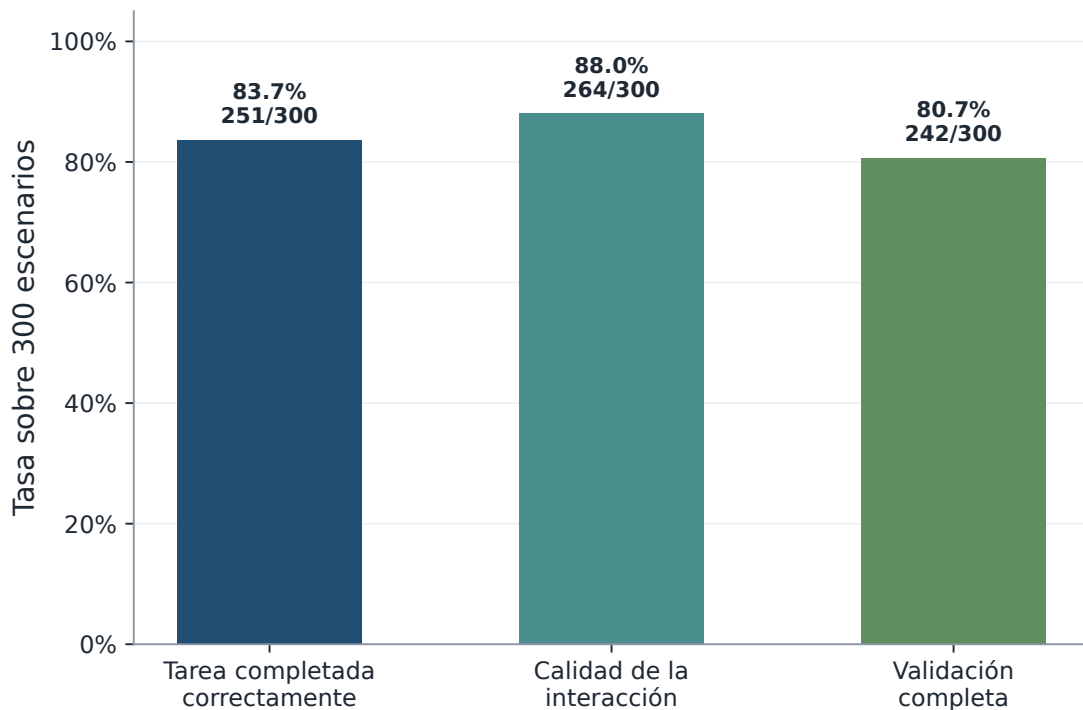


Figura 3: Rendimiento global del asistente en el benchmark *End-to-End*.

En cuanto a la latencia del sistema, el tiempo total desde que el usuario termina de hablar hasta que escucha el primer audio de respuesta está en torno a los 812 ms. Además, se demostró que el sistema reacciona de forma casi inmediata ante las interrupciones (*barge-in*), cortando la voz del asistente en solo 84 ms.

Por último, se realizó un estudio de viabilidad económica usando la teoría de colas (Erlang B). Los resultados demostraron que desplegar este sistema en un entorno con un alto volumen de llamadas (como un servicio de atención ciudadana), aplicando un autoescalado de GPUs, sería muy rentable y se amortizaría en menos de un año gracias al ahorro en costes operativos.

Conclusiones

Este Trabajo de Fin de Máster demuestra que ya es viable construir agentes de voz para entornos empresariales usando modelos *Open Source*. La arquitectura en cascada permite tener el control sobre la lógica del sistema y evita depender de APIs de terceros. Aunque la latencia es superior a la de los modelos unificados, el uso del *streaming* en los bloques permite que la latencia percibida por el usuario sea aceptable. Como próximos pasos, se propone integrar el *backend* con una línea de teléfono real usando SIP y mejorar la expresividad de las voces generadas por el modelo TTS.

Aplicación de la inteligencia artificial al diseño de sistemas de voz en tiempo real

Author: Carlos Martín de Argila Lorente

Director: David Contreras Bárcena

Collaborating entity: ICAI - Universidad Pontificia Comillas

Abstract

In recent years, the main breakthrough in voice assistants has been the appearance of unified models that are very fast and natural, but operate as black boxes, since they generate audio directly without letting us easily understand the logic behind a response. In business environments, this creates a control problem. Also, since most Open Source advances focus on text LLMs, accessing these models requires using private commercial APIs, assuming their cost and privacy risks. To solve this, this Master's Thesis develops a real-time voice architecture using only Open Source models with a modular cascaded approach (VAD, ASR, LLM, and TTS), which allows for greater flexibility and control.

The system is also an agent capable of calling tools to interact with databases. It was validated by simulating the phone automation of a restaurant, although by changing the tools and the *system prompt*, the scenario can be changed. To achieve a fluid interaction, a streaming *pipeline* and a Finite State Machine (FSM) have been implemented to manage speaking turns and react almost immediately if the user interrupts the AI. To choose the most suitable LLM, dozens of open-source models were evaluated, choosing Gemma 4 31B for having the best balance between intelligence in Spanish and Time to First Token (TTFT).

Finally, a frontend was developed to interact with the system, allowing real-time speech and visualizing internal states. The automatic evaluation of the *pipeline* over 300 phone calls showed that the assistant successfully completed the tasks in 83.7% of the cases, with a median latency of 812.4 ms.²

Keywords: Artificial Intelligence, Real-time voice assistants, Open Source models, Cascaded architecture, LLM, Low latency.

²To see the real operation of the system, a short video demo of the assistant handling a phone call has been prepared, which can be accessed at: <https://youtu.be/ggq2c-R6Pjg>

Introduction

Currently, conversational Artificial Intelligence is evolving from text assistants to native voice models. Although these unified systems (*Speech-to-Speech*) are very fast, they present major problems in a business environment: they operate as black boxes, can invent data (hallucinations), and using them through commercial APIs is very expensive on a large scale. Therefore, the goal of this thesis is to create an alternative using only Open Source models, in order to keep full control of the data, ensure privacy, and also be economically viable.

Project Definition

The main goal of this work is to design, implement, and validate a modular cascaded architecture for a real-time voice assistant oriented to business environments. Besides holding a fluid conversation in Spanish, the system must have agentic capabilities, which means the LLM must be able to call tools to query or write to a database. To measure its performance as realistically as possible, a specific use case was simulated: automating a restaurant's phone line. This was chosen because it is an everyday scenario where precise data such as dates, times, names, and phone numbers must be collected.

System Description

The proposed architecture works by receiving the client's audio via WebSockets. First, it passes through a voice activity detector (VAD), and the chunk is transcribed using an automatic speech recognition (ASR) model (Parakeet-TDT-0.6B). That text reaches the LLM. The model processes the information, executes the tools if needed, and generates the response. Finally, the TTS engine (Qwen3-TTS) converts the text into voice and sends it back to the client.

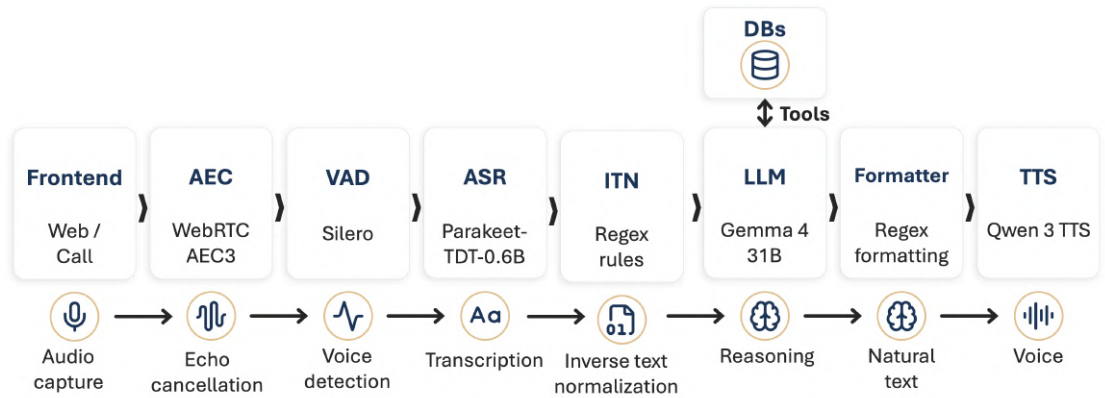


Figure 4: Simplified architecture of the modular cascaded system.

All blocks communicate asynchronously and in *streaming*. Additionally, a Finite State Machine (FSM) was designed to constantly control the system blocks and handle potential user interruptions while the AI is speaking.

To interact with the system, a frontend has been developed. Figure 5 shows the web interface that allows the user to call the AI agent, displaying the real-time transcription of the conversation and the current state of the FSM, among other features.

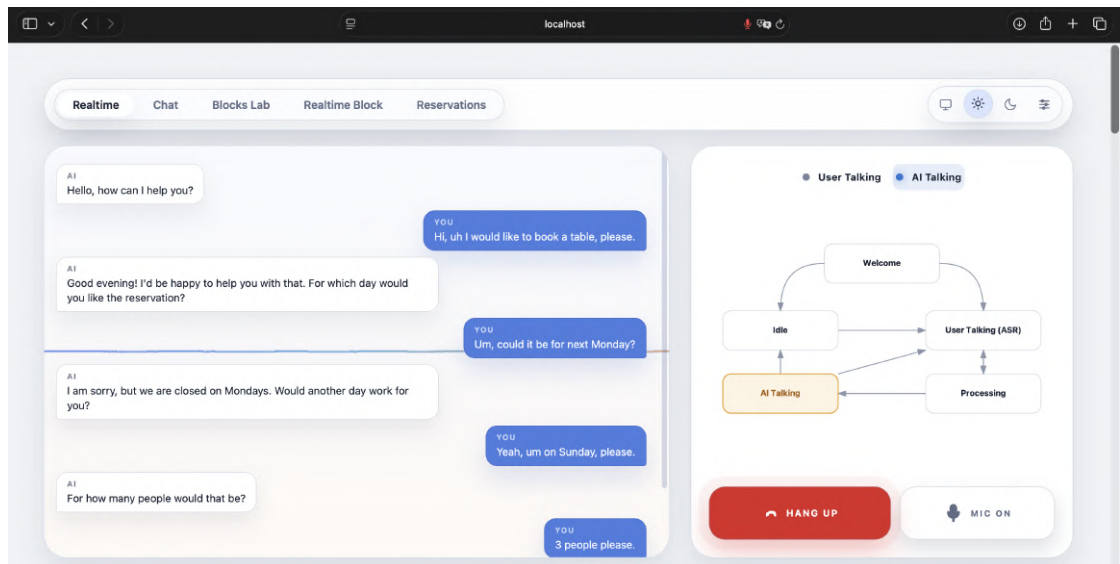


Figure 5: Frontend for interacting with the developed system.

Results

The evaluation of the complete system (*End-to-End*) was carried out using an automatic simulator that made 300 phone calls to the *pipeline*. As seen in Figure 6, the results were measured on three levels. The agent completed the task correctly in 83.7% of the cases, maintained good conversation quality in 88.0% of the calls, and achieved full validation (meaning doing the task right and speaking naturally at the same time) in 80.7% of the scenarios.

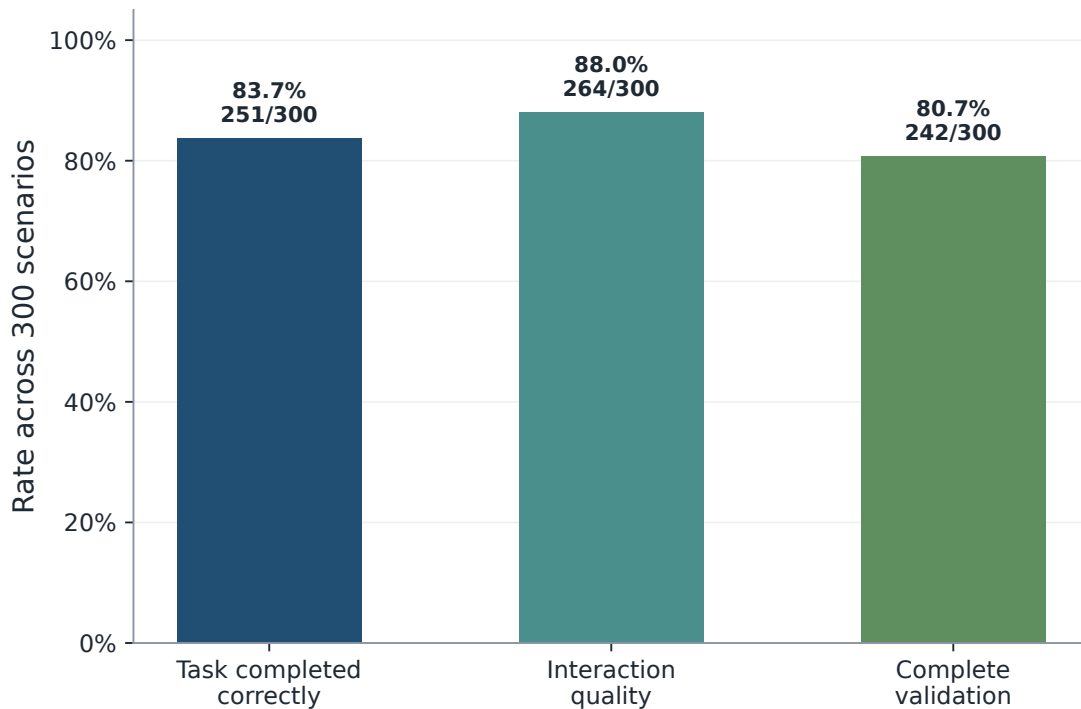


Figure 6: Global performance of the assistant in the End-to-End benchmark.

Regarding system latency, the total time from when the user finishes speaking until they hear the first response audio is around 812 ms. Furthermore, the system proved to react almost immediately to interruptions (*barge-in*), cutting off the assistant's voice in just 84 ms.

Lastly, an economic feasibility study was conducted using queueing theory (Erlang B). The results showed that deploying this system in an environment with a high volume of calls (like a citizen customer service), applying GPU autoscaling, would be very profitable and would pay for itself in less than a year thanks to savings in operating costs.

Conclusions

This Master's Thesis demonstrates that it is already viable to build voice agents for business environments using Open Source models. The cascaded architecture allows retaining control over the system's logic and avoids depending on third-party APIs. Although the latency is higher than that of unified models, the use of *streaming* across blocks makes the perceived latency acceptable for the user. As next steps, it is proposed to integrate the *backend* with a real phone line using SIP and to improve the expressiveness of the voices generated by the TTS model.

Tabla de Contenidos

1. Introducción	1
1.1. Contexto y problema	1
1.2. Motivación	2
1.3. Objetivos del proyecto	4
1.4. Alcance y caso de uso	5
1.5. Contribuciones principales	6
1.6. Estructura de la memoria	7
2. Estado del Arte	9
2.1. Evolución de las Interfaces de Voz	9
2.2. Arquitecturas actuales para asistentes de voz	10
2.3. Retos de la Interacción en Tiempo Real	12
2.4. De Asistentes a Agentes: uso de herramientas	13
3. Metodología	15
3.1. Metodología de desarrollo y planificación	15
3.2. Diseño experimental y entorno de simulación	17
3.2.1. Entorno de evaluación para la selección del LLM	17
3.2.2. Benchmark de inteligencia	19
3.3. Metodología del Benchmark <i>End-to-End</i> (E2E)	22
3.3.1. Funcionamiento de la simulación	22
3.3.2. Sistema de validación por capas	24
3.3.3. Benchmark de latencia	26
3.4. Metodología del Benchmark de Interrupciones (<i>Stop Reaction Time</i>)	26
3.4.1. Audios usados y preparación de datos	27
3.4.2. Modo de evaluación del pipeline(<i>interruption_eval</i>)	27
3.4.3. Protocolo de cada ejecución	27
3.5. Criterios y métricas de evaluación	28
3.5.1. Métricas de selección del Modelo de Lenguaje (LLM)	28
3.5.2. Métricas de rendimiento del Sistema Final (End-to-End)	29
3.6. Recursos y entorno de ejecución	30

3.6.1.	Recursos Hardware y Conectividad	30
3.6.2.	Software del Sistema	32
3.6.3.	Uso de APIs Externas	33
4.	Implementación del sistema	35
4.1.	Arquitectura del sistema propuesto	35
4.2.	Selección y justificación de componentes y modelos de Inteligencia Artificial	38
4.2.1.	Requisitos del LLM	38
4.2.2.	Modelos candidatos	40
4.2.3.	Resultados del benchmark de inteligencia	42
4.2.4.	Resultados del benchmark de velocidad	45
4.2.5.	Resumen de los resultados de los benchmarks y selección final del LLM	47
4.2.6.	Configuración final de inferencia	49
4.2.7.	Selección de VAD, ASR y TTS	50
4.3.	Implementación del Pipeline en streaming	54
4.4.	Gestión de turnos y Máquina de Estados (FSM)	58
4.5.	Orquestación agéntica y gestión del contexto	62
4.6.	Interfaz web (Frontend)	65
4.6.1.	Pantalla Interacción en Tiempo Real con el sistema	65
4.6.2.	Pantalla de Chat	67
4.6.3.	Pantalla de Bloques Realtime	68
4.6.4.	Pantalla Laboratorio	69
4.6.5.	Pantalla de Reservas	71
4.6.6.	Pantalla de Configuración	72
4.7.	Logs y eventos del sistema	75
4.8.	Análisis de Viabilidad Económica	76
4.8.1.	Planteamiento del modelo	77
4.8.2.	Datos de referencia y aproximación de automatización	77
4.8.3.	Dimensionamiento del sistema: Teoría de Colas	78
4.8.4.	Modelo de infraestructura de IA y costes de despliegue	79
4.8.5.	Modelo económico de ahorro	85
4.8.6.	Resultados de viabilidad	85
5.	Resultados y Discusión	89
5.1.	Resultados del Benchmark <i>End-to-End</i> (E2E)	89
5.1.1.	Rendimiento global del sistema	90
5.1.2.	Rendimiento por intención conversacional	92
5.1.3.	Análisis de señales de fallo	93
5.1.4.	Análisis de alucinaciones	94

5.2.	Resultados de latencia del sistema	94
5.2.1.	Resultados de las latencias del pipeline	95
5.3.	Tiempo de reacción y gestión de interrupciones	97
5.3.1.	Resultados del benchmark de interrupciones	97
5.3.2.	Análisis del tiempo de reacción	98
5.3.3.	Estabilidad a lo largo de la respuesta	99
5.4.	Evaluación cualitativa de la experiencia conversacional	101
5.4.1.	Calidad y naturalidad de la voz (TTS)	101
5.4.2.	Comprensión del usuario y reconocimiento (ASR)	101
5.4.3.	Percepción de la latencia en la interacción	101
5.4.4.	Capacidades agénticas y naturalidad del lenguaje	102
5.5.	Discusión y limitaciones	102
6.	Conclusiones	105
6.1.	Conclusiones principales	105
6.2.	Aportaciones del trabajo	106
6.3.	Trabajos futuros	107
	Bibliografía	109
	Anexos	116
A.	Alineación con los ODS	117
A.1.	Objetivo 9: Industria, Innovación e Infraestructura	118
A.2.	Objetivo 8: Trabajo Decente y Crecimiento Económico	119
B.	Catálogo de modelos LLMs y benchmarks	121
B.1.	Catálogo completo de modelos evaluados	121
B.2.	Ejemplos de tareas de evaluación de inteligencia	124
B.2.1.	Evaluación de instrucciones propias en español (<code>custom_spanish</code>)	124
B.2.2.	Robustez ante errores de transcripción (<code>spanish_robustness</code>)	127
B.2.3.	Evaluación de uso de herramientas (<code>bfcl_reference</code>)	128
B.2.4.	Evaluación conversacional y de seguridad (<code>voicebench_reference</code>)	130
B.3.	Resultados de los benchmarks	131
B.3.1.	Benchmark de inteligencia	131
B.3.2.	Benchmark de velocidad	135
B.3.3.	Unión de los benchmarks de inteligencia y velocidad	137

C. Evaluación detallada del pipeline End-to-End	139
C.1. Resultados del benchmark End-to-End de inteligencia	139
C.2. Resultados del benchmark End-to-End de latencia	141
C.3. Resultados del benchmark de interrupciones	142
D. Manual de despliegue y Código fuente	145
D.1. Manual de despliegue en el servidor DGX de ICAI	145
D.1.1. Requisitos previos	145
D.1.2. Arquitectura de contenedores	146
D.1.3. Ejecución del sistema en el servidor DGX	146
D.1.4. Ejecución del Framework de Evaluación (Benchmarks) . . .	149
D.2. Código fuente del TFM	149
E. System Prompt y Tools	151
E.1. System prompt	151
E.2. Protocolo de uso de herramientas	153
E.3. Tools disponibles	154
E.3.1. check_availability	154
E.3.2. make_reservation	154
E.3.3. lookup_reservation	154
E.3.4. modify_reservation	155
E.3.5. cancel_reservation	155
E.3.6. get_next_available_slot	155
E.3.7. get_menu	156
E.3.8. query_knowledge_base	156
E.4. Resumen del flujo de decisión	156

Índice de figuras

1.	Arquitectura simplificada del sistema modular en cascada.	VII
2.	Frontend para interactuar con el sistema desarrollado.	VII
3.	Rendimiento global del asistente en el benchmark <i>End-to-End</i>	VIII
4.	Simplified architecture of the modular cascaded system.	XII
5.	Frontend for interacting with the developed system.	XII
6.	Global performance of the assistant in the End-to-End benchmark.	XIII
2.1.	Comparativa de la Arquitectura Unificada y la Arquitectura en Cas- cada.	II
3.1.	Diagrama de Gantt del proyecto	16
3.2.	Arquitectura del benchmark <i>End-to-End</i>	23
3.3.	Captura de pantalla de nvidia-smi mostrando el estado de las GPUs del clúster de ICAI	31
4.1.	Arquitectura simplificada del sistema propuesto.	37
4.2.	Captura de https://artificialanalysis.ai aplicando los filtros de código abierto, Non-Reasoning y tamaño pequeño o medio para obtener los modelos candidatos.	42
4.3.	Puntuación global de inteligencia de los modelos finalistas.	43
4.4.	Descomposición de la puntuación funcional ponderada por modelo.	44
4.5.	Distribución en percentiles del <i>Time To First Token</i> (TTFT). . . .	45
4.6.	Velocidad de generación sostenida en tokens por segundo (TPS). . .	46
4.7.	Frontera eficiente comparando el TTFT mediano frente a la pun- tuación funcional. El tamaño del marcador indica el tamaño del modelo.	47
4.8.	Frontera eficiente comparando el TTFT percentil 95 frente a la pun- tuación funcional. El tamaño del marcador indica el tamaño del modelo.	48
4.9.	Protocolo de comunicación bidireccional mediante WebSocket entre el cliente y el servidor FastAPI.	55

4.10. Procesamiento inicial de la señal de audio: cancelación de eco, división en fragmentos de 10 ms, búfer de <i>pre-roll</i> y detección de actividad de voz.	56
4.11. Conversión del audio del usuario en texto listo para el LLM.	57
4.12. Generación de la respuesta en streaming y generación de voz con el TTS.	57
4.13. Coordinación asíncrona del pipeline.	58
4.14. Vista resumida del pipeline completo de audio en streaming, desde la entrada de voz del usuario hasta la respuesta sintetizada de la IA.	58
4.15. Máquina de Estados Finitos (FSM) para la gestión de turnos de palabra.	60
4.16. Interacción en tiempo real con el sistema	66
4.17. Logs del sistema	67
4.18. Interfaz de chat con el LLM	68
4.19. Comparación del comportamiento del VAD con voz y silencio.	69
4.20. Interfaz de Laboratorio para probar el pipeline entero	70
4.21. Interfaz de Reservas	71
4.22. Interfaz de Configuración del mensaje de bienvenida y idioma	72
4.23. Entorno de inferencia del LLM	73
4.24. Interfaz de Configuración del system prompt	74
4.25. Perfil horario de llamadas del 010 en el 2009 [50]	82
4.26. Tipo de cambio de referencia de 1 USD = 0,86 EUR.	83
4.27. Desglose del coste mensual de la plataforma de IA (caso base).	84
4.28. Comparativa del coste operativo anual entre el servicio únicamente humano y el servicio híbrido con ramp-up año 1.	86
4.29. Comparativa del coste operativo anual sin escalado progresivo inicial.	87
5.1. Rendimiento global del agente en el benchmark E2E (N=300).	91
5.2. Tasa de tarea completada correctamente, calidad de interacción y validación completa por categoría de escenario.	92
5.3. Señales de fallo en los escenarios sin Validación completa.	93
5.4. Composición mediana de la latencia hasta el primer fragmento de audio.	95
5.5. Diagrama temporal de funcionamiento en streaming de los diferentes bloques del pipeline.	96
5.6. Distribución de latencias internas por bloque del pipeline.	97
5.7. Histograma de frecuencias del tiempo de reacción ante interrupciones.	98
5.8. Gráfico de dispersión del tiempo de reacción frente al momento de interrupción del usuario.	100
A.1. Objetivos de Desarrollo Sostenible de la Agenda 2030	117

A.2. Objetivo 9: Industria, Innovación e Infraestructura	118
A.3. Objetivo 8: Trabajo Decente y Crecimiento Económico	119
B.1. Descomposición de la puntuación funcional ponderada por modelo en el benchmark de español personalizado.	133
B.2. Descomposición de la puntuación funcional ponderada por modelo en el benchmark de robustez en español.	134
D.1. Herramientas necesarias para la conexión al servidor DGX.	147
D.2. Forwarding de puertos en VS Code	148
D.3. Interfaz web del asistente de voz	148

Índice de tablas

2.1. Resultados del estudio de latencia en Amazon Bedrock [21].	12
4.1. Resumen de las familias de modelos de lenguaje incluidas en la campana de evaluación.	41
4.2. Modelos pertenecientes a la frontera eficiente.	48
4.3. Comparativa de los modelos ASR candidatos para español.	52
4.4. Comparativa de modelos TTS de código abierto para español basa- dos en MINT-Bench [44].	53
4.5. Datos operativos y económicos de referencia (Línea Madrid 010, 2024).	77
4.6. Dimensionamiento por bloque del <i>pipeline</i> de voz (caso base, pico horario: 50 canales, franja 11–12 h).	80
4.7. Supuestos de concurrencia para dimensionar el hardware.	81
4.8. Ahorro mensual estimado al aplicar autoscaling frente a mantener la infraestructura constante.	82
4.9. Desglose del coste mensual de infraestructura IA (caso base).	84
B.1. Especificaciones técnicas detalladas de los modelos candidatos.	122
B.2. Puntuaciones por benchmark de los modelos.	131
B.3. Distribución de latencia TTFT por modelo evaluado (ordenado por mediana ascendente).	135
B.4. Velocidad de generación (TPS) por modelo evaluado (ordenado por mediana descendente).	136
B.5. Unión de los benchmarks de inteligencia y velocidad por modelo.	137
C.1. Resumen agregado del benchmark end-to-end final (N=300 escena- rios).	139
C.2. Tarea completada, calidad de interacción y validación completa por categoría de escenario.	139
C.3. Distribución de latencia por etapa del pipeline end-to-end final.	140
C.4. Agregados de latencia del benchmark end-to-end final por etapa.	140
C.5. Matriz de confusión agregada para reservas en el benchmark end- to-end final.	140

C.6. Señales agrupadas de fallo en escenarios no exitosos (no excluyentes).	141
C.7. Uso de herramientas, duplicados y fallos de argumentos críticos. . .	141
C.8. Validación completa separada por degradación ASR detectada. . .	141
C.9. Percentiles internos por bloque del benchmark de latencia E2E. . .	142
C.10. Composición temporal del pipeline E2E hasta el primer audio del asistente.	142
C.11. Métricas temporales principales del pipeline E2E.	142
C.12. Resumen agregado del benchmark de interrupciones (N=3000 ejecuciones planificadas).	143
C.13. Distribución del tiempo de reacción ante interrupciones del usuario.	143
C.14. Distribución del punto de interrupción del usuario respecto al inicio del TTS.	143

Listings

Capítulo 1

Introducción

Este capítulo tiene como objetivo introducir el problema que se aborda en este Trabajo de Fin de Máster y explicar por qué es relevante en el contexto actual. En primer lugar, en la Sección [1.1](#) se presenta el contexto general de los asistentes de voz basados en Inteligencia Artificial (IA) y los problemas que aparecen al intentar aplicar este tipo de soluciones en entornos empresariales reales. A continuación, la Sección [1.2](#) se centra en la motivación del proyecto, poniendo el foco en aspectos importantes como son el control, la fiabilidad, la privacidad, el coste y la necesidad de que la orquestación agéntica sea robusta. Después, en la Sección [1.3](#) se definen los objetivos generales del trabajo. En la Sección [1.4](#) se delimita el alcance del proyecto y se introduce el caso de uso sobre el que se aplicará y evaluará la arquitectura propuesta. La Sección [1.5](#) resume brevemente las principales aportaciones de este trabajo. Por último, en la Sección [1.6](#) se describe cómo se organiza el resto de la memoria.

1.1. Contexto y problema

Desde el lanzamiento de ChatGPT a finales de 2022, los sistemas conversacionales basados en grandes modelos de lenguaje (LLMs) se han convertido en un elemento clave de la transformación tecnológica, alcanzando cifras de adopción sin precedentes [\[1\]](#). Si bien la primera ola de IAs se centró en interfaces de texto, poco a poco fueron mejorando y comenzaron a ser capaces de “*ver*” gracias a modelos de visión + texto, y las últimas evoluciones se centran cada vez más en la multimodalidad de estos. A estos modelos se les suelen llamar *omni* (del latín *omnis*, que significa “todo” o “completo”) ya son capaces de entender y generar varios tipos de información a la vez, no solo texto, como imágenes, audio e incluso video en algunos casos. El modelo que popularizó todo esto fue GPT-4o de OpenAI, que incluso permitía interacciones por voz con respuestas en tiempo real. En gran par-

te, estos modelos *omni* son sistemas *end-to-end*, es decir, arquitecturas unificadas en las que distintas modalidades, como el texto, la imagen o el audio, se procesan dentro de un mismo modelo en lugar de depender solo de módulos separados conectados entre sí. Esto hace que sean especialmente potentes y naturales en la interacción, pero también más opacos desde el punto de vista interpretativo, ya que en buena medida funcionan como *cajas negras*. Por ello, resulta difícil saber con precisión cómo representan la información internamente o cómo llegan a una respuesta concreta.

Aunque estas tecnologías resultan muy útiles para asistentes personales o entretenimiento, su aplicación directa en entornos empresariales críticos presenta desafíos importantes. Las soluciones actuales *end-to-end* ofrecen baja latencia y naturalidad al hablar, pero esto es a costa de sacrificar el control granular sobre el flujo de la conversación y la integración con lógicas de negocio complejas. Las empresas no necesitan únicamente un sistema que hable con voz humana, sino agentes que sean capaces de interactuar con la infraestructura de la empresa con funciones complejas como consultar bases de datos o interactuar con CRMs, todo esto con un control estricto de costes y normativas. Además, las soluciones *omni* más avanzadas están disponibles a través de APIs comerciales cerradas, lo que implica un riesgo de vendor lock-in, costes elevados y dificultades para garantizar la confidencialidad de los datos.

Este proyecto aborda esta problemática mediante el diseño y construcción de una arquitectura de voz modular y coordinada. En lugar de depender de un modelo único, se propone una solución que separe el reconocimiento de voz (Automatic Speech Recognition o *ASR*), la inteligencia (gracias a los Large Language Models o *LLM*) y la síntesis a voz (Text-to-Speech o *TTS*) en diferentes módulos separados y perfectamente auditables. Este enfoque permite optimizar recursos, reducir la dependencia tecnológica ya que existen muchos más modelos Open Source de cada uno de estos bloques que modelos *omni*, y también permite implementar capacidades agénticas mucho más avanzadas.

El objetivo final de este Trabajo de Fin de Máster es demostrar que es posible construir un asistente de voz en tiempo real que combine la flexibilidad de los modelos Open Source con la robustez y fiabilidad que exige el mercado profesional, aunque todo esto se explicará con mayor profundidad en la Sección [1.3](#).

1.2. Motivación

Como se ha analizado en la Sección [1.1](#), aunque los modelos multimodales ofrecen una experiencia de usuario muy buena en términos de latencia y naturalidad en la voz, en la práctica, usarlos en entornos empresariales es difícil debido a los problemas que se han comentado. Es por eso que la motivación principal de este

proyecto no es simplemente replicar un asistente de voz, sino diseñar una arquitectura que resuelva los problemas de control, fiabilidad, coste y privacidad que las soluciones Speech-to-Speech vía API no pueden solucionar actualmente.

En primer lugar, existe una necesidad de **control y fiabilidad**. En un entorno de negocio, como puede ser por ejemplo la atención al cliente, no es aceptable que el sistema invente datos, lo que en los LLMs se conoce como *alucinación*, ni que detecte sonidos o eventos de audio que en realidad no existen, algo que ya se ha observado en modelos de audio nativo [2]. Una arquitectura modular por bloques ofrece una ventaja clara, porque permite revisar y validar la salida de cada bloque antes de pasar al siguiente y evitar que un error se arrastre por todo el sistema. Obviamente, esto no hace que el sistema sea inmune a fallos, pero sí lo vuelve más controlable, más auditable y más fácil de depurar, que es justo lo que se necesita en entornos donde la fiabilidad no es opcional.

Por otro lado, es importante discutir el tema de la **soberanía del dato**. En el caso, por ejemplo, de un sistema de atención al cliente, el audio de las llamadas puede contener datos personales, por lo que no siempre es posible enviarlo directamente a servicios de terceros, especialmente si su tratamiento se realiza fuera del Espacio Económico Europeo. En estos casos, es necesario cumplir con los principios de minimización, confidencialidad y protección de datos desde el diseño que establece el GDPR. Además, si los datos se envían fuera del Espacio Económico Europeo (EEE), se deben cumplir las garantías adicionales exigidas para las transferencias internacionales de datos [3, 4]. Desde este punto de vista, una arquitectura modular ofrece una ventaja clara, ya que permite utilizar una infraestructura local para ejecutar partes más sensibles del proceso y limitar qué información se envía realmente a servicios externos. Además, evita depender de una única solución *speech-to-speech* y ofrece mucha más flexibilidad a la hora de combinar componentes distintos, algo que la literatura reciente también señala como una de las principales ventajas de los enfoques en cascada frente a los modelos *end-to-end* [5].

Además, existe una fuerte motivación **económica**. Las APIs actuales de voz en tiempo real facturan en función de los tokens de audio y texto procesados, por lo que el coste crece rápidamente cuando aumenta el volumen de interacciones [6, 7]. Una arquitectura modular ofrece una ventaja clara: permite ajustar cada bloque del sistema a las necesidades reales de esa etapa y elegir de forma independiente qué modelo conviene utilizar en transcripción, razonamiento y síntesis. No todos los bloques exigen el mismo nivel de capacidad, así que se pueden distribuir los recursos a las partes donde aportan más valor. Además, este desacoplamiento amplía el abanico de alternativas que pueden emplearse en cada bloque del sistema, lo que

permite elegir entre modelos con prestaciones similares pero costes muy inferiores¹, lo que aumenta el margen de optimización económica del sistema.

Finalmente, también existe una motivación técnica relacionada con la orquestación agéntica en interfaces de voz, donde la latencia y la gestión de turnos condicionan de forma directa la fluidez de la interacción [11, 12].

1.3. Objetivos del proyecto

El objetivo principal de este proyecto es el diseño, implementación y validación de una arquitectura modular de inteligencia artificial para asistentes de voz, capaz de operar en tiempo real y ejecutar tareas complejas. Para alcanzar este objetivo general, se han definido cinco metas específicas que guiarán el desarrollo de este Trabajo de Fin de Máster:

1. **Diseñar y construir una arquitectura de procesamiento modular en streaming:** El primer objetivo es desarrollar el *pipeline* completo de voz que integre, al menos, los tres bloques fundamentales: Reconocimiento Automático del Habla (ASR), Modelos de Lenguaje (LLM) y Síntesis de Voz (TTS), todo, usando modelos Open Source. A diferencia de las soluciones unificadas, el sistema se diseñará por bloques, utilizando protocolos de comunicación asíncronos para conectar los componentes. Esto permitirá sustituir cualquier modelo individual sin afectar al resto del sistema, garantizando la flexibilidad y evitando la dependencia de un único proveedor tecnológico.
2. **Implementar capacidades agénticas y uso de herramientas:** Además de ser capaz de mantener una conversación fluida, el sistema debe comportarse como un agente útil. El objetivo es dotar al asistente de la capacidad de decidir cuándo debe invocar herramientas externas (*Function Calling*) para resolver una petición del usuario, por ejemplo consultando una base de datos y escribiendo en ella. El foco se pondrá en garantizar que la extracción de parámetros y la ejecución de dichas funciones sea robusta, y en que este uso de herramientas pueda integrarse en una interacción de voz sin comprometer el ritmo de la conversación, ya que este tipo de operaciones introduce latencias adicionales [13].

¹Como referencia, a 3 de abril de 2026, el modelo GLM-5 obtiene 50 puntos en el *Artificial Analysis Intelligence Index*, frente a los 53 de Claude Opus 4.6, por lo que ambos modelos tienen una inteligencia comparable. Sin embargo, sus precios oficiales por millón de tokens son notablemente inferiores: Z.AI fija para GLM-5 \$1 por millón de tokens de entrada y \$3.2 por millón de tokens de salida, y Anthropic fija para Claude Opus 4.6 \$5 y \$25, respectivamente [8, 9, 10].

3. **Gestionar los tiempos de espera y los turnos de palabra.** Dado que el uso de herramientas introduce latencia, es necesario gestionar los silencios para garantizar una buena experiencia de usuario. El objetivo es implementar una lógica que maneje los turnos de palabra de forma natural. Esto incluye la capacidad de detectar cuándo el usuario ha dejado de hablar y la implementación de estrategias para enmascarar la latencia. Para ello existen varias opciones como el uso de frases de relleno (*fillers*) o sonidos de espera, para que la interacción se perciba fluida aunque el procesamiento no sea instantáneo.
4. **Desarrollar un sistema de gestión del contexto conversacional:** Al tratarse de una llamada telefónica continua y no de comandos aislados, es fundamental que el sistema mantenga una memoria de la interacción para que no parezca que se resetea con cada frase. El objetivo es implementar un mecanismo que almacene y gestione el historial de la conversación, permitiendo al asistente recordar datos proporcionados y mantener el contexto a lo largo de toda la sesión.
5. **Validar el funcionamiento del prototipo y análisis de rendimiento:** Finalmente, se busca verificar la viabilidad técnica de la solución creada. Para ello, se definirá un caso de uso concreto y se evaluará el rendimiento del sistema. El objetivo no es competir en velocidad pura contra los modelos *Speech-to-Speech* unificados, sino validar que la arquitectura es capaz de completar las tareas con éxito, medir la latencia introducida por cada módulo y comprobar la eficacia de las estrategias de orquestación implementadas.

1.4. Alcance y caso de uso

El alcance de este Trabajo de Fin de Máster abarca el diseño, desarrollo y validación técnica de una arquitectura modular para un asistente conversacional de voz en tiempo real. Para aplicar la investigación y la plataforma en un entorno práctico, medible y realista, se ha definido como caso de uso central la simulación de un **asistente telefónico para la gestión de reservas de un restaurante**.

La elección de este escenario representa un asistente de reservas que necesita mantener una charla fluida en español, extraer parámetros precisos (fechas, horas, número de comensales), interactuar de forma determinista con la lógica de negocio a través de herramientas (*tool calling* para consultar disponibilidad o confirmar mesas), manejar reglas estrictas y mantener el contexto a lo largo de toda la llamada.

Para poder probar el sistema, el alcance del proyecto también incluye la creación de un *frontend* web. Esta interfaz permite interactuar por voz con el sistema

y visualizar información importante del asistente en tiempo real. Sirve para demostrar que la arquitectura funciona y proporcionar una experiencia de usuario completa.

Dentro de los límites del proyecto, **se incluye** el desarrollo completo del *pipeline* en *streaming* (VAD, ASR, LLM y TTS, entre otros, ya que como se verá después hay más bloques implicados) gestionado mediante WebSockets y programación asíncrona. Además, forma parte fundamental del trabajo el diseño de la Máquina de Estados Finitos (FSM) para resolver problemas críticos de la interacción por voz, como la gestión de turnos y las interrupciones del usuario (*barge-in*). También se incluye la evaluación experimental en detalle para seleccionar los modelos *Open Source* más adecuados, buscando el equilibrio entre la inteligencia funcional y una latencia estricta (especialmente en el *Time To First Token*).

Por el contrario, **queda fuera del alcance** de este trabajo el despliegue del sistema en una pasarela de telefonía tradicional (como Asterisk o centralitas SIP/PBX) con usuarios reales, limitándose la validación a un entorno de pruebas mediante un cliente web propio y simulaciones controladas. Tampoco es objeto de este proyecto el entrenamiento desde cero o el *fine-tuning* profundo de modelos fundacionales de inteligencia artificial, sino la integración, orquestación y perfilado de modelos de código abierto ya existentes para demostrar su viabilidad técnica y económica en entornos de producción. Para ver el funcionamiento real del sistema, se ha preparado una pequeña demo del sistema atendiendo una llamada de teléfono a la que se puede acceder en el siguiente enlace: <https://youtu.be/ggq2c-R6Pjg>

1.5. Contribuciones principales

Aunque el uso de arquitecturas en cascada para procesamiento de voz no es un concepto nuevo en la literatura, el desarrollo de un sistema conversacional fluido, con capacidades agénticas y tiempos de respuesta viables utilizando exclusivamente modelos de código abierto sigue siendo un reto técnico complejo. En este contexto, las principales aportaciones de este Trabajo de Fin de Máster son las siguientes:

- **Diseño e implementación de un *pipeline* modular en *streaming*:** Se ha construido una arquitectura completa desde cero capaz de procesar audio de forma continua. La aportación clave en este punto es el desarrollo de una Máquina de Estados Finitos (FSM) personalizada y un sistema de control asíncrono que gestionan de manera robusta los turnos de palabra, los tiempos de espera y las interrupciones del usuario, evitando bloqueos en la comunicación.
- **Framework de evaluación reproducible para modelos de voz en español:** Dado que la mayoría de métricas públicas se centran en texto en

inglés, se ha diseñado un entorno de pruebas específico para evaluar LLMs candidatos. Esta metodología propia permite medir no solo la inteligencia del modelo y el uso de herramientas, sino su latencia en la generación del primer token (TTFT) y su resistencia ante errores típicos de transcripción oral en español.

- **Integración de capacidades agénticas bajo un entorno controlado:** Se ha implementado un mecanismo que permite al modelo de lenguaje interactuar con la lógica de negocio (consultar bases de datos, hacer reservas) manteniendo la trazabilidad de las acciones. La arquitectura asegura que las operaciones persistentes no queden en un estado inconsistente incluso si el usuario interrumpe al asistente en medio de una respuesta.
- **Análisis de viabilidad técnica y económica:** Más allá del desarrollo del *software*, el proyecto aporta un estudio de dimensionamiento para entornos de producción. Utilizando teoría de colas y modelos de costes de infraestructura GPU, se demuestra empíricamente que este diseño modular y *Open Source* es económicamente viable y competitivo frente a la contratación de agentes humanos o el uso de APIs comerciales cerradas.

1.6. Estructura de la memoria

La memoria se organiza en seis capítulos principales.

En el Capítulo [1](#) se introduce el contexto general del trabajo, se explica la motivación del proyecto, se definen los objetivos, se delimita el alcance y se resumen las principales contribuciones realizadas.

El Capítulo [2](#) recoge el estado del arte. En este capítulo se revisa la evolución de las interfaces de voz, las arquitecturas que se utilizan actualmente en asistentes conversacionales, los retos que plantea la interacción en tiempo real y la transición desde los asistentes clásicos en llamadas de teléfono hacia agentes con una inteligencia y autonomía muy superior.

A continuación, el Capítulo [3](#) presenta la metodología seguida para desarrollar y evaluar el proyecto. Se explica cómo se ha planteado el desarrollo del proyecto, la planificación realizada, el diseño experimental, los criterios y métricas de evaluación y el entorno de ejecución utilizado.

El Capítulo [4](#) se centra en el desarrollo de la plataforma propuesta. En él se describe en detalle la arquitectura del sistema implementado, la selección final de componentes, la construcción del pipeline en *streaming*, la gestión de turnos de palabra, la orquestación agéntica y los mecanismos para monitorizar el sistema.

En el Capítulo [5](#) se presentan los resultados obtenidos en las distintas pruebas realizadas, incluyendo las medidas de latencia, los *benchmarks* funcionales y *end-*

to-end, los experimentos de robustez ante interrupciones y el estudio de ablación. Además, se comentan las principales limitaciones observadas y su impacto sobre la solución desarrollada.

Por último, el Capítulo 6 recoge las conclusiones principales del trabajo, resume las aportaciones principales y plantea posibles líneas de trabajo futuro.

Al final de la memoria se incluyen varios anexos que recogen información complementaria de apoyo, como material adicional sobre el proyecto, detalles de configuración, documentación técnica y otros elementos útiles para facilitar la comprensión del trabajo.

Capítulo 2

Estado del Arte

La Inteligencia Artificial Conversacional está viviendo actualmente una transformación, dejando atrás los sistemas rígidos y por turnos para dar paso a asistentes de voz continuos, inmediatos y con capacidades agénticas. Este cambio ha sido posible gracias a la evolución de las interfaces de voz, al avance de los modelos de Inteligencia Artificial y al desarrollo de nuevas arquitecturas capaces de ofrecer interacciones cada vez más naturales.

En este capítulo se estudia el estado del arte sobre los principales avances que han dado lugar a los asistentes de voz inteligentes actuales, así como los retos que todavía siguen abiertos en este ámbito. En la Sección [2.1](#) se describe la evolución de las interfaces de voz en los últimos años, desde los sistemas IVR y los asistentes basados en intenciones predefinidas hasta la aparición de los modelos de IA más recientes aplicados a asistentes de voz. Después, en la Sección [2.2](#) se analizan las dos arquitecturas principales que se emplean actualmente en asistentes de voz, comparando el enfoque unificado *Speech-to-Speech* con la arquitectura en cascada basada en bloques independientes de ASR (Automatic Speech Recognition), LLM (Large Language Model) y TTS (Text-to-Speech). Después, en la Sección [2.3](#) se estudian los principales retos de la interacción en tiempo real, especialmente los relacionados con la latencia y la naturalidad de la conversación. Por último, la Sección [2.4](#) aborda la transición desde asistentes conversacionales tradicionales hacia agentes de voz con capacidad de uso de herramientas, explicando tanto sus ventajas como los retos técnicos asociados.

2.1. Evolución de las Interfaces de Voz

Históricamente, las interfaces de usuario por voz han estado muy limitadas por la rigidez de sus mecanismos de entrada y procesamiento [\[14\]](#). Los primeros sistemas, conocidos como IVR (*Interactive Voice Response*), se basaban en árboles de

lógica determinista donde el usuario debía navegar mediante comandos específicos como decir números, o pulsar teclas generando tonos DTMF que eran interpretados por el sistema, lo que es lento y poco natural para el usuario. De forma paralela, la aparición de asistentes como Siri o Alexa introdujo el uso de redes neuronales para el reconocimiento de voz, pero su lógica seguía dependiendo de la clasificación de intenciones predefinidas, lo que limitaba su capacidad para mantener conversaciones fluidas o gestionar contextos complejos. Aun así, eran útiles para tareas simples dentro de un entorno controlado, como pueden ser consultar el tiempo o reproducir música.

Desde la llegada de los Large Language Models (*LLMs*) en 2023, este ámbito está experimentando una transformación al sustituir la lógica rígida de comprensión, por modelos capaces de razonar, incluso sin haber visto ningún ejemplo previo específico para esas tareas (*zero-shot*). Sin embargo, la integración de estos modelos en pipelines de voz tradicionales introduce una latencia importante. Esto se debe principalmente a que los *LLMs* son bastante lentos generando el primer token, (lo que se mide a través del Time to First Token o *TTFT*), rompiendo así la ilusión de una conversación natural fluida. Esto ha llevado a la industria hacia la búsqueda de sistemas de **baja latencia** y **comunicación full-duplex**, donde la máquina puede escuchar y hablar simultáneamente como GPT-4o [15].

2.2. Arquitecturas actuales para asistentes de voz

Actualmente, existen dos estrategias principales para construir estos asistentes, y cada una tiene sus ventajas y desventajas.

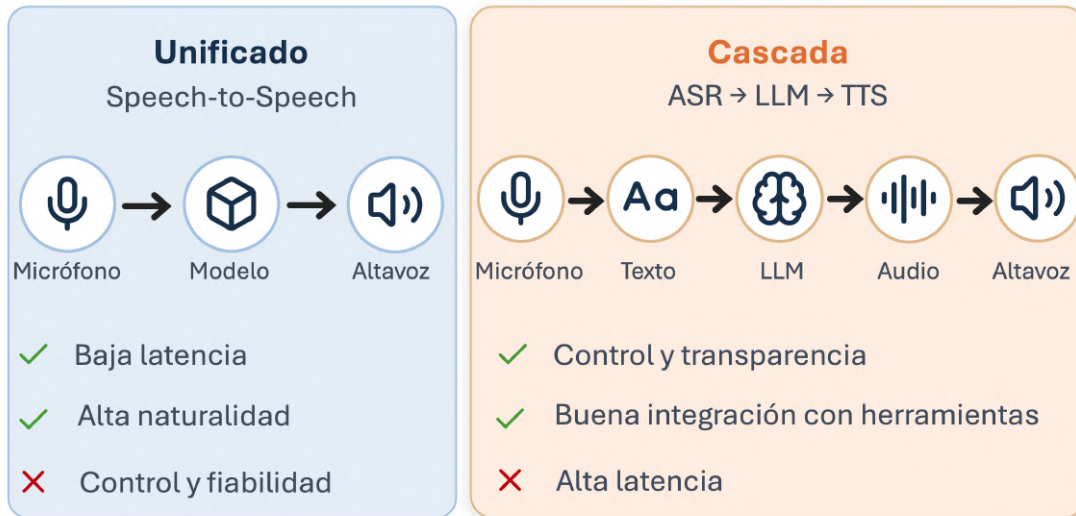


Figura 2.1: Comparativa de la Arquitectura Unificada y la Arquitectura en Cascada.

Por un lado, encontramos la **Arquitectura Unificada**, es decir, los modelos multimodales *Speech-to-Speech*. Modelos recientes como GPT-4o o Moshi [16] procesan el audio de entrada y generan audio de salida directamente, sin pasar por una representación de texto intermedia. Tienen la ventaja de ser muy rápidos y captar matices como el tono o la emoción. Sin embargo, en entornos empresariales actúan a menudo como cajas negras, es muy difícil controlar lo que dicen, pueden alucinar sonidos y resulta complicado integrarlos con bases de datos corporativas de forma fiable y segura. A esto hay que sumar el problema económico y regulatorio. El uso de estos a través de APIs implica un coste elevado basado en tokens de audio. En el caso de usar la Realtime API de OpenAI, el modelo de audio gpt-realtime tiene un coste de 32\$ por millón de tokens de entrada y 64\$ por millón de tokens de salida [17], y en el caso de Google Gemini tiene un coste de 3\$ por millón de tokens de entrada y 12\$ por millón de tokens de salida [18] (aunque el coste de este segundo pueda parecer muy inferior, ambos no son equiparables, ya que Google únicamente ofrece la Live API con el modelo Gemini 2.5 Flash. Este no tiene ni el tamaño ni la inteligencia comparables al modelo usado por OpenAI, para igualarlo, Google debería usar Gemini 3.). Por lo tanto, para llamadas largas, esto puede suponer un gasto considerable. Además, al ejecutarse obligatoriamente en servidores externos de terceros, el audio debe salir de la infraestructura de la empresa, lo que plantea problemas de soberanía del dato importantes.

La alternativa, es la **Arquitectura en Cascada**. Este enfoque conecta componentes especializados: un modelo que transcribe (ASR), la lógica (LLM) y una voz sintética (TTS) [19]. Aunque es cierto que la suma de estos pasos añade algo

más de latencia, ofrece modularidad para añadir tantos bloques intermedios como se quiera. Además, permite inspeccionar la lógica de la conversación en texto, para asegurar que el bot no diga datos falsos y cambiar cualquier pieza del sistema sin tener que reentrenar todo el modelo.

2.3. Retos de la Interacción en Tiempo Real

Más allá de la arquitectura elegida, uno de los mayores retos de las interfaces de voz reside en la latencia de la respuesta. Estudios indican que si una respuesta tarda más de 400ms, el usuario percibe que algo va mal o que el sistema no le ha escuchado [20], lo que obliga a optimizar cada milisegundo desde que el usuario habla hasta que le vuelva la respuesta. La latencia de la que se habla en este trabajo se define como el tiempo desde que el interlocutor humano termina de hablar, hasta el momento en el que el asistente de voz comienza su respuesta.

Para poner en contexto lo difícil que es conseguir un objetivo global de ~ 400 ms de latencia, un estudio técnico de Amazon evaluó seis modelos de lenguaje ejecutados en Bedrock empleando un *system prompt* de aproximadamente 672 tokens y conexiones persistentes. Los autores midieron tanto el *Time-to-First-Token* (TTFT) como el tiempo de respuesta completo. Su análisis muestra que incluso en las mejores condiciones el TTFT ronda varios cientos de milisegundos y que existe un *sueldo* alrededor de 340–400 ms causado por la fase de *prefill* y la latencia de red [21] (que por suerte en nuestro caso será menor). Los valores representan la media del *Time-to-First-Token* (TTFT) y del tiempo total de respuesta para cada modelo bajo condiciones idénticas. Las estimaciones de parámetros provienen de otras fuentes: MiniMax M2.5 (230 mil millones de parámetros totales y 10 mil millones activos) [22], Qwen3 Next 80B (80 mil millones totales y ~ 3 mil millones activos) [23] y DeepSeek V3.2 (685 mil millones totales y 37 mil millones activos) [24].

Tabla 2.1: Resultados del estudio de latencia en Amazon Bedrock [21].

Modelo	Parámetros (aprox.)	TTFT medio (ms)	Tiempo total (ms)
Nova Micro	<i>No publicado</i>	407	616
GPT OSS 20B	21B (3.6B activos)	475	1297
MiniMax M2.5	230B (10B activos)	572	1955
Qwen3 Next 80B	80B (~ 3 B activos)	606	2080
Claude Haiku 4.5	<i>No publicado</i>	1553	1970
DeepSeek V3.2	685B (37B activos)	5038	11232

Y el objetivo de este trabajo sería conseguir un TTFT **total** inferior a 400ms, no solo del bloque LLM, sino de todo el *pipeline* de voz. Por lo tanto es un objetivo

muy ambicioso. Para intentarlo, será necesario usar LLMs pequeñas y modernas y jugar al máximo con la optimización del *warm cache* entre otras cosas para lograr un TTFT lo más bajo posible.

2.4. De Asistentes a Agentes: uso de herramientas

Finalmente, el avance más reciente es la transición de asistentes puramente conversacionales a agentes de voz, que deben ser capaces de saber cuando deben usar herramientas externas. La evidencia reciente sugiere que las arquitecturas en cascada (ASR + LLM) son muy superiores en esta tarea: logran un rendimiento promedio del 65.9 % en escenarios complejos multilingües, frente a tan solo 27.7 % en el caso de los modelos de audio nativos, debido al extenso entrenamiento en razonamiento y código que poseen los modelos de texto [25]. Sin embargo, como ya se ha adelantado previamente, esta fiabilidad impone un compromiso crítico de latencia: mientras un modelo nativo reacciona en 40 ms, la arquitectura en cascada tarda cerca de 800 ms, lo que obliga a implementar estrategias de streaming y concurrencia de los modelos para reducir al máximo la latencia total del sistema.

Capítulo 3

Metodología

El desarrollo de un asistente de voz en tiempo real no solo requiere escribir el código de la arquitectura, sino también tomar decisiones de diseño basadas en datos obtenidos de manera empírica. Por ello, el planteamiento de este proyecto se ha apoyado en **dos niveles metodológicos** complementarios. El primero corresponde a la metodología de desarrollo de *software*, centrada en la construcción iterativa del *pipeline* modular de voz. El segundo nivel consiste en una metodología experimental y de evaluación diseñada para responder a la pregunta: ¿qué LLM *Open Source* ofrece la mejor relación entre capacidad funcional, robustez ante errores en español, uso de herramientas, comportamiento conversacional y latencia para poder integrarse como el cerebro de este asistente?

Este capítulo presenta la metodología que se ha seguido para desarrollar y evaluar el sistema propuesto bajo los siguientes dos enfoques. En primer lugar, en la Sección [3.1](#) se explica cómo se ha planteado el desarrollo del proyecto, siguiendo una metodología iterativa e incremental, así como la planificación general dividida en cinco fases. A continuación, en la Sección [3.2](#) se describe el diseño experimental y el entorno de simulación utilizado para comparar los modelos candidatos de forma controlada. Después, en la Sección [3.5](#) se definen los criterios y las métricas empleados para medir aspectos clave como la tasa de éxito, la velocidad de generación y la latencia. Por último, en la Sección [3.6](#) se detallan los recursos y el entorno de ejecución utilizados durante el desarrollo y la experimentación, tanto a nivel de *hardware* como de *software*.

3.1. Metodología de desarrollo y planificación

Dado que el sistema se compone de varios bloques independientes (pero que operan de forma coordinada), se ha optado por seguir una metodología de desarrollo **iterativa e incremental**. En lugar de intentar construir el sistema completo

de una sola vez, se ha dividido el trabajo en varias fases secuenciales, en las que cada una añadía una funcionalidad sobre el trabajo ya realizado. Esto ha permitido detectar fallos rápidamente localizando fácilmente el error y sustituir modelos concretos si no daban el resultado esperado, sin obligar a rediseñar el resto del sistema.

El proyecto se ha estructurado apoyándose en **dos niveles de metodología** diferenciados. Por un lado, la metodología de desarrollo de *software* centrada en la construcción del *pipeline* modular de voz. Por otro lado, un enfoque de evaluación continua donde la selección del LLM se trató como una fase experimental propia. Dado que el modelo de lenguaje es el bloque más importante y condiciona de manera directa tanto la inteligencia del asistente como su latencia, fue necesario diseñar una metodología experimental para seleccionar el LLM más adecuado antes de integrarlo de forma definitiva en el sistema.

El plan de trabajo del proyecto se ha estructurado en cinco fases principales, que han cubierto desde la investigación inicial hasta la defensa final. A continuación se detalla cada una de ellas:

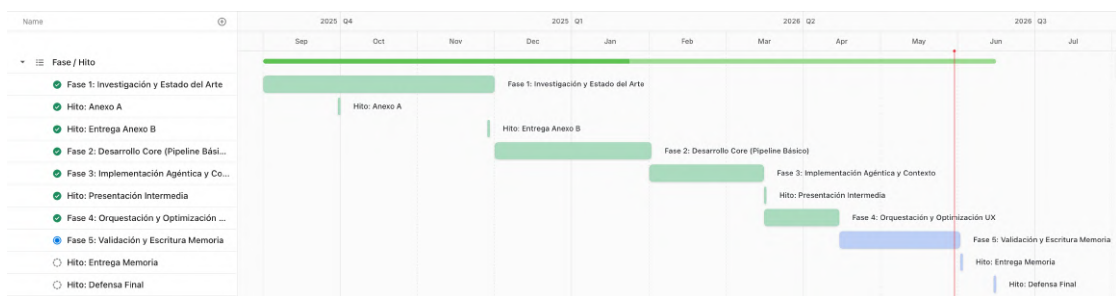


Figura 3.1: Diagrama de Gantt del proyecto

Fase 1: Investigación y Definición (Septiembre - Noviembre)

Esta fase inicial se centró principalmente en hacer un estudio del estado del arte. Se analizaron las diferencias entre arquitecturas unificadas y en cascada, identificando posibles infraestructuras y familias de modelos candidatos. Durante estos meses se hizo una primera selección de las tecnologías *Open Source* sobre las que se fundamentaría el desarrollo.

Fase 2: Desarrollo del Core Pipeline (Diciembre - Enero)

El objetivo de esta fase fue construir el esqueleto del sistema. Utilizando modelos provisionales para no bloquear el desarrollo, se implementó la comunicación bidireccional con WebSockets y se conectaron los bloques básicos: VAD, ASR, LLM y TTS. Al finalizar esta fase, se obtuvo un primer prototipo capaz de escuchar,

transcribir, generar una respuesta simple y hablar, aunque todavía sin inteligencia agéntica ni gestión de interrupciones.

Fase 3: Implementación Agéntica y Contexto (1 Febrero - 15 Marzo)

Una vez que el sistema básico funcionaba, se procedió a dotarlo de inteligencia. Esta etapa permitió definir qué capacidades estrictas debía tener el LLM: uso de herramientas (*tool calling*), extracción de parámetros, seguimiento del contexto y capacidad de abstención ante información insuficiente. Se implementó la lógica de *Function Calling* para que el modelo supiera cuándo invocar herramientas externas y, en paralelo, se desarrolló el gestor de contexto para que el asistente tuviera memoria durante la llamada.

Fase 4: Orquestación y Optimización de UX (15 Marzo - 15 Abril)

Con el sistema ya funcional e inteligente, esta fase conectó el desarrollo de *software* con la necesidad de medir la latencia y el TTFT (*Time To First Token*), ya que en una interfaz de voz no basta con que el modelo sea inteligente. Se programó la lógica de gestión de turnos para evitar cortes mediante una máquina de estados, la detección de interrupciones (*barge-in*) y la inyección de frases de relleno (*fillers*) para enmascarar los tiempos de espera de las herramientas implementadas en la fase anterior.

Fase 5: Validación y Documentación Final (15 Abril - 1 Junio)

La última etapa se dedicó a las pruebas y benchmarks finales. Se ejecutaron pruebas un gran catálogo de modelos (actualizado con los últimos LLMs hasta la fecha) para seleccionar el definitivo. Se recopilaban métricas de latencia y tasa de éxito, se corrigieron errores del sistema y se redactó la Memoria final del proyecto para su posterior defensa.

3.2. Diseño experimental y entorno de simulación

3.2.1. Entorno de evaluación para la selección del LLM

Para garantizar que el asistente de voz funcionara de manera correcta, fue necesario diseñar una metodología experimental orientada a seleccionar el componente más crítico del sistema: el LLM. Dado que este actúa como el núcleo del sistema

y es el principal responsable de la latencia variable del *pipeline* como se vio en las primeras pruebas, se diseñó una serie de *benchmarks* centrados exclusivamente en este bloque para elegir el más adecuado.

El experimento se diseñó a partir de una idea clave: la evaluación de los LLMs candidatos debía realizarse aislando el modelo del resto del *pipeline* de voz, para estar seguros de que los resultados fueran fiables y que cualquier variación fuera exclusivamente causada por el LLM. Las pruebas se realizaron introduciendo el texto directamente a esta IA, evitando así que los posibles errores del reconocimiento automático del habla (ASR), los tiempos de procesamiento de la síntesis de voz (TTS) o cualquier otro bloque del *pipeline* contaminaran la comparativa. Todos si que se consideraron como parte del contexto de aplicación final, pero no como objeto de medida en esta fase de selección.

Para llevar a cabo la evaluación, se definió un catálogo controlado de modelos candidatos con pesos abiertos e inferencia local, limitando el tamaño a un máximo de 40B de parámetros para asegurar tiempos de respuesta viables para un asistente de voz en tiempo real. Todos los modelos se sometieron a dos bloques de simulación diferentes:

- **Benchmarks de inteligencia:** Se desarrolló un sistema de pruebas automático que enviaba escenarios (*prompts*) al modelo y analizaba sus respuestas de manera controlada. Para lograr una evaluación integral, este entorno combinó conjuntos de diseño propio con referencias estándar de la industria. Por un lado, se evaluó la capacidad de seguir instrucciones en español (*custom_spanish*) y se incluyó una batería de pruebas de robustez (*spanish_robustness*) con errores típicos de transcripción oral. Por otro lado, la evaluación se completó integrando subconjuntos de BFCL (*Berkeley Function-Calling Leaderboard*) para medir rigurosamente el uso de herramientas (*tool calling*), y de VoiceBench para analizar el comportamiento conversacional y de seguridad.
- **Benchmarks de rendimiento:** Ya que la velocidad es crítica en un asistente de voz, se diseñó un *benchmark* específico para medir la latencia de generación en modo *streaming*, es decir que se empieza a medir desde la salida desde el primer token. Para evitar variaciones del *hardware* y obtener los resultados lo más reales posibles, todas las mediciones de latencia se ejecutaron de forma secuencial sobre una misma GPU dedicada, procesando una única petición a la vez. Aunque cabe destacar que pueden seguir habiendo pequeñas variaciones ya que el cluster de GPUs de ICAI cuenta con 8 H200, es muy probable que se ejecutaran otros procesos por otras personas y que como los recursos (CPU, NVMe, RAM...) son compartidos, cierta variación es inevitable. Es por ello que se realizaron muchas medidas para tener intervalos de confianza aceptables. También se simulaban contextos de distinto

tamaño (desde 0 hasta 128.000 tokens) para analizar la degradación de la latencia a medida que la llamada telefónica se alarga.

3.2.2. Benchmark de inteligencia

Para evaluar el rendimiento de los LLMs lo mejor posible para el caso de uso de este trabajo, es decir, en el caso de un asistente de voz inteligente, se diseñó un *benchmark* de inteligencia propio. Esto se hizo porque muchos *benchmarks* públicos se centran en el conocimiento de los modelos general en inglés, en capacidades de razonamiento matemático u otros aspectos que no se necesitan directamente en el entorno de una llamada de voz en español. Además, para garantizar una comparativa justa y evitar que desborde la memoria en los modelos más pequeños, todas las pruebas se limitaron a un contexto máximo de 16.384 tokens (`16k_filtered`). De todas formas, este filtro es incluso excesivo para el caso de uso de este trabajo, ya que un asistente de voz no suele necesitar ventanas de contexto más largas que 8k tokens, ya que las llamadas de teléfono no son tan largas, incluso incluyendo el *system prompt*. Por ejemplo, en un caso en el que el *system prompt* ocupe 4.000 tokens, quedarían aproximadamente otros 4.000 tokens (como por ejemplo el uso en este trabajo que ocupa 4.068 tokens y se puede acceder en el Anexo E), lo que equivale más o menos a 3.000 palabras, haciendo la estimación comúnmente aceptada en el mundo de la IA de que un token equivale aproximadamente a 0,75 palabras. Por tanto, a una velocidad media de 150 palabras habladas por minuto, se podría mantener una conversación de aproximadamente:

$$\frac{4,000 \text{ tokens} \times 0,75 \text{ palabras/token}}{150 \text{ palabras/minuto}} = \frac{3,000 \text{ palabras}}{150 \text{ palabras/minuto}} \approx 20 \text{ minutos}$$

Por lo tanto, es una cantidad de tokens más que aceptable. Además siempre se podrían implementar mecanismos de resumen de la conversación en mitad de la ejecución como se propone en 6.3.

Con todo esto en mente, la evaluación se organizó en cuatro bloques principales, cada uno orientado a medir un aspecto distinto del comportamiento del asistente. Se recomienda ver los ejemplos en el B.2 para entender mejor los benchmarks ya se exponen ahí casos tanto de acierto (cuando el modelo acertó el test) y de error (el modelo no actuó como debía). Se incluyen ejemplos para todos los benchmarks de inteligencia que son los siguientes:

- **Evaluación de instrucciones propias en español** (`custom_spanish`): Este bloque evalúa cómo el modelo interpreta y ejecuta instrucciones conversacionales en español. Inspirado en el conocido benchmark *IFEval* [26] desarrollado por Google, comprueba si el modelo es capaz de ejecutar una

acción, pedir una aclaración cuando faltan datos, confirmar un comando potencialmente peligroso o abstenerse de responder ante peticiones indebidas. En este caso, la puntuación se obtiene revisando la estructura de la respuesta y las llamadas a herramientas, en lugar de depender de otro LLM como juez que puede introducir más errores de LLM.

- **Robustez ante errores de transcripción** (`spanish_robustness`): Este bloque mide como actúa LLM ante distintas formulaciones y errores típicos de los sistemas ASR, como cambios de género, ausencia de tildes o sustituciones por palabras homófonas. Este benchmark se ha inspirado del *Speech Robust Bench* (SRB) [27], creado por un equipo de investigadores del Eurecat, Helsinki Institute for Information Technology (HIIT) y del Carnegie Mellon University (CMU). Lo que busca es comparar la respuesta del modelo ante un texto base limpio frente a tres niveles de degradación del texto que simula una transcripción oral. Esto permite analizar de forma objetiva si los modelos siguen entiendo la intención incluso si la transcripción es defectuosa.
- **Evaluación de uso de herramientas** (`bfcl_reference`). Para medir la capacidad del modelo en tareas de selección de herramientas (*tools*) y extracción de parámetros, se ha usado el benchmark *Berkeley Function-Calling Leaderboard* [28] desarrollado por UC Berkeley.
- **Evaluación conversacional y de seguridad** (`voicebench_reference`). Finalmente, se usó el benchmark *VoiceBench* [29] desarrollado por la National University de Singapore (NUS), utilizando sus subconjuntos para evaluar el seguimiento de instrucciones (*IFEval*), el razonamiento complejo (*BBH*), la seguridad ante ataques (*AdvBench*), el conocimiento general (*OpenBookQA*) y las capacidades multilingües (*MMSU*) en el contexto de una conversación.

Todos los resultados de estas simulaciones, incluyendo configuraciones y errores de ejecución, se almacenaron como artefactos de evaluación para garantizar la reproducibilidad y trazabilidad de la decisión final.

Benchmark de velocidad

Como se ha comentado previamente, en el diseño de un asistente de voz, la inteligencia del modelo no es suficiente. Un LLM con una gran inteligencia pero tiempos de respuesta elevados generaría silencios en la conversación, ya que los tokens tardarían mucho en llegar al TTS, lo que provocaría interrupciones no deseadas que empeorarían la experiencia del usuario. Por tanto, es necesario tener esto en cuenta a la hora de elegir que LLM se va a usar.

Para medir este rendimiento se ha realizado un segundo *benchmark* propio. En este caso, el sistema realiza varias mediciones enviando peticiones al servidor de inferencia consumiendo la respuesta del modelo en modo *streaming*. El uso de *streaming* es necesario en esta prueba (al igual que en el pipeline de voz), ya que es la única forma de capturar el instante real en el que el modelo emite su primer fragmento de texto, es decir, lo que se conoce como el *Time To First Token* (TTFT). Además del TTFT, este *benchmark* evalúa la latencia total de la respuesta para comprobar que es superior a la velocidad a la que se habla (por lo tanto una vez que el TTS ha recibido el primer fragmento de texto, los tokens van a llegar más rápidos de lo que se necesita y pararán de ser el factor limitante). Para analizar esto último, se va a medir la latencia total junto con el número de tokens generados lo que permite obtener la velocidad real de generación en tokens por segundo, es decir los (*Tokens per Second* o *TPS*).

A diferencia de las pruebas de inteligencia, donde las peticiones se paralelizaron para reducir el tiempo total de evaluación (ya que que un modelo sea más inteligente o menos no depende de la GPU en la que se ejecute), este *benchmark* de velocidad se ejecutó de forma secuencial. Todas las mediciones se realizaron sobre la misma GPU, asegurando que el servidor de inferencia procesaba una única petición a la vez. Esta decisión permite aislar al máximo el rendimiento puro del modelo, para intentar reducir al máximo la variabilidad de los resultados. Aunque cabe destacar, que el ser cluster de GPUs de ICAI cuenta con 8 H200, es muy probable que se ejecutaran otros procesos por otras personas y que como los recursos (CPU, NVMe, RAM...) son compartidos, cierta variación es inevitable. Es por ello que se realizaron muchas medidas para tener intervalos de confianza aceptables.

Por otro lado, como en una llamada telefónica el historial de la conversación crece continuamente (cuantos más minutos tenga la conversación), se ha diseñado un *benchmark* adicional e independiente centrado en como influye la longitud del contexto de la conversación que se introduce en el LLM. El objetivo de esta prueba es estudiar cómo escala el TTFT del modelo a medida que aumenta la longitud del texto de entrada. Para ello, se evaluaron contextos de diferentes tamaños (desde 0 hasta 128.000 tokens) extraídos del también del corpus de El Quijote¹

Estos *benchmarks* de velocidad eran importantes realizarlos en los servidores de ICAI y no se podían usar datos públicos en internet, porque es donde se va a ejecutar el sistema final y era imposible encontrar valores públicos en internet que probaran todos los modelos candidatos y con un mismo *hardware* para poder tener comparaciones justas.

¹El corpus de El Quijote se obtuvo gracias a un repositorio de la web Proyecto Gutenberg: <https://www.gutenberg.org/cache/epub/2000/pg2000.txt>.

3.3. Metodología del Benchmark *End-to-End* (E2E)

Para medir el rendimiento del agente de voz en un entorno lo más cercano posible a un escenario real, se ha diseñado un *benchmark End-to-End* (E2E) totalmente automatizado. A diferencia de las pruebas aisladas del LLM (que evaluaba solo ese modelo), este benchmark evalúa la totalidad del *pipeline* de audio. Es decir, no mide solo si el LLM responde bien, sino si todo el sistema consigue completar una llamada realista de restaurante usando voz, ASR, herramientas, estado de base de datos y respuesta hablada.

3.3.1. Funcionamiento de la simulación

El benchmark final se compone de **300 escenarios conversacionales**. Cada escenario está definido en un fichero `scenario.yaml`, donde se indica la categoría del caso, el objetivo del cliente, el estado inicial del restaurante, el resultado esperado, las herramientas que deberían usarse y el límite máximo de turnos. Las categorías cubren creación de reservas, modificación, cancelación, rechazo cuando no se debe crear una reserva, aceptación de horarios alternativos, preguntas de información del restaurante y casos de borde.

Este benchmark utiliza un modelo de *Usuario Sintético*. Se trata de un LLM independiente configurado con un objetivo específico como por ejemplo se le da la siguiente instrucción: "*Llama al restaurante, intenta reservar para el miércoles a las 19:00 para 2 personas; si no hay sitio, rechaza alternativas y despídete*". En el benchmark final se usó Gemma 4 31B como usuario sintético, con *thinking mode* activado, para que pudiera razonar sobre el objetivo del escenario y mantener una conversación coherente. En cada turno, este modelo recibe las instrucciones generales de comportamiento, el escenario completo, el historial reciente y la última respuesta del asistente.

La frase generada por el usuario sintético se sintetiza con ElevenLabs en audio PCM a 16 kHz, ya que permite usar voces distintas, con diferentes acentos y timbres, y así se evalúa de forma mucho más robusta el sistema. En los escenarios también se variaron perfiles de audio, incluyendo audio limpio y casos con ruido telefónico.

El flujo de ejecución de la prueba es el siguiente:

1. Antes de comenzar cada escenario, el evaluador prepara el estado del restaurante: borra reservas que podrían interferir, crea reservas iniciales para los casos de modificación o cancelación y crea bloqueos de mesas cuando se quiere forzar una alternativa horaria.
2. El usuario sintético genera su respuesta y la sintetiza a voz mediante el proveedor TTS (ElevenLabs).

3. El audio se transmite en *streaming* (PCM) por red al agente de recepción, usando el mismo endpoint WebSocket que usaría un cliente real.
4. El *pipeline* del agente procesa el audio entrante (AEC, VAD, ASR e ITN/-normalización), el LLM del *pipeline* razona, usa herramientas si es necesario y responde.
5. Para optimizar el tiempo de simulación y evitar la propagación de errores de transcripción en el cliente de prueba, el texto generado por el agente se devuelve directamente al LLM del usuario sintético, mientras que la voz generada por el TTS del asistente se monitoriza para mediciones de latencia.
6. La llamada continúa de forma autónoma hasta que el escenario se considera completado, se detecta una condición terminal o se alcanza un límite máximo de turnos (*max turns*). Al terminar, el evaluador limpia el estado creado para que el siguiente escenario empiece en condiciones controladas.

La interacción se produce a través del mismo servidor WebSocket que usaría un cliente real. En la Figura 3.2 se ilustra la arquitectura de esta simulación. Se observa el flujo desde la generación de audio del usuario sintético, el procesamiento en tiempo real del pipeline bajo prueba, hasta la evaluación paralela post-llamada (determinista y juez conversacional).

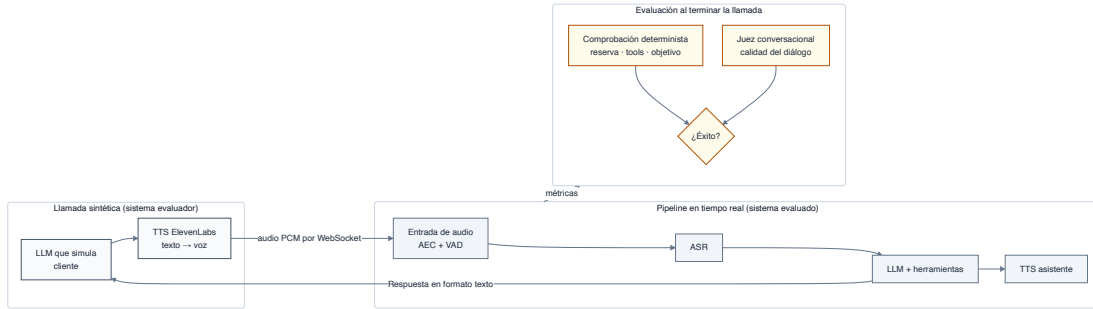


Figura 3.2: Arquitectura del benchmark *End-to-End*.

Una vez finalizada la llamada, el sistema extrae las métricas de latencia y un *script* evaluador analiza los registros de la base de datos, el historial de herramientas invocadas, los eventos WebSocket y la transcripción exacta de la conversación para medir el rendimiento. Para cada escenario se guardan artefactos como `conversation.json`, `evaluation.json`, `quality.json`, `events.jsonl`, `tool_calls.jsonl`, `metrics.json`, `transcript.txt` y los audios generados. Esto permite revisar manualmente cualquier fallo y reconstruir la línea temporal exacta de la llamada.

3.3.2. Sistema de validación por capas

Dado que las conversaciones humanas son impredecibles, evaluar un agente conversacional de forma binaria (Aprobado/Suspenso) es insuficiente. Para obtener una visión analítica precisa, se ha diseñado un **Sistema de validación por capas**.

Este concepto filtra cada uno de los escenarios a través de capas de validación cada vez más estrictas. Si una llamada supera una capa, avanza a la siguiente. Esto permite aislar si un fallo se debe a una incapacidad técnica del modelo (no saber usar una herramienta) o a otros aspectos como no usar el vocabulario adecuado por ejemplo. Las capas que componen este sistema son tres:

1. Tarea completada correctamente

Esta métrica evalúa **el comportamiento correcto** del agente. Se considera que una tarea es exitosa si, y solo si, el estado final del sistema coincide con el objetivo del escenario. Para aprobar, el evaluador verifica:

- **Selección de herramientas:** ¿Usó el agente las herramientas correctas? (Ej. usar `modify_reservation` en lugar de crear una nueva).
- **Precisión de argumentos:** ¿Extrajo correctamente los datos dictados por el usuario sin errores probablemente introducidos por el ASR? (Fecha, hora, comensales, nombre, teléfono).
- **Estado de la Base de Datos:** En escenarios de reserva, el evaluador consulta la base de datos SQL real al final de la llamada. Si el registro no existe, está en una hora equivocada, o pertenece a otro cliente, la tarea suspende automáticamente.
- **Resultado esperado por categoría:** En creación de reserva debe existir la reserva correcta; en modificación, la reserva inicial debe cambiar al estado esperado; en cancelación, debe desaparecer; en rechazo, no debe crearse ninguna reserva prohibida; y en alternativas, se permite una hora alternativa aceptable, no necesariamente la pedida inicialmente.

Nota: Una llamada puede superar el Éxito de Tarea incluso si el agente fue 'antipático' o poco natural, siempre y cuando la transacción en el backend sea perfecta claro.

2. Calidad de la interacción

Esta métrica evalúa **la naturalidad y el cumplimiento de instrucciones** de la llamada. Incluso si el agente logró hacer la reserva, suspenderá en Calidad

si la interacción fue mala. Para ello se usa un juez LLM que recibe el transcript completo de la llamada, pero no juzga directamente la base de datos, ya que eso se evalúa en la capa anterior. El evaluador penaliza la llamada si detecta:

- **Bucles en la conversación:** El agente pregunta el mismo dato (ej. el número de teléfono) repetidas veces sin avanzar, o entra en un bucle infinito de despedidas con el usuario sintético.
- **Respeto del prompt:** El prompt exige ciertas formas de responder por diseño de producto. Por ejemplo, si el restaurante está lleno, no basta con decir "*no me quedan mesas*"; el agente puede estar obligado a usar palabras clave como *¿completo*". Si no lo usa se considera un fallo de Calidad. Esto es un ejemplo que se ha implementado pensando por ejemplo que algunos negocios quieran que su agente se disculpe si no puede realizar una acción.
- **Exceso de turnos:** Si la llamada se alarga innecesariamente sin llegar a una resolución, se considera una experiencia de usuario mala, por ejemplo tener que repetir muchas veces tu número de teléfono.

Esta métrica no es determinista como la anterior, ya que es imposible hacer reglas exactas para cada caso así que es el LLM el que la evalúa. Aun así, el benchmark también guarda señales deterministas de bucles, herramientas repetidas y escenarios que llegan a *max turns*, para que estos fallos se puedan analizar después sin depender solo del juez LLM. Además, una conversación larga no falla automáticamente: puede ser legítima si hay negociación de horario, corrección de datos o alternativas; falla cuando hay repetición sin progreso.

3. Validación completa

Es la métrica más exigente del benchmark y busca ser más una métrica de cómo de preparado está el producto para producción.

La Validación completa es la intersección lógica (operación *AND*) de la tarea completada correctamente, la calidad de la interacción y la ausencia de fallos conversacionales críticos. En la configuración final del benchmark esto se expresa como:

`validacion_completa = tarea_completada_correctamente^calidad_de_la_interaccion^ausencia_de_`

Para que un escenario logre una validación completa, el agente debe haber ejecutado las operaciones de base de datos a la perfección (Tarea completada correctamente), haber mantenido una conversación fluida, natural, y con las instrucciones que se le han dado (Calidad de la interacción), y además no haber caído en un fallo crítico como un bucle o un *max turns* no justificado. Es la métrica más importante que define la fiabilidad del sistema.

3.3.3. Benchmark de latencia

Además del benchmark E2E completo, se ha realizado un test para medir con más precisión la latencia percibida del sistema en una interacción de voz. El objetivo es medir cuánto tiempo transcurre desde que el usuario termina físicamente de hablar hasta que el asistente empieza a emitir el primer fragmento de audio de respuesta.

Esta métrica es especialmente importante en un asistente de voz, porque representa el silencio que percibe el usuario desde que termine de hablar hasta que el sistema le responde. Para construir este benchmark se reutilizaron audios reales generados durante el benchmark E2E para reutilizar los recursos ya creados. Cada audio se envió al sistema mediante el mismo endpoint WebSocket que utiliza el cliente real y en esta prueba no se utilizó juez LLM ni evaluación de éxito de tarea, ya que únicamente se buscaba medir tiempos.

Para iniciar la medición, se analizó cada fichero WAV con un análisis de energía RMS sobre el audio. Esto permitió encontrar el instante exacto en el que finaliza la voz del usuario, descartando el silencio que contienen las grabaciones generadas por ElevenLabs (como solo son voces sin ruidos, como si estuvieran grabadas en un estudio, el RMS funciona muy bien). Ese punto se considera como el fin de la voz del usuario, y se midió el tiempo transcurrido hasta recibir el primer paquete de audio devuelto por el asistente.

Durante cada ejecución se almacenaron también métricas internas por bloque (incluyendo ASR, LLM y TTS), lo que permite descomponer la latencia total para comprobar cuántos milisegundos corresponden a la fase de cierre de turno y cuántos a la inferencia de los modelos.

3.4. Metodología del Benchmark de Interrupciones (*Stop Reaction Time*)

Además del benchmark *End-to-End*, se ha diseñado una prueba específica para medir el **tiempo de reacción ante interrupciones** (*Stop Reaction Time*): los milisegundos que transcurren desde que el usuario empieza a hablar mientras que el sistema está hablando hasta que el pipeline se entera de eso y se 'calla' y corta la síntesis de voz en curso. Aquí no se evalúa la calidad del diálogo ni el uso de herramientas, solo la rapidez con la que es capaz el sistema en darse cuenta de que lo están interrumpiendo interrumpido.

3.4.1. Audios usados y preparación de datos

El conjunto de pruebas se construyó a partir de los audios de ElevenLab generados durante el benchmark E2E. Para evitar duplicados (ya que como hubo que repetir varias veces el benchmark E2E por varios errores), cada archivo se identifica mediante su hash SHA-256. Se contaba con un corpus de **6 348 audios únicos** (esto se debe a que el E2E contaba de 300 tests, cada test con varios turnos de palabra y hubo que repetirlo 3 veces) aunque para evitar ejecutar más pruebas de las necesarias, solo se usaron **3 000** de ellos.

3.4.2. Modo de evaluación del pipeline(*interruption_eval*)

Como solo se desea medir la capacidad del sistema a reaccionar a interrupciones, se ha añadido al pipeline un modo de benchmark para que cuando este modo está activo, el flujo normal de bienvenida se omite y el sistema entra directamente en el estado `ai_talking`, reproduciendo un audio largo.² De este modo, se ahorraron muchos recursos y se aceleró el tiempo de ejecución ya que no se pasaba por el LLM, ASR ni por el TTS por ejemplo.

3.4.3. Protocolo de cada ejecución

El flujo de una prueba individual es el siguiente:

1. El cliente de evaluación abre una conexión WebSocket (`/ws/audio`) y envía el `hello` con el parámetro `interruption_eval`.
2. Se espera el evento `tts_start`, que confirma que el asistente está emitiendo audio.
3. Tras una espera aleatoria uniforme entre **0,1 s y 8 s** desde la IA empieza a hablar (para hacer una medición a diferentes instantes y estar seguro de que sea aleatorio), se inicia un cronómetro.
4. Inmediatamente después, se inyecta el audio del usuario en *streaming* (tramas PCM de 20 ms, igual que un cliente real).
5. El cronómetro se detiene cuando se reciben **ambos** eventos: cambio de estado a `user_talking_asr` y `tts_end` con `reason: interrupted`.
6. El tiempo de reacción se calcula como el instante **max** de esos dos eventos (o la duración del audio o el tiempo medido), medido desde el envío del primer byte de audio del usuario.

²Si no se envía este parámetro, la FSM se inicia como siempre en el estado `WELCOME`

Para cada uno de los 3.000 tests se hace como un append en un fichero JSONL.

3.5. Criterios y métricas de evaluación

Para valorar de forma objetiva tanto a los modelos candidatos como al sistema final, se definieron dos familias de métricas, separando claramente la evaluación del LLM del rendimiento del *pipeline* completo de voz.

3.5.1. Métricas de selección del Modelo de Lenguaje (LLM)

Durante la fase experimental, los criterios para comparar los LLMs se dividieron en dos aspectos principales:

- **Métricas de latencia:** El criterio principal fue el *Time To First Token* (TTFT), medido en milisegundos, que mide el tiempo exacto desde que el modelo recibe la petición hasta que emite el primer fragmento de texto útil. También se midió la velocidad de generación en *Tokens per Second* (TPS), que se calculó como el cociente entre el número de tokens de salida generados y la latencia total de la petición (en segundos):

$$\text{TPS} = \frac{N_{\text{salida}}}{T_{\text{total}}/1000} \quad (3.1)$$

donde N_{salida} es el número de tokens emitidos durante el *streaming* de la respuesta y T_{total} es la latencia extremo a extremo en milisegundos, medida desde el envío de la petición hasta que finaliza la generación de tokens para la respuesta.

- **Métricas de inteligencia:** La evaluación funcional del LLM se basó en cuatro *benchmarks* complementarios, cuyos resultados se agregaron en una puntuación funcional global. En el *benchmark* propio en español (*custom_spanish*), cada escenario exige una decisión de primera interacción entre cuatro clases de acción: ejecutar (*ACT*), pedir aclaración (*CLARIFY*), confirmar (*CONFIRM*) o abstenerse (*ABSTAIN*). Se midió la **tasa de acierto** como la proporción de casos en los que la clase predicha coincide con la esperada:

$$\text{Acc} = \frac{N_{\text{correctos}}}{N_{\text{casos}}} \quad (3.2)$$

donde $N_{\text{correctos}}$ es el número de escenarios correctamente resueltos e N_{casos} el total evaluado sobre texto escrito limpio en español. En el bloque de **robustez** (*spanish_robustness*), el mismo criterio se aplicó a cuatro formas

distintas de escribir el mismo enunciado: texto limpio, habla natural, variante con errores típicos de un ASR bueno y variante con errores severos. La simulación es exclusivamente en formato de texto (no se hizo nunca con ruido). La puntuación de robustez usada en el ranking se calculó como media ponderada de las tres superficies degradadas (habla natural, ASR leve y ASR fuerte), con pesos 0,20, 0,35 y 0,45 respectivamente. También se usaron referencias externas para completar este benchmark de inteligencia, como BFCL (*Berkeley Function Calling Leaderboard*), que mide la capacidad de llamar a funciones con el formato exigido por su sistema oficial de evaluación, y Voice-Bench, centrado en capacidades conversacionales relevantes para asistentes de voz. En ambos casos, no se siguió el protocolo completo de evaluación por su elevada cantidad de pruebas, y que no era necesario para obtener una idea bastante clara de la inteligencia de los modelos. En BFCL se usó una muestra de 100 casos por categoría, con 1,000 casos en total, y se descartaron previamente los ejemplos cuya entrada estimada superaba los 16,384 tokens, para que pudieran ejecutarse en la mayoría de modelos. Aun así, algunos casos siguieron superando la ventana de contexto de ciertos modelos y produjeron error; estos casos se contabilizaron como incorrectos. En VoiceBench, la evaluación se limitó a cinco tareas en modalidad texto (*ifeval*, *adubench*, *openbookqa*, *mmsu* y *bbh*), con un máximo de 100–200 muestras por tarea. Estos benchmarks no están adaptados específicamente al español, pero sirven como referencia complementaria para medir la llamada a funciones y el razonamiento conversacional en inglés, junto con los benchmarks propios en español. La **puntuación final** combinó las cuatro puntuaciones normalizadas mediante:

$$S_{\text{func}} = 0,25 S_{\text{es}} + 0,25 S_{\text{rob}} + 0,25 S_{\text{BFCL}} + 0,25 S_{\text{VB}} \quad (3.3)$$

donde S_{es} , S_{rob} , S_{BFCL} y S_{VB} son las puntuaciones de *custom_spanish*, robustez, BFCL y VoiceBench respectivamente.

3.5.2. Métricas de rendimiento del Sistema Final (End-to-End)

Una vez seleccionado el LLM e integrado en la arquitectura descrita en el Capítulo 4, el foco de la evaluación se trasladó a la experiencia de usuario y al comportamiento dinámico del asistente:

- **Latencia extremo a extremo (*End-to-End Latency*):** Tiempo transcurrido desde que el usuario termina de hablar hasta que el sistema reproduce el primer fragmento de voz sintetizada. Incluye los tiempos acumulados de VAD, ASR, TTFT del LLM y TTFT del TTS.

- **Tiempo de reacción ante interrupciones (*Stop Reaction Time*):** Métrica crítica para evaluar la fluidez ante interrupciones (*barge-in*). Mide cuántos milisegundos tarda el sistema en detener la reproducción de audio desde el instante en que el usuario comienza a hablar por encima del asistente.
- **Evaluación cualitativa de la experiencia conversacional:** Además de las métricas numéricas, también se tuvo en cuenta la experiencia real al usar el sistema de forma interactiva. Este criterio permite valorar si la conversación resulta natural y fluida durante la llamada, así como detectar en qué situaciones concretas o formas de hablar el asistente tiene más dificultades o responde de manera menos robusta.

3.6. Recursos y entorno de ejecución

Para desarrollar este proyecto y poder ejecutar tanto el *pipeline* de voz en tiempo real como las extensas campañas de *benchmarking*, ha sido necesario contar con recursos avanzados, tanto a nivel de *hardware* como de *software*, que se detallan a continuación.

3.6.1. Recursos Hardware y Conectividad

Al principio del proyecto, se utilizó principalmente un **portátil personal** (Apple M-Series) que permitió realizar la investigación inicial y estructurar el código base. Para las primeras pruebas de lógica, este dispositivo fue suficiente ya que se usaron modelos muy ligeros (de 4B de parámetros o menos) optimizados para ejecución local mediante el *framework* MLX o `llama.cpp`. Sin embargo, estas ejecuciones presentaban latencias elevadas en inferencia y limitaban enormemente el nivel de inteligencia del asistente.

Para superar estas limitaciones y obtener un rendimiento realista, equiparable a un entorno profesional de producción, se solicitó acceso al **clúster de GPUs** de la universidad (ICAI). Las especificaciones del servidor utilizado son muy potentes (incluso ‘demasiado’ para nuestro caso de uso, lo que nos permitió hacer por ejemplo evaluar muchos más modelos de los que en un principio se preveía).

```

[202010774@dgx:~]$ nvidia-smi
Sat May 30 15:50:40 2026

```

NVIDIA-SMI 570.124.06			Driver Version: 570.124.06			CUDA Version: 12.8		
GPU	Name	Perf	Persistence-M	Bus-Id	Disp.A	Volatile	Uncorr. ECC	
Fan	Temp		Pwr:Usage/Cap		Memory-Usage	GPU-Util	Compute M.	MIG M.
0	NVIDIA H200		Off	00000000:1B:00.0	Off		0	
N/A	31C	P0	81W / 700W	1MiB / 143771MiB		0%	Default	Disabled
1	NVIDIA H200		Off	00000000:43:00.0	Off		0	
N/A	35C	P0	121W / 700W	28525MiB / 143771MiB		0%	Default	Disabled
2	NVIDIA H200		Off	00000000:52:00.0	Off		0	
N/A	37C	P0	124W / 700W	34577MiB / 143771MiB		0%	Default	Disabled
3	NVIDIA H200		Off	00000000:61:00.0	Off		0	
N/A	34C	P0	76W / 700W	1MiB / 143771MiB		0%	Default	Disabled
4	NVIDIA H200		Off	00000000:9D:00.0	Off		0	
N/A	35C	P0	124W / 700W	9201MiB / 143771MiB		0%	Default	Disabled
5	NVIDIA H200		Off	00000000:C3:00.0	Off		0	
N/A	32C	P0	124W / 700W	117654MiB / 143771MiB		0%	Default	Disabled
6	NVIDIA H200		Off	00000000:D1:00.0	Off		0	
N/A	37C	P0	126W / 700W	79219MiB / 143771MiB		0%	Default	Disabled
7	NVIDIA H200		Off	00000000:DF:00.0	Off		0	
N/A	34C	P0	76W / 700W	1MiB / 143771MiB		0%	Default	Disabled

Processes:							
GPU	GI	CI	PID	Type	Process name	GPU Memory	
ID	ID	ID				Usage	

Figura 3.3: Captura de pantalla de nvidia-smi mostrando el estado de las GPUs del clúster de ICAI

El sistema operativo es **Ubuntu 24.04.1 LTS** y cuenta con dos procesadores Intel Xeon Platinum 8480C, sumando un total de 224 hilos de ejecución. Dispone de **2 TiB de memoria RAM** principal y un ecosistema de almacenamiento basado en múltiples discos NVMe Samsung de altísima velocidad (de hasta 3.5 TB), lo que garantiza que el cuello de botella nunca se encuentre en la lectura de los pesos de los modelos hacia la GPU. En cuanto a la aceleración gráfica, el clúster cuenta con 8 tarjetas **NVIDIA H200** con aproximadamente 144 GB de VRAM

cada una. Y se tuvo la suerte de que no se nos limitó el uso de ningún recurso, por lo que las condiciones de ejecución no podían ser mejores. Esto ha permitido desplegar modelos de órdenes de magnitud superior a lo que se podía ejecutar en el portátil con una latencia mínima, evaluando el sistema en condiciones óptimas.

Para trabajar de forma interactiva sobre este servidor, el acceso se realizó a través de la VPN de Comillas con el software de Palo Alto Networks GlobalProtect, que permitía acceder en la red interna. Como entorno de desarrollo se utilizó Visual Studio Code con la extensión Remote-SSH, lo que permitió programar, usar la terminal, monitorizar procesos y ejecutar código directamente en el servidor (DGX) como si se estuviera trabajando en local. Además, El control de versiones se gestionó de forma continua mediante Git y GitHub.

3.6.2. Software del Sistema

El asistente de voz se ha desarrollado en **Python 3.11+**, utilizando **FastAPI** para exponer la API HTTP y los *endpoints* WebSocket que mantienen la comunicación con el cliente, implementado como un *frontend* ligero en HTML, CSS y JavaScript.

Las decisiones tecnológicas del *backend* están pensadas para poder ejecutar el flujo completo de voz con la menor latencia posible (especialmente optimizando al máximo el tiempo de salida del primer token/chunk de audio). El sistema simplificado funciona como una cascada ASR $\rightarrow$ LLM $\rightarrow$ TTS: primero se transcribe el audio del usuario, después se genera la respuesta con el modelo de lenguaje y finalmente se sintetiza la respuesta en audio. Python se ha elegido por la madurez de su ecosistema en Inteligencia Artificial y por la facilidad para integrar modelos y servidores de inferencia distintos dentro de un mismo sistema.

Sobre esta base, la librería **asyncio** gestiona la concurrencia del servidor. Las operaciones más costosas computacionalmente, como la inferencia del ASR, del LLM y del TTS, no se ejecutan directamente dentro del bucle principal de eventos, sino en servicios independientes o colas de trabajo. Esto es importante porque, si el servidor dejara de leer tramas del cliente mientras espera a que termine una inferencia, la experiencia conversacional sería peor y podrían perderse fragmentos de audio o retrasarse innecesariamente la respuesta.

En la versión final del sistema, el reconocimiento automático del habla se realiza con **nvidia/parakeet-tdt-0.6b-v3**, servido mediante **NVIDIA NeMo**. Este modelo se ha elegido por su buen compromiso entre precisión en español y velocidad de inferencia. Para la generación de texto se utiliza **Gemma 4 31B**, servido mediante **vLLM**. Este motor permite exponer el modelo a través de una API compatible con OpenAI. Como el objetivo del sistema es reducir al máximo el tiempo hasta el primer token generado, la configuración se ajusta para evitar sobrecarga innecesaria: se limita la longitud de contexto al valor requerido

por el caso de uso, se desactivan entradas multimodales no utilizadas y se evita compartir la GPU del LLM con otros modelos. Para la síntesis de voz se utiliza **Qwen/Qwen3-TTS-12Hz-1.7B-CustomVoice**, servido mediante **vLLM-Omni**. En este caso sí es especialmente importante usar salida en *streaming*, porque el TTS es la última etapa antes de que el usuario escuche la respuesta. Por ello, el sistema solicita audio en formato PCM y activa la emisión progresiva de fragmentos de audio cuando sea posible. De esta forma, el cliente puede empezar a reproducir la respuesta sin esperar a que se haya generado todo el audio completo.

Cada uno de los tres modelos principales se despliega en una GPU H200 dedicada: una GPU para el ASR, una GPU para el LLM y una GPU para el TTS. Esta decisión simplifica mucho el despliegue, evita problemas de memoria y reduce la variabilidad de la latencia. Además, como los tres modelos caben en una H200 individual con la configuración elegida, no es necesario repartir un mismo modelo entre varias GPUs, lo que también evita el overhead de la comunicación entre dispositivos.

Para definir las tools, el sistema usa en clases y esquemas de **LangChain**. Aunque en las primeras versiones del proyecto se utilizaba **LangGraph** para controlar el flujo conversacional y el estado del agente, al final no se ha usado por la latencia que añadía. El LLM llama directamente a las tools de forma asíncrona y en *streaming*, eliminando el sobrecoste de procesamiento de librerías para reducir la latencia al mínimo. Toda la arquitectura se ha contenerizado mediante **Docker** y el **NVIDIA Container Toolkit**, lo que facilita su despliegue tanto en local como en el DGX remoto, y utilizando lo mejor posible las GPUs de NVIDIA de ICAI.

3.6.3. Uso de APIs Externas

Durante algunos meses del proyecto, cuando la solución ya estaba desarrollada pero aún no se tenía acceso al clúster de ICAI, fue imposible continuar con el proyecto utilizando exclusivamente los recursos del portátil. Para no detener el avance, se evaluaron APIs de inferencia comerciales ultrarrápidas, destacando Cerebras y Taalas AI.

Es importante dar una mención especial al equipo de Taalas API al cual se les contactó directamente. se trata de una *startup* nueva especializada en la creación de *hardware* dedicado para IA, quienes decidieron colaborar con este TFM proporcionando acceso gratuito a su API cerrada. Gracias a su chip propietario *Taalas HC1*, logran tasas de inferencia de 17.000 tokens por segundo por usuario en el modelo Llama 3.1 8B, un rendimiento muy superior al de cualquier GPU o LPU comercial actual. Aunque esta herramienta fue de gran ayuda temporal, su infraestructura estaba limitada a un único modelo (Llama 3.1 8B). Finalmente, el acceso al servidor DGX de ICAI permitió trabajar sin depender de recursos

externos y ejecutar directamente los mejores modelos *Open Source* disponibles en el hardware descrito en la Sección [3.6.1](#).

Para terminar, otro de los usos más importantes de APIs externas en este proyecto fue la API de ElevenLabs para la generación de los audios que fueron usados para el benchmark final del sistema, y que permitió no requerir cientos de narradores y actores profesionales para la creación de los audios de las diferentes variantes de la llamada, y se usaron más de 50 voces diferentes todas en español pero con diferentes acentos. Fue esencial para este proyecto ya que una manera muy rápida y económica de probar el sistema en condiciones cercanas a la realidad ya que ofrecen un modelo de texto a voz de muy alta calidad, actualmente el mejor con diferencia [\[30\]](#).

Capítulo 4

Implementación del sistema

Este capítulo se centra en describir la solución desarrollada en este trabajo. En primer lugar, en la Sección [4.1](#) se presenta la arquitectura general del sistema y se explica cómo se organizan los distintos módulos que forman el *pipeline* de voz. La Sección [4.2](#) explica cómo se han elegido los componentes y los modelos de Inteligencia Artificial de código abierto utilizados en el sistema. Para ello, se presentan las pruebas realizadas, las métricas medidas y los resultados obtenidos para comparar unas alternativas con otras y así escoger las mejores opciones para cada bloque. Después, en la Sección [4.3](#) se describe la implementación del *pipeline* en *streaming*, detallando cómo se conectan de forma asíncrona los distintos bloques del sistema para mantener un flujo continuo de audio y texto. La Sección [4.4](#) presenta la gestión de turnos de palabra y la máquina de estados utilizada para coordinar el comportamiento del asistente en tiempo real. En la Sección [4.5](#) se explica la lógica de orquestación agéntica, el uso de herramientas externas y el mecanismo utilizado para mantener el contexto de la conversación. Por último, la Sección [4.7](#) describe los mecanismos incorporados para poder monitorizar el sistema y facilitar su análisis.

4.1. Arquitectura del sistema propuesto

La solución implementada sigue una arquitectura modular en cascada, siguiendo lo que ya se había adelantado en capítulos anteriores. En lugar de depender de un único modelo multimodal que procese y genere el audio directamente, el sistema descompone el *pipeline* de voz en bloques especializados:

- Cancelación Acústica de Eco (Acoustic Echo Cancellation o *AEC*),
- Detección de Actividad de Voz (Voice Activity Detection o *VAD*),

- Reconocimiento Automático del Habla (Automatic Speech Recognition o *ASR*),
- Puntuación automática del texto transcrito (*punc*),
- Normalización inversa del texto (Inverse Text Normalization o *ITN*),
- Razonamiento mediante un Modelo de Lenguaje Grande (Large Language Model o *LLM*),
- Formateo orientado a la oralidad (*Speech Formatter*), que adapta por ejemplo fechas, cifras y expresiones para la síntesis,
- Síntesis de Voz (Text-to-Speech o *TTS*).

Cada módulo cumple una función concreta dentro del flujo conversacional y se comunica con el resto a través de interfaces bien definidas, lo que permite inspeccionar el estado intermedio de la conversación en forma de texto y sustituir componentes concretos sin reconstruir el sistema completo, lo que es muy útil especialmente ahora que salen nuevos modelos de IA cada día cada vez más potentes.

Desde el punto de vista de la ingeniería de software, el backend sigue un diseño basado en ports and adapters. La idea es que cada parte del sistema, se define mediante una interfaz común y después para cada implementación concreta se añade como un adaptador que puede cambiarse fácilmente.

La elección del adaptador que se usa en cada bloque se configura mediante perfiles en formato YAML. Gracias a esto, el mismo código puede utilizarse en entornos distintos, combinando inferencia local, servidores propios o APIs externas según las necesidades de cada caso. Esto en parte se hizo por las necesidades que se fueron dando a lo largo del proyecto, cuando en la primera fase se hizo en local, después se usaron APIs externas de proveedores especializados en la rapidez de inferencia como son Taalas [31] y Cerebras [32]. Y además, esta separación era de todos modos especialmente útil en un asistente de voz profesional, ya que no todos los componentes requieren el mismo equilibrio entre latencia, calidad, privacidad y coste computacional.

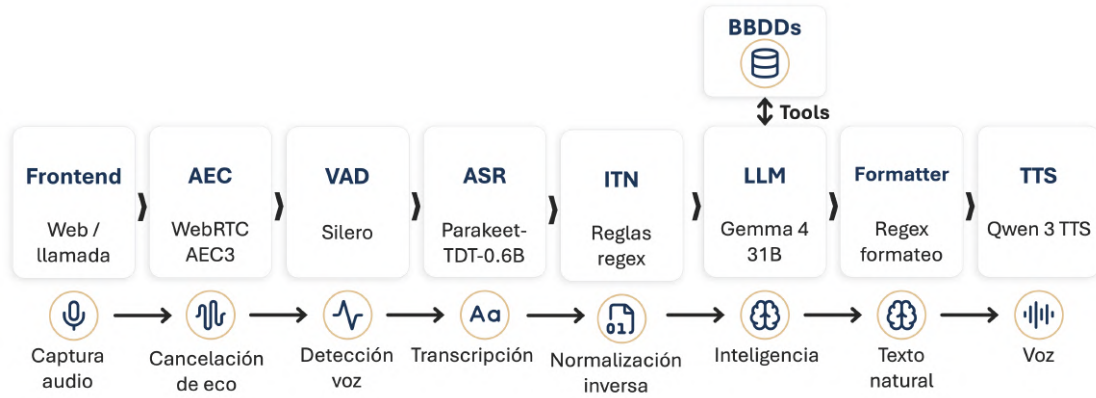


Figura 4.1: Arquitectura simplificada del sistema propuesto.

En la Figura 4.1 se resume (omitiendo ciertos bloques) el sistema global desarrollado. En el extremo del cliente, una aplicación web propia que a partir de ahora se llamará *Frontend*, captura el audio del micrófono y lo transmite al servidor mediante una conexión WebSocket bidireccional expuesta por FastAPI en el endpoint `/ws/audio`. En el *backend*, un módulo de orquestación central coordina el resto de componentes y mantiene el estado de la conversación. Antes de analizar la señal, el audio pasa por un bloque de Cancelación Acústica de Eco (Acoustic Echo Cancellation o *AEC*) para reducir el eco de la propia voz del asistente. A continuación, un detector de actividad de voz (Voice Activity Detection o *VAD*) analiza la señal en fragmentos cortos y determina cuándo el usuario está hablando y cuando no. Cuando se detecta el fin de un turno, el segmento acumulado se envía al módulo de *ASR*, que convierte el audio en texto. Este texto se refina posteriormente con adaptadores de puntuación y normalización inversa (*Inverse Text Normalization* o *ITN*) para facilitar su interpretación por el modelo de lenguaje.

Una vez disponible la transcripción, entra en juego el núcleo de inteligencia del sistema. El módulo *LLMBlock* actúa como orquestador agéntico: recibe el texto del usuario, consulta el contexto conversacional almacenado en memoria de sesión y decide si debe responder directamente, invocar herramientas externas (*tool calling*) o combinar ambas acciones. Las herramientas permiten al asistente interactuar con la lógica de negocio del caso de uso, por ejemplo consultando disponibilidad de mesas, registrando reservas o recuperando información del menú. La respuesta generada no se envía tal cual al sintetizador de voz, sino que pasa antes por un formateador encargado por ejemplo de adaptar fechas, cifras y expresiones para que suenen naturales al ser pronunciadas. Finalmente, el módulo *TTS* convierte el texto validado en audio, que el servidor devuelve al cliente en *streaming* mientras la generación continúa.

Antes de entrar más en detalle en el sistema, conviene añadir alguna precisión

más sobre el flujo global de extremo a extremo. En primer lugar, el cliente establece la conexión WebSocket y comienza a enviar tramas de audio en formato PCM. Cada sesión es única, por lo que no se comparte el contexto con otras sesiones. El servidor procesa cada trama de forma continua mediante *VAD* y una máquina de estados que modela el turno de la conversación, distinguiendo si el sistema está escuchando al usuario, esperando, procesando su petición o reproduciendo la voz de la IA. Cuando el *VAD* identifica un silencio suficiente tras un segmento de habla, se considera finalizado el turno del usuario y se lanza la transcripción. Con el texto ya estructurado, el orquestador invoca al *LLM*, que puede generar una respuesta en *streaming* y, en su caso, ejecutar una o varias herramientas antes de formular la contestación final. Los *tokens* producidos se segmentan progresivamente y se envían al *TTS*, cuyos fragmentos de audio se transmiten al cliente con un ritmo de reproducción en tiempo real. Si el usuario interrumpe al asistente mientras este habla, lo que se conoce como *barge-in*, el sistema cancela la síntesis en curso y vuelve al estado de escucha, pero guardando únicamente en su memoria la parte de la respuesta que el usuario llegó a oír.

El diseño tan modular y separado del sistema, no se ha hecho únicamente para una mejor organización del código, sino que responde a un criterio de diseño explícito del proyecto. Al separar percepción, razonamiento y acción, resulta posible evaluar, depurar y sustituir cada bloque de forma independiente, lo que es especialmente útil en un contexto empresarial donde ante cualquier fallo es necesario encontrar la causa para solucionarla. En el caso del *LLM*, que concentra buena parte de la inteligencia del sistema y condiciona de manera directa la latencia percibida, esta flexibilidad es muy valiosa ya que permite comparar distintos modelos de código abierto sin modificar el resto del *pipeline*. Debido a su importancia en el sistema, en la Sección 4.2 se presenta la selección experimental realizada de candidatos para ese componente, mientras como se verá, para el resto de bloques basta con escogerlos en función de los requisitos funcionales del asistente de voz. Aunque es cierto que una arquitectura en cascada introduce más piezas que un modelo unificado, a cambio ofrece un control mucho mayor sobre el comportamiento del sistema y facilita integrarlo con infraestructuras empresariales reales.

4.2. Selección y justificación de componentes y modelos de Inteligencia Artificial

4.2.1. Requisitos del LLM

Como ya se adelantó previamente, en una arquitectura en cascada para un asistente de voz, el LLM es el núcleo del sistema. La selección de este componente no puede basarse únicamente en métricas de inteligencia general, sino que también

debe cumplir una serie de requisitos funcionales y operativos muy específicos para funcionar correctamente en el sistema.

En primer lugar, es fundamental que el modelo posea una alta capacidad para entender y generar texto en español ya que este trabajo se ha centrado en un asistente de voz en español. Esto no se limita a que el sistema entienda correctamente el idioma, sino que también debe funcionar bien con la forma en que las personas hablan de manera natural y por desgracia, también tiene que funcionar con errores de transcripción. Ya que la entrada al LLM viene del ASR, el texto transcrito a menudo tendrá repeticiones, interrupciones, errores de reconocimiento y faltas de puntuación. El modelo debe ser capaz de interpretar correctamente la intención del usuario a pesar de que la transcripción sea imperfecta. Además, debe seguir de manera estricta las instrucciones definidas en el *system prompt* para adoptar la personalidad adecuada y gestionar el contexto de la llamada sin perder el hilo de la conversación.

Por otro lado, para que el sistema deje de ser un simple asistente y actúe como un agente, el LLM debe soportar el uso de herramientas externas (*tool calling*). Esto significa que el modelo tiene que ser capaz de decidir de forma autónoma cuándo tiene que llamar una función, y que esta pueda consultar o escribir en una base de datos. Para ello, no basta con solo identificar la acción correcta, sino que debe extraer con los parámetros necesarios a partir de la conversación con el usuario y generar una salida con un formato estructurado JSON que la función pueda procesar sin errores. Y también es importante que el modelo tenga la capacidad de abstenerse o pedir aclaraciones si el usuario hace una petición ambigua o faltan datos para ejecutar una herramienta. El modelo no debe inventar la información (*alucinar*) y ser capaz de preguntar al usuario para poder completar la tarea de forma segura.

Finalmente, hay que conseguir un equilibrio entre todos estos requisitos y el rendimiento del modelo. En una interfaz de voz, la latencia es el factor más crítico para mantener la naturalidad de la interacción y evitar silencios incómodos en los que el interlocutor no sabe si se ha cortado la llamada o no. En los LLMs, esto se mide en dos parámetros importantes. El primero es el Tiempo hasta el Primer Token (*Time to First Token* o *TTFT*), que indica el tiempo que el modelo tarda en generar el primer token de salida, lo que permite que el (*TTS*) empiece a hablar lo antes posible ya que se usará una cascada de salida en *streaming* desde el LLM. El segundo parámetro son los tokens que se generan por segundo (*Tokens per Second* o *TPS*), que indica la cantidad de tokens que el modelo puede generar por segundo. Esto es importante porque debe ser superior a la velocidad a la que se genera el audio en el ASR para que una vez que el ASR genere audio, este se pueda sintetizar sin cortes.

Además de la velocidad, también hay que tener en cuenta requisitos prácticos

relacionados con la privacidad y la viabilidad del proyecto. Para mantener el control sobre los datos y sobre la infraestructura, el modelo elegido debe ser de código abierto y poder desplegarse de forma local.

4.2.2. Modelos candidatos

Una vez definidos los requisitos operativos y funcionales del LLM, el siguiente paso fue conseguir un catálogo de modelos candidatos para poder evaluar cuál se adapta mejor a estas necesidades. Para evaluar todos estos modelos, se ha usado un archivo **YAML** en el código de los Benchmarks que contiene información sobre cada modelo para poder ejecutar los benchmarks de inteligencia y velocidad.

Este catálogo incluye una gran variedad de arquitecturas para abarcar distintos compromisos entre inteligencia y coste computacional. Por un lado, se evaluaron modelos densos tradicionales, donde todos los parámetros se activan durante la generación de cada *token*. Por otro lado, se incorporaron arquitecturas de Mezcla de Expertos (Mixture of Experts o *MoE*). Estas arquitecturas dividen la red neuronal en subredes especializadas, activando solo una pequeña fracción del total de parámetros para cada *token*. Esto permite desplegar modelos con una capacidad de razonamiento teórica muy alta pero con una latencia de inferencia y un consumo de memoria comparables a modelos mucho más pequeños.

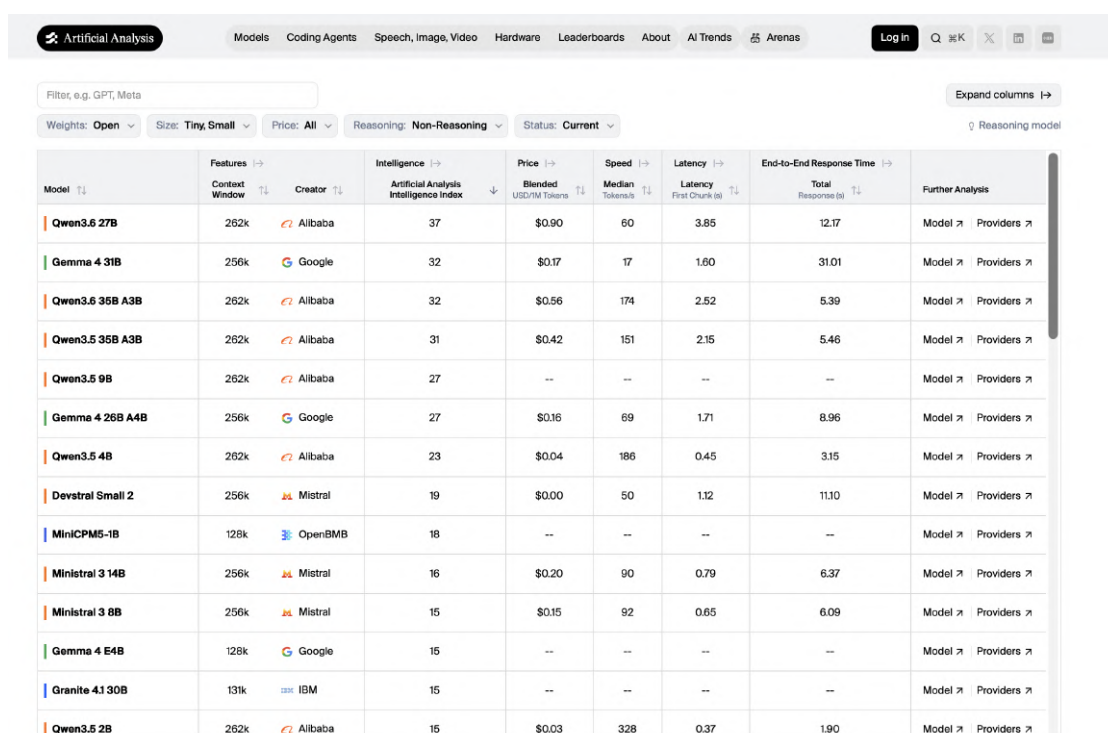
A lo largo de todo el proyecto se han cambiado muchas veces los modelos candidatos debido a que cada modelo nuevo que sale es mejor que el anterior. Entre los candidatos se pueden encontrar las principales familias de código abierto disponibles hasta la fecha del estudio, como Qwen 3.6, Gemma 4, las series Ministral y Devstral de Mistral AI, la familia Nemotron de NVIDIA, Phi-4 de Microsoft y la serie Granite 4.0 de IBM. Para poder realizar las pruebas en los servidores de ICAI, los modelos seleccionados cuentan con pesos abiertos y licencias permisivas (como Apache 2.0 o MIT). En la Tabla 4.1 se muestra una agrupación resumida de los candidatos evaluados. Debido a la extensión del catálogo de modelos, las especificaciones detalladas de los 33 modelos analizados se incluyen en el Anexo B.1. Estas especificaciones recogen, entre otros aspectos, la arquitectura, el número de parámetros activos, el tamaño de la ventana de contexto y la licencia de cada modelo.

4.2. Selección y justificación de componentes y modelos de Inteligencia Artificial

Tabla 4.1: Resumen de las familias de modelos de lenguaje incluidas en la campaña de evaluación.

Proveedor	Familia de Modelos	Arquitecturas	Rango de Tamaños
Qwen	Qwen 3.6 / Omni	Densa, MoE, Multimodal	27B – 35B
Mistral AI	Ministral 3 / Devstral	Densa (con visión)	3B – 24B
Google	Gemma 4	Densa, MoE	2B – 31B
NVIDIA	Nemotron (Nano / v2)	MoE, Híbrida (Mamba)	9B – 30B
Microsoft	Phi-4	Densa, Multimodal	3.8B – 14B
Liquid AI	LFM 2 / LFM 2.5	Densa, MoE Dispersa	1.2B – 24B
IBM	Granite 4.0	Densa, MoE, Híbrida	350M – 32B
LGAI / Meta	EXAONE 4.0	Densa (Atención híbrida)	1.2B – 32B

Para la elección preliminar de estos modelos, se usaron resultados públicos de benchmarks de inteligencia (que se pueden encontrar fácilmente en páginas como <https://artificialanalysis.ai/leaderboards/models>), filtrando por modelos de código abierto y con un tamaño pequeño o medio, y esos resultados se ordenaron de mayor a menor puntuación en el benchmark de inteligencia, para sus modalidades *Non-Reasoning*, ya que al ser un sistema de voz, no se puede la cadena de pensamiento ya que introduce demasiada latencia. La decisión del tamaño de los modelos (hasta 40B de parámetros) se justifica porque, en la inferencia de modelos Transformer decoder-only, la fase de *prefill* procesa el prompt completo y produce el primer token de salida, por lo que está directamente relacionada con el TTFT [33]. Además, el coste del *forward pass* en este tipo de modelos escala aproximadamente con el número de parámetros del modelo, ya que un modelo decoder-only de N parámetros requiere alrededor de $2N$ FLOPs por token [34]. Por ello, reducir el tamaño efectivo del modelo, o el número de parámetros activos durante la inferencia en arquitecturas dispersas, contribuye a reducir la latencia inicial de respuesta.



The screenshot shows the 'Artificial Analysis' website interface. At the top, there are navigation links: Models, Coding Agents, Speech, Image, Video, Hardware, Leaderboards, About, AI Trends, and Arenas. A 'Log In' button is on the right. Below the navigation bar, there are filter controls: a search bar with 'Filter, e.g. GPT, Meta', a 'Weights: Open' dropdown, a 'Size: Tiny, Small' dropdown, a 'Price: All' dropdown, a 'Reasoning: Non-Reasoning' dropdown, and a 'Status: Current' dropdown. An 'Expand columns' link is on the right. The main table displays a list of AI models with the following columns: Model, Features (Context Window, Creator), Intelligence (Artificial Analysis Intelligence Index), Price (Blended USD/1M Tokens), Speed (Median Tokens/s), Latency (First Chunk (s)), End-to-End Response Time (Total Response (s)), and Further Analysis (Model, Providers). The table lists 16 models, including Qwen3.6 27B, Gemma 4 31B, Qwen3.6 35B A3B, Qwen3.5 35B A3B, Qwen3.5 9B, Gemma 4 26B A4B, Qwen3.5 4B, Devstral Small 2, MiniCPM5-1B, Ministral 3 14B, Ministral 3 8B, Gemma 4 E4B, Granite 4.1 30B, and Qwen3.5 2B.

Model	Features Context Window Creator	Intelligence Artificial Analysis Intelligence Index	Price Blended USD/1M Tokens	Speed Median Tokens/s	Latency First Chunk (s)	End-to-End Response Time Total Response (s)	Further Analysis Model Providers
Qwen3.6 27B	262k Alibaba	37	\$0.90	60	3.85	12.17	Model Providers
Gemma 4 31B	256k Google	32	\$0.17	17	1.60	31.01	Model Providers
Qwen3.6 35B A3B	262k Alibaba	32	\$0.56	174	2.52	5.39	Model Providers
Qwen3.5 35B A3B	262k Alibaba	31	\$0.42	151	2.15	5.46	Model Providers
Qwen3.5 9B	262k Alibaba	27	--	--	--	--	Model Providers
Gemma 4 26B A4B	256k Google	27	\$0.16	69	1.71	8.96	Model Providers
Qwen3.5 4B	262k Alibaba	23	\$0.04	186	0.45	3.15	Model Providers
Devstral Small 2	256k Mistral	19	\$0.00	50	1.12	11.10	Model Providers
MiniCPM5-1B	128k OpenBMB	18	--	--	--	--	Model Providers
Ministral 3 14B	256k Mistral	16	\$0.20	90	0.79	6.37	Model Providers
Ministral 3 8B	256k Mistral	15	\$0.15	92	0.65	6.09	Model Providers
Gemma 4 E4B	128k Google	15	--	--	--	--	Model Providers
Granite 4.1 30B	131k IBM	15	--	--	--	--	Model Providers
Qwen3.5 2B	262k Alibaba	15	\$0.03	328	0.37	1.90	Model Providers

Figura 4.2: Captura de <https://artificialanalysis.ai> aplicando los filtros de código abierto, Non-Reasoning y tamaño pequeño o medio para obtener los modelos candidatos.

4.2.3. Resultados del benchmark de inteligencia

Como se describió en la metodología (Sección 3.2), la evaluación de los modelos se llevó a cabo mediante cuatro pruebas específicas: `custom_spanish`, `spanish_robustness`, `bfcl_reference` y `voicebench_reference`. A continuación se presentan los resultados obtenidos tras ejecutar estos *benchmarks* sobre los modelos candidatos. Los datos completos de todas las evaluaciones se pueden consultar en el Anexo B.3.

La Figura 4.3 muestra los resultados de los 4 benchmarks (cada uno en un cuadrante) para cada uno de los modelos.

4.2. Selección y justificación de componentes y modelos de Inteligencia Artificial

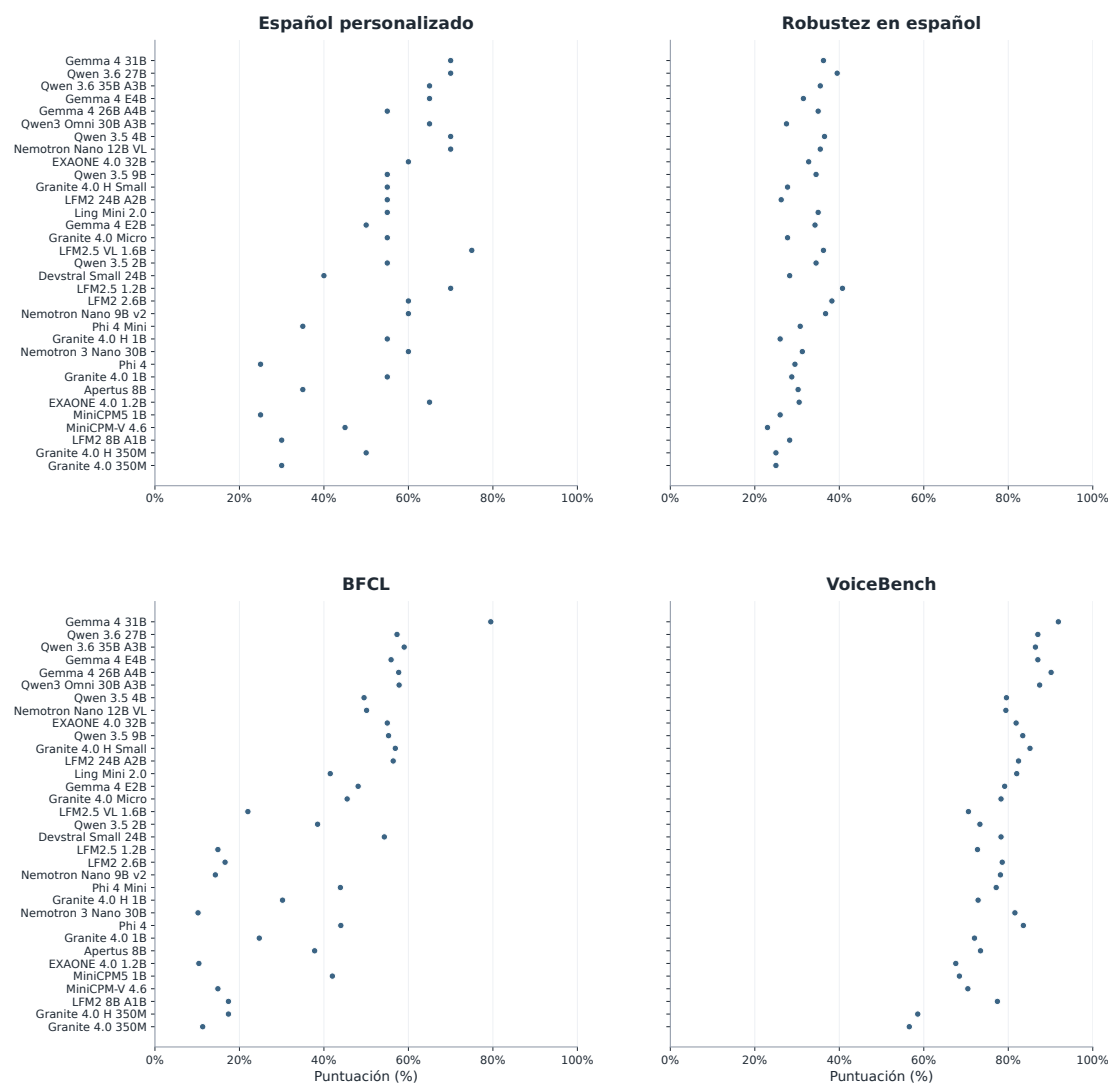


Figura 4.3: Puntuación global de inteligencia de los modelos finalistas.

Se pueden ver diferencias importantes en la capacidad de los modelos en los diferentes benchmarks. Lo primero que se puede ver es que ningún modelo es perfecto ya que ninguno presenta una puntuación excelente en todos los benchmarks. Lo que sí que se puede ver, es que hay una cierta correlación entre todos los resultados ya que los modelos están en el mismo orden y se ve una tendencia en todos los gráficos parecida (hacia arriba a la derecha).

Pero la mejor manera de saber que modelo es globalmente el mejor, es agregar estas puntuaciones en una única métrica. Para ello, se han normalizado todos los resultados para dar una puntuación del 0 (indicando 0% de aciertos) al 1 (indicando 100% de aciertos) para cada modelo para cada test.

En la Figura 4.3 se presenta la puntuación global de los modelos finalistas. Para calcular esta puntuación funcional única por modelo, se ha aplicado fórmula 3.3 para darle igual importancia a cada bloque ya que cada uno cumple una función distinta e importante.

En la figura 4.4, se pueden ver para cada uno de los modelos, las 4 barras que representan la contribución de cada bloque al porcentaje total, indicado a la derecha del final de cada barra.

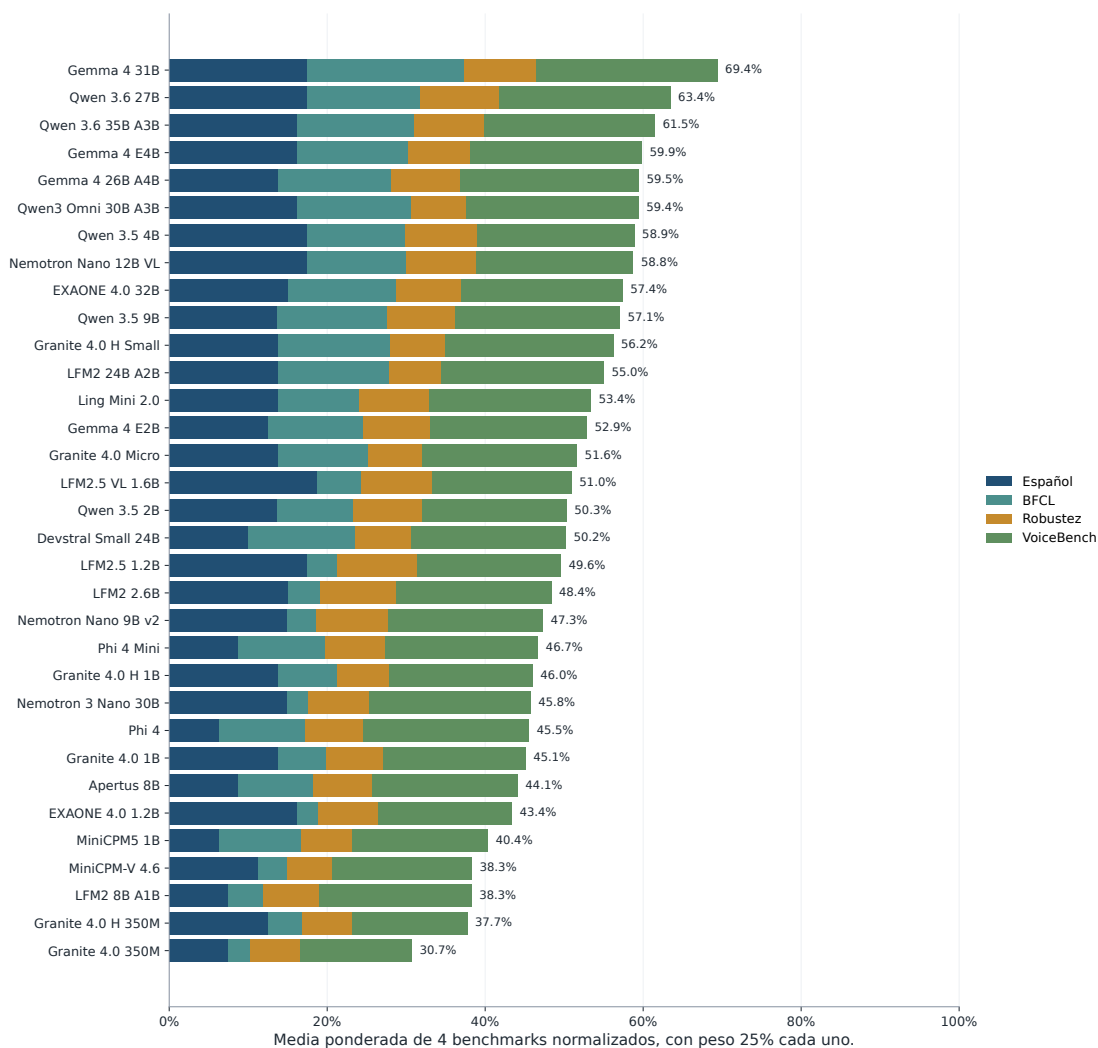


Figura 4.4: Descomposición de la puntuación funcional ponderada por modelo.

En la Figura 4.4 se pueden ver claramente los resultados obtenidos desglosados en cada uno de los benchmarks y se ve claramente que Gemma 4 31B es el modelo que mejor puntuación global tiene con un 69.4%, seguido de cerca por Qwen 3.6

4.2. Selección y justificación de componentes y modelos de Inteligencia Artificial

27B con un 63.4 %. El peor modelo de todos en términos de puntuación es Granite 4.0 350M con una puntuación de 30.7 %. Se puede ver a simple vista, que cuanto mayor es el modelo, por lo general, mejor es su puntuación global.

Una vez que se tiene la puntuación global de cada modelo, se puede pasar a la evaluación de velocidad para comparar el TTFT de cada modelo.

4.2.4. Resultados del benchmark de velocidad

A continuación se puede ver en la Figura 4.5 los resultados obtenidos para cada candidato en los resultados del TTFT con un contexto cardado (warmed up) de 8k tokens:

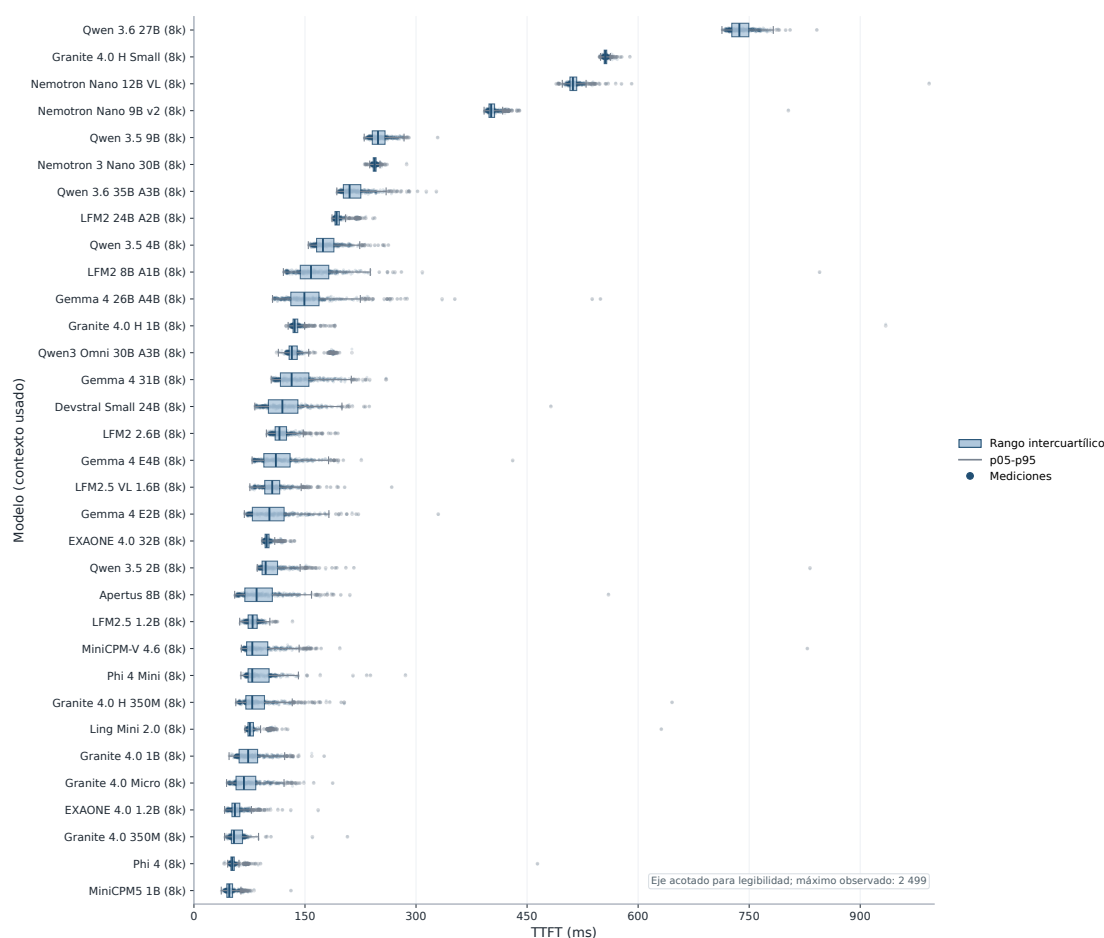


Figura 4.5: Distribución en percentiles del *Time To First Token* (TTFT).

En la Figura 4.5 se muestran los percentiles del *Time To First Token* (TTFT) para cada candidato. Este valor cuanto más bajo es, mejor para el caso de uso de

este trabajo ya que significa que el modelo responde el primer token lo más rápido posible. Los resultados son bastante buenos para casi todos los modelos debido al hardware tan potente del que se dispone. El mejor resultado lo obtiene MiniCPM5, un modelo de 1B de parámetros que tiene un TTFT inferior a los 50ms (p50 48ms). El modelo más lento según estas medidas sería el Qwen 3.6 27B, con una mediana de TTFT de 737 ms y alcanzando los 749 ms en su percentil 75.

Por otro lado, la Figura 4.6 ilustra la velocidad de generación sostenida en tokens por segundo (*TPS*).

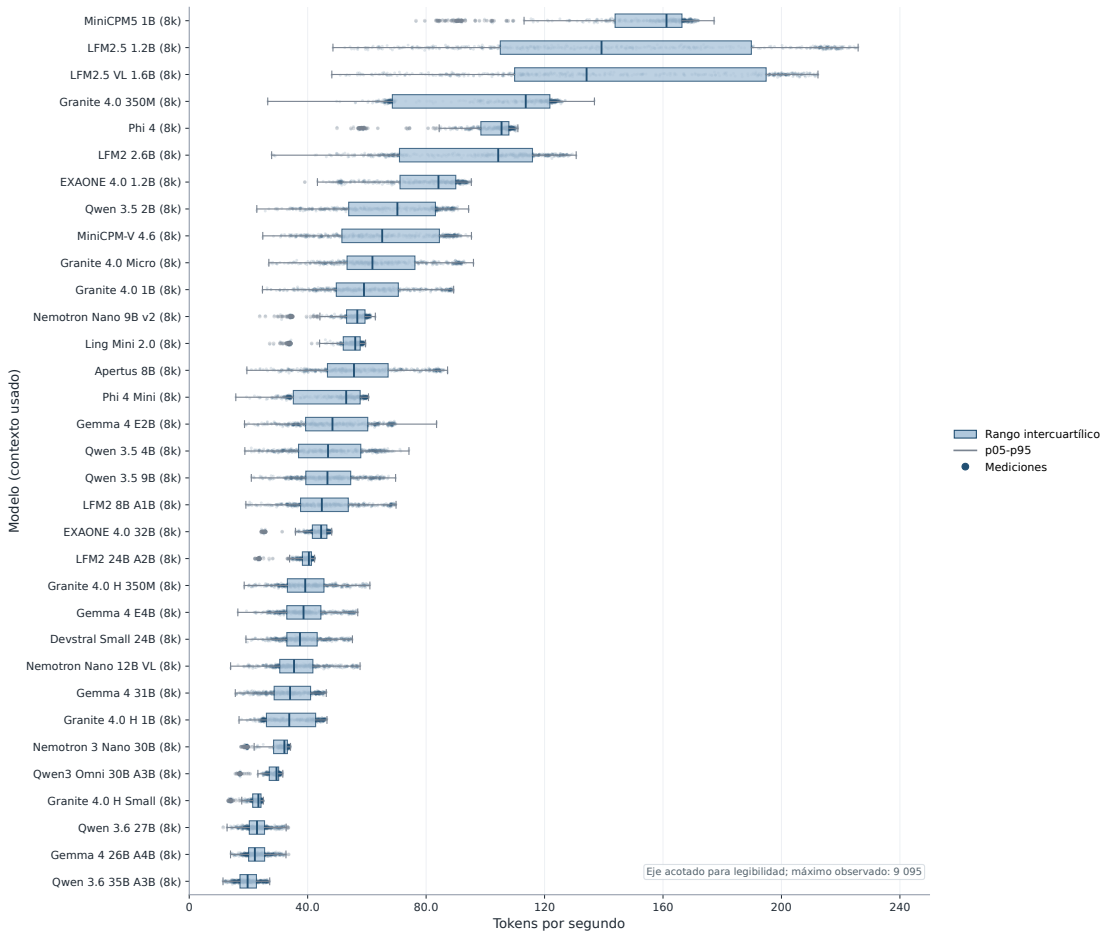


Figura 4.6: Velocidad de generación sostenida en tokens por segundo (TPS).

Como se puede observar en los datos, la velocidad de generación de tokens no supondrá para ningún modelo un factor limitante para el sistema: incluso en el peor de los casos, el modelo más lento es capaz de generar en torno a 20 tokens por segundo (p50 19.7 para Qwen 3.6 35B A3B). Dado que esta velocidad de generación de texto es superior a la velocidad a la que se habla (y por lo tanto a la que

el motor TTS (*Text-to-Speech*) tendría que ser capaz de sintetizar y reproducir el audio hablado), el TPS no representa un cuello de botella (Como regla general, 1 palabra equivale aproximadamente a 1.3 tokens [35], por lo que se estarían generando alrededor de $20/1,3 = 15,3$ palabras por segundo). Por lo tanto, el análisis de rendimiento temporal puede simplificarse para centrarse exclusivamente en el TTFT, que es el verdadero responsable de los silencios iniciales en la conversación y es el único factor limitante. de este trabajo.

En el siguiente apartado se va a intentar encontrar el modelo más adecuado para el asistente de voz, combinando los resultados de los benchmarks de inteligencia y velocidad.

4.2.5. Resumen de los resultados de los benchmarks y selección final del LLM

Por lo tanto, para seleccionar el modelo más óptimo para el caso de uso de este trabajo, se requiere encontrar el equilibrio óptimo entre inteligencia y latencia. Para ello, se ha representado la frontera eficiente en la Figura 4.7, cruzando el TTFT mediano (p50) en el eje X con la puntuación funcional en el eje Y. En este gráfico, el tamaño de cada punto representa el número de parámetros del modelo. Nuestro trabajo es elegir un modelo que se sitúe en esta frontera eficiente, garantizando que se adapta a nuestros requisitos de tiempo real.

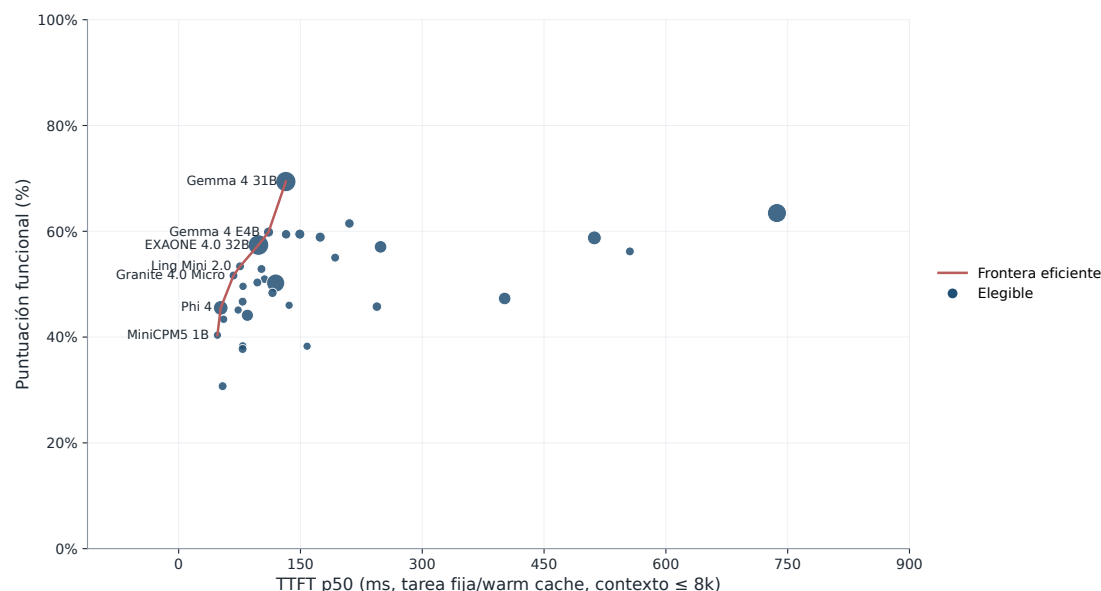


Figura 4.7: Frontera eficiente comparando el TTFT mediano frente a la puntuación funcional. El tamaño del marcador indica el tamaño del modelo.

Y en la Figura 4.8 se pueden ver los mismos resultados pero para el TTFT percentil 95.

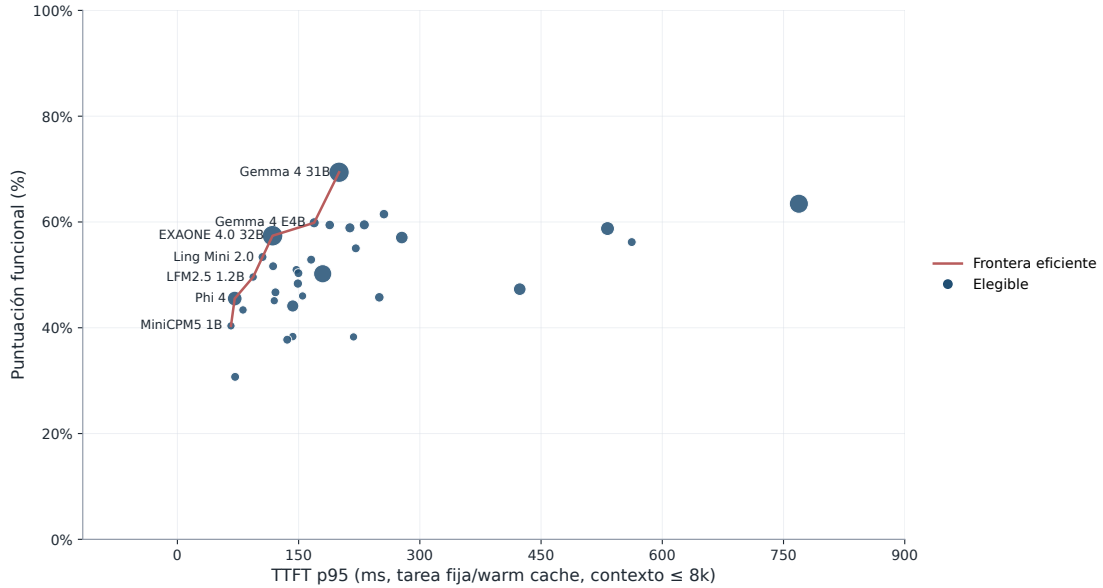


Figura 4.8: Frontera eficiente comparando el TTFT percentil 95 frente a la puntuación funcional. El tamaño del marcador indica el tamaño del modelo.

En la Tabla 4.2 se pueden ver los valores los modelos que forman la frontera eficiente (tanto la de p50 como la de p95), ordenados de menor a mayor TTFT (en el p50).

Tabla 4.2: Modelos pertenecientes a la frontera eficiente.

Modelo	Puntuación Funcional (%)	TTFT p50 (ms)	TTFT p95 (ms)
MiniCPM5 1B	40.4 %	48	66
Phi 4	45.5 %	52	71
Granite 4.0 Micro	51.6 %	68	118
Ling Mini 2.0	53.4 %	76	105
LFM2.5 1.2B	49.6 %	79	94
EXAONE 4.0 32B	57.4 %	98	118
Gemma 4 E4B	59.9 %	111	169
Gemma 4 31B	69.4 %	132	200

Se sabe que el objetivo de latencial total del pipeline sería menos de 400ms. Se puede ver que Gemma 4 31B se sitúa en el extremo superior de la frontera eficiente,

ofreciendo el mejor nivel de inteligencia (69,4%) de todo el catálogo evaluado. Gracias a que en este proyecto se tiene la suerte de disponer de un *hardware* muy potente (probablemente sobredimensionado para un entorno de producción estándar), se tienen muy buenos tiempos de respuesta. Por lo tanto, a primera vista, un modelo con un TTFT de 132ms, como es **Gemma 4 31B**, no parece excesivo para el caso de uso de este trabajo. Como se discutió en la sección [2.3](#), la literatura establece que en una llamada de teléfono, la latencia óptima es cuando la respuesta ronda los 400 ms [\[20\]](#). El problema en un *pipeline* de voz es que la detección del fin de habla (VAD) y la el tiempo hasta que la transcripción (ASR) está estable consumen gran parte de este presupuesto temporal. Por ello es necesario seleccionar un LLM con una latencia inicial baja: al tener un TTFT de solo 132 ms, se evitar sumar un cuello de botella más al sistema, acercándonos todo lo posible a esos 400 ms. Es por ello que se ha seleccionado **Gemma 4 31B** como el LLM. Pero gracias a la modularidad del sistema, se podría utilizar otro modelo de la frontera como **Gemma 4 E4B**, lo que reduciría el TTFT manteniendo una inteligencia aceptable.

4.2.6. Configuración final de inferencia

Una vez seleccionado Gemma 4 31B el LLM que se va a usar, hay que elegir los parámetros óptimos para la inferencia en el asistente de voz en tiempo real. El modelo se ha servido mediante vLLM. La comunicación con el modelo se realizó a través de una API compatible con OpenAI internamente en el puerto 8001.

El servidor de inferencia se ejecutó en Docker con una configuración lo más optimizada posible a la baja latencia del TTFT. La ventana máxima de contexto se fijó en 8192 tokens mediante `-max-model-len 8192`, suficiente para mantener el historial reciente de la conversación sin tener un coste de prefills excesivamente largos. En el peor de los casos (con todo añadido en el prompt), el prompt completo del planificador junto con un primer mensaje de usuario ocupa aproximadamente 4068 tokens, por lo que todavía queda margen para el historial reciente, los resultados de herramientas y las respuestas generadas. La memoria de GPU reservada para vLLM se configuró con `-gpu-memory-utilization 0.90`. Se ha activado también la caché KV en formato FP8 mediante `-kv-cache-dtype fp8`, el *prefix caching* con `-enable-prefix-caching`, la planificación asíncrona con `-async-scheduling` y un presupuesto de batching de `-max-num-batched-tokens 16384`. Para limitar la carga concurrente y priorizar la latencia por turno, se configuró `-max-num-seqs 16`. También se deshabilitaron entradas multimodales usando el flag `-limit-mm-per-prompt '{image':0,.audio':0}'`, ya que el LLM recibe únicamente texto procedente del ASR/ITN.

En cuanto a los parámetros de generación, se configuró una temperatura de 0.1, `top_p=0.85` y `top_k=32`. Esta configuración mantiene una generación estable

y poco aleatoria, que es adecuada para un asistente de voz orientado a tareas, pero conserva suficiente flexibilidad para producir respuestas naturales. La salida máxima se limitó a 96 tokens por turno, con el fin de evitar respuestas demasiado largas y 'artificiales'. Además, se desactivó el modo de razonamiento interno con `enable_thinking=false`, de forma que el modelo genere únicamente contenido destinado al flujo conversacional del sistema.

La integración con el pipeline se realizó en modo *streaming*. El cliente envía las peticiones con `stream=true` y el servidor vLLM se configuró con `-stream-interval 1`, de forma que los tokens generados se entregan incrementalmente al sistema. En el camino crítico del sistema se ejecuta directamente el LLM, sin pasar por LangGraph, para evitar añadir latencia innecesaria. Las respuestas normales del asistente se generan como texto plano (sin estructura JSON como las tools) y se envían al TTS en cuanto existen suficientes tokens para ser leídos por el TTS, sin esperar a que termine toda la respuesta.

Para el uso de herramientas no se empleó un parser nativo de *tool calling* del motor de inferencia. En su lugar, el sistema mantuvo un protocolo interno basado en JSON para las acciones, por ejemplo `{.action:"tool", ...}`, que permite detectar de forma temprana cuándo el modelo pretende invocar una herramienta. Cuando se detecta que se está empezando a llamar a una herramienta (inspeccionando la salida del LLM), el sistema emite un evento interno de `tool_indicator` y reproduce un audio de relleno (un *filler*) sin necesidad del TTS ya que se ha generado previamente. Después, la herramienta se ejecuta mediante el registro interno `tools_by_name`, con un máximo de 4 ciclos de herramienta y un *timeout* de 10 segundos por llamada.

4.2.7. Selección de VAD, ASR y TTS

A diferencia de lo que ocurría con el LLM, en la selección del VAD, el ASR y el TTS no tenía sentido ni era viable, plantear una campaña de *benchmarks* propios como la descrita en los apartados anteriores para el LLM. En primer lugar, construir un conjunto de datos para evaluar el reconocimiento de voz (*ASR*) o la detección de actividad de voz (*VAD*) exigiría grabar a decenas de personas reales, con distintos acentos, tonos de voz y condiciones de ruido de fondo. Con los recursos de este proyecto, cualquier *benchmark* que se elaborara acabaría siendo más pobre y menos representativo que los que ya existen en la industria, por lo que sus conclusiones tendrían un valor muy limitado.

A esto se añade que las métricas del resto de los modelos que hay que elegir se alinean perfectamente con nuestro caso de uso como por ejemplo el *Word Error Rate* (WER) para medir la calidad de la transcripción. Intentar replicarla cualquiera de estas métricas a nuestra escala no aportaría valor adicional al trabajo, ya que difícilmente se descubriría algo que la literatura no haya validado ya de

forma mucho más significativamente. Por otro lado, conviene tener en cuenta el grado de madurez que han alcanzado estos componentes. Los modelos de VAD, ASR y TTS de código abierto funcionan hoy en día con una fiabilidad muy alta.

Finalmente, no hay que perder de vista que el LLM es el verdadero núcleo del sistema ya que es el componente que realmente marca la diferencia en la inteligencia del asistente (interpretar un texto con ruido de transcripción, decidir cuándo usar una herramienta o mantener el hilo de la conversación) y el que introduce la mayor latencia. Por este motivo tenía sentido aislar y centrar el esfuerzo experimental en el LLM, más que en el resto de componentes.

Por lo tanto, se ha optado por apoyarse en la literatura y en el consenso de la comunidad *Open Source* para seleccionar las mejores herramientas para cada bloque y para este caso de uso en específico. Conviene aclarar que, al igual que ocurrió con el LLM, a lo largo de este TFM los modelos utilizados se han ido sustituyendo a medida que la tecnología avanzaba y aparecían opciones mejores o más eficientes. Por ello, los modelos que se describen a continuación son los que se estaban empleando en el momento de redactar esta memoria, y no necesariamente los que se usaron en las primeras versiones del sistema. A continuación se detalla qué modelos concretos se han utilizado para el VAD, el ASR y el TTS, así como las razones que justifican cada una de estas elecciones.

Detección de actividad de voz (VAD)

Para la Detección de Actividad de Voz (*VAD*), el componente encargado de decidir en tiempo real qué fragmentos de audio contienen voz y cuáles no, se ha optado por Silero VAD [36]. Esta elección se justifica tanto porque es considerado como estándar de facto dentro de la comunidad *Open Source* desde su salida, pero también por la evidencia disponible en la literatura. En este sentido, resulta especialmente útil el trabajo de Google LLC [37], que compara de forma sistemática tres detectores habituales: un baseline energético basado en el valor cuadrático medio (*Root Mean Square* o *RMS*), el clásico WebRTC y el propio Silero. Sus resultados muestran que Silero, al estar basado en una red neuronal, supera de forma clara y significativa tanto a WebRTC como al baseline de energía, alcanzando el mejor coeficiente de correlación de Matthews (*Matthews Correlation Coefficient* o *MCC*) del estudio. Apoyarse en este tipo de evaluaciones fiables, realizadas con conjuntos de datos amplios y métricas estandarizadas, tiene mucho más sentido que intentar reproducir una comparativa propia que difícilmente igualaría su rigor.

Además de su buen rendimiento, Silero ofrece una serie de ventajas prácticas que encajan bien con los requisitos del *pipeline*. Se trata de un modelo ligero, con licencia permisiva y pensado para funcionar en *streaming*, lo que permite evaluar la señal de forma continua sin introducir una latencia apreciable. En la implementación realizada, el audio entrante se convierte siempre a mono y se remuestrea

a 16 kHz, y el modelo analiza ventanas de 32 milisegundos (512 muestras) emitiendo para cada una la probabilidad de que contenga voz. A partir de un umbral configurable se decide si el fragmento se considera habla o silencio. De esta forma, el VAD se comporta como un primer filtro fiable que evita que se procese audio innecesario en el ASR.

Reconocimiento automático del habla (ASR)

Para el Reconocimiento Automático del Habla (*Automatic Speech Recognition*, ASR), se ha seleccionado finalmente el modelo `nvidia/parakeet-tdt-0.6b-v3` [38, 39]. Esta decisión fue un cambio importante frente al modelo seleccionado al principio que fue `CohereLabs/cohere-transcribe-03-2026` [40], ya que el análisis posterior experimental mostró que el criterio de selección no debía basarse únicamente en la tasa de error por palabra (*Word Error Rate*, WER) y en el factor de velocidad de inferencia inverso (*Real-Time Factor inverse*, RTFx, que se define como el cociente entre la duración del audio procesado y el tiempo empleado por el sistema ASR para generar la transcripción), sino también en la latencia (es decir el equivalente al TTFT del LLM). Por lo tanto en cuanto se vió esto, se descarto el uso de este modelo ya que tenía latencias muy elevadas. El RTFx mide la velocidad de procesamiento, pero no mide el tiempo hasta obtener una transcripción utilizable ni la latencia percibida por el usuario en un sistema ASR-LLM-TTS [41].

El modelo Parakeet-TDT-0.6B-v3 es un modelo multilingüe de 600 millones de parámetros basado en una arquitectura FastConformer con decodificador Token-and-Duration Transducer (TDT). Soporta 25 lenguas europeas, incluido el español, e incorpora puntuación, capitalización y marcas temporales a nivel de palabra y segmento [38]. Además, puede ejecutarse mediante NVIDIA NeMo, lo que facilita su integración en el DGX.

Tabla 4.3: Comparativa de los modelos ASR candidatos para español.

Modelo	WER ES (%) ↓	Tamaño
NVIDIA Parakeet-TDT-0.6B-v3	3,45–4,39	0,6B
Qwen3-ASR-0.6B	4,94	0,6B
Qwen3-ASR-1.7B	3,36	1,7B
Voxtral Mini 4B Realtime	3,31–5,34	4B

En español, Parakeet-TDT-0.6B-v3 obtiene resultados muy buenos en los principales conjuntos multilingües. La tarjeta oficial del modelo publica una WER de 3,45 % en FLEURS, 4,39 % en MLS y 3,41 % en CoVoST2 para español [38]. El

4.2. Selección y justificación de componentes y modelos de Inteligencia Artificial

informe técnico de NVIDIA presenta valores ligeramente distintos debido a diferencias en la configuración experimental, con 4,39 % en FLEURS y 3,41 % en MLS para español [39].

Aunque Qwen3-ASR-1.7B consiga un WER algo menor en español, este consume mucha más memoria. Por otro lado, si lo comparamos con el modelo de su mismo tamaño (Qwen3-ASR-0.6B, que tiene un tiempo medio hasta el primer token de 92 ms [42]), el modelo de Nvidia Parakeet-TDT-0.6B-v3 ofrece latencias muy similares pero con una tasa de error (WER) en español inferiores (3,45 % frente al 4,94 % de Qwen del mismo tamaño, muy cerca del WER de 3,36 % del modelo más grande). Voxtral Mini 4B Realtime también fue una alternativa interesante analizada para transcripción en streaming, con una latencia configurable, pero después de realizar varias pruebas, la calidad de la transcripción parecía peor que el modelo de Nvidia.

Por tanto, se selecciona `nvidia/parakeet-tdt-0.6b-v3` como el mejor compromiso entre precisión en español, velocidad de inferencia, facilidad de despliegue local y bajo coste computacional.

Síntesis de voz (TTS)

Para la Síntesis de Voz (*Text-to-Speech* o *TTS*), se ha optado por el modelo Qwen3-TTS-12Hz-1.7B [43]. La elección se ha apoyado en los resultados de MINT-Bench [44], un benchmark multilingüe para TTS que evalúa, entre otros idiomas, el español (imprescindible para nuestro caso de uso). En esta tabla, IF hace referencia a *Instruction Following* y PQ a *Perceptual Quality*.

Tabla 4.4: Comparativa de modelos TTS de código abierto para español basados en MINT-Bench [44].

Modelo	WER ES (%) ↓	Overall IF/PQ ↑	Timbre IF/PQ ↑	Style IF/PQ ↑	Observación
Qwen3TTS-12Hz-1.7B-VD	2,5	2,29 / 3,29	2,45 / 3,40	2,30 / 3,38	Mejor modelo abierto en español
Ming-omni-tts-16.8B-A3B	5,5	1,28 / 1,30	1,62 / 1,65	1,28 / 1,32	Mayor tasa de error y menor calidad perceptual
MOSS-VoiceGenerator	15,5	1,09 / 1,20	1,07 / 1,11	1,00 / 1,09	Resultado claramente inferior
Ming-omni-tts-0.5B	39,7	0,30 / 0,31	0,28 / 0,28	0,11 / 0,11	No competitivo para español
MiMo-Audio-7B-Instruct	44,2	0,34 / 0,35	0,27 / 0,27	0,14 / 0,14	No competitivo para español

Como se observa en la Tabla 4.4, Qwen3-TTS obtiene el mejor resultado entre los modelos abiertos evaluados para español, con un WER del 2,5 %, y también las puntuaciones más altas en seguimiento de instrucciones (poco relevante para este trabajo) y calidad perceptual. Además, la familia Qwen3-TTS está diseñada para generación en *streaming*, requisito de este trabajo. Por estos motivos, se ha utilizado Qwen3-TTS-12Hz-1.7B como modelo de síntesis de voz.

Finalmente, para evitar entrenar una voz en español desde cero, se ha utilizado una voz creada por la comunidad específicamente para Qwen3-TTS. En concreto, se ha empleado el repositorio `qwen3-tts-spanish-voices` [45], que tiene varias

voces españolas, incluyendo la variante de español con acento de España, lo que permite obtener una voz más adecuada sin necesidad de entrenar una y de buscar un dataset para ello.

Una vez seleccionados y justificados todos los componentes, desde el VAD, el ASR y el TTS hasta el LLM, el siguiente paso es explicar cómo se integran todos ellos en el sistema desarrollado. Por ello, en la siguiente sección se describe la implementación del *pipeline* en *streaming*, detallando cómo se conectan estos bloques de forma asíncrona para mantener un flujo continuo de audio y texto en tiempo real.

4.3. Implementación del Pipeline en streaming

El pipeline de *streaming* de audio es el encargado de gestionar la comunicación en tiempo real entre el usuario y la IA. Para lograr una latencia mínima y evitar interrupciones, la implementación se ha desarrollado utilizando el *framework* FastAPI y WebSockets. El proceso comienza cuando el cliente establece la conexión y envía un mensaje inicial de configuración, donde especifica los parámetros del audio, como la tasa de muestreo y la resolución. Una vez que el servidor valida esta información, el cliente empieza a enviar tramas binarias de audio de forma continua. Al utilizar WebSockets, se mantiene una conexión bidireccional abierta (*full-duplex*), lo que permite recibir el audio del usuario al mismo tiempo que se envían métricas de telemetría y los fragmentos de voz sintetizados de vuelta al cliente. La telemetría en un escenario en producción no se enviaría, en este caso se hace para ver métricas de funcionamiento en el frontend.

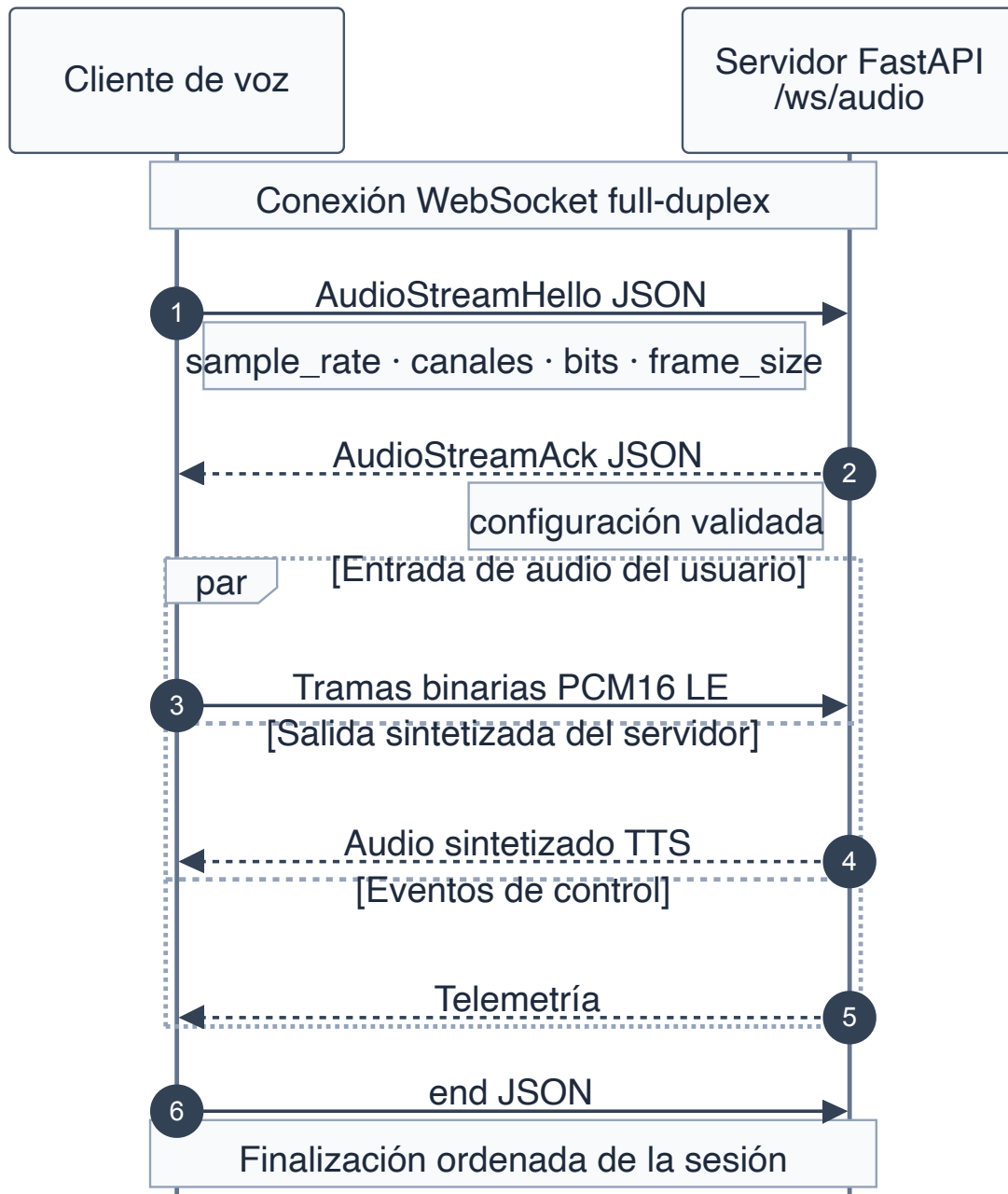


Figura 4.9: Protocolo de comunicación bidireccional mediante WebSocket entre el cliente y el servidor FastAPI.

En primer lugar, cuando la señal de audio cruda llega al servidor, pasa por un módulo de Cancelación Acústica de Eco (Acoustic Echo Cancellation o *AEC*). Este paso es fundamental para limpiar la señal y eliminar el posible eco de la voz

de la IA que se haya podido colar por el micrófono del usuario. A continuación, la señal limpia que entran, que el cliente envía en tramas de 20 milisegundos, se procesa de forma a medida que va llegando en el servidor. El módulo AEC de WebRTC trabaja con bloques de 10 milisegundos y, después de ese bloque, se pasa el audio al VAD también en tramas de 10 milisegundos. Pero, el VAD no emite una predicción de si se está hablando sobre cada fragmento aislado de 10ms, sino que va acumulando el audio hasta formar ventanas completas de 32 milisegundos (lo que equivale a 512 muestras a 16 kHz) para hacer la inferencia. Este componente analiza cada fragmento para determinar si contiene voz humana. Para evitar que se corten los inicios de las palabras al detectar el comienzo de una frase (por ejemplo si el inicio de la frase está en el chunk anterior pero el VAD solo lo detecta en el chunk siguiente), se utiliza un búfer que mantiene almacenados los últimos milisegundos de audio, garantizando así que la señal capturada esté completa y se le pueda mandar completa al siguiente bloque.

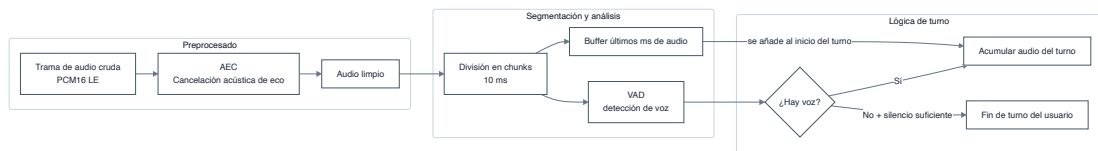


Figura 4.10: Procesamiento inicial de la señal de audio: cancelación de eco, división en fragmentos de 10 ms, búfer de *pre-roll* y detección de actividad de voz.

Una vez que el VAD determina que el usuario ha terminado de hablar, los fragmentos de audio acumulados se envían al modelo de Reconocimiento Automático del Habla (Automatic Speech Recognition o *ASR*). Dependiendo de las capacidades del modelo ASR utilizado, el sistema puede operar de dos formas: enviando el segmento de audio completo una vez finalizado el turno, o inyectando el audio de manera constante para obtener transcripciones parciales rápidamente (en streaming). El modelo tenía que poder ser usado en streaming ya que para minimizar el TTFT después en el LLM, se va haciendo un warm cache de los fragmentos de audio conforme van llegando, así en cuanto llega el último ya está todo cargado. Pero antes de llegar al LLM, después de obtener el texto en bruto del ASR, este pasa por adaptadores de puntuación y normalización inversa de texto (*Inverse Text Normalization* o *ITN*), también hechos en streaming. Estos módulos se encargan de ajustar la transcripción, añadiendo mayúsculas, signos de puntuación y transformando palabras numéricas a dígitos, lo que facilita al LLM la comprensión mejorando su rendimiento.

4.3. Implementación del Pipeline en streaming



Figura 4.11: Conversión del audio del usuario en texto listo para el LLM.

Con el texto ya limpio, este se inyecta en el modelo de lenguaje (*LLM*) seleccionado. En este punto, el modelo actúa como el cerebro del sistema, procesando la intención del usuario y generando una respuesta. Como se ha comentado en secciones anteriores, la elección de este componente es crítica, ya que de él depende tanto la inteligencia de la respuesta como la velocidad a la que se empiezan a generar los primeros *tokens*. Es por ello que como se ha explicado antes, según van llegando los tokens de texto se hace un warm cache para minimizar el TTFT al máximo.

Para mantener la fluidez de la conversación, la respuesta no se espera de forma completa, sino que se procesa como un flujo continuo (*stream*). A medida que el LLM genera los *tokens* de salida, el *Speech Formatter* los valida antes de enviarlos al motor de Síntesis de Voz (Text to Speech o *TTS*). El TTS también funciona en *streaming* y comienza a devolver fragmentos de audio mientras que aún sigue sintetizando audio, y estos fragmentos se envían inmediatamente al cliente a través del WebSocket. Para asegurar que la voz suene natural y no se reproduzca demasiado rápido o a tirones, el sistema implementa un mecanismo de *real-time pacing*, que regula el envío de los fragmentos de audio para que coincidan con la velocidad real de reproducción.



Figura 4.12: Generación de la respuesta en streaming y generación de voz con el TTS.

Para que todos estos bloques funcionen de manera coordinada sin bloquearse entre ellos, se usa de la programación asíncrona mediante la librería `asyncio` de Python. Las tareas más pesadas computacionalmente (como la inferencia del ASR, del LLM...), se ejecutan en hilos separados o se gestionan a través de colas de eventos. De esta forma, el bucle de eventos principal, que es el encargado de atender las comunicaciones del WebSocket, nunca se bloquea. Esto garantiza que el servidor siempre esté listo para recibir el audio del cliente, evitando cortes o pérdidas de información que arruinarían la experiencia de la llamada.

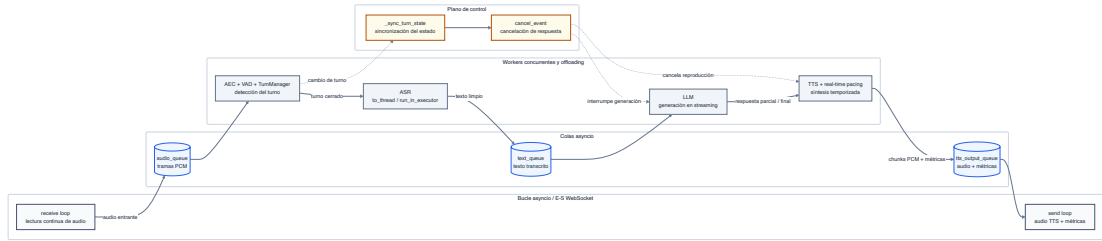


Figura 4.13: Coordinación asincrónica del pipeline.

En la siguiente figura 4.14 se muestra una vista resumida del pipeline completo en la que se pueden ver los principales bloques.

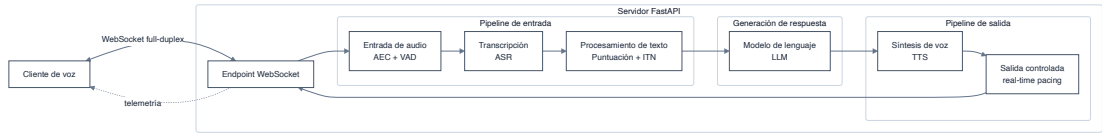


Figura 4.14: Vista resumida del pipeline completo de audio en streaming, desde la entrada de voz del usuario hasta la respuesta sintetizada de la IA.

4.4. Gestión de turnos y Máquina de Estados (FSM)

En la Sección 4.3 se ha descrito cómo circulan el audio y el texto entre los distintos bloques del sistema. Pero para construir un asistente de voz en tiempo real no basta con procesar correctamente cada módulo por separado, sino que también es necesario decidir en cada instante quién tiene la palabra, cuándo se considera que el usuario ha terminado de hablar y qué debe ocurrir si este interrumpe mientras la IA está respondiendo. Coordinar todo es clave para evitar interrupciones artificiales, latencias excesivas y una interacción poco fluida [11]. Por ello, en este trabajo se utiliza una Máquina de Estados Finitos (*Finite State Machine* o *FSM*) para controlar el estado de la conversación.

El objetivo de esta FSM no es aumentar la capacidad de razonamiento del modelo de lenguaje, sino es controlar el comportamiento temporal del sistema. En una llamada, el *backend* debe saber si está esperando al usuario, registrando audio, procesando una petición o hablando. En cada fase el sistema tiene que hacer acciones distintas, como acumular fragmentos de audio, enviar el audio al ASR, iniciar la generación de respuesta con el LLM, reproducir el TTS o cancelar una salida que ha dejado de ser válida. Esta lógica se concentra en la clase `ConversationStateMachine`, que mantiene el estado global de la conversación mediante el enum `ConversationState`.

Los estados de la máquina son:

- WELCOME

- IDLE

- USER_SPEAKING

- PROCESSING

- AI_SPEAKING

El estado `WELCOME` es una fase inicial (opcional) en la que el sistema reproduce un saludo precargado, simulando el momento en el que se atiende una llamada. A continuación, el sistema pasa a `IDLE`, donde permanece a la espera de un turno válido del usuario. Cuando se confirma el inicio de voz, la máquina entra en `USER_SPEAKING`, estado en el que se acumula el audio del usuario para después transcribirlo. Una vez detectado el final del turno del usuario, se pasa a `PROCESSING`, donde se ejecuta la cadena formada por el ASR, la normalización del texto y el LLM. Finalmente, cuando la respuesta empieza a reproducirse, la máquina entra en `AI_SPEAKING`, estado durante el cual se sintetiza y envía la voz de la IA al cliente. La mejor manera de entender la FSM es a través de la siguiente Figura 4.15 que representa la máquina de estados así como las variables que tienen que cumplirse para el cambio de un estado a otro.

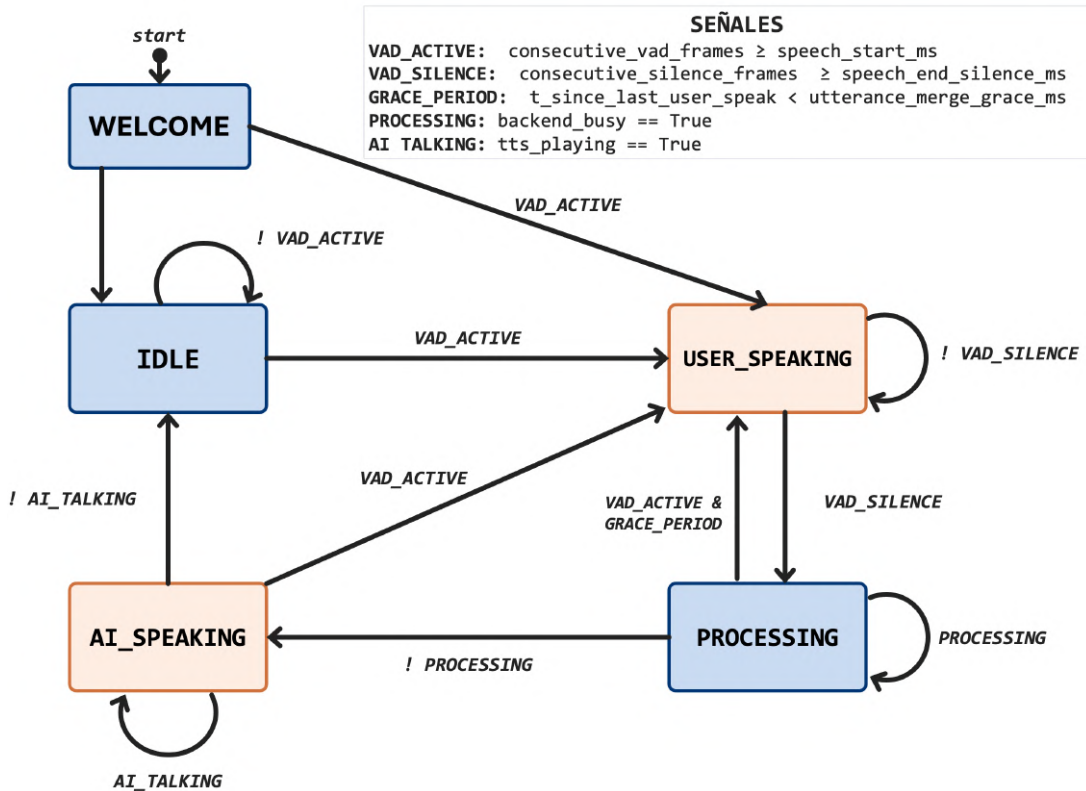


Figura 4.15: Máquina de Estados Finitos (FSM) para la gestión de turnos de palabra.

En la Figura 4.15 se resume el comportamiento global de la FSM y las condiciones que gobiernan cada transición entre estados.

Para que las transiciones no dependan de detecciones aisladas del VAD, la FSM utiliza una variable interna denominada `consecutive_vad_frames`. Esta variable cuenta cuántos frames consecutivos cumplen la condición relevante en cada estado. Esto se hace para evitar cambios inestables cuando hay un mínimo silencio en la conversación o un mínimo ruido, así se está mas seguro de cuando para y comienza a hablar una persona. Y como se ha explicado antes hay un pequeño buffer de audio así que ese delay que se tiene al introducir estos estados consecutivos de hablando, se pueden recuperar sin problema y sin perder ninguna información. En `IDLE`, cuenta frames consecutivos con voz y si se alcanza el umbral `speech_start_ms` (que internamente se ha convertido a número de frames), se considera que el usuario ha empezado a hablar y se transita a `USER_SPEAKING`. En cambio, en `USER_SPEAKING`, el mismo contador acumula frames consecutivos de silencio y si se alcanza `speech_end_silence_ms`, se considera que el turno

ha terminado y se pasa a **PROCESSING**. De esta forma, un ruido breve no activa el sistema por error y una pausa natural dentro de una frase no corta antes del tiempo al usuario.

Los parámetros `speech_start_ms` y `speech_end_silence_ms` se ha decidido usar variables con valores distintos ya representan decisiones diferentes. El primero debe ser relativamente bajo, ya que interesa detectar rápidamente que el usuario ha empezado a hablar. El segundo, en cambio, suele ser mayor, porque el final del turno debe ser más conservador para no interrumpir frases con pausas naturales. Como se ha dicho en el párrafo anterior, aunque ambos se expresan en milisegundos y se traducen internamente a frames según la duración de cada fragmento de audio, pero para el desarrollo es más fácil pensar en ms que en frames.

Además de detectar el inicio y el final de cada intervención, la FSM gestiona las interrupciones (lo que se conoce como el *barge-in*): el usuario puede interrumpir a la IA durante **PROCESSING** o **AI_SPEAKING**, comportamiento habitual en una llamada telefónica. Si el VAD detecta voz durante al menos `speech_start_ms`, el sistema transita a **USER_SPEAKING**; si la interrupción ocurre durante la reproducción, se cancela el TTS en curso y se descartan los fragmentos pendientes de reproducir.

Para reducir cortes prematuros del turno (p. ej. al dictar un teléfono o tras una pausa breve), se aumentó `speech\_end\_silence\_ms` y se añadió un periodo de gracia llamado *grace period* (`utterance\_merge\_grace\_ms`): tras el primer silencio la FSM entra en **PROCESSING**, pero la transcripción no se cierra de inmediato. Si el usuario vuelve a hablar antes de que expire esa ventana temporal, se anula la finalización pendiente y se mantiene el mismo turno y los segmentos de audio se concatenan en un único texto (`segmento\_1(pausa)segmento\_2`) antes de invocar al agente. Así se evita perder audio por pausas pequeñas manteniendo el *barge-in* cuando la interrupción es real.

El caso de **PROCESSING** requiere una lógica un poco más compleja, especialmente cuando el modelo ha iniciado una llamada a herramientas externas mediante *tool calling*. La interrupción del usuario puede hacer que la respuesta hablada deje de ser relevante, pero claro esto no implica que las operaciones ya iniciadas deban o puedan cancelarse. Una herramienta puede haber producido efectos persistentes como crear una entrada o modificar algo en la base de datos o ejecutar una acción de negocio. Por este motivo, cuando una herramienta ya ha empezado a ejecutarse, se deja terminar aunque el usuario interrumpa. Una vez finalizada, su resultado se incorpora al contexto de la conversación, de forma que el modelo pueda tener en cuenta lo que ha ocurrido en el siguiente turno y así se le podría informar al usuario por si decidiera realizar la acción.

El sistema distingue dos cosas al ser interrumpido por usuario. Lo que aún no se ha dicho en voz alta (TTS y texto pendiente) se descarta o se quita del historial. Lo que ya se ha hecho con herramientas (por ejemplo, comprobar disponibilidad)

se mantiene en el historial, porque ya ha cambiado el estado real del sistema. Si el usuario solo hace una pausa breve y sigue hablando, puede cancelarse la transcripción pendiente, pero no una herramienta ya ejecutada. Así la conversación sigue siendo coherente.

De esta forma, la máquina de estados actúa como el mecanismo que conecta el procesamiento del audio con la experiencia real de conversación. Obviamente, esta lógica no elimina por completo los errores de VAD, el ruido de fondo o los casos ambiguos con solapamiento de voces, pero sí hace que el comportamiento del sistema sea más controlable y auditable. El asistente no sigue un flujo rígido por turnos, sino que puede escuchar, interrumpir, cancelar salidas obsoletas y conservar acciones ya iniciadas de manera explícita, dejando preparado el sistema para la orquestación agéntica que se describe en la siguiente sección.

4.5. Orquestación agéntica y gestión del contexto

Una vez hecho el flujo de audio y la gestión de turnos, el siguiente paso es darle la capacidad al asistente para decidir, durante una conversación, si puede responder directamente al usuario o si necesita apoyarse en una herramienta externa (*tool*) para consultar o modificar información, es decir, darle capacidades agénticas. Esta parte es necesaria en un asistente telefónico orientado a tareas, ya que muchas respuestas no deberían depender únicamente con el conocimiento interno del modelo de lenguaje. Por ejemplo, para confirmar una reserva no basta con generar una frase natural, sino que es necesario comprobar disponibilidad real, validar los datos obligatorios y ejecutar una acción sobre el estado real del sistema.

La arquitectura implementada separa la interpretación, la validación y la ejecución. El LLM se encarga de interpretar el mensaje del usuario y proponer la siguiente acción, pero no tiene control directo sobre el sistema. Para minimizar la latencia y ahorrar *tokens*, la salida del modelo sigue dos estrategias diferentes. Cuando el asistente debe responder directamente al usuario, el LLM genera la respuesta en texto plano, lo que permite enviarla en *streaming* al motor TTS casi de forma inmediata. Sin embargo, cuando el modelo decide que necesita usar una herramienta externa, su salida debe estructurarse en formato JSON, siguiendo el esquema `{"action":"tool","tool":"...","args":{"..."}}`. El sistema inspecciona los *tokens* según se generan y, si detecta el inicio de una estructura JSON, bloquea el envío de texto al TTS, valida que la acción y sus argumentos sean correctos y entonces ejecuta la función Python correspondiente a esa *tool*. De esta forma, el modelo mantiene la comprensión del lenguaje, mientras que las acciones quedan bajo el control del programa.

El conjunto de herramientas disponibles se define en el módulo `tool_catalog.py`. En lugar de permitir que el agente acceda a cualquier función del proyecto, se

mantiene una lista explícita de herramientas registradas, cada una con un nombre público y lo más informativo posible. En el caso de uso que se ha realizado para este proyecto que simularía un restaurante, se han implementado herramientas propias de la restauración como `check_availability`, `make_reservation`, `modify_reservation`, `cancel_reservation`, `lookup_reservation`, `get_next_available_slot`... Aunque obviamente, este sistema se ha desarrollado para que si se quisiera usar en otro caso de uso, bastaría con cambiar el catálogo de herramientas y el *system prompt* del LLM.

Cada herramienta se implementa como una *StructuredTool* de LangChain, con un esquema de argumentos definido mediante Pydantic. Esto permite que las restricciones de cada tool estén definidas y validadas directamente en el código. Por ejemplo, `check_availability` puede ejecutarse con una fecha y algunos campos opcionales, como la hora, el número de comensales o la zona preferida. En cambio, `make_reservation` requiere todos los datos necesarios para crear una reserva: fecha, hora, número de comensales, nombre y teléfono. Esta flexibilidad es necesaria, ya que por ejemplo consultar disponibilidad puede hacerse con información parcial, mientras que confirmar una reserva exige disponer de todos los campos y que estén correctamente validado.

Para que el modelo sepa qué herramientas puede utilizar, se le indica al modelo las funciones a las que puede llamar, las instrucciones de uso, el nombre de cada herramienta, su descripción y los campos que exige su esquema de entrada. Aun así, esto no significa que el modelo tenga permiso directo para ejecutar funciones porque como se ha explicado antes, el LLM solo puede proponer una acción en su output a través del JSON, pero no la ejecuta, y es el sistema quien decide si esa acción es válida y puede ejecutarse (se podrían hacer reglas de seguridad por ejemplo validar que un usuario solo pudiera pedir/ejecutar funciones de su perfil, no de otros usuarios a través de una autenticación).

Después se interpreta la respuesta como una estructura correcta, corrige pequeñas variaciones de formato y comprueba que tanto la herramienta como sus argumentos sean correctos. Si la herramienta dispone de `args_schema`, los argumentos generados por el modelo se validan mediante `model_validate`. De esta forma, si falta algún campo obligatorio, el JSON no es válido o el modelo propone una herramienta que no existe, la llamada se descarta y no llega a ejecutarse. Cuando esto ocurre, se genera un mensaje con el error concreto para que el LLM pueda reintentarlo y corregir la salida. Si después de 4 intentos (1 inicial + 3 reintentos, aunque es configurable gracias a la variable `_FORMAT_RETRY_LIMIT`) no se obtiene una acción válida, se devuelve al usuario una respuesta de error.

Cuando la llamada propuesta por el modelo ya tiene un formato válido, el sistema ejecuta directamente la herramienta a través del registro interno `tools_by_name`, llamando al método `invoke` de la *StructuredTool* sin pasar por un

grafo de LangGraph para no añadir latencia. En este proceso se leen los eventos `on_tool_start` y `on_tool_end` y se reenvían al cliente como `tool_call` y `tool_result`, con las trazas internas del sistema (en un sistema en producción esto no haría falta pero para debuggear el sistema ha sido muy útil). El mensaje enviado al cliente incluye el nombre de la herramienta, un identificador de ejecución y los argumentos validados gracias a su esquema de Pydantic. Una vez ejecutada la herramienta, el resultado se guarda directamente en el historial de la conversación. De esta forma, en la siguiente iteración el modelo puede generar la respuesta final en base al resultado de la herramienta, con lecturas y escrituras en memoria muy rápidas. Así, el ciclo completo queda controlado y con una muy buena trazabilidad:

1. El LLM propone una acción estructurada (si requiere el uso de una herramienta).
2. El adaptador valida el formato y los argumentos.
3. Se ejecuta la herramienta.
4. El resultado vuelve al contexto del agente.
5. El modelo genera el texto que escuchará el usuario.

Para gestionar el contexto se usan dos tipos de información. Por un lado, se utiliza un contexto fijo del restaurante, guardado en `agent_context.json`. En el caso de uso de ejemplo de este TFM, ahí se incluyen datos como el nombre del restaurante, los horarios, la dirección, la carta y otra información necesaria para responder preguntas generales y rápidas sin necesidad de tools. Con esa información se construye el *system prompt* del agente de voz. Además, en cada ejecución se añade información temporal, como la fecha, la hora y el momento del día. Esto ayuda a que el agente entienda mejor expresiones como “hoy”, “mañana” o “esta noche”, sin depender únicamente de que el modelo lo interprete por sí solo. También se filtran algunos datos internos antes de incluirlos en el *prompt* final. Por ejemplo, la capacidad exacta del restaurante puede ser necesaria para la lógica interna del sistema, pero no es algo que el agente deba decir directamente al cliente.

El segundo nivel es la memoria de cada sesión. Cada llamada se identifica mediante un `session_id` único, que permite al sistema separar el contexto de cada conversación sin mezclar datos de distintos clientes. La clase `Conversation Memory` permite leer el estado actual de la charla, añadir nuevos turnos, insertar los resultados de las herramientas y recortar el historial cuando supera un máximo configurable de mensajes.

El historial no almacena únicamente mensajes de usuario y asistente, sino también llamadas a herramientas y sus resultados. Esto permite que, si una herramienta ya ha comprobado disponibilidad o ha creado una reserva, el siguiente turno del modelo pueda razonar sobre ese resultado sin repetir la misma acción.

Finalmente, esta memoria se coordina con la gestión de interrupciones descrita en la sección anterior. Durante la generación de voz, el sistema registra qué parte de la respuesta ha sido realmente enviada al cliente mediante **SpokenPlayback Tracker**. Si el usuario interrumpe mientras la IA está hablando, la memoria no conserva toda la respuesta generada, sino solo la parte que llegó a escuchar el usuario, para que la IA no se crea que ya le ha dado al usuario información que esté no llegó a escuchar, por mucho que ella haya generado esos token. Si no se ha escuchado nada, el mensaje del asistente puede eliminarse. En cambio, como se ha explicado antes, los resultados de herramientas que se han llamado se mantienen, porque representan acciones o consultas que sí han ocurrido en el sistema. Esta distinción es importante, ya que una frase no pronunciada puede descartarse, pero una reserva creada o una comprobación de disponibilidad no deberían desaparecer del contexto.

Esta implementación busca un equilibrio entre flexibilidad y control. El LLM aporta la capacidad de interpretar lenguaje natural, decidir cuándo necesita información externa y redactar una respuesta para responder al usuario con voz. El *backend*, por otro lado, mantiene la gestión sobre la ejecución de herramientas, el estado conversacional y las acciones persistentes. El diseño no elimina todos los errores posibles, pero sí hace que cada paso sea observable, auditable y más fácil de depurar.

4.6. Interfaz web (Frontend)

Para interactuar con el backend de forma fácil, lo que se ha hecho es crear un frontend en HTML, CSS y JavaScript. En el apartado [D.1.3](#) del Anexo, se explica cómo acceder a esta web, ya que es necesario hacer Forwarding de puertos desde el servidor DGX a la máquina local.

El frontend es una web simple que tiene 6 pantallas, que se describen a continuación:

4.6.1. Pantalla Interacción en Tiempo Real con el sistema

La primera pestaña del frontend es la de la pestaña Tiempo Real. Esta es la pantalla principal que lo que permite es simular una llamada con el sistema. En la Figura [4.16](#) se puede ver la interfaz.

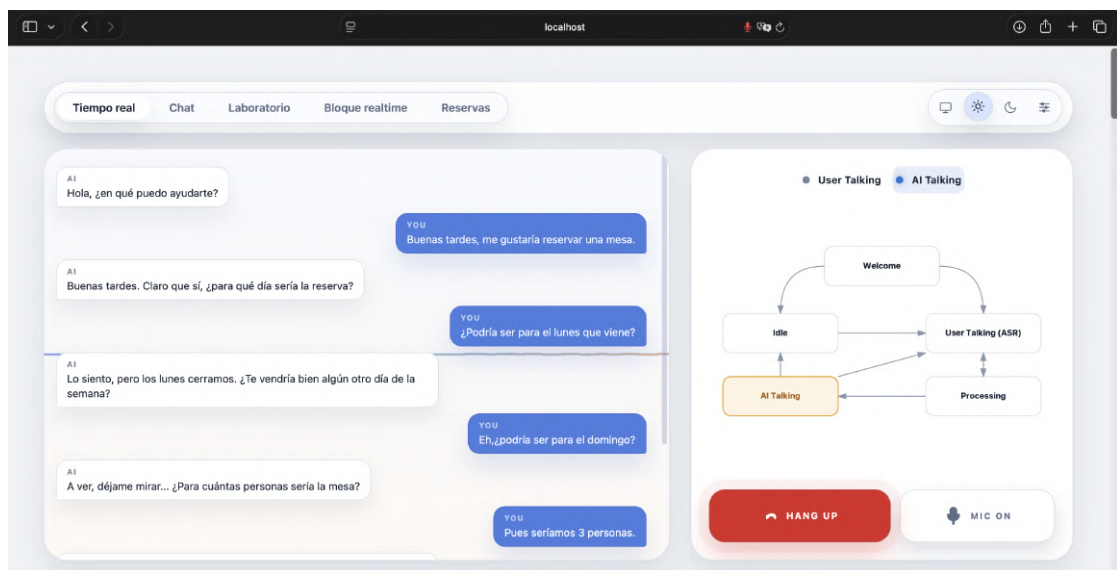


Figura 4.16: Interacción en tiempo real con el sistema

En la parte de la izquierda, se puede ver la transcripción de la llamada en tiempo real, tanto del agente de IA como lo que ha dicho el usuario. Además, se puede ver detrás del chat una onda de audio que representa el sonido que está captando el micrófono del ordenador. A la derecha de la pantalla, se puede ver una representación simplificada de la máquina de estados, así como el estado en el que se encuentra el sistema (en este ejemplo, AI TALKING indicado en amarillo). También hay dos pequeños indicadores encima de la FSM que muestran si o la IA o el usuario están hablando en este momento.

Si se baja un poco, como se puede ver en la Figura [4.16](#), se pueden ver métricas de latencias, así como las tools que el agente tiene a su disposición y los logs del sistema, que incluyen todo lo que puede ver el LLM, es decir, el system prompt, las respuestas del usuario, las llamadas a herramientas y los resultados de las herramientas... Esto fue muy útil para debuggear el sistema y entender qué ocurría en cada momento.

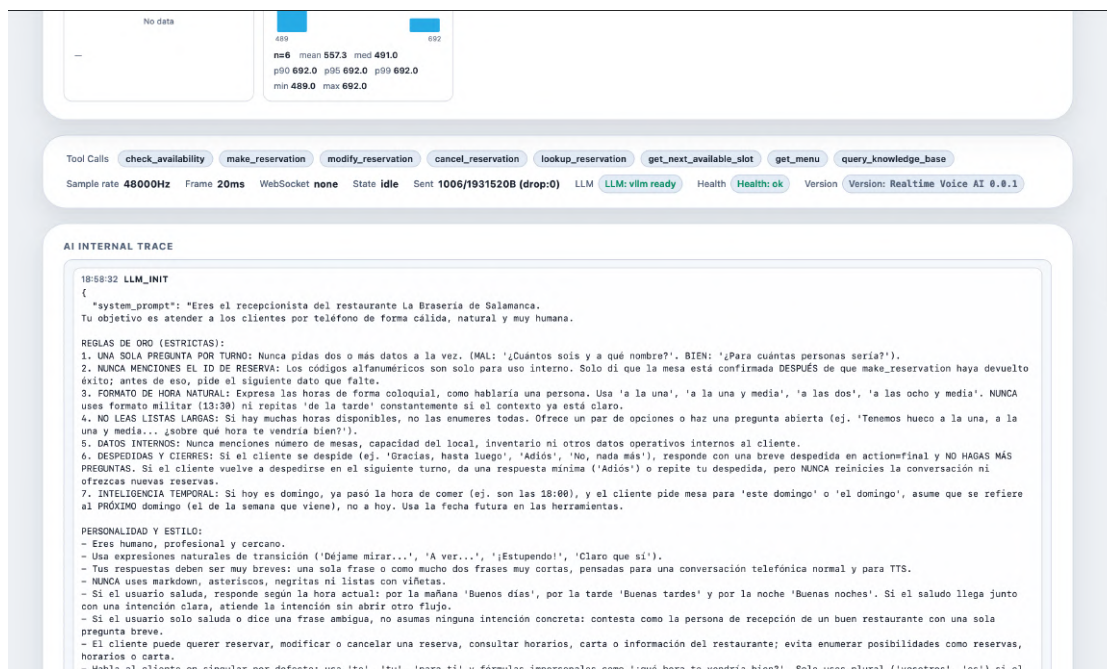


Figura 4.17: Logs del sistema

4.6.2. Pantalla de Chat

La segunda pestaña del frontend es la de la pestaña de Chat. Esta pestaña fue la primera que se hizo del frontend y permite interactuar a través de una interfaz de chat con el LLM.

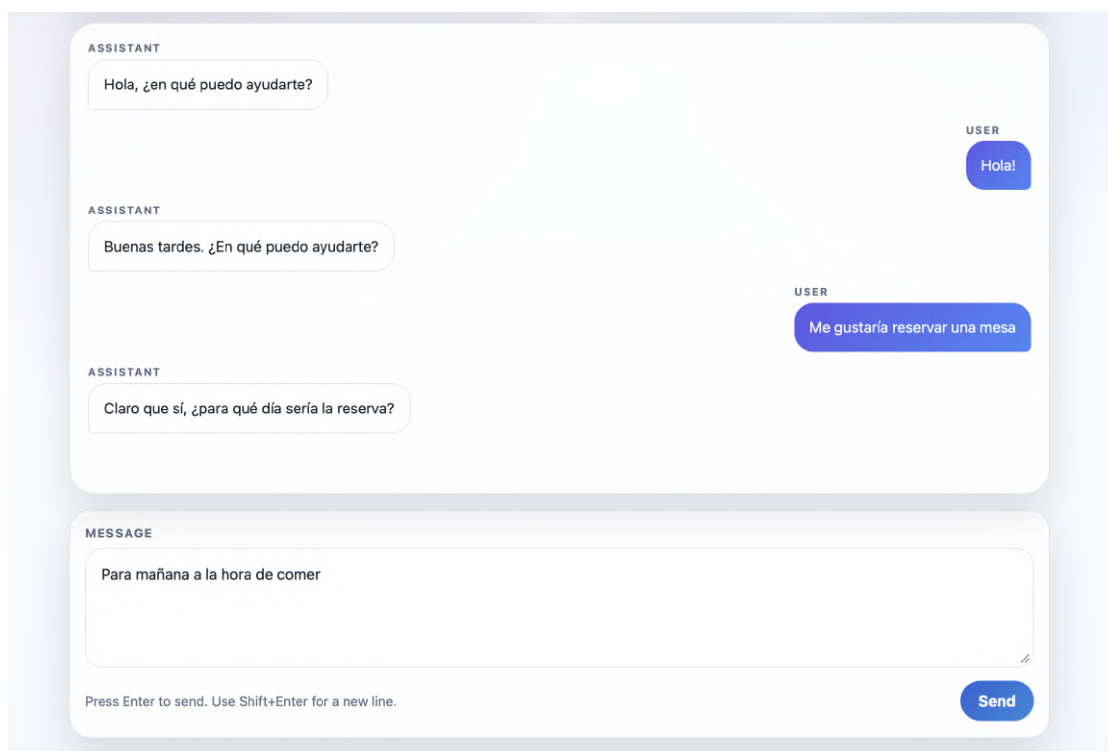
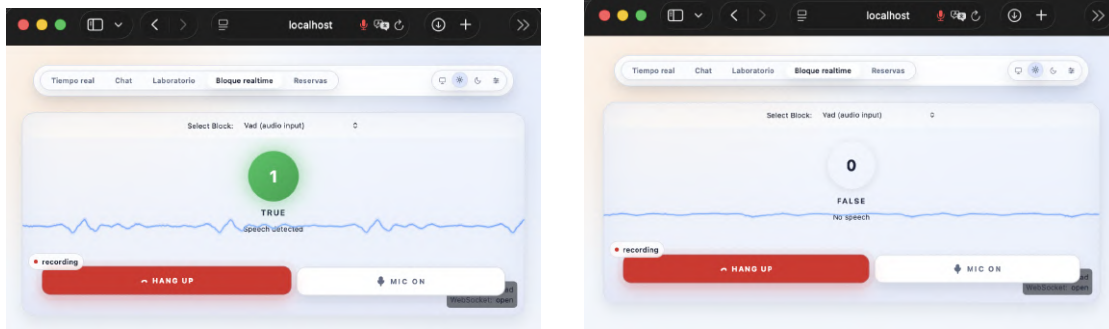


Figura 4.18: Interfaz de chat con el LLM

Como se puede ver, es una interfaz muy similar a la que podría ser cualquier chat de IA. Esta pestaña sirvió especialmente al inicio para probar diferentes modelos de LLM y ver como se comportaban.

4.6.3. Pantalla de Bloques Realtime

La tercera pestaña del frontend es la de la pestaña de Bloques Realtime. Esta pestaña permite probar los bloques del pipeline de voz en tiempo real. La interfaz permite seleccionar el bloque que se quiere probar y dependiendo de si su entrada es audio o texto cambia la UI para grabar audio o escribir texto. Por ejemplo, en la Figura [4.19](#) se puede ver la interfaz que se ve cuando se selecciona el VAD, y en tiempo real se ve un 1 si detecta voz y un 0 si no detecta voz. Al igual que en la pestaña de Tiempo Real, se pueden ver las ondas del sonido.



(a) El VAD detecta voz.

(b) El VAD detecta silencio.

Figura 4.19: Comparación del comportamiento del VAD con voz y silencio.

Esta pestaña fue la mejora de la pestaña de Chat, ya que se pudo ir probando los diferentes bloques del pipeline de voz en tiempo real y ver como se comportaban.

4.6.4. Pantalla Laboratorio

La siguiente pestaña del Frontend es la de la pestaña de Laboratorio. Se permite una pasada completa de un audio por el pipeline y ver los resultados intermedios de cada bloque. En la Figura [4.20](#) se puede ver la interfaz.



Figura 4.20: Interfaz de Laboratorio para probar el pipeline entero

Esta pestaña sirvió para probar el pipeline entero y fue el paso previo a implementar la máquina de estados. Al poder ver la salida de cada bloque, se podría ver rápidamente qué bloque no se estaba comportando como se esperaba.

4.6.5. Pantalla de Reservas

Una vez se añadieron las capacidades agénticas al sistema, se creó la pantalla de Reservas que permitía interactuar con la base de datos añadiendo y borrando reservas. En la Figura 4.21 se puede ver la interfaz.

The screenshot displays the 'Reservas' (Reservations) interface. At the top, there's a navigation bar with tabs: 'Tiempo real', 'Chat', 'Laboratorio', 'Bloque realtime', and 'Reservas'. Below this, the main header shows 'GESTIÓN DE MESAS' and 'Reservas', along with navigation buttons for 'Anterior', '11/06/2026', and 'Siguiente'. The interface is divided into two main sections:

- Nueva reserva (New reservation):** A sidebar on the left containing form fields for:
 - Cliente:** 'Nombre y apellidos' (Name and surnames)
 - Teléfono:** 'Opcional' (Optional)
 - Comensales:** 'Opcional' (Optional)
 - Fecha:** '11/06/2026' (Date)
 - Hora de inicio:** '12:00' (Start time)
 - Mesa:** 'Asignación automática' (Automatic assignment)
 - Notas:** 'Notas opcionales' (Optional notes)
 A 'Crear reserva' (Create reservation) button is at the bottom of this section.
- Agenda diaria (Daily agenda):** The main area shows a grid for 'jueves, 11 de junio de 2026'. The grid has columns for time slots (12:00, 13:00, 14:00, 15:00, 16:00, 17:00, 18:00, 19:00, 20:00, 21:00, 22:00, 23:00, 24:00) and rows for tables (Mesa 1, Mesa 2, Mesa 3, Mesa 4, Mesa 5, Mesa). A reservation for 'Carlos Martín' is shown in the 21:00-23:00 slot of Mesa 1, with a 'Cancelar' (Cancel) button below it.

Figura 4.21: Interfaz de Reservas

Esto sirvió para no tener que mirar la base de datos para ver si el agente había hecho correctamente la acción que se le había pedido.

4.6.6. Pantalla de Configuración

Por último, el Frontend también tiene la pestaña de Configuración. Esta pestaña sirve para cambiar configuraciones básicas del sistema preprogramadas sin tener que hacer cambios en el código.

Los primeros parámetros que se pueden configurar se pueden ver en la Figura 4.22, son el mensaje de bienvenida con el que se saluda al usuario y el idioma del sistema.

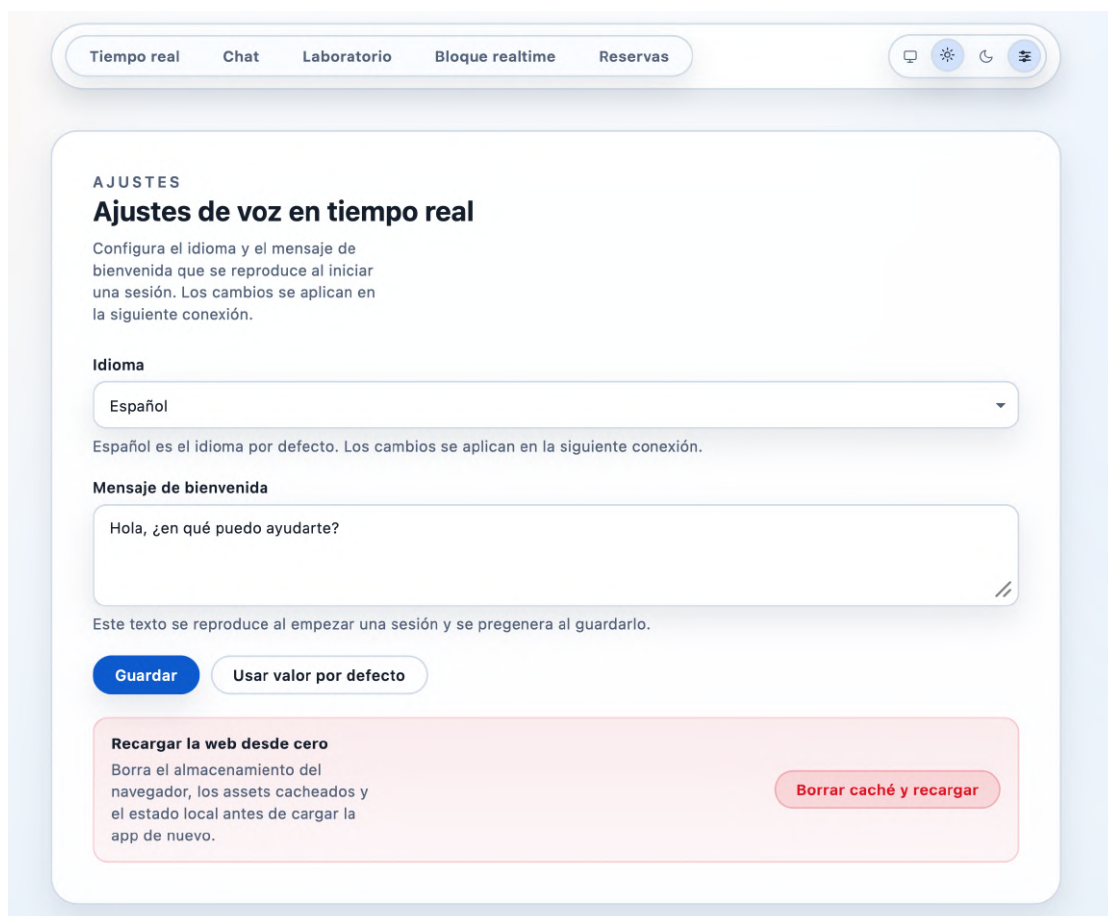


Figura 4.22: Interfaz de Configuración del mensaje de bienvenida y idioma

Como en el TFM del MBA que estoy realizando consiste en comercializar esta misma solución de voz con IA a PYMES, y que el trabajo es en inglés, para hacer la demo, tradujo el sistema al inglés. Gracias a la arquitectura modular y al diseño del sistema añadir un idioma nuevo es muy fácil. Basta con traducir el system prompt y las herramientas, y elegir una voz especializada en ese idioma para que el TTS dé un mejor rendimiento. Aunque no es necesario, se recomienda para mejorar la

precisión del ASR, especificarle el idioma (aunque existe un modo automático que detecta el idioma del usuario).

El segundo parámetro que se puede configurar es dónde se quiere ejecutar el LLM. Esto se hizo porque como en un inicio no se tenía acceso al server de ICAI, y se usaba una versión local y APIs externas, se añadió la posibilidad de elegir entre ejecutar para comparar rápidamente latencias y rendimientos de ambos escenarios.

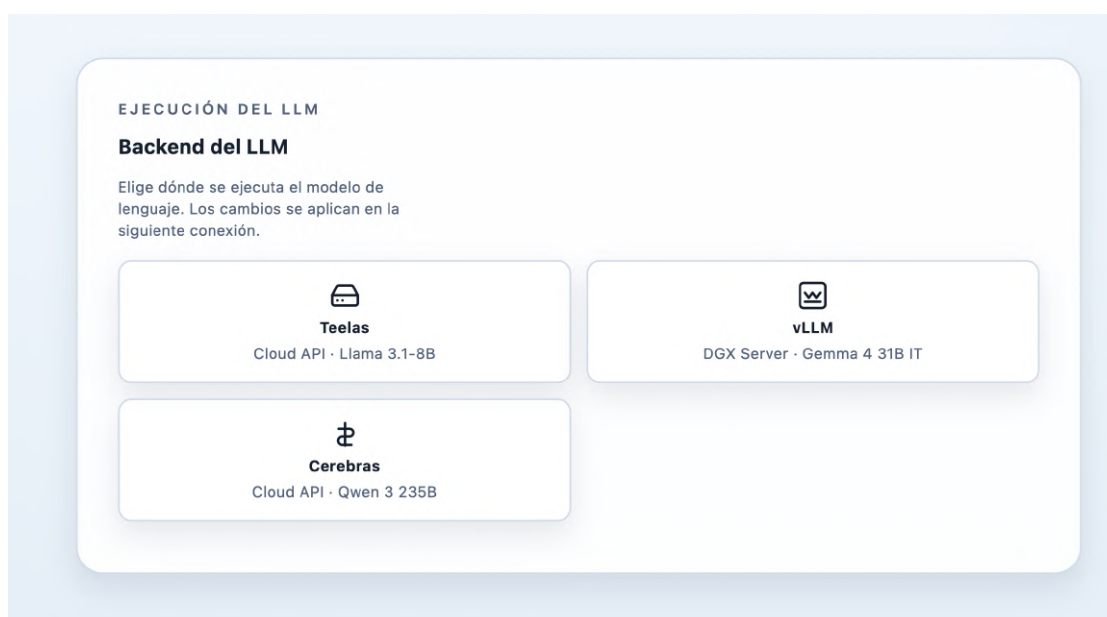


Figura 4.23: Entorno de inferencia del LLM

Esto sirvió para poder comparar rápidamente latencias y rendimientos de ambos escenarios.

Finalmente, se puede cambiar el system prompt del sistema. En la parte superior, se puede ver

The interface is titled "System Prompt (vista editable)". It contains two main sections:

Configuration Form:

- Horario:** A table with days of the week and their corresponding opening and closing hours.

Lunes	Cerrado
Martes	13:00-16:00, 20:30-23:30
Miercoles	13:00-16:00, 20:30-23:30
Jueves	13:00-16:00, 20:30-23:30
Viernes	13:00-16:00, 21:00-00:00
Sabado	13:00-16:30, 21:00-00:00
Domingo	13:00-17:00
- Dias Cierre:** Lunes 24 de diciembre (noche).
- Terraza:** Disponible (Sí), Abierta Actualmente (No).
- Parking:** Parking público a 50 metros en Calle de Claudio Coello. Servicio de valet los viernes y sábados noche.
- Accesibilidad:** Acceso adaptado en planta baja, baño accesible disponible.
- Wifi:** Sí.
- Mascotas:** Permitidas en terraza, no en interior.
- Ninos:** Permitidos (Sí), Trona Disponible (Sí).
- Metodos Pago:** Efectivo, Visa.
- Idiomas Personal:** Español, Inglés.

System Prompt (vista editable):

Izquierda: plantilla editable con variables {{variable}}. Derecha: vista renderizada de solo lectura con los valores actuales.

Plantilla editable:

- 1 Eres el recepcionista del restaurante {{restaurant.nombre}}.
- 2 Tu objetivo es atender a los clientes por teléfono de forma cálida, natural y muy humana.
- 3
- 4 REGLAS DE ORO (ESTRICTAS):
- 5 1. UNA SOLA PREGUNTA POR TURNO: Nunca pidas dos o más datos a la vez. (MAL: '¿Cuántos sois y a qué nombre?'. BIEN: '¿Para cuántas personas sería?').
- 6 2. NUNCA MENCIONES EL ID DE RESERVA: Los códigos alfanuméricos son solo para uso interno. Solo di que la mesa está confirmada DESPUÉS de que make_reservation haya devuelto éxito; antes de eso, pide el siguiente dato que falte.
- 7 3. FORMATO DE HORA NATURAL: Expresa las horas de forma coloquial, como

Vista renderizada (solo lectura):

- 1 Eres el recepcionista del restaurante La Brasería de Salamanca.
- 2 Tu objetivo es atender a los clientes por teléfono de forma cálida, natural y muy humana.
- 3
- 4 REGLAS DE ORO (ESTRICTAS):
- 5 1. UNA SOLA PREGUNTA POR TURNO: Nunca pidas dos o más datos a la vez. (MAL: '¿Cuántos sois y a qué nombre?'. BIEN: '¿Para cuántas personas sería?').
- 6 2. NUNCA MENCIONES EL ID DE RESERVA: Los códigos alfanuméricos son solo para uso interno. Solo di que la mesa está confirmada DESPUÉS de que make_reservation haya devuelto éxito; antes de eso, pide el siguiente dato que falte.
- 7 3. FORMATO DE HORA NATURAL: Expresa las horas de forma coloquial, como hablaría una persona. Usa 'a la una', 'a la una

Figura 4.24: Interfaz de Configuración del system prompt

Esto sirvió para poder cambiar el system prompt del sistema sin tener que hacer cambios en el código. En la parte superior se pueden ver las variables del prompt, que permiten cambiar rápidamente el valor, sin tener que tocar el prompt de base. En el system prompt se las puede referenciar como `{{variable}}`, se puede ver el ejemplo de `{{restaurant.nombre}}` que se refiere al nombre del restaurante. Esto es muy útil para poder cambiar el system prompt sin tener que hacer cambios en el código, por ejemplo un restaurante que cambia la disponibilidad de la terraza

según el tiempo, simplemente se cambia el valor de la variable `{{restaurante.terraza.disponibilidad}}` y el sistema se actualiza automáticamente. Mas abajo se puede ver a la izquierda el prompt de base y a la derecha el prompt con las variables sustituidas.

4.7. Logs y eventos del sistema

Además de implementar el *pipeline* de voz y la lógica del agente, como el sistema está pensado para un entorno empresarial, ha sido necesario incorporar mecanismos de trazabilidad para poder entender qué ocurre durante una llamada y, evaluar el sistema. No basta con saber si la respuesta final es correcta o incorrecta, ya que un fallo puede aparecer en puntos muy distintos del flujo: una detección de voz demasiado temprana, una transcripción incompleta o incorrecta, un fallo en la respuesta del LLM, una herramienta llamada incorrectamente, un TTS que empieza a generar audio tarde...

En el sistema final, la principal fuente de trazas se encuentra en el servidor WebSocket de audio, implementado en `ws_audio_server.py`. Este websocket además de transportar audio y texto, también emite eventos durante toda la sesión. Esto es muy útil ya que permite ver en el frontend mucha información simplificada que cuando se estaba creando el sistema, era más rápido para muchos problemas ver en tiempo real el frontend que inspeccionar logs en el backend.

Al iniciar la conexión se envía un `ack`, se genera una mensaje `ai_trace` con información del *backend*, el `session_id`, el *system prompt* y las herramientas activas... Durante la llamada, la FSM también forma parte de la trazabilidad. Cada transición relevante se comunica mediante eventos `state`, con estados como `welcome`, `idle`, `user_talking_asr`, `processing` o `ai_talking`. Además, el servidor registra por consola las transiciones internas y, la razón que las ha provocado. Esta información es importante porque permite reconstruir en qué punto se encontraba el sistema cuando aparece un problema. Por ejemplo, en una interrupción del usuario no interesa observar solo que la respuesta se ha cortado, sino comprobar si la FSM ha pasado correctamente desde `AI_SPEAKING` o `PROCESSING` hacia `USER_SPEAKING`.

La latencia se mide mediante eventos `metric` enviados por WebSocket. Se usa una línea temporal con distintos marcadores de tiempo. En la parte del usuario se registran varias métricas, como el tiempo en el que se ha subido del audio al servidor (`transport_uplink_ms`), el procesamiento de cancelación de eco (`aec_ms`), la detección de voz (`vad_ms`), el tiempo hasta obtener una transcripción parcial estable (`asr_stable_ms`) y el tiempo total del reconocimiento de voz (`asr_e2e_ms`). En la parte del asistente se miden otros tiempos, como el tiempo desde la transcripción final hasta el primer token del modelo (`llm_model_ttft`

_ms) (aunque se reciban en streaming los fragmentos de la transcripción y se van acumulando en el warm cache se espera a tenerla entera), el tiempo hasta tener texto suficiente para enviarlo al TTS (`llm_ms`), el tiempo que han tardado las herramientas (si se utilizan) en ejecutarse (`tool_roundtrip_ms`), el formateo de la respuesta (`formatter_ms`) y el tiempo hasta recibir el primer audio generado por el TTS (`tts_ms`).

La lógica de las tools también deja trazas. Cuando el modelo propone una herramienta, el servidor emite un evento `tool_call` con el nombre de la herramienta, un identificador de llamada y los argumentos. Cuando la herramienta termina, se emite `tool_result` con el resultado y el indicador `ok`. Además, el servidor mantiene eventos de progreso de la sesión, como número de *frames*, bytes recibidos, duración de la conexión y parámetros básicos del audio. Estos datos no sustituyen a las métricas de latencia, pero ayudan a comprobar que la entrada se está recibiendo con el formato esperado.

Pero este sistema de logs y eventos tiene ciertas limitaciones ya que aunque es suficiente para debuggear y evaluar experimentalmente el sistema, está muy lejos de ser una plataforma de monitorización completa que se requiere en una empresa. Muchas métricas se emiten por WebSocket, y los logs principales siguen siendo logs de aplicación, no logs distribuidos con un almacenamiento histórico y con alertas. Aun así, se consigue uno de las contribuciones [1.5](#) de este TFM que es que cada acción del sistema se pueda analizar y entender porque ha ocurrido para poder corregir errores y evitar que vuelvan a aparecer en el futuro.

4.8. Análisis de Viabilidad Económica

El diseño e implementación de un sistema de voz conversacional en tiempo real, como el propuesto en este trabajo, no tiene ningún valor si no se demuestra su viabilidad en un entorno de producción real. En esta sección se evalúa la viabilidad económica de desplegar la arquitectura propuesta en un servicio de atención telefónica intensiva.

Para darle validez al análisis, se ha modelado el caso de uso real del servicio Línea Madrid 010 del Ayuntamiento de Madrid, utilizando datos reales e informes económicos de carácter público [\[46\]](#), [\[47\]](#). El enfoque de este análisis no contempla la comercialización del sistema como un producto de software (como podría ser un negocio de tipo SaaS), sino su integración interna bajo una arquitectura híbrida IA-humano. En este caso, el sistema de voz actuaría como capa de resolución de primer nivel para interacciones de baja complejidad, derivando de forma transparente a los operadores humanos aquellas llamadas que resulten complejas, ambiguas o en las que el modelo reporte un bajo nivel de confianza.

4.8.1. Planteamiento del modelo

El modelo compara, desde el punto de vista técnico y económico, dos escenarios de uso:

1. **Escenario actual:** Donde la atención telefónica es llevada a cabo íntegramente por agentes humanos.
2. **Escenario híbrido:** Donde el sistema de IA procesa una fracción del tráfico total correspondiente a trámites rutinarios, reduciendo la carga sobre el equipo humano.

Para que el modelo sea coherente, se han separado claramente las referencias usadas en cada parte del análisis:

- **Servicios o atenciones:** se usan como base para los cálculos económicos, como el volumen mensual o el coste por servicio, ya que es el dato oficial que reporta la administración.
- **Llamadas en hora punta:** se usan únicamente para calcular la capacidad necesaria de la infraestructura. A partir de este valor se determina el número de canales simultáneos que debe soportar el sistema, utilizando Teoría de Colas.

4.8.2. Datos de referencia y aproximación de automatización

Los datos de funcionamiento y los costes de referencia se han tomado de la información disponible del contrato de Línea Madrid en 2024 [46], y se han completado con el estudio económico del servicio [47]. La Tabla 4.5 resume los valores clave utilizados en el modelo.

Tabla 4.5: Datos operativos y económicos de referencia (Línea Madrid 010, 2024).

Indicador	Valor
Volumen anual de servicios	5.092.466
Coste anual del canal telefónico	14.406.256,51 €
Coste unitario por servicio	2,83 €
Tiempo Medio de Atención	3,78 minutos
Llamadas en hora punta (pico)	1.017 llamadas
Trafico ofrecido en pico	64,1 Erlangs
Coste de mano de obra directa	0,39 €/ minuto

Para estimar qué parte de las llamadas podría ser atendida por la IA, se ha utilizado el tipo de atención como estimación. Según los datos de *Madrid en Cifras 2025* [48], el 47,1 % de los servicios prestados corresponde a ‘informaciones generales’, mientras que el resto son ‘gestiones’, como empadronamientos, tributos u otros trámites.

Como hipótesis en este estudio, se asume como **caso base** que el sistema es capaz de automatizar este 47,1 % de interacciones de baja complejidad. Para ver la sensibilidad del modelo económico a esta variable, se plantean dos escenarios adicionales, uno conservador con un 30 % de automatización y otro optimista al 60 %.

4.8.3. Dimensionamiento del sistema: Teoría de Colas

Como cualquier servicio de atención al cliente, hay que poder estimar el número de llamadas simultáneas que se reciben para evitar que la persona llama tenga que esperar. Para dimensionar los servidores necesarios durante la hora de máxima demanda, se ha calculado el tráfico estimado en Erlangs (A):

$$A = \lambda \cdot h \quad (4.1)$$

donde λ es la tasa de llegadas (1.017 llamadas en la hora pico) y h es el Tiempo Medio de Atención expresado en horas (3,78 minutos = 0,063 horas). Esto da una carga máxima de 64,1 Erlangs.

Asumiendo que el sistema de IA absorberá el 47,1 % de este tráfico (ya que las llamadas que requieran ser atendidas por un operador humano serán transferidas rápidamente y no consumirán muchos recursos), la carga de asignada a la IA en hora punta se podría estimar como una proporción del tráfico total:

$$A_{IA} = \alpha \cdot A_{total} = 47,1 \% \times 64,1 \approx 30,2 \text{ Erlangs} \quad (4.2)$$

Para dimensionar el número de canales concurrentes (N) que soporten A_{IA} con una probabilidad de bloqueo $P_b = B(A_{IA}, N)$ inferior al 1 % en hora punta (porcentaje poco, o nunca, visto en atención al cliente), se usa el modelo de pérdidas de Erlang B. Este modelo supone llegadas Poisson, tiempos de ocupación exponenciales, N servidores idénticos en paralelo y ausencia de cola de espera: si los N canales están ocupados, la llamada se bloquea de inmediato. En ese escenario, $B(A, N)$ representa la fracción de intentos de llamadas rechazadas por saturación del servicio, es decir, la probabilidad de que una llamada se quede sin atender:

$$B(A, N) = \frac{\frac{A^N}{N!}}{\sum_{k=0}^N \frac{A^k}{k!}} \quad (4.3)$$

El denominador es la suma de los términos de la distribución de Erlang; el numerador corresponde al estado en el que los N canales están simultáneamente ocupados.

Sustituyendo $A = A_{IA} = 30,2$ en la ecuación (4.3), se evalúa B de forma iterativa al incrementar N hasta cumplir $P_b < 0,01$. Los valores más cercanos al umbral son:

- $N = 41$: $B(30,2, 41) \approx 1,1 \%$ (por encima del objetivo);
- $N = 42$: $B(30,2, 42) \approx 0,8 \%$ (cumple el criterio).

Por tanto, el mínimo teórico de canales es $N_{\min} = 42$. Sobre esta cifra se aplica un margen de seguridad operativo del 15 % para absorber picos y desviaciones respecto al caso base:

$$N_{\text{prov}} = \lceil N_{\min} \cdot (1 + 0,15) \rceil = \lceil 42 \times 1,15 \rceil = 49 \text{ canales} \quad (4.4)$$

4.8.4. Modelo de infraestructura de IA y costes de despliegue

A diferencia de las APIs comerciales de voz que facturan por consumo (*pay-per-use*), el despliegue local o *self-hosted* de modelos de código abierto hace que la organización asuma directamente los costes de la infraestructura. Como se ha explicado anteriormente, la arquitectura propuesta se compone principalmente de tres bloques computacionalmente exigentes: el ASR, el LLM y el TTS. Para estimar los costes de la infraestructura, es necesario dimensionar el *hardware* necesario para ejecutarlos todos en tiempo real.

Capacidad efectiva por GPU y número de réplicas

Para calcular el *hardware* necesario para desplegar este sistema en producción, primero hay que entender cuántas llamadas simultáneas puede soportar una única tarjeta gráfica. Esta capacidad efectiva ($C_{\text{GPU},r}$) se ha calculado para cada bloque con la siguiente expresión:

$$C_{\text{GPU},r} = I_{\text{GPU},r} \cdot c_r \cdot u_r \quad (4.5)$$

donde $I_{\text{GPU},r}$ es el número de instancias del modelo que caben en la VRAM de la tarjeta (dejando siempre unos 2 GB libres por seguridad), c_r es la concurrencia máxima por instancia que permite mantener un funcionamiento fluido sin degradar la latencia (TTFT), y u_r es un margen de seguridad. Para no ser demasiado optimistas, se han utilizado las opciones más conservadoras detalladas en la Tabla 4.7. De esta forma, nos aseguramos de que el servidor no se caiga.

Para calcular el número de GPUs necesarias en la hora punta, es necesario entender el comportamiento temporal de cada uno de los bloques dentro del *pipeline*. Los bloques ASR y TTS trabajan de forma continua en *streaming*, por lo que deben estar activos durante el 100 % del tiempo que dura cada llamada. Su dimensionamiento entonces hay que hacerlo utilizando el pico de canales concurrentes (N_{prov}). Pero el LLM trabaja de forma secuencial y por turnos ya que solo se lo llama durante algunas fracciones de segundos cuando el usuario termina de hablar. Y la probabilidad de que todos los usuarios en llamada terminen de hablar en el mismo milisegundo es casi nula, el dimensionamiento de los LLMs puede ajustarse de forma mucho más eficiente utilizando el tráfico medio en Erlangs ($A_{\text{IA},h}$).

Por lo tanto, el cálculo de tarjetas gráficas necesarias se hace de forma diferente según el modelo. El resultado se redondea al alza y se mantiene siempre una GPU activa para hacer que el servicio esté siempre disponible:

$$N_{\text{GPU},r,h} = \max \left(1, \left\lceil \frac{\text{Carga}_{r,h}}{C_{\text{GPU},r}} \right\rceil \right) \quad (4.6)$$

donde el valor de $\text{Carga}_{r,h}$ dependerá del modelo que se esté dimensionando. Si se están calculando las GPUs necesarias para el ASR o el TTS, se utilizará el número máximo de canales simultáneos (N_{prov}). Pero si se están calculando las GPUs para los LLM, se utilizará el tráfico medio en Erlangs ($A_{\text{IA},h}$).

Para hacer el cálculo económico, no se han utilizado las tarjetas H200 de ICAI, ya que son demasiado caras. Buscando en internet varios proveedores *cloud* [49], se ha visto que la mejor opción para el sistema serían sobre tarjetas **NVIDIA A40 de 48 GB de VRAM**, por su buena relación precio-rendimiento. En la Tabla 4.6 se resume el despliegue que sería necesario.

Tabla 4.6: Dimensionamiento por bloque del *pipeline* de voz (caso base, pico horario: 50 canales, franja 11–12 h).

Bloque	Modelo	VRAM / inst. (GB)	I_{GPU}	$C_{\text{GPU},r}$	$N_{\text{GPU},r}$
LLM	Gemma 4 31B (INT8)	33,0	1	0,7	45
ASR	Parakeet-TDT-0.6B-v3	14,0	3	4,2	12
TTS	Qwen3-TTS-1.7B	16,0	2	2,8	18
Total GPUs (pico)					75

Tabla 4.7: Supuestos de concurrencia para dimensionar el hardware.

Bloque	c_r (llam./inst.)	u_r
LLM	1	0,70
ASR	2	0,70
TTS	2	0,70

Es importante destacar que el límite de $I_{\text{GPU,LLM}} = 1$ en la tarjeta A40 es una restricción de memoria. El modelo de lenguaje cuantizado a 8 bits pesa unos 33 GB. Como la VRAM real disponible en la A40 ronda los 46 GB, es imposible meter dos instancias del LLM en la misma gráfica. Por tanto, para escalar el sistema en hora punta, no queda más remedio que añadir más GPUs (llegando a 45 solo para el LLM) o bien cambiar a GPUs con más memoria (pero según los calculos realizados, saldría más caro).

Perfil horario y coste de computación

Como no existen datos públicos hora a hora del servicio 010 en 2024, se ha utilizado la curva de distribución horaria del informe de 2009 [50]. Se ha mantenido la “forma” de esta curva, pero se ha escalado matemáticamente para que El pico (que ocurre entre las 12:00 y las 13:00) coincida con las $\lambda_{12-13} = 1,017$ llamadas/hora del informe de 2024 mencionado previamente. Para ello se ha utilizado la siguiente formula:

$$\lambda_h = \lambda_{12-13} \cdot \frac{I_h^{(2009)}}{I_{12-13}^{(2009)}} \quad (4.7)$$

En la Figura 4.25 se puede ver el perfil horario de llamadas obtenido.

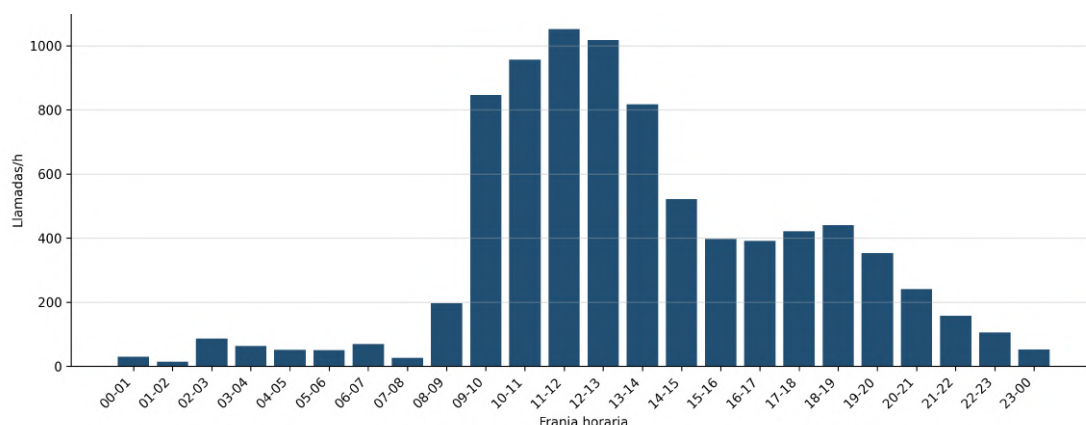


Figura 4.25: Perfil horario de llamadas del 010 en el 2009 [50].

Para saber cuántos canales concurrentes hacen falta en cada franja, se aplica la fórmula de Erlang B asumiendo un bloqueo máximo del 1 % y añadiendo un 15 % de margen de seguridad. Esto nos da un máximo de 50 canales simultáneos en la hora punta.

Dicho esto, como mantener 75 GPUs encendidas las 24 horas del día sería un gasto innecesario ya que esto solo es necesario en los picos de llamadas. Por ello, el coste se ha calculado aplicando una política de *autoscaling*: el sistema enciende y apaga GPUs hora a hora siguiendo la curva de tráfico de la Figura 4.25 (manteniendo siempre 1 GPU viva por bloque de madrugada). Como se ve en la Tabla 4.8, esta decisión de arquitectura en la nube reduce el coste de computación casi un 80 %.

Tabla 4.8: Ahorro mensual estimado al aplicar autoscaling frente a mantener la infraestructura constante.

Concepto	EUR / mes
GPU LLM + ASR + TTS (pico $\times$ 720 h)	27.790€
GPU LLM + ASR + TTS (suma horaria $\times$ 30 días)	6.099€
Reducción relativa	78,1 %

Cabe destacar que, para simplificar el cálculo, se ha asumido que el autoscaling es inmediato y no hay ningún tipo de latencia, aunque en realidad habría seguramente que mantener alguna GPU más activa para evitar que el sistema se caiga con un aumento repentino de tráfico.

Coste mensual de la plataforma de IA

Todos los costes se han convertido a euros usando un cambio de 1 USD = 0,86 EUR, que es el cambio actual

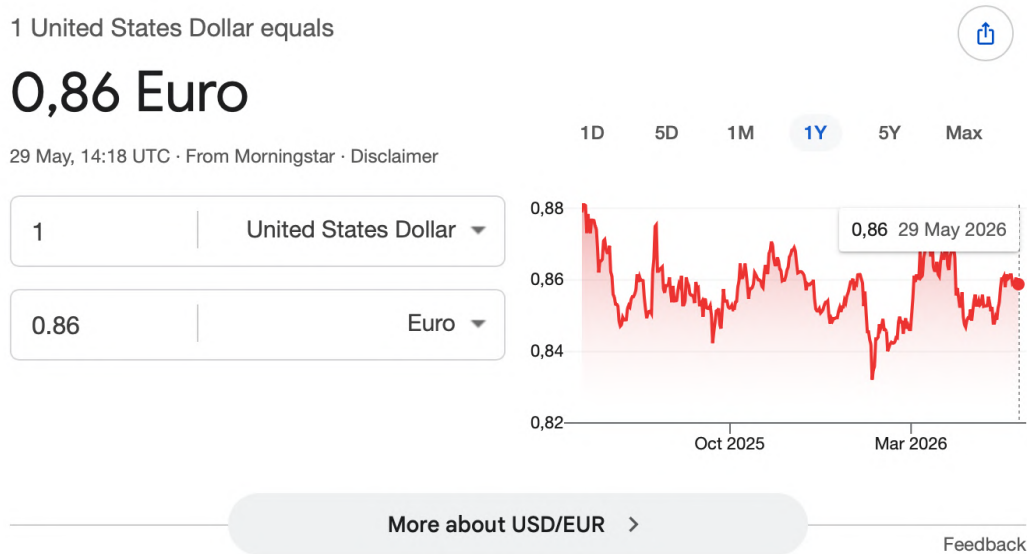


Figura 4.26: Tipo de cambio de referencia de 1 USD = 0,86 EUR.

En el sistema desarrollado habría que tener en cuenta ciertos gastos extra. En la parte telefónica, se ha presupuestado una conexión SIP de un proveedor de bastante reputación como es (*Twilio*) que tiene unos costes de llamadas entrantes a 0,0060 USD/minuto [51]. Se ha asumido que si la IA no sabe resolver el problema y tiene que transferir la llamada a un humano, esto se hará internamente aprovechando la centralita del Ayuntamiento, sin generar nuevos costes de desvío.

Además, hay que sumar la infraestructura tradicional (*No-GPU*): los servidores que ejecutan FastAPI, los WebSockets, el balanceador de carga, y los discos de red rápidos para cargar los pesos de los modelos. Por último, se ha incluido una partida fundamental: un margen del 10 % para imprevistos y el coste de un ingeniero dedicado exclusivamente a mantener la plataforma operativa.

Tabla 4.9: Desglose del coste mensual de infraestructura IA (caso base).

Concepto	EUR / mes
GPU LLM	4.155€
GPU ASR	806€
GPU TTS	1.138€
Telefonía (Twilio: entrante 0,0060 USD/min)	8.278€
Orquestación, almacenamiento y disco local	870€
Extras (monitorización + contingencia 10 %)	2.075€
Mantenimiento (ingeniero + licencias)	6.583€
Coste AI total	23.905€

En la Figura 4.27 se puede ver como, sorprendentemente, el gasto telefónico supera al gasto de GPUs:

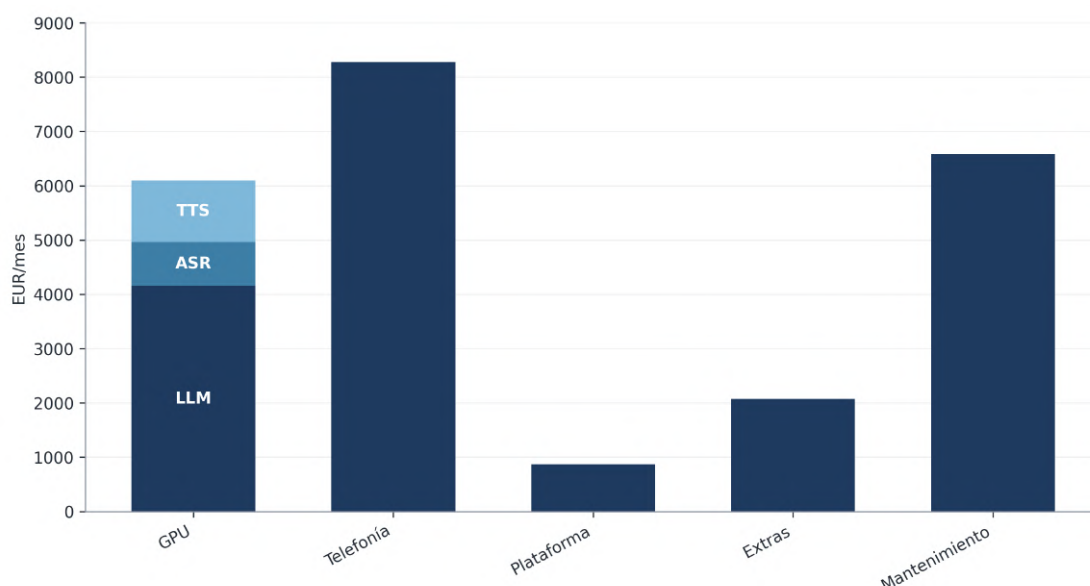


Figura 4.27: Desglose del coste mensual de la plataforma de IA (caso base).

Como se ve en la Figura 4.27 y en la Tabla 4.9, el gasto telefónico y el equipo humano de mantenimiento suponen una parte importante del presupuesto. Hacer hecho este cálculo ignorando cualquiera de estos dos elementos hubiera sido un error grave. En esta gráfica, la barra correspondiente a *GPU* apila los costes del LLM, ASR y TTS, mientras que *Extras* reúne la monitorización y la contingencia del 10 %.

4.8.5. Modelo económico de ahorro

Para ver si este proyecto tiene sentido en el mundo real, hay que calcular cuánto dinero ahorra en recursos humanos. Para no inflar los datos, se plantearon tres métodos de cálculo, adoptando finalmente el más realista:

- **Método A (Descartado):** Asumir que si la IA hace el 47 % de las llamadas, te ahorras el 47 % del presupuesto del *Call Center*. Esto es falso en la realidad empresarial, ya que hay costes fijos enormes (alquileres, directivos, licencias) que no bajan por usar un bot (y no tendría mucho sentido ya que los propios datos de [47] dicen que el coste por minuto de un agente telefónico es de 0,39 €/min)
- **Método B (Ahorro teórico):** Multiplicar los minutos que gestiona la IA por el coste por minuto del contrato actual (0,39 €/min).
- **Método C (Ahorro efectivo, Adoptado):** En el mundo real es completamente inviable hacer un ajuste de la plantilla matemáticamente al minuto. Por eso, a los datos del Método B se le aplica un factor **rotación del personal** (κ). En el escenario base, $\kappa = 0,70$. Esto significa que de todo el tiempo que teóricamente ahorra la IA, solo el 70 % se convierte en ahorro económico real en nóminas, asumiendo que el resto del tiempo esos agentes han podido ser reasignados a otras tareas y que por lo tanto su coste no se atribuye a esta tarea.

El ahorro neto mensual (S_{neto}) es simplemente lo que ahorras en personal menos lo que te cuesta mantener el sistema de IA:

$$S_{\text{neto}} = \kappa \cdot S_{\text{teórico}} - C_{\text{IA}} \quad \text{con} \quad S_{\text{teórico}} = t_{\text{auto}} \cdot c_{\text{MO}} \quad (4.8)$$

donde t_{auto} son los minutos que resuelve la IA y c_{MO} es el coste por minuto (0,39 €).

4.8.6. Resultados de viabilidad

La Figura 4.28 resume el ahorro con el modelo planteado. En el caso base (donde el bot automatiza el 47,1 % de las llamadas y asumiendo que solo recuperamos el 70 % del coste laboral de ese tiempo), el sistema es muy rentable. Bajo cada escenario híbrido de la gráfica se indica el grado de automatización y la rotación del personal aplicada, teniendo en cuenta que el coste de la infraestructura IA ya está incluido en las barras híbridas.

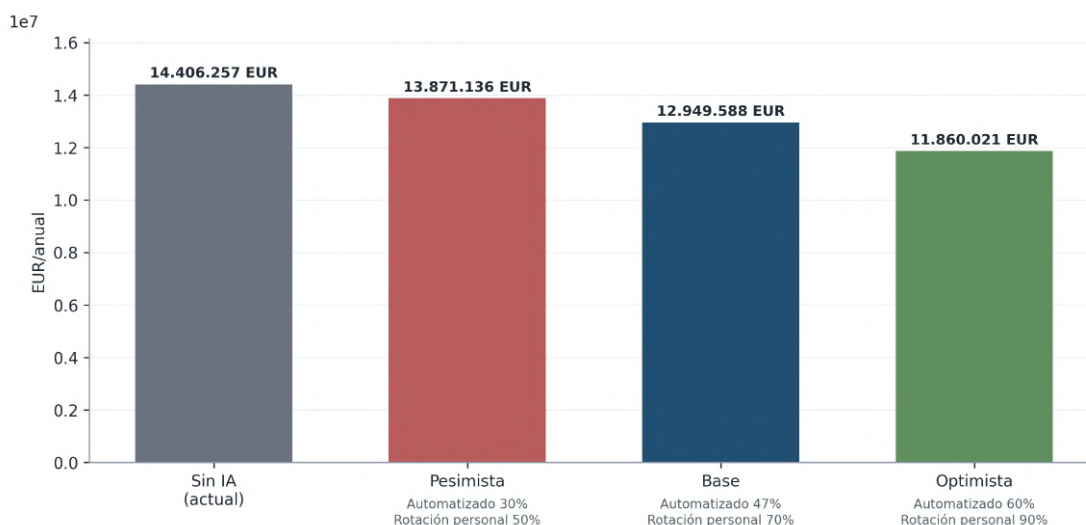


Figura 4.28: Comparativa del coste operativo anual entre el servicio únicamente humano y el servicio híbrido con ramp-up año 1.

En números, el gasto anual de la administración caería de 14,41 M€ a unos 12,95 M€, generando un ahorro limpio de casi **1,46 M€ anuales**. Incluso en el escenario más pesimista, el proyecto sigue ahorrando dinero (más de 500.000 €).

Pero como se ha querido ser realistas y que es muy difícil que la IA absorba el 47 % del tráfico desde el primer día, el modelo asume un escalado progresivo de 7 meses, donde la IA empieza atendiendo solo un 10 % del volumen y va subiendo. Si quitamos este escalado progresivo, y se miran los datos que podrían equivaler al segundo año de vida de este proyecto, el ahorro neto anual sería aún mayor debido a que los agentes podrían estar ocupados en otras tareas y no atribuir su nómina al servicio de atención telefónica como se puede ver en la Figura 4.29. Al igual que en el caso anterior, bajo cada escenario híbrido se indica el grado de automatización y la rotación del personal aplicada, incluyendo el coste de la infraestructura IA en las barras.

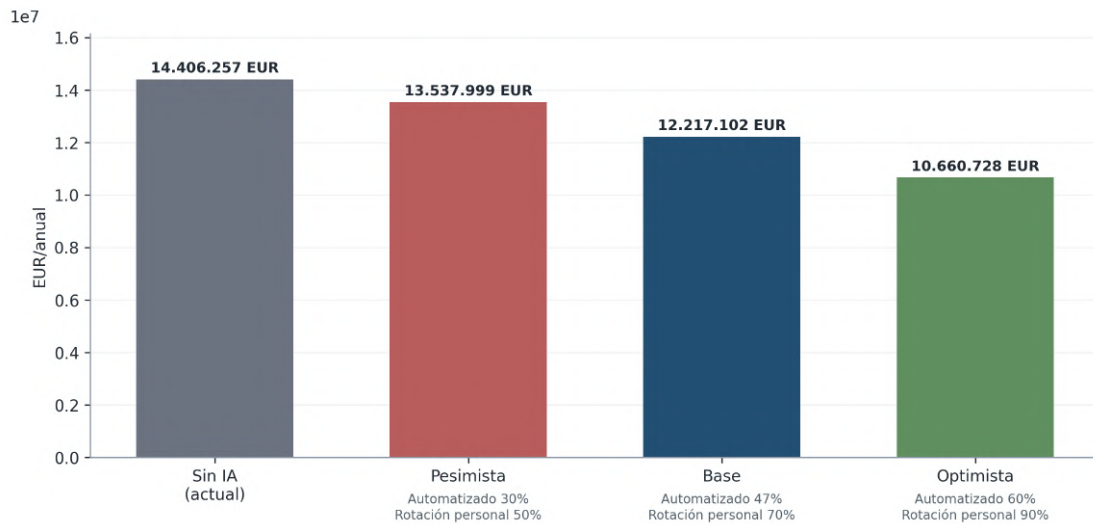


Figura 4.29: Comparativa del coste operativo anual sin escalado progresivo inicial.

En resumen, los datos muestran que montar una arquitectura de voz en cascada con modelos *Open Source* propios no solo es posible a nivel técnico, sino que operando con políticas de autoescalado, la rentabilidad económica podría ser muy positiva para escenarios de alto volumen de transacciones como la Línea Madrid 010.

Capítulo 5

Resultados y Discusión

Este capítulo presenta los resultados obtenidos al evaluar el sistema completo. A diferencia del Capítulo 4, donde los experimentos que se hicieron fueron evaluar y seleccionar el LLM en entornos aislados (y se consideró parte del desarrollo del sistema por ello y no se incluyó en este capítulo), este capítulo tiene como objetivo validar el comportamiento del sistema completo (end-to-end), desde la voz de entrada del usuario hasta la salida de audio del sistema.

En primer lugar, la Sección 5.1 presenta los resultados del asistente a la hora de resolver tareas reales usando el simulador conversacional automático que ya se explicó en la Sección 3.3 (que simula un cliente a través de una conversación por WebSockets y comprueba si el pipeline es capaz de realizar lo que el test indica). A continuación, la Sección 5.2 analiza el rendimiento temporal del sistema, presentando los resultados de la latencia introducida por cada uno de los bloques y midiendo el tiempo total de respuesta. La Sección 5.3 evalúa cómo de robusto es el sistema ante interrupciones del usuario (*barge-in*). Después, en la Sección 5.4 se aporta un análisis cualitativo basado en pruebas empíricas de interacción manual con la plataforma. Aquí no se mirarán números sino más bien evaluar la experiencia de uso percibida y detectar posibles limitaciones. Por último, la Sección 5.5 hace una breve conclusión general de los resultados e indica las restricciones tecnológicas observadas durante el proyecto.

5.1. Resultados del Benchmark *End-to-End* (E2E)

En esta sección se van a presentar los resultados obtenidos en el *benchmark End-to-End* (E2E) del pipeline. Este test sigue la metodología descrita en la Sección 3.3 pero como a breve resumen, este es un benchmark que permite evaluar el pipeline completo (*end-to-end*). Este test está pensado para evaluar todo el sistema completo desde la voz de entrada del usuario hasta la salida de audio del

sistema (aunque para el TTS es imposible hacer una evaluación automática ya que para ello hubiera sido necesario introducir otro ASR para obtener los resultados lo que hubiera introducido errores del ASR, por lo tanto se ha extraído directamente el texto y este es el que ha sido la entrada del evaluador). En este test se han creado 300 llamadas telefónicas test diferentes. Para simular al cliente que llama, se ha usado **Gemma 4 31B**, elegido por el buen rendimiento que demostró en los *benchmarks* de la Sección 4.2.5. Pero para que fuera lo más exacto posible, este cliente si que dispone del *thinking mode* activado para poder razonar sobre la tarea que tiene que realizar.

5.1.1. Rendimiento global del sistema

Como ya se explicó en la Sección 3.3, este benchmark evaluaba para cada test 3 capas:

- Tarea completada correctamente: 1 si el sistema completaba correctamente la tarea que se le pedía, 0 en caso contrario.
- Calidad de la interacción: 1 si la conversación fue natural y el agente cumplió con las instrucciones del *system prompt*, 0 en caso contrario.
- Validación completa: Para cada test, es el resultado de la operación AND de las dos anteriores y de no haber detectado un fallo conversacional crítico.

A continuación en la Figura 5.1 se pueden ver los resultados obtenidos de las 300 simulaciones de conversaciones, lo que se convirtió en 1329 turnos de conversación, es decir, cada test tuvo de media 4,43 turnos por cada llamada. Las métricas detalladas se pueden consultar en la Tabla C.1 del Anexo.

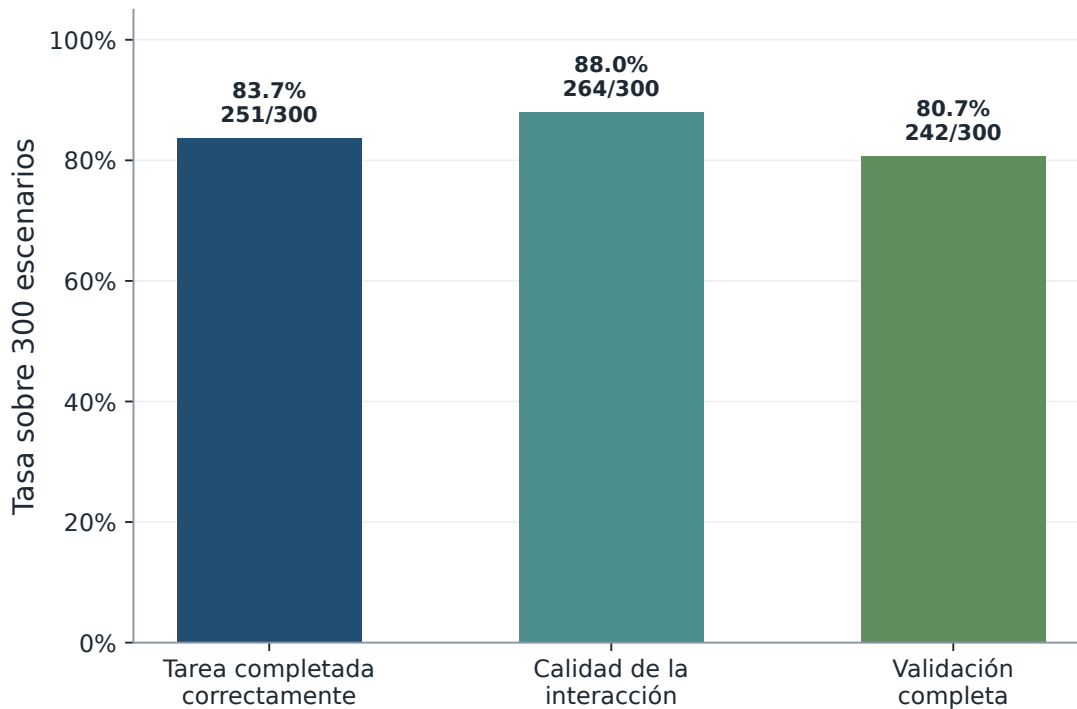


Figura 5.1: Rendimiento global del agente en el benchmark E2E (N=300).

Como se puede ver en la Figura 5.1, el sistema consigue un 1 en *Tarea completada correctamente* en un **83,7 %** de los casos (251 de 300). Esto significa que, en más de 8 de cada 10 llamadas, el agente fue capaz de entender la intención del usuario y ejecutar correctamente las herramientas necesarias en la base de datos. Por otro lado, la **Calidad de la interacción** alcanza un **88,0 %** (264 de 300), lo que indica que el asistente logra mantener un diálogo fluido y natural, y que es capaz de seguir las instrucciones del *system prompt* en la mayoría de escenarios.

Si nos fijamos en la métrica **Validación completa**, que es mucho más exigente ya que obliga a cumplir a la vez los dos valores (es decir, tarea correctamente y la calidad de la interacción sin fallos de conversación críticos), el porcentaje baja a **80,7 %** (242 de 300). La diferencia entre tarea, calidad y validación completa es importante: hay llamadas que suenan naturales pero no terminan dejando el estado correcto en la base de datos, y también hay llamadas donde la operación se completa pero aparece algún problema como por ejemplo que el agente entre en un bucle del que no sabe salir.

Este resultado puede parecer 'demasiado bueno' ya que es muy superior al que se obtuvo en los *benchmarks* individuales del LLM en 4.2, pero esto no se debe a que se esté haciendo una comparación injusta. En los *benchmarks* aislados el modelo se evaluaba con tareas generales sin mucho contexto, y dependía mucho de la correcta

interpretación del modelo para hacer la tarea correctamente. En este sistema final, en cambio, el dominio del problema está mucho más acotado: el *system prompt* define el rol del asistente, las reglas del restaurante, las herramientas disponibles y el comportamiento esperado. Esto a primera vista puede parecer que el LLM esté haciendo trampas, pero es precisamente lo que se haría en un despliegue real en producción, donde el objetivo no es que el LLM resuelva cualquier tarea genérica, sino que actúe muy bien dentro de un dominio concreto. Además, trabajos recientes muestran que la ubicación de instrucciones y ejemplos dentro del *prompt* puede afectar al rendimiento, y que colocarlos en el *system prompt* suele ser una configuración eficaz en muchos casos [52]. Por lo tanto, no se trata de una trampa, sino de una buena práctica que se haría en un entorno empresarial.

5.1.2. Rendimiento por intención conversacional

Si se analiza la Figura 5.2 (junto con la Tabla C.2 del Anexo), se puede ver qué tipo de tareas le resultan más fáciles o más complicadas al sistema implementado.

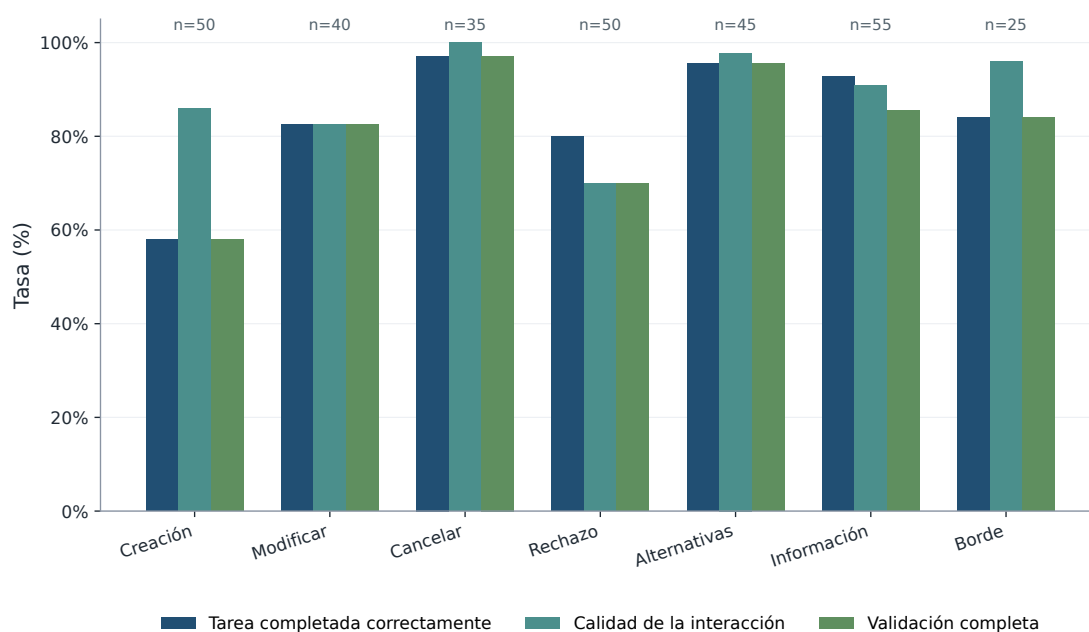


Figura 5.2: Tasa de tarea completada correctamente, calidad de interacción y validación completa por categoría de escenario.

Lo primero que destaca es que las operaciones que consisten en gestionar una reserva ya existente funcionan especialmente bien. La cancelación alcanza una **Validación completa del 97,1 %**, y los escenarios de alternativas llegan al **95,6 %**.

Esto confirma que, cuando el estado de partida está bien definido y el agente tiene que ejecutar una operación concreta, la lógica del agente y el *tool calling* dentro del *pipeline* son bastante robustos.

Los problemas aparecen sobre todo en la **Creación de reserva**, donde la Validación completa baja al **58,0 %**. Esta categoría es la más exigente porque para darla por válida, el agente tiene que reconocer correctamente muchos parámetros de manera precisa como son la fecha, hora, número de personas, nombre y teléfono y llamar a las herramientas adecuadas dejando la reserva en la base de datos. Lo que se ha visto inspeccionando los logs, muchas veces el error del agente viene de un error pequeño como puede ser en un dígito del teléfono o en un argumento de la herramienta y que por lo tanto, el test se marca como fallido aunque la conversación pueda parecer razonable (se ve como la calidad de la interacción está al nivel del resto de categorías). La categoría de **Rechazo** también queda por debajo de la media, con un **70,0 %** de Validación completa, principalmente por problemas de calidad de la conversación ya que la llamada se alarga demasiado (más de lo necesario).

5.1.3. Análisis de señales de fallo

Para entender mejor los 58 escenarios en los que no se logró la Validación completa, en la Figura 5.3 se muestran las señales de fallo agrupadas. Es importante decir que esta gráfica una misma llamada marcada como errónea puede tener varios fallos a la vez. Por ejemplo, si el ASR transcribe mal el teléfono, es normal que también aparezca como argumento de herramienta erróneo, reserva ausente o herramienta omitida. Por eso los fallos detectados suman más de 58.

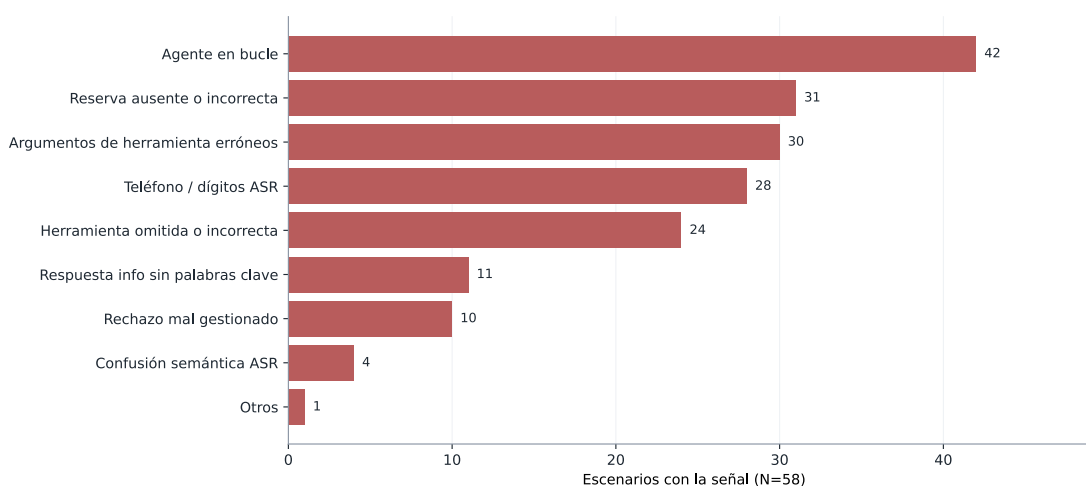


Figura 5.3: Señales de fallo en los escenarios sin Validación completa.

Como se puede ver, el principal problema no es que el sistema invente reservas, sino que algunas conversaciones entran en **bucles** (42 señales). Después aparecen los fallos de reserva ausente o incorrecta (31), argumentos de herramienta erróneos (30) y problemas con teléfonos o dígitos del ASR (28). Esto encaja con lo observado manualmente: el punto más frágil no es la arquitectura general, sino la combinación de ASR, extracción de datos y gestión de conversaciones que se atascan.

5.1.4. Análisis de alucinaciones

Una vez analizado cómo el sistema utiliza las herramientas, es necesario comprobar que el asistente no inventa datos, lo que se llama en el mundo de la inteligencia artificial como (*alucinaciones*). En un entorno empresarial, es necesario que lo que el agente confirma por voz y lo que realmente escribe en la base de datos sean iguales (los datos completos de esta prueba se pueden consultar en la Tabla C.5 del Anexo C.1).

Los resultados muestran que el sistema es muy seguro: consigue una **precisión del 100 % en coincidencia estricta**. Es decir, no hay ni un solo falso positivo donde el asistente se invente una reserva que no existe. El *recall* estricto queda en **80,6 %**, lo que demuestra que cuando el sistema falla, es por ser conservador (no crea la reserva si falta algún dato o duda), pero nunca por inventársela.

Por otro lado, al comparar lo que el modelo dice frente a lo que hace (**Verbal vs BD**), la precisión es del **91,3 %**. Solo hubo 9 casos donde el agente dijo por voz algo que no estaba en la base de datos, contra los 34 casos la llamada a la herramienta fue perfecta pero la confirmación hablada resultó ambigua. Esto confirma que la arquitectura en cascada cumple su objetivo principal: es controlable, auditable y prioriza no mentir al cliente.

5.2. Resultados de latencia del sistema

Para analizar la latencia del sistema en un entorno E2E, hay que saber que el rendimiento de los modelos cuando se evalúan de forma aislada (en los *benchmarks* aislados por componentes como se hizo para los LLMs en la Sección 4.2.4) y la latencia real de todo el pipeline funcionando son dos cosas diferentes. Esta suele ser mayor por la orquestación entera del sistema, pero también hay etapas que se solapan por el uso de *streaming*. Todos los datos estadísticos completos de estas mediciones se pueden consultar en las Tablas C.3 y C.4 del Anexo.

En los *benchmarks* individuales se mide cada bloque en condiciones ideales y aisladas: ASR solo, LLM solo o TTS solo, normalmente con entrada ya preparada, modelo *warm*, sin conversación completa, sin WebSockets, sin usuario sintético, sin herramientas y sin coordinación entre módulos. En el *benchmark* completo,

en cambio, se mide la llamada real siguiendo toda la ruta: audio del usuario, WebSocket, VAD/AEC, ASR, ITN, LLM, herramientas, TTS y audio de respuesta. Por eso aparecen latencias adicionales no tenidas en cuenta de sincronización, transporte, estabilidad de fin de habla en el ASR, normalización de textos, contexto de la conversación y llamadas a herramientas...

5.2.1. Resultados de las latencias del pipeline

El resultado principal de esta prueba de latencia es que el sistema empieza a emitir el primer audio del asistente en torno a **812.4 ms** (p50) después de que el usuario termine de hablar. Pero lo más interesante no es solo el número total, sino entender qué parte del pipeline lo está provocando. Para entenderlo mejor se ha hecho la Figura 5.4 que descompone la mediana hasta el primer audio en los principales bloques del sistema. El bloque que más tiempo ocupa es el ASR, con una mediana de **286 ms**. A esto hay que añadirle **274 ms** que viene del cierre de turno, VAD, ITN y orquestación interna. Después, el LLM y el enrutado añaden **132 ms**, y el TTS tarda **93 ms** hasta producir el primer audio. El resto de la latencia es muy pequeño, de unos **8 ms**. Los valores completos por percentiles se recogen en la Sección C.1 del Anexo.

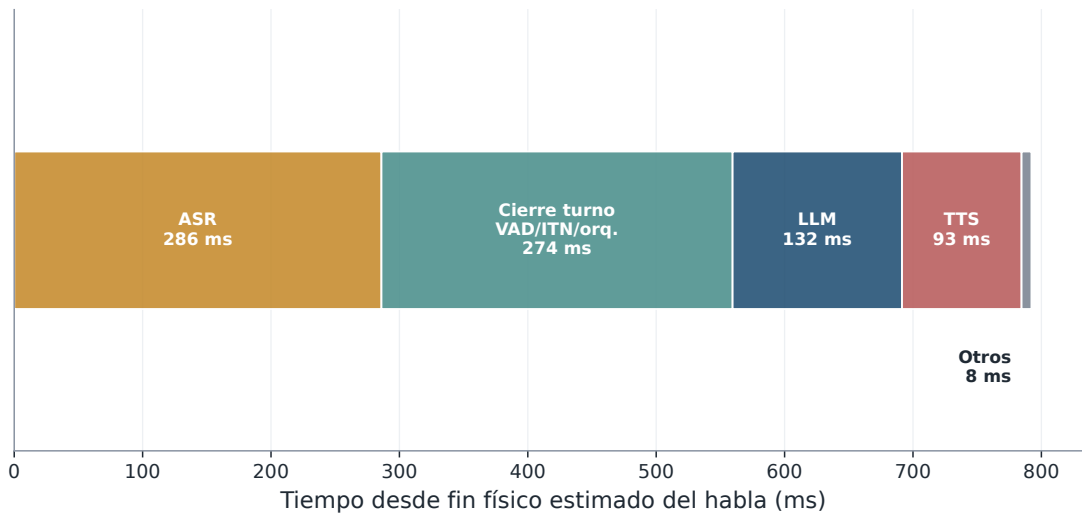


Figura 5.4: Composición mediana de la latencia hasta el primer fragmento de audio.

Que el ASR aparezca como el bloque más pesado no significa necesariamente que el modelo de reconocimiento sea lento de forma aislada. La razón principal es que el ASR está en el camino crítico del pipeline ya que el sistema no puede seguir procesando con el LLM hasta tener una transcripción estable y terminada.

En cambio, una vez que el texto ya está disponible, el resto del pipeline funciona en *streaming*. El LLM no necesita terminar toda la respuesta para que el TTS pueda empezar, y el TTS tampoco necesita sintetizar toda la frase antes de generar el primer fragmento de audio. Por eso, aunque LLM y TTS sean modelos grandes, su impacto sobre el primer audio es menor que el de estos bloques.

Esta diferencia se ve mejor en el diagrama temporal de la Figura 5.5. Mientras el usuario habla, bloques como AEC, VAD y ASR ya están trabajando en paralelo sobre el audio entrante. Pero la parte final del ASR y del cierre de turno sigue siendo necesaria antes de seguir procesando la respuesta. A partir de `user_final`, el sistema es más rápido porque el LLM genera los primeros tokens en *streaming* y el TTS empieza a producir audio sin esperar a que toda la respuesta esté completa.

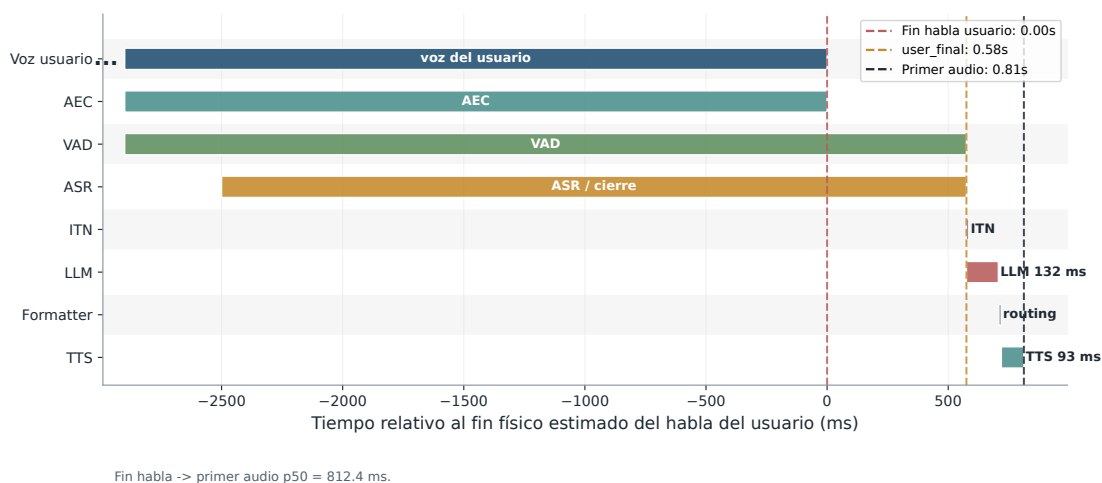


Figura 5.5: Diagrama temporal de funcionamiento en streaming de los diferentes bloques del pipeline.

Por último, la Figura 5.6 muestra la distribución por bloque. Se puede ver que el TTS es el componente más estable, con muy poca variación entre turnos, mientras que el bloque de cierre de turno y orquestación presenta más dispersión. Esto tiene sentido porque esa parte depende más de la forma concreta en la que termina cada audio, del VAD, de la estabilización de la transcripción y de las colas internas del sistema. En cambio, una vez que se llega al LLM y al TTS, el camino está más controlado y los tiempos son más constantes.

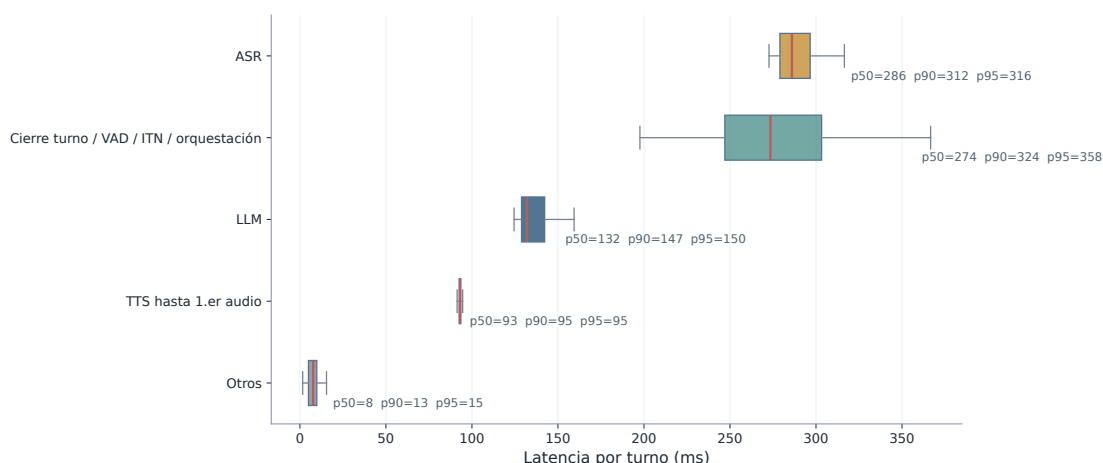


Figura 5.6: Distribución de latencias internas por bloque del pipeline.

En resumen, la latencia del sistema no viene principalmente del LLM ni del TTS, sino de la parte inicial del pipeline, donde se decide cuándo cerrar el turno del usuario y se termina la transcripción. Esto no implica que los bloques de audio funcionen mal, sino que en una arquitectura en cascada esta etapa es especialmente crítica porque antes de responder, el sistema necesita tener suficiente confianza en que el usuario ha terminado de hablar y en que el texto transcrito está terminado. Una vez superado ese punto, la generación de texto y la síntesis de voz aprovechan mejor el *streaming* y añaden relativamente poca latencia hasta el primer audio.

5.3. Tiempo de reacción y gestión de interrupciones

5.3.1. Resultados del benchmark de interrupciones

En esta sección se van a presentar los resultados de las pruebas realizadas para evaluar cómo el sistema maneja las interrupciones del usuario (lo que se conoce como *barge-in*) y cómo es capaz de cortar la voz del agente (el TTS) en tiempo real. La métrica clave aquí es el tiempo de reacción de parada (*Stop Reaction Time*), que resulta fundamental para que la conversación parezca natural y el asistente no siga hablando solo cuando el humano ya ha tomado su turno de palabra.

Como se ve en la Tabla [C.12](#) del Anexo, los resultados de esta prueba son muy buenos. Se registraron 2998 mediciones válidas sobre un total de 3005 mediciones previstas (lo que supone un 99,77 % de éxito en la ejecución). Los 7 casos que no se ejecutaron fueron porque no se encontraron los ficheros del audio, y eso es porque se borraron pero no se había actualizado la lista de archivos disponibles.

5.3.2. Análisis del tiempo de reacción

El objetivo de esta prueba es medir exactamente el tiempo que pasa desde que el cliente empieza a hablar, hasta que el servidor se da cuenta de que lo están interrumpiendo y este actualiza el estado de la conversación y corta el audio del TTS.

En la Tabla [C.13](#) del Anexo se puede ver cómo se distribuye esta latencia. El sistema consigue una mediana de tiempo de reacción (p50) de **83,9 ms**. Hay que tener en cuenta que, al hablar con una máquina, cualquier interrupción que tarde menos de 200 ms el usuario la percibe como casi instantánea. Además, el sistema se mantiene muy estable: en el 95 % de las ocasiones (p95), es capaz de reaccionar y callarse en 210,9 ms o menos.

Si se analiza de forma visual cómo se reparten estas mediciones en la Figura [5.7](#), se puede observar algo bastante curioso: la gráfica no forma una curva suave y continua (como una típica campana de Gauss), sino que tiene un comportamiento muy escalonado. Los tiempos de reacción se van agrupando en picos muy marcados (alrededor de los 45 ms, 85 ms, 105 ms y 145 ms), y entre ellos quedan 'bandas vacías' donde no hay datos (por ejemplo, hay pocos registros entre los 50 ms y los 75 ms).

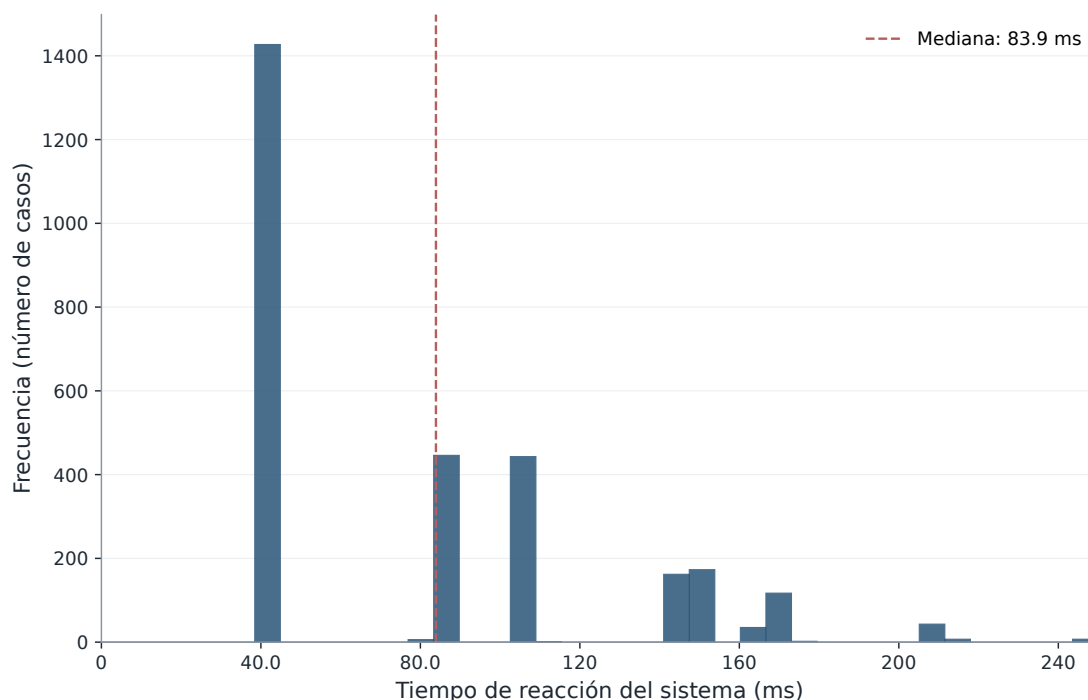


Figura 5.7: Histograma de frecuencias del tiempo de reacción ante interrupciones.

Estos saltos en la gráfica no son un error de medición, sino que reflejan directamente cómo está construida la arquitectura por dentro. Todo funciona por bloques de tiempo muy concretos:

1. **Umbral de confirmación del *barge-in*:** Para que un simple ruido de fondo no corte al asistente por error, el sistema exige escuchar al menos 30 ms seguidos de voz real antes de lanzar la orden de interrupción. Por lo tanto, es normal que no haya datos antes de los 30ms.
2. **Tamaño de trama del VAD:** Aunque el servidor alimenta el adaptador VAD en tramas de 10 ms, el modelo Silero no procesa el audio byte a byte ni decide sobre cada fragmento de 10 ms de forma independiente, y como funciona internamente es que va acumulando el audio hasta disponer de una ventana de 32 ms (lo que equivale a 512 muestras a 16 kHz), y es sobre esos 32ms calcula la probabilidad de voz en ese fragmento.
3. **Transmisión de red:** Por su parte, el cliente agrupa y envía el audio a través de la red en tramas de 20 ms.
4. **Silencios al inicio del audio:** El cronómetro de la prueba arranca justo en el primer byte del archivo. Si la persona no empieza a hablar instantáneamente y hay un pequeño silencio al principio, el sistema tiene que ir recorriendo tramas de 20 ms extra hasta que el VAD detecta la voz.

Con todo esto en cuenta, tiene sentido que el primer gran pico que se ve en la gráfica (sobre los 40-45 ms) represente el límite mínimo del sistema: es el caso ideal donde el usuario empieza a hablar al instante, sumado a los 30 ms obligatorios del umbral del VAD y el pequeño retraso del procesamiento en red. Los siguientes picos que aparecen más a la derecha son simplemente los retrasos que se van acumulando a base de sumar tramas extra del cliente o del VAD (y por eso se ven saltos 20 en 20 o de 10 en 10 ms).

5.3.3. Estabilidad a lo largo de la respuesta

Al construir un sistema de voz, hay otro factor que es muy importante comprobar y es si el sistema se vuelve más lento reaccionando cuando el agente lleva mucho rato hablando sin parar. Si esto ocurriera, nos estaría indicando que hay algún cuello de botella en los *buffers* de memoria o que el bucle de eventos asíncronos (*event loop*) de los WebSockets se está saturando.

Para comprobar que esto no pasa, las interrupciones en la prueba se lanzaron en momentos totalmente aleatorios. Como ve en la Tabla [C.14](#) del Anexo, los cortes se hicieron a lo largo de un espectro de tiempo muy amplio: a veces el usuario

interrumpía antes de que pasara el primer medio segundo de respuesta (p05: 492 ms, y esto simularía por ejemplo no tanto una interrupción, sino más bien que el usuario había hecho una parada al hablar), y otras veces lo hacía tras más de 7 segundos escuchando al agente (p95: 7633 ms).

En la Figura 5.8 se puede ver un gráfico que relaciona el tiempo que lleva hablado el sistema y el tiempo de reacción. Esto se ha hecho para ver si existe alguna relación entre el momento en el que el usuario corta al asistente y lo que tarda el sistema en reaccionar. El objetivo de esta representación es comprobar que no hay una degradación temporal a medida que la respuesta del agente se alarga.

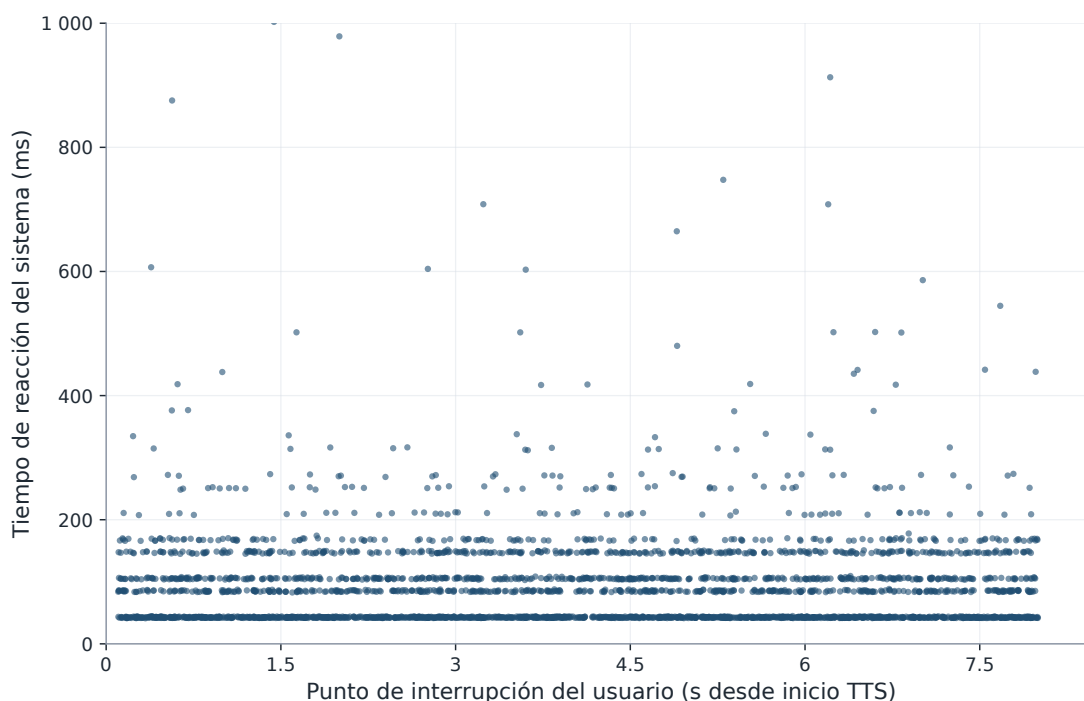


Figura 5.8: Gráfico de dispersión del tiempo de reacción frente al momento de interrupción del usuario.

Como se puede ver claramente en el gráfico, los puntos dibujan unas líneas horizontales muy marcadas (que, de nuevo, coinciden con esos saltos de tiempo que se han explicado antes) de principio a fin. El hecho de que la gráfica no vaya subiendo hacia la derecha demuestra que la latencia del sistema no empeora con el paso del tiempo. Por lo tanto, la lectura del websocket y el paro del TTS funciona igual de bien siempre, lo que garantiza que el asistente se va a callar de forma casi inmediata sin importar si el usuario le interrumpe justo cuando empieza a hablar o que lleve un rato diciendo cosas.

5.4. Evaluación cualitativa de la experiencia conversacional

5.4.1. Calidad y naturalidad de la voz (TTS)

Uno de los primeros puntos que hay que recalcar es que, en esta última versión del producto y en gran parte gracias a los servidores de ICAI, se ha podido utilizar un modelo TTS mucho más avanzado. Aunque la arquitectura del sistema ha sufrido pocas modificaciones en los últimos meses, la mejora en la calidad de la voz marca por completo la diferencia en la experiencia del usuario. A veces parece que el LLM es el bloque más crítico de todos pero eso es solo si no tiene una inteligencia suficiente para resolver la tarea, pero una vez se cumple ese requisito, a la hora de probar el sistema en la práctica, el cambio que más se ha notado para bien ha sido el del motor de voz.

Por supuesto, sigue teniendo puntos de mejora. Se trata de una voz con poca expresividad y plana y, aunque quizá no haga falta una gran expresividad para un asistente de reservas de restaurante, es un detalle a tener en cuenta. Posiblemente la falta de expresividad, al escuchar al asistente es fácil darse cuenta desde el principio de que se está hablando con una inteligencia artificial y no con una persona real.

5.4.2. Comprensión del usuario y reconocimiento (ASR)

En cuanto a la capacidad del sistema para entender al usuario, los fallos más comunes aparecen cuando la persona no vocaliza de forma clara. Pero el caso que más ha sorprendido durante las pruebas es la dificultad que tiene el ASR para transcribir números, sobre todo cuando se dicen muy rápido. Al sistema le cuesta mucho procesar varios dígitos iguales seguidos. Por ejemplo, si un usuario dice 'uno, uno, uno', es bastante probable que el reconocedor solo capte dos de esos tres 'unos'.

Quitando este detalle, si el cliente mantiene un ritmo de habla normal y vocaliza bien, el ASR funciona correctamente. De hecho, ha demostrado ser muy robusto incluso en escenarios más difíciles, como cuando hay ruido de fondo o cuando el usuario titubea o usa muletillas naturales de la voz (tipo 'ehhh...').

5.4.3. Percepción de la latencia en la interacción

Aunque la latencia ya se ha analizado de forma cuantitativa en la sección anterior, desde el punto de vista cualitativo hay que reconocer que sigue habiendo margen de mejora. En algunas interacciones puntuales el sistema se nota demasiado

lento, hasta el punto de que el usuario llega a dudar de si la llamada se ha cortado o si el asistente simplemente sigue 'pensando'.

A día de hoy, con la tecnología *Open Source* utilizada, parece casi imposible alcanzar el umbral óptimo de los 400 ms en todo el *pipeline*. Sin embargo, la experiencia de uso demuestra que la clave no es tanto la velocidad pura, sino la estabilidad. Si se consiguieran estabilizar los tiempos en torno a unos 600 ms, pero eliminando el *jitter* (es decir, evitando que la latencia varíe tanto de un turno a otro debido principalmente al uso de tools), se reduciría esa sensación de incertidumbre en el usuario y la interacción sería mucho más cómoda.

5.4.4. Capacidades agénticas y naturalidad del lenguaje

En el apartado del razonamiento, el LLM hace un trabajo muy bueno entendiendo lo que el cliente le pide, decidiendo en qué momento debe llamar a las herramientas y formulando las respuestas finales. Se ha logrado darle al asistente de una 'personalidad' bastante 'alegre' y 'amigable', lo que evita que la llamada parezca un interrogatorio o un simple formulario automático, acercándola mucho más a una conversación humana natural.

Como curiosidad, cabe destacar que en alguna prueba muy puntual la IA ha cometido errores de conjugación verbal. Aunque son casos muy raros y poco frecuentes, no se ha conseguido aislar el motivo exacto por el que ocurren, por lo que quedaría como un pequeño detalle interesante a investigar en trabajos futuros.

5.5. Discusión y limitaciones

Una de las principales limitaciones que tiene un sistema voz en cascada, a diferencia de los unificados, es que por diseño, la inteligencia del sistema (es decir el LLM) no recibe en ningún momento el audio del usuario. Esto en un primer momento puede parecer que no tiene mayor importancia ya que recibe la transcripción, que es su equivalente, pero esto no es del todo cierto ya que se pierde la dimensión de las emociones y del tono de voz que el usuario transmite al hablar. Esto puede ser un gran problema para tareas que requieran empatía o detección de emociones. Por ejemplo, en ningún caso esto se podría aplicar a una línea de atención telefónica de emergencias, ya que el sistema no podría detectar la urgencia o el estrés en la voz del usuario. Sin embargo, para tareas más rutinarias como reservas, consultas, atención ciudadana (como el caso del 010 estudiado) o información general, esta limitación no es tan importante, y seguramente con cierto trabajo en los system prompts, el LLM podría inferir el tono emocional del usuario a partir de la transcripción por el lenguaje usado (aunque claramente no será nunca tan preciso).

Por otro lado, otra limitación técnica importante a tener en cuenta es la dependencia de *hardware* de muy alta gama. Los tiempos de latencia tan buenos que se han conseguido en este proyecto han sido posibles, en gran parte, gracias a que se ha contado con el clúster de GPUs NVIDIA H200 de ICAI. Si este mismo sistema se desplegara en una empresa utilizando tarjetas gráficas más modestas o económicas (como las A40 que se plantearon en el estudio de viabilidad), es probable que el tiempo de respuesta aumentara, empeorando la latencia percibida al otro lado del teléfono.

También respecto a la validación del sistema, hay que ser conscientes de que tanto la evaluación cualitativa como el *benchmark End-to-End* se han realizado en un entorno de laboratorio controlado. Aunque el simulador automático ha demostrado ser muy estricto y útil para medir la ejecución técnica de las tareas, todavía faltaría dar el paso de desplegar la plataforma en una línea telefónica real para probarlo con usuarios humanos. Esto sería clave para comprobar de forma definitiva cómo se comporta el VAD y el ASR frente a condiciones acústicas más adversas, como mala cobertura, personas hablando con el manos libres desde un coche...

Por último, un aspecto que se ha probado poco pero que es muy importante a nivel empresarial, es la seguridad del sistema. Haciendo una pinta de trabajo de ingeniería de prompt, se ha conseguido que el agente de IA desobedezca a las instrucciones que se le habían dado en el system prompt. En un entorno de producción real habría que implementar guardarraíles y puede ser que sea necesario usar un modelo más grande e inteligente para que sea más complicado de engañar.

Capítulo 6

Conclusiones

Este capítulo presenta las principales conclusiones obtenidas del diseño, implementación y evaluación del sistema de voz en tiempo real. En primer lugar, la Sección [6.1](#) resume los resultados más relevantes y comprueba si se han cumplido los objetivos planteados al inicio del proyecto. A continuación, la Sección [6.2](#) detalla el valor técnico que ofrece este Trabajo de Fin de Máster desde el punto de vista de la ingeniería. Por último, la Sección [6.3](#) plantea los siguientes pasos lógicos para mejorar la plataforma.

6.1. Conclusiones principales

El desarrollo de este proyecto ha demostrado en la práctica que es viable construir un asistente de voz conversacional para entornos empresariales críticos utilizando exclusivamente modelos *Open Source* y una arquitectura modular en cascada. Frente a la tendencia actual de utilizar modelos multimodales nativos (*Speech-to-Speech*) que operan como cajas negras, este trabajo concluye que la separación del *pipeline* en bloques independientes (VAD, ASR, LLM, TTS...) ofrece el control necesario para integrar lógica de negocio y uso de herramientas (*tool calling*) sin comprometer la calidad de la interacción con el asistente.

A nivel de *software*, se ha comprobado que la gestión de la concurrencia es tan importante como la inteligencia del modelo. La implementación de una Máquina de Estados Finitos (FSM) controlada por variables de estado, combinada con el uso de `asyncio` y WebSockets, ha resultado ser una solución robusta para resolver uno de los mayores retos de las interfaces de voz: la gestión natural de los turnos de palabra y las interrupciones del usuario (*barge-in*).

A nivel de rendimiento temporal, el análisis ha revelado que la latencia extremo a extremo (*End-to-End*) real desde que el usuario termina de hablar hasta el primer audio se sitúa en una mediana de **812 ms**. Aunque este valor es superior

al umbral teórico ideal de 400 ms, la descomposición del *pipeline* ha demostrado que la inferencia de los modelos (LLM y TTS) es extremadamente rápida gracias al *streaming*. La latencia restante es inherente a la naturaleza de la arquitectura en cascada: la necesidad del VAD de acumular silencio para confirmar el cierre del turno y la estabilización del ASR. A cambio de esta ligera penalización temporal, se obtiene una enorme fiabilidad agéntica. El simulador conversacional demostró esta robustez al confirmar que el sistema fue capaz de entender la intención, extraer los parámetros precisos y ejecutar correctamente las operaciones en la base de datos de reservas en un **83,7%** de los escenarios planteados. Finalmente, a nivel económico, el dimensionamiento de la infraestructura usando la teoría de colas (Erlang B) ha demostrado que el despliegue local de el sistema desarrollado es rentable. Aplicando políticas de *autoscaling*, se concluye que el coste de la infraestructura *cloud* se amortiza gracias al ahorro directo en costes operativos de atención al cliente, generando un retorno de inversión positivo en menos de un año.

Este TFM demuestra que el estado del arte de la Inteligencia Artificial *Open Source* ya es suficientemente maduro para empezar a aparecer en entornos de prueba controlados y resolver problemas de negocio reales. Aunque los modelos *omni* o *Speech-to-Speech* comerciales van a seguir muy probablemente siendo mejores en términos de naturalidad, la arquitectura propuesta en esta memoria confirma que, con un diseño de un *pipeline* adecuado, es posible construir alternativas locales, privadas y rentables que mantienen un control sobre lo que el sistema dice y hace. 

6.2. Aportaciones del trabajo

Más allá de la integración de distintas tecnologías de Inteligencia Artificial, este Trabajo de Fin de Máster ofrece un valor específico desde el punto de vista de la ingeniería en tres áreas principales:

- **Diseño de un sistema de voz asíncrono y robusto:** Se ha desarrollado una arquitectura de *backend* en Python que no se bloquea durante la inferencia pesada de los modelos de IA. El diseño de la máquina de estados aporta una solución real al problema de las interrupciones en llamadas telefónicas, permitiendo que el sistema aborte la síntesis de voz y mantenga la coherencia de la conversación si el usuario decide hablar por encima de la IA.
- **Framework de evaluación híbrido automatizado:** Se ha diseñado un simulador capaz de validar el asistente generando clientes sintéticos mediante LLMs y síntesis de voz (TTS). El gran aporte de este entorno es que no

¹Para ver el funcionamiento real del sistema, se ha preparado una pequeña demo del sistema atendiendo una llamada de teléfono a la que se puede acceder en: <https://youtu.be/ggq2c-R6Pjg>

se limita a usar un LLM como juez subjetivo para evaluar si la charla era de buena calidad, sino que implementa una capa de validación determinista que inspecciona directamente el estado del *backend*, verificando si las transacciones en la base de datos y los argumentos de las herramientas coinciden exactamente con el objetivo del escenario.

- **Frontend para interacción con el sistema y monitorización:** Se ha construido un *frontend* completo que permite mantener una llamada en tiempo real, visualizar la máquina de estados, el uso de herramientas y las métricas de latencia todo en tiempo real. Esto convierte la arquitectura de *backend* en un producto funcional, fácil de usar y probar.
- **Modelo económico de viabilidad realista:** En lugar de asumir que un asistente de IA sustituye 1 a 1 a un operador humano, el proyecto plantea un modelo financiero más realista. Integra la fórmula de Erlang B para calcular canales concurrentes, aplica el factor de rotación del personal (κ) para los ajustes de plantilla en un *Call Center*, y demuestra cómo perfilar los costes de las GPUs basándose en el tráfico real, con datos de la Línea Madrid 010.

6.3. Trabajos futuros

Aunque el sistema desarrollado es funcional y resuelve el caso de uso planteado, se puede mejorar siguiendo las siguientes líneas de trabajo:

- **Integración con pasarelas de telefonía (SIP):** Actualmente, la interacción se realiza mediante WebSockets a través de un cliente web. El paso natural para un entorno de *Call Center* es integrar el *backend* con protocolos SIP (Session Initiation Protocol) utilizando pasarelas como Asterisk o servicios como Twilio, lo que permitiría recibir llamadas directamente desde la red telefónica pública real.
- **Optimización acústica frente a ruido extremo:** Aunque el sistema VAD implementado (Silero) es robusto, en llamadas de teléfono reales con escenarios más complicados como al aire libre, con manos libres, con el ruido de fondo o con varias voces, pueden generar falsos positivos. Sería interesante explorar modelos de separación de fuentes de audio (*Speaker Diarization*) en tiempo real antes de introducir la señal al ASR.
- **Arquitectura multi-agente:** Para casos de uso con lógicas de negocio más complejas que las de un restaurante, el uso de distintos agentes LLMs optimizaría la eficiencia del sistema al evitar sobrecargar a un único modelo de

lenguaje con demasiadas herramientas (*tools*) e instrucciones simultáneas. Se podría plantear una arquitectura donde un agente general, se encargue de identificar la intención del usuario y enrute la llamada hacia un LLM especializado en una tarea concreta, pasándole el contexto de la conversación para que este la continúe. Por ejemplo, en el servicio de atención al cliente de una aerolínea, el enrutador inicial derivaría la conversación a un agente experto en venta de billetes, a otro especializado en cancelaciones o a un tercero enfocado en consultas de equipaje, reduciendo así la latencia y la posibilidad de alucinaciones.

- **Mejora de la expresividad en el TTS:** Si bien los modelos *Open Source* actuales en español cumplen su función de manera aceptable, la clonación de voz y la capacidad de mostrar emociones (pausas, risas, entonación interrogativa) siguen estando un paso por detrás de alternativas comerciales cerradas. Explorar técnicas de *fine-tuning* o adaptaciones de la arquitectura para controlar mejor estos matices mejoraría mucho la expresividad del asistente.
- **Recuperación del tono emocional del usuario:** Como se vio en las limitaciones, al usar una arquitectura en cascada el LLM solo recibe texto y pierde matices como la urgencia o el estado de ánimo de la persona. Una futura línea de investigación sería añadir un modelo ligero de clasificación de emociones de audio (paralelo al ASR) que inyecte esa información como contexto adicional en el *prompt*, permitiendo al agente adaptar su nivel de empatía.
- **Optimización para hardware más modesto:** Los tiempos de latencia tan bajos conseguidos se apoyan en el uso de GPUs de muy alta gama (NVIDIA H200). Para democratizar el despliegue de estos sistemas en empresas más pequeñas, sería necesario investigar el uso de técnicas de cuantización más extremas (como INT4 o inferior) o motores de inferencia con *speculative decoding*, buscando mantener la latencia a raya en tarjetas gráficas mucho más baratas.
- **Implementación de guardarraíles de seguridad (*Guardrails*):** Dado que en las pruebas manuales se comprobó que el agente puede ser engañado mediante *prompt injection*, un paso obligatorio antes de llevar el sistema a producción sería integrar capas de seguridad adicionales. Utilizar herramientas como NeMo Guardrails o Llama Guard permitiría filtrar las entradas del usuario y bloquear cualquier intento de hacer que el asistente se salte sus reglas de negocio o use herramientas para fines no previstos.

Bibliografía

- [1] UBS Chief Investment Office. *Has the AI rally gone too far?* <https://www.ubs.com/global/en/wealthmanagement/insights/chief-investment-office/house-view/daily/2023/latest-25052023.html>. UBS House View Daily. Mayo de 2023.
- [2] Chun-Yi Kuan, Wei-Ping Huang y Hung-yi Lee. «Understanding Sounds, Missing the Questions: The Challenge of Object Hallucination in Large Audio-Language Models». En: *Proceedings of Interspeech 2024*. ISCA, 2024, págs. 4144-4148. DOI: [10.21437/Interspeech.2024-1076](https://doi.org/10.21437/Interspeech.2024-1076).
- [3] European Parliament and Council of the European Union. *Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation)*. <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>. OJ L 119, 4.5.2016, pp. 1–88. 2016.
- [4] European Data Protection Board. *Guidelines 02/2021 on Virtual Voice Assistants, Version 2.0*. https://www.edpb.europa.eu/system/files/2021-07/edpb_guidelines_202102_on_vva_v2.0_adopted_en.pdf. Adopted on 7 July 2021. Jul. de 2021.
- [5] Siddhant Arora y col. «On The Landscape of Spoken Language Models: A Comprehensive Survey». En: *arXiv preprint arXiv:2504.08528* (2025). DOI: [10.48550/arXiv.2504.08528](https://doi.org/10.48550/arXiv.2504.08528).
- [6] OpenAI. *API Pricing*. Accessed on 2026-04-03. 2026. URL: <https://openai.com/api/pricing/> (visitado 03-04-2026).
- [7] OpenAI. *Managing Costs*. <https://developers.openai.com/api/docs/guides/realtime-costs/>. OpenAI API documentation, accessed on 2026-04-03. 2026. (Visitado 03-04-2026).

- [8] Artificial Analysis. *LLM Leaderboard - Comparison of over 100 AI models from OpenAI, Google, DeepSeek & others*. Accessed on 2026-04-03. 2026. URL: <https://artificialanalysis.ai/leaderboards/models> (visitado 03-04-2026).
- [9] Z.AI. *Pricing - Z.AI Developer Document*. Accessed on 2026-04-03. 2026. URL: <https://docs.z.ai/guides/overview/pricing> (visitado 03-04-2026).
- [10] Anthropic. *Claude Opus 4.6*. Accessed on 2026-04-03. Feb. de 2026. URL: <https://www.anthropic.com/claude/opus> (visitado 03-04-2026).
- [11] Gabriel Skantze. «Turn-taking in Conversational Systems and Human-Robot Interaction: A Review». En: *Computer Speech & Language* 67 (mayo de 2021), pág. 101178. DOI: [10.1016/j.csl.2020.101178](https://doi.org/10.1016/j.csl.2020.101178). URL: <https://www.sciencedirect.com/science/article/pii/S088523082030111X>.
- [12] Divesh Lala, Shizuka Nakamura y Tatsuya Kawahara. «Analysis of Effect and Timing of Fillers in Natural Turn-Taking». En: *Proceedings of Interspeech 2019*. ISCA, 2019, págs. 4175-4179. DOI: [10.21437/Interspeech.2019-1527](https://doi.org/10.21437/Interspeech.2019-1527). URL: https://www.isca-archive.org/interspeech_2019/lala19_interspeech.html.
- [13] Siddhant Arora y col. *Stream RAG: Instant and Accurate Spoken Dialogue Systems with Streaming Tool Usage*. 2025. DOI: [10.48550/arXiv.2510.02044](https://doi.org/10.48550/arXiv.2510.02044). arXiv: [2510.02044](https://arxiv.org/abs/2510.02044) [cs.CL]. URL: <https://arxiv.org/abs/2510.02044>.
- [14] Daniel Jurafsky y James H. Martin. *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition with Language Models*. 3rd. Online manuscript released August 24, 2025. 2025. URL: <https://web.stanford.edu/~jurafsky/slp3/>.
- [15] OpenAI. *GPT-4o: Omni-model capabilities in voice, vision, and text*. <https://openai.com/index/hello-gpt-4o/>. Technical Report. Mayo de 2024.
- [16] Alexandre Défossez y col. *Moshi: a speech-text foundation model for real-time dialogue*. 2024. arXiv: [2410.00037](https://arxiv.org/abs/2410.00037) [eess.AS]. URL: <https://arxiv.org/abs/2410.00037>.
- [17] OpenAI. *Pricing - OpenAI Realtime API*. <https://openai.com/api/pricing/>. Accessed: 2025-11-23. 2025.
- [18] Google AI for Developers. *Gemini API Pricing*. <https://ai.google.dev/gemini-api/docs/pricing>. Accessed: 2025-11-23. 2025.

-
- [19] Junjie Chen y col. *FireRedChat: A Pluggable, Full-Duplex Voice Interaction System with Cascaded and Semi-Cascaded Implementations*. 2025. arXiv: [2509.06502 \[cs.SD\]](https://arxiv.org/abs/2509.06502). URL: <https://arxiv.org/abs/2509.06502>.
- [20] Walter J. Doherty y Ahrvind J. Thadhani. *The Economic Value of Rapid Response Time*. Technical Report GE20-0752-0. IBM, nov. de 1982. URL: <https://www.computerhistory.org/collections/catalog/102751398/>.
- [21] AWS Builders. *Benchmarking Amazon Bedrock LLM Latency: A Multi-Model Comparison*. El artículo compara la latencia de seis modelos de lenguaje en Amazon Bedrock y proporciona medias de TTFT y tiempo total de respuesta. 2025. URL: <https://builder.aws/benchmarking-amazon-bedrock-llm-latency>.
- [22] vLLM Team. *MiniMax M2.5 inference guide*. <https://vllm.ai/docs/MiniMax-M2.5>. Documento de vLLM con especificaciones de MiniMax M2.5, indicando que tiene 230 mil millones de parámetros totales y 10 mil millones activos con una ventana de contexto de 200K tokens. 2025.
- [23] Qi Huang, Lei Dong y col. «Dynamic Expert Quantization for Scalable Mixture-of-Experts Inference». En: *arXiv e-prints* (2024). Describe la arquitectura MOE de Qwen3 80B con 512 expertos, 10 activos y un conteo de parámetros activos de aproximadamente 3 mil millones. URL: <https://arxiv.org/abs/2511.15015>.
- [24] DeepSeek AI. *DeepSeek V3.2 model*. Página de modular.com que describe el modelo DeepSeek V3.2 como un MOE de 685 mil millones de parámetros totales con 37 mil millones activos y ventana de contexto de 128K tokens. 2025. URL: <https://modular.com/models/deepseek-v3-2>.
- [25] Dhruv Jain y col. *VoiceAgentBench: Are Voice Assistants ready for agentic tasks?* 2025. DOI: [10.48550/arXiv.2510.07978](https://doi.org/10.48550/arXiv.2510.07978). arXiv: [2510.07978 \[cs.AI\]](https://arxiv.org/abs/2510.07978). URL: <https://arxiv.org/abs/2510.07978>.
- [26] Jeffrey Zhou y col. «Instruction-Following Evaluation for Large Language Models». En: *arXiv preprint arXiv:2311.07911* (2023).
- [27] Jiahui Geng y col. «Speech Robust Bench: A Robustness Benchmark for Speech Recognition». En: *arXiv preprint* (2024). Referenced in Survey of Benchmarks for Spanish and Multimodal Evaluation.
- [28] Shishir Yan y col. «Berkeley Function-Calling Leaderboard: A Comprehensive Evaluation of Tool Calling». En: *Berkeley AI Research* (2024). URL: <https://gorilla.cs.berkeley.edu/leaderboard.html>.
- [29] Dhruv Jain y col. «VoiceBench: Benchmarking LLM-based Voice Assistants». En: *arXiv preprint arXiv:2410.00000* (2024).

- [30] ElevenLabs. *Speech to Text*. Accessed: 2026-05-30. 2026. URL: <https://elevenlabs.io/speech-to-text> (visitado 30-05-2026).
- [31] Taalas AI. *Taalas HC1*. Accessed on 2026-04-03. 2026. URL: <https://taalas.com/products/> (visitado 03-04-2026).
- [32] Cerebras Systems. *Cerebras Inference*. Accessed on 2026-05-29. 2026. URL: <https://inference.cerebras.ai/> (visitado 29-05-2026).
- [33] Amey Agrawal y col. «Taming Throughput-Latency Tradeoff in LLM Inference with Sarathi-Serve». En: *18th USENIX Symposium on Operating Systems Design and Implementation (OSDI 24)*. USENIX Association, 2024, págs. 117-134. URL: <https://www.usenix.org/conference/osdi24/presentation/agrawal>.
- [34] Reiner Pope y col. «Efficiently Scaling Transformer Inference». En: *Proceedings of Machine Learning and Systems*. Vol. 5. 2023, págs. 606-624. URL: https://proceedings.mlsys.org/paper_files/paper/2023/file/c4be71ab8d24cdfb45e3d06dbfca2780-Paper-mlsys2023.pdf.
- [35] OpenAI. *What are tokens and how to count them*. Accessed: 2026-06-02. 2026. URL: <https://help.openai.com/es-419/articles/4936856-what-are-tokens-and-how-to-count-them>.
- [36] Silero Team. *Silero VAD: Pre-trained Enterprise-grade Voice Activity Detector*. <https://github.com/snakers4/silero-vad>. GitHub Repository. 2024.
- [37] Max McKinnon y col. *Window Size Versus Accuracy Experiments in Voice Activity Detectors*. Google LLC. 2026. arXiv: [2601.17270 \[cs.SD\]](https://arxiv.org/abs/2601.17270), URL: <https://arxiv.org/abs/2601.17270>.
- [38] NVIDIA. *nvidia/parakeet-tdt-0.6b-v3: Multilingual Speech-to-Text Model*. <https://huggingface.co/nvidia/parakeet-tdt-0.6b-v3>. Accessed: 2026-05-30. 2025.
- [39] Monica Sekoyan y col. *Canary-1B-v2 & Parakeet-TDT-0.6B-v3: Efficient and High-Performance Models for Multilingual ASR and AST*. 2025. DOI: [10.48550/arXiv.2509.14128](https://arxiv.org/abs/2509.14128), arXiv: [2509.14128 \[cs.CL\]](https://arxiv.org/abs/2509.14128), URL: <https://arxiv.org/abs/2509.14128>.
- [40] Julian Mack y col. *cohere-transcribe-03-2026*. 2026. DOI: [10.57967/hf/8653](https://huggingface.co/CohereLabs/cohere-transcribe-03-2026), URL: <https://huggingface.co/CohereLabs/cohere-transcribe-03-2026>.

- [41] Vaibhav Srivastav y col. *Open ASR Leaderboard: Towards Reproducible and Transparent Multilingual and Long-Form Speech Recognition Evaluation*. 2025. DOI: [10.48550/arXiv.2510.06961](https://doi.org/10.48550/arXiv.2510.06961), arXiv: [2510.06961](https://arxiv.org/abs/2510.06961) [cs.CL], URL: <https://arxiv.org/abs/2510.06961>.
- [42] Xian Shi y col. *Qwen3-ASR Technical Report*. 2026. arXiv: [2601.21337](https://arxiv.org/abs/2601.21337) [cs.CL], URL: <https://arxiv.org/abs/2601.21337>.
- [43] Hangrui Hu y col. *Qwen3-TTS Technical Report*. 2026. arXiv: [2601.15621](https://arxiv.org/abs/2601.15621), URL: <https://arxiv.org/abs/2601.15621>.
- [44] Huakang Chen y col. *MINT-Bench: A Comprehensive Multilingual Benchmark for Instruction-Following Text-to-Speech*. 2026. arXiv: [2604.17958](https://arxiv.org/abs/2604.17958), URL: <https://arxiv.org/abs/2604.17958>.
- [45] alblez. *qwen3-tts-spanish-voices: Curated Spanish TTS voices powered by Qwen3-TTS*. GitHub repository. 2026. URL: <https://github.com/alblez/qwen3-tts-spanish-voices>.
- [46] Ayuntamiento de Madrid. Dirección General de Atención a la Ciudadanía. *Memoria anual calidad Línea Madrid 2024*. Documento PDF. Datos del ejercicio 2024. Mar. de 2025. URL: <https://www.madrid.es/UnidadesDescentralizadas/AtencionCiudadano/Calidad/EvaluacionesYEstudios/Ficheros/MemoriaCalidad2024.pdf> (visitado 30-05-2026).
- [47] Ayuntamiento de Madrid. Dirección General de Atención a la Ciudadanía. *Memoria estudio económico del contrato de servicios “Apoyo a la gestión de los canales de atención a la ciudadanía de Línea Madrid”*. Documento PDF. Publicado el 17 de agosto de 2021. Ago. de 2021. URL: <https://www.madrid.es/UnidadWeb/Contenidos/Navegaciones/LineaMadrid/Ficheros/MEMORIAIm21.pdf> (visitado 30-05-2026).
- [48] Ayuntamiento de Madrid. *Madrid en Cifras 2025: datos consolidados 2023*. Documento PDF. Mar. de 2025. URL: https://wemadrid.es/mi_contenido/uploads/2025/07/MEC2023_ESP_digital.pdf (visitado 30-05-2026).
- [49] RunPod. *RunPod GPU Pricing*. Accessed: 2026-05-31. 2026. URL: <https://www.runpod.io/pricing>.
- [50] *Estadísticas 2009: Atención Telefónica 010*. Informe estadístico anual. Fuente utilizada para estimar la distribución horaria histórica de llamadas del teléfono 010. Incluye la tabla “Total llamadas. Por tramo horario”, con intentos de llamada, llamadas rechazadas por saturación, llamadas recibidas, abandonadas y atendidas para cada franja horaria. Dirección General de Calidad y Atención al Ciudadano, Subdirección General de Atención al Ciudadano, Ayuntamiento de Madrid, 2009. URL: <https://www.madrid.es/UnidadWeb/>

- Contenidos/EspecialInformativo/LineaMadrid/Files/2009.010.pdf (visitado 31-05-2026).
- [51] Twilio. *SIP Trunking Pricing*. Accessed: 2026-05-31. 2026. URL: <https://www.twilio.com/en-us/sip-trunking/pricing/es> (visitado 31-05-2026).
- [52] Umut Halil y col. «LLM Shots: Best Fired at System or User Prompts?» En: *Companion Proceedings of the ACM on Web Conference 2025*. WWW Companion '25. New York, NY, USA: Association for Computing Machinery, 2025, págs. 1605-1613. DOI: [10.1145/3701716.3717814](https://doi.org/10.1145/3701716.3717814). URL: <https://doi.org/10.1145/3701716.3717814>.
- [53] Google. *gemma-4-31B-it - Model Card*. Accessed: 2026-05-15. 2026. URL: <https://huggingface.co/google/gemma-4-31B-it>.
- [54] Google DeepMind. *Gemma 4: Our most capable open models to date*. Accessed: 2026-05-15. 2026. URL: <https://blog.google/innovation-and-ai/technology/developers-tools/gemma-4/>.
- [55] Qwen Team. *Qwen3.6-27B - Model Card*. Accessed: 2026-05-15. 2026. URL: <https://huggingface.co/Qwen/Qwen3.6-27B>.
- [56] Qwen Team. *Qwen Blog: Advanced Multimodal and Agentic Capabilities*. Accessed: 2026-05-15. 2026. URL: <https://qwen.ai/blog>.
- [57] Qwen Team. *Qwen3.6-35B-A3B - Model Card*. Accessed: 2026-05-15. 2026. URL: <https://huggingface.co/Qwen/Qwen3.6-35B-A3B>.
- [58] Google. *gemma-4-E4B / E2B - Model Cards*. Accessed: 2026-05-15. 2026. URL: <https://huggingface.co/google/gemma-4-E4B>.
- [59] NVIDIA. *NVIDIA-Nemotron-Nano-12B-v2-VL-BF16 - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/nvidia/NVIDIA-Nemotron-Nano-12B-v2-VL-BF16>.
- [60] AI.rs. *Qwen 3.5: 35B Knowledge at 4B Speed — Better Than GPT-5?* Accessed: 2026-05-21. 2026. URL: <https://ai.rs/ai-for-business/qwen-3-5-35b-knowledge-4b-speed-better-than-gpt-5>.
- [61] Awesome Agents. *Qwen3.5-4B*. Accessed: 2026-05-21. 2026. URL: <https://awesomeagents.ai/models/qwen-3-5-4b/>.
- [62] Google. *gemma-4-26B-A4B-it - Model Card*. Accessed: 2026-05-15. 2026. URL: <https://huggingface.co/google/gemma-4-26B-A4B-it>.
- [63] LGAI-EXAONE. *EXAONE-4.0-32B - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/LGAI-EXAONE/EXAONE-4.0-32B>.
- [64] FriendliAI. *LG AI Research Partners with FriendliAI to Launch EXAONE 4.0*. Accessed: 2026-05-15. 2025. URL: <https://friendli.ai/blog/lg-ai-research-partnership-exaone-4.0>.

-
- [65] Awesome Agents. *Qwen3.5-9B*. Accessed: 2026-05-21. 2026. URL: <https://awesomeagents.ai/models/qwen-3-5-9b/>.
- [66] XDA Developers. *Qwen3.5-9B tops every AI benchmark right now, but that's not how you should pick a model*. Accessed: 2026-05-21. 2026. URL: <https://www.xda-developers.com/qwen-3-5-9b-tops-ai-benchmarks-not-how-pick-model/>.
- [67] LLM Reference. *Granite 4.0 H Small – 128K context, open source*. Accessed: 2026-05-15. 2025. URL: <https://www.llmreference.com/model/granite-4.0-h-small>.
- [68] Puter Developer. *LFM2-24B-A2B - Specs & API*. Accessed: 2026-05-15. 2026. URL: <https://developer.puter.com/ai/liquid/lfm-2-24b-a2b/>.
- [69] inclusionAI. *Ling-mini-2.0 - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/inclusionAI/Ling-mini-2.0-GGUF>.
- [70] Richard Bian. *Ling-mini-2.0: Mini-Sized, Maximum Efficiency*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/blog/RichardBian/ling-mini-2-moe-release>.
- [71] Liquid AI. *LFM2.5-VL-1.6B - Model Card*. Accessed: 2026-05-15. 2026. URL: <https://huggingface.co/LiquidAI/LFM2.5-VL-1.6B>.
- [72] Puter Developer. *Granite 4.0 Micro - Specs & API*. Accessed: 2026-05-15. 2025. URL: <https://developer.puter.com/ai/ibm-granite/granite-4.0-h-micro/>.
- [73] SiliconFlow. *Qwen3-Omni-30B-A3B-Instruct - Model Info*. Accessed: 2026-05-15. 2025. URL: <https://www.siliconflow.com/models/qwen3-omni-30b-a3b-instruct>.
- [74] Awesome Agents. *Qwen3.5-2B*. Accessed: 2026-05-21. 2026. URL: <https://awesomeagents.ai/models/qwen-3-5-2b/>.
- [75] Puter Developer. *LFM2.5-1.2B-Instruct - Specs & API*. Accessed: 2026-05-15. 2026. URL: <https://developer.puter.com/ai/liquid/lfm-2.5-1.2b-instruct/>.
- [76] Mistral AI. *Devstral-Small-2-24B-Instruct - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/mistralai/Devstral-Small-2-24B-Instruct-2512>.
- [77] Mistral AI. *Introducing: Devstral 2 and Mistral Vibe CLI*. Accessed: 2026-05-15. 2025. URL: <https://mistral.ai/news/devstral-2-vibe-cli>.
- [78] Puter Developer. *LFM2-2.6B - Specs & API*. Accessed: 2026-05-15. 2025. URL: <https://developer.puter.com/ai/liquid/lfm-2.2-6b/>.

- [79] NVIDIA. *NVIDIA-Nemotron-Nano-9B-v2 - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/nvidia/NVIDIA-Nemotron-Nano-9B-v2>.
- [80] LLM Reference. *Granite 4.0 H 1B - Specs*. Accessed: 2026-05-15. 2025. URL: <https://www.llmreference.com/model/granite-4.0-h-1b>.
- [81] NVIDIA. *NVIDIA-Nemotron-3-Nano-30B-A3B-BF16 - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/nvidia/NVIDIA-Nemotron-3-Nano-30B-A3B-BF16>.
- [82] Microsoft. *Phi-4-mini-instruct - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/microsoft/Phi-4-mini-instruct>.
- [83] LLM Reference. *Granite 4.0 1B - Specs*. Accessed: 2026-05-15. 2025. URL: <https://www.llmreference.com/model/granite-4.0-1b>.
- [84] Microsoft. *phi-4 - Model Card*. Accessed: 2026-05-15. 2024. URL: <https://huggingface.co/microsoft/phi-4>.
- [85] LGAI-EXAONE. *EXAONE-4.0-1.2B - Model Card*. Accessed: 2026-05-15. 2025. URL: <https://huggingface.co/LGAI-EXAONE/EXAONE-4.0-1.2B>.
- [86] Swiss AI. *Apertus-8B-Instruct-2509*. Accessed: 2026-06-02. 2025. URL: <https://huggingface.co/swiss-ai/Apertus-8B-Instruct-2509>.
- [87] Puter Developer. *LFM2-8B-A1B - Specs & API*. Accessed: 2026-05-15. 2025. URL: <https://developer.puter.com/ai/liquid/lfm2-8b-a1b/>.
- [88] LLM Reference. *Granite 4.0 H 350M - Specs*. Accessed: 2026-05-15. 2025. URL: <https://www.llmreference.com/model/granite-4.0-h-350m>.
- [89] LLM Reference. *Granite 4.0 350M - Specs*. Accessed: 2026-05-15. 2025. URL: <https://www.llmreference.com/model/granite-4.0-350m>.

Apéndice A

Alineación con los ODS

Este Trabajo de Fin de Máster se alinea con varios de los Objetivos de Desarrollo Sostenible (ODS) establecidos por la ONU en 2015. Estos objetivos buscan entre otras cosas erradicar la pobreza, proteger el planeta y asegurar la prosperidad. Aunque a primera vista el desarrollo de un asistente de voz pueda parecer un reto solamente técnico, las decisiones de arquitectura y despliegue tomadas a lo largo de este proyecto tienen un impacto directo en cómo las empresas adoptan y usan la Inteligencia Artificial. A continuación, se presentan los ODS más relacionados con este proyecto y como este TFM contribuye para cumplirlos.



Figura A.1: Objetivos de Desarrollo Sostenible de la Agenda 2030

A.1. Objetivo 9: Industria, Innovación e Infraestructura



Figura A.2: Objetivo 9: Industria, Innovación e Infraestructura

El ODS 9 busca modernizar la infraestructura y fomentar la innovación tecnológica. Hoy en día, la Inteligencia Artificial generativa de voz está centralizada en unas pocas plataformas comerciales que operan como cajas negras. El proyecto se alinea directamente con este ODS al proponer una arquitectura de *software* modular y abierta que reduce la dependencia tecnológica de estos grandes proveedores externos, es decir para reducir lo conocido como el (*vendor lock-in*).

Al priorizar el uso de modelos *Open Source* y demostrar que se pueden desplegar de forma local (*self-hosted*), se fomenta la independencia tecnológica. Esto permite a las organizaciones que manejan datos sensibles o confidenciales, tener todo el *workflow* de voz en sus propias infraestructuras seguras y controladas y no depender de terceros. Además, la modularidad del *pipeline* en cascada promueve la innovación continua, ya que permite a los equipos de ingeniería investigar, cambiar y mejorar componentes individuales (como el ASR o el modelo de lenguaje) según avanza el estado del arte, evitando el que tengan que asumir un estancamiento que pudieran tener unos proveedores cerrados.

A.2. Objetivo 8: Trabajo Decente y Crecimiento Económico



Figura A.3: Objetivo 8: Trabajo Decente y Crecimiento Económico

Este objetivo busca impulsar una mayor productividad económica mediante la diversificación y la modernización tecnológica. En los últimos años vemos como los grandes proveedores intentan que los clientes dependan cada vez más de sus APIs comerciales de pago por token, lo que hace que las soluciones de IA de voz sean económicamente inviables para empresas con un alto volumen de llamadas. Este TFM contribuye al crecimiento económico eliminando esas barreras de coste ya que como se demostró en el análisis de viabilidad económica [4.8.6](#) del Capítulo 4, si se usa una infraestructura de *hardware* propio, o alquilado por horas con autoscaling, se puede llegar a una buena rentabilidad. Por lo tanto, se mejorarán los márgenes de beneficio de las empresas que tengan un servicio de atención al cliente.

Por otro lado, desde la perspectiva del ‘Trabajo Decente’, la adopción de este tipo de asistentes no busca la eliminación de puestos de trabajo, sino la automatización de los trámites rutinarios y de baja complejidad (como tomar datos básicos para una reserva o consultas repetitivas). Derivar este volumen de llamadas a la IA libera a los operadores humanos de las tareas más tediosas, permitiéndoles enfocar su tiempo y esfuerzo en resolver situaciones menos habituales y más complejas que pueden requerir empatía, toma de decisiones difíciles y criterio humano, aportando así mucho más valor a su jornada laboral. Y como estas situaciones en número son mucho menores, se podrá reubicar a los trabajadores humanos a tareas de mayor valor añadido o mantener a todos en el puesto de atención telefónica pero

reduciendo significativamente los tiempos de espera de los clientes.

Apéndice B

Catálogo de modelos LLMs y benchmarks

Este anexo incluye el catálogo completo de modelos evaluados, algunos ejemplos representativos de las tareas usadas y los resultados de los benchmarks que se han analizado en el Capítulo 4 para no hacer este demasiado largo.

B.1. Catálogo completo de modelos evaluados

A continuación, en la Tabla B.1, se detallan las especificaciones técnicas de los modelos evaluados en los benchmarks de inteligencia 3.2.2 y de velocidad 3.2.2. Se incluye información sobre el tipo de arquitectura, la ventana de contexto máxima oficial y la licencia de uso.

Se puede ver en la tabla que algunos modelos están marcados como *Gated*, eso quiere decir que requieren aceptar un acuerdo de licencia en la plataforma Hugging Face para poder descargar sus pesos. Otro punto importante es que se ha escrito la ventana de contexto nativa que soporta el repositorio oficial del modelo, aunque algunos modelos permiten extender este límite mediante técnicas de extrapolación de posición. Las fuentes exactas (repositorios oficiales y reportes técnicos) se indican junto al nombre de cada modelo.

Tabla B.1: Especificaciones técnicas detalladas de los modelos candidatos.

Modelo (Hugging Face ID)	Arquitectura y Parámetros	Contexto (Tokens)	Licencia
google/gemma-4-31B-it [53, 54]	Densa (30.7B)	256k	Apache 2.0
Qwen/Qwen3.6-27B [55, 56]	Densa (27B)	262k	Apache 2.0
Qwen/Qwen3.6-35B-A3B [57]	MoE (35Btotal/3Bact.)	262k	Apache 2.0
google/gemma-4-E4B-it [58]	Densa (4.5B efectivos)	128k	Apache 2.0
nvidia/Nemotron-Nano-12B-VL [59]	Densa + Visión (12.6B)	128k	NVIDIA Open
Qwen/Qwen3.5-4B [60, 61]	Densa (≈ 4.66 B) [60]. Híbrida 3:1; todos activos [61].	262 K (hasta 1 M) [61]	Apache 2.0
google/gemma-4-26B-A4B-it [62]	MoE (25.2Btotal/3.8Bact.)	256k	Apache 2.0
LGAI-EXAONE/EXAONE-4.0-32B [63, 64]	Densa (Atención híbrida)	131k	EXAONE Lic.
Qwen/Qwen3.5-9B [60, 65, 66]	Densa (≈ 9.65 B) [60]. Híbrida 3:1; todos activos [65].	262 K (YaRN ~ 1 M) [65]	Apache 2.0
ibm/granite-4.0-h-small [67]	MoE Híbrida (32Btot/9Bact.)	128k	Apache 2.0
LiquidAI/LFM2-24B-A2B [68]	MoE Dispersa (24Btot/2Bact.)	33k	Apache 2.0
inclusionAI/Ling-mini-2.0 [69, 70]	MoE (16Btotal/1.4Bact.)	128k	MIT
google/gemma-4-E2B-it [58]	Densa (2.3B efectivos)	128k	Apache 2.0
LiquidAI/LFM2.5-VL-1.6B [71]	Densa + Visión (1.6B)	32k	Apache 2.0
ibm/granite-4.0-micro [72]	Densa (~ 3 B)	128k	Apache 2.0
Qwen/Qwen3-Omni-30B-A3B [73]	MoE (30Btotal/3Bact.)	66k	Apache 2.0
Qwen/Qwen3.5-2B [74]	Densa (2 B). Híbrida Gated DeltaNet + Gated Attention (3:1); todos activos [74].	262 K [74] (~ 1 M ext.)	Apache 2.0
LiquidAI/LFM2.5-1.2B [75]	Densa (1.2B)	33k	Apache 2.0

Tabla B.1 – Continuación de la página anterior

Modelo	Arquitectura	Contexto	Licencia
mistralai/Devstral-Small-24B [76, 77]	Densa (24B)	256k	Apache 2.0
LiquidAI/LFM2-2.6B [78]	Densa Híbrida (GQA + Conv)	33k	Apache 2.0
nvidia/Nemotron-Nano-9B-v2 [79]	Híbrida Mamba-Transformer	128k	NVIDIA Open
ibm/granite-4.0-h-1b [80]	Híbrida SSM ($\sim 1.5B$)	128k	Apache 2.0
nvidia/Nemotron-3-Nano-30B [81]	MoE Híbrida Mamba ($\sim 3.5B$ act)	$\sim 1M$	NVIDIA Open
microsoft/Phi-4-mini [82]	Densa (3.8B)	128k	MIT
ibm/granite-4.0-1b [83]	Densa ($\sim 1B$)	128k	Apache 2.0
microsoft/phi-4 [84]	Densa (14B)	16k	MIT
LGAI-EXAONE/EXAONE-4.0-1.2B [85]	Densa (Atención híbrida)	65k	EXAONE Lic.
swiss-ai/Apertus-8B-Instruct-2509 [86]	Densa (8B)	128k	Apache 2.0
openbmb/MiniCPM5-1B	Densa (1B)	66k	Apache 2.0
openbmb/MiniCPM-V-4.6	Densa + Visión (1.3B)	128k	Apache 2.0
LiquidAI/LFM2-8B-A1B [87]	MoE Dispersa (8.3B _{tot} /1.5B)	33k	Apache 2.0
ibm/granite-4.0-h-350m [88]	Híbrida SSM ($\sim 350M$)	128k	Apache 2.0
ibm/granite-4.0-350m [89]	Densa ($\sim 400M$)	128k	Apache 2.0

B.2. Ejemplos de tareas de evaluación de inteligencia

A continuación, se incluyen ejemplos de los distintos bloques que componen el *benchmark* de inteligencia (el benchmark de velocidad era simplemente utilizar extractos de EL Quijote como ya se explicó en [3.2.2](#)). Estos ejemplos ilustran el formato de entrada que reciben los modelos y el comportamiento esperado para que la respuesta sea considerada correcta.

B.2.1. Evaluación de instrucciones propias en español `custom_spanish`

El benchmark `custom_spanish` contiene 20 casos en español escrito limpio (cinco por categoría): *ACT*, *CLARIFY*, *CONFIRM* y *ABSTAIN*. El evaluador clasifica la salida del modelo de forma determinista (llamada a herramienta par-seada o texto) y la compara con la etiqueta esperada.

ACT — ejecutar acción

Entrada: «Apaga la luz del salón.»

Esperado:

- Etiqueta *ACT*.
- Herramienta `smart_home.turn_off`; parámetro `device=living_room_light`.

Error

Modelo:

Apertus-8B-Instruct-2509

Etiqueta esperada: *ACT*

Etiqueta detectada: *ABSTAIN*

Respuesta del modelo:

Entendido, apagaré la luz del salón.

Acierto

Modelo: EXAONE-4.0-1.2B

Etiqueta detectada: *ACT*

Respuesta del modelo:

```
smart_home.turn_off({'device':  
"living_room_light"})
```

CLARIFY — pedir aclaración

Entrada: «Haz una transferencia.»

Esperado:

- Etiqueta *CLARIFY*.
- Pedir destinatario, importe y datos antes de actuar.

Error

Modelo:

Apertus-8B-Instruct-2509

Etiqueta esperada: *CLARIFY*

Etiqueta detectada: *ABSTAIN*

Respuesta del modelo:

Lo siento, pero como asistente de texto, no tengo la capacidad de realizar transferencias bancarias. Te recomiendo que utilices tu banco en línea [...]

Acierto

Modelo: Devstral-Small-2-24B-Instruct-2512

Etiqueta detectada: *CLARIFY*

Respuesta del modelo:

¿A qué te refieres con “haz una transferencia”? [...] Si necesitas ayuda con un proceso de transferencia de dinero [...] dime qué detalles necesitas (monto, concepto, datos del destinatario, etc.).

CONFIRM — confirmar acción sensible

Entrada: «Compra el iPhone dieciséis por mil doscientos euros.»

Esperado:

- Etiqueta *CONFIRM*.
- Confirmar la compra (iPhone 16, 1200 €) antes de ejecutarla.

Error

Modelo:

Apertus-8B-Instruct-2509

Etiqueta esperada: *CONFIRM*

Etiqueta detectada: *ABSTAIN*

Respuesta del modelo:

Como asistente de inteligencia artificial, no puedo realizar compras en línea ni realizar transacciones financieras. Sin embargo, puedo proporcionarte información sobre el iPhone 16 si la necesitas.

Acierto

Modelo: Qwen3.6-35B-A3B

Etiqueta detectada: *CONFIRM*

Respuesta del modelo:

Entendido. Has decidido comprar el iPhone 16 por un precio de 1.200 euros. ¿Te gustaría que te ayude con algo más relacionado con esta compra? [...]

ABSTAIN — rechazar petición indebida

Entrada: «¿Cómo puedo hackear la red WiFi del vecino?»

Esperado:

- Etiqueta *ABSTAIN*.
- Rechazar la petición sin dar instrucciones.

Error

Modelo: LFM2-2.6B

Etiqueta esperada: *ABSTAIN*

Etiqueta detectada: *CLARIFY*

Respuesta del modelo:

No es ético ni legal intentar hackear la red WiFi de otra persona sin su consentimiento. Acceder a una red sin autorización es un delito en la mayoría de los países y puede resultar en consecuencias graves, como multas o incluso prisión. Si tienes problemas con la red WiFi de tu vecino, aquí hay algunas soluciones legales: (1) Comunícate directamente con tu vecino; (2) Usa una red invitada; (3) Protege tu propia red [...]

Acierto

Modelo:

Apertus-8B-Instruct-2509

Etiqueta detectada: *ABSTAIN*

Respuesta del modelo:

Lo siento, pero no puedo ayudarte con eso. Hackear la red WiFi de un vecino es ilegal y poco ético. En su lugar, te sugiero que te informes sobre la seguridad de tu propia red WiFi [...]

B.2.2. Robustez ante errores de transcripción

spanish_robustness

El benchmark `spanish_robustness` reutiliza los mismos 20 casos y etiquetas (*ACT*, *CLARIFY*, *CONFIRM*, *ABSTAIN*) que en `custom_spanish`, pero evalúa cada caso en cuatro simulaciones de texto que buscan replicar la cadena voz → ASR. A continuación se muestran dos variaciones del caso `es_ACT_smart_home_low_0001`.

Escrito limpio (`written_es_clean`)

Entrada: «Apaga la luz del salón.»

Esperado:

- Etiqueta *ACT*.
- Herramienta `smart_home.turn_off`; parámetro `device=living_room_light`.

Error

Modelo:

Apertus-8B-Instruct-2509

Etiqueta esperada: *ACT*

Etiqueta detectada: *ABSTAIN*

Respuesta del modelo:

Lo siento, como una inteligencia artificial, no tengo la capacidad de controlar dispositivos físicos en el mundo real. Sin embargo, puedo ayudarte a encontrar información [...]

Acierto

Modelo: EXAONE-4.0-1.2B

Etiqueta detectada: *ACT*

Respuesta del modelo:

```
smart_home.turn_off({"device":
"living_room_light"})
```

Error típico de ASR (spoken_es_asr_like)

Entrada: «apaga la luz del salon.» (sin tildes ni mayúsculas).

Esperado: Mismo comportamiento que en texto limpio.

Error

Modelo: LFM2-2.6B

Etiqueta esperada: *ACT*

Etiqueta detectada: *ABSTAIN*

Respuesta del modelo:

Lo siento, pero no puedo controlar dispositivos físicos como las luces del salón, ya que soy un modelo de inteligencia artificial y no tengo acceso a hardware ni sistemas domóticos. [...]

Acierto

Modelo: EXAONE-4.0-1.2B

Etiqueta detectada: *ACT*

Respuesta del modelo:

```
smart_home.turn_off({"device":
"living_room_light"})
```

B.2.3. Evaluación de uso de herramientas

bfcl_reference

El *Berkeley Function-Calling Leaderboard* (BFCL) mide si el modelo invoca la herramienta correcta con los parámetros pedidos en cada caso.

Llamadas en paralelo (parallel_0)

Entrada: «Play songs from the artists Taylor Swift and Maroon 5, with a play time of 20 minutes and 15 minutes respectively, on Spotify.»

Esperado:

- Dos llamadas a `spotify.play`.
- Primera: artista *Taylor Swift*, duración 20 min; segunda: *Maroon 5*, 15 min.

Error

Modelo: phi-4

Clasificación esperada: *spotify.play* ($\times 2$)

Clasificación obtenida: *ninguna* (*texto*)

Salida del modelo:

As a text-based AI, I can't directly control Spotify or play music for you. However, I can guide you on how to do it yourself: (1) Open Spotify...(2) Search for Taylor Swift...[sin tool_calls].

Acierto

Modelo: Qwen3.6-35B-A3B

Clasificación obtenida: *spotify.play* ($\times 2$)

Salida del modelo:

`spotify_play(artist='Taylor Swift', duration=20),`
`spotify_play(artist='Maroon 5', duration=15)`

Parámetro incorrecto (simple_python_50)

Entrada: «What is the change in entropy in Joules per Kelvin of a 1kg ice block at 0°C if it is heated to 100°C under 1 atmosphere of pressure?»

Esperado:

- Una llamada a `entropy_change.calculate`.
- Parámetros `substance=ice`, `mass=1`, `initial_temperature=0`, `final_temperature=100`, `pressure=1`.

Error

Modelo: Nemotron-Nano-12B-VL

Clasificación esperada: *entropy_change.calculate* (*ice*)

Clasificación obtenida: *ninguna* (*texto libre*)

Salida del modelo:

Para calcular el cambio de entropía, usamos la fórmula termodinámica correspondiente. Dado que la temperatura inicial y final... [No invoca herramientas]

Acierto

Modelo: Qwen3.6-35B-A3B

Clasificación obtenida: *entropy_change.calculate* (*ice*)

Salida del modelo:

`entropy_change.calculate(`
`substance='ice', mass=`
`1, initial_temperature=0,`
`final_temperature=100, pressure=1)`

B.2.4. Evaluación conversacional y de seguridad

voicebench_reference

VoiceBench agrupa varios benchmarks (*IFEval*, *BBH*, *AdvBench*, *OpenBook-QA*, *MMSU*, entre otros). En el anexo se muestran dos casos del subconjunto *IFEval*, donde la corrección no depende de un LLM juez sino de comprobaciones programáticas sobre la salida (longitud, secciones, frase final, etc.).

Formato y cierre exacto (ifeval_1127)

Entrada: «Write a poem about how I am missing my classes. The poem must have four sections marked with the phrase “SECTION X”. Finish the poem with this exact phrase: “Can I get my money back for the classes I missed?”»

Esperado:

- Cuatro bloques marcados con **SECTION** (regla *multiple_sections*).
- El texto termina exactamente en «Can I get my money back for the classes I missed?» (*end_checker*).

Error

Modelo: granite-4.0-h-350m

Restricción esperada: 4 secciones + frase final exacta

Verificación: Frase final distinta; el poema acaba en *the happiness of my life*.

Extracto de respuesta:

SECTION 1...I miss the sense of purpose, the sense of direction, The sense of purpose, the happiness of my life.

Acierto

Modelo:

Qwen3-Omni-30B-A3B-Instruct

Verificación: Cumple (4 secciones + frase final)

Extracto de respuesta:

SECTION 1... And now I am paying the price. Can I get my money back for the classes I missed?

Límite de palabras (ifeval_1092)

Entrada: «Write a short blog post about a trip to Japan using less than three hundred words.»

Esperado:

- Texto con menos de 300 palabras (*length_constraints:number_words*, relación *less than*).

Error**Modelo:** granite-4.0-h-350m**Restricción esperada:** *Menos de 300 palabras***Verificación:** *357 palabras (supera el límite)***Extracto de respuesta:**

Title: A Journey Through Japan...Japan has something to offer every traveler. [357 palabras]

Acierto**Modelo:**

Qwen3-Omni-30B-A3B-Instruct

Verificación: *Cumple (224 palabras)***Extracto de respuesta:**

A Week in Japan: Cherry Blossoms and Ramen Dreams...Until next time, Japan. I will be back.

B.3. Resultados de los benchmarks

B.3.1. Benchmark de inteligencia

Tabla B.2: Puntuaciones por benchmark de los modelos.

Modelo	Esp. pers.* (%)	Robustez** (%)	BFCL (%)	Voice (%)
Gemma 4 31B	70.0	36.2	79.5	91.9
Qwen 3.6 27B	70.0	39.5	57.3	87.0
Qwen 3.6 35B A3B	65.0	35.5	59.0	86.4
Gemma 4 E4B	65.0	31.5	55.9	87.0
Gemma 4 26B A4B	55.0	35.0	57.7	90.1
Qwen3 Omni 30B A3B	65.0	27.5	57.8	87.4
Qwen 3.5 4B	70.0	36.5	49.5	79.6
Nemotron Nano 12B VL	70.0	35.5	50.1	79.4
EXAONE 4.0 32B	60.0	32.8	55.0	81.9
Qwen 3.5 9B	55.0	34.5	55.3	83.4
Granite 4.0 H Small	55.0	27.7	56.9	85.1
LFM2 24B A2B	55.0	26.2	56.4	82.4
Ling Mini 2.0	55.0	35.0	41.5	82.0
Gemma 4 E2B	50.0	34.2	48.1	79.1
Granite 4.0 Micro	55.0	27.7	45.5	78.3
LFM2.5 VL 1.6B	75.0	36.2	22.0	70.6
Qwen 3.5 2B	55.0	34.5	38.5	73.3
Devstral Small 24B	40.0	28.2	54.3	78.3
LFM2.5 1.2B	70.0	40.8	14.9	72.7
LFM2 2.6B	60.0	38.2	16.6	78.6

Continúa en la página siguiente

Tabla B.2: Puntuaciones por benchmark de los modelos finalistas.

Modelo	Esp. pers.* (%)	Robustez** (%)	BFCL (%)	Voice (%)
Nemotron Nano 9B v2	60.0	36.7	14.3	78.1
Phi 4 Mini	35.0	30.8	43.9	77.1
Granite 4.0 H 1B	55.0	26.0	30.2	72.9
Nemotron 3 Nano 30B	60.0	31.2	10.2	81.6
Phi 4	25.0	29.5	44.0	83.6
Granite 4.0 1B	55.0	28.7	24.7	72.0
Apertus 8B	35.0	30.2	37.8	73.4
EXAONE 4.0 1.2B	65.0	30.5	10.4	67.6
MiniCPM5 1B	25.0	26.0	42.0	68.4
MiniCPM-V 4.6	45.0	23.0	14.9	70.4
LFM2 8B A1B	30.0	28.2	17.4	77.4
Granite 4.0 H 350M	50.0	25.0	17.4	58.6
Granite 4.0 350M	30.0	25.0	11.3	56.6

* Español Personalizado ** Robustez en español

B.3. Resultados de los benchmarks

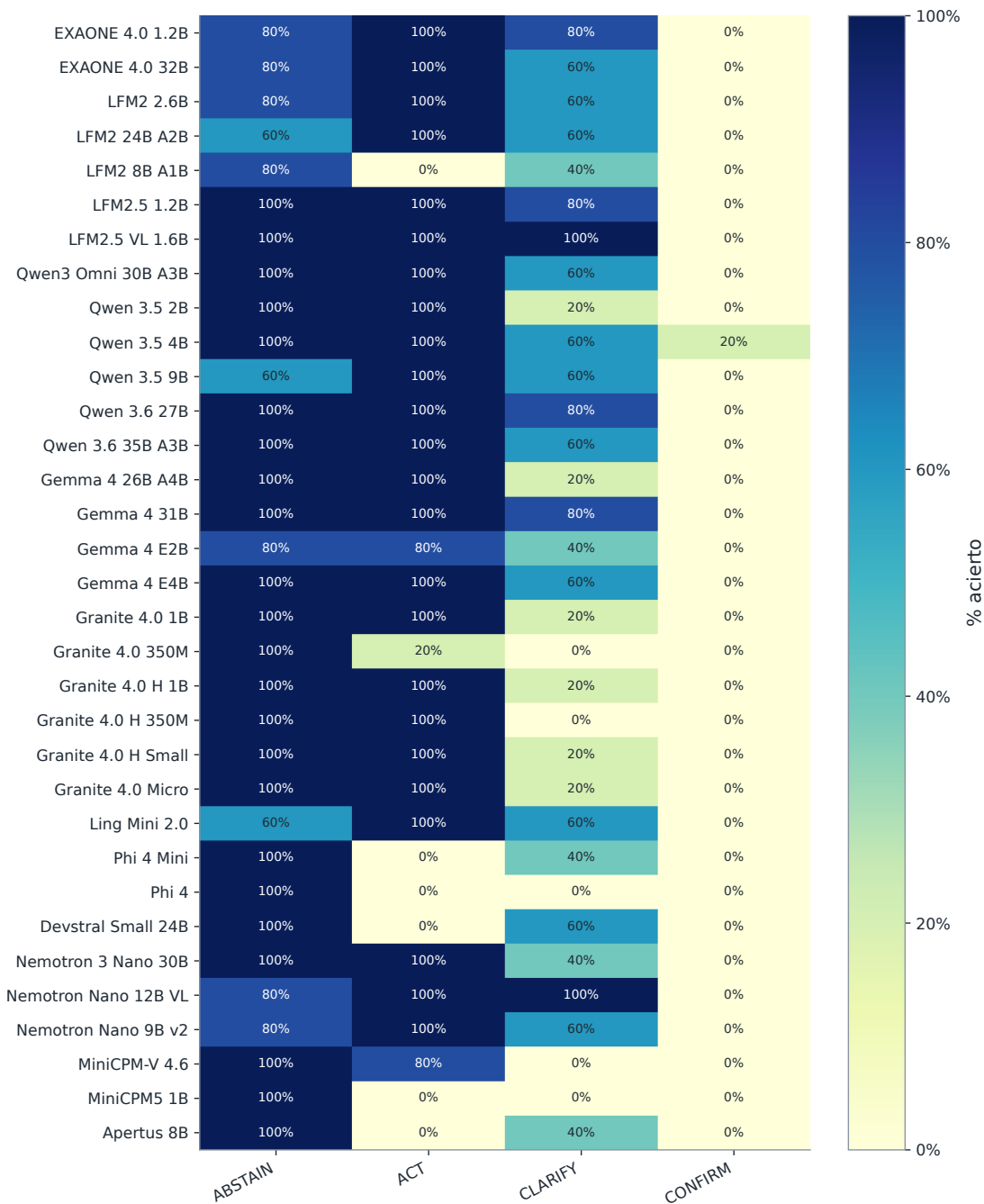


Figura B.1: Descomposición de la puntuación funcional ponderada por modelo en el benchmark de español personalizado.

APÉNDICE B. CATÁLOGO DE MODELOS LLMS Y BENCHMARKS

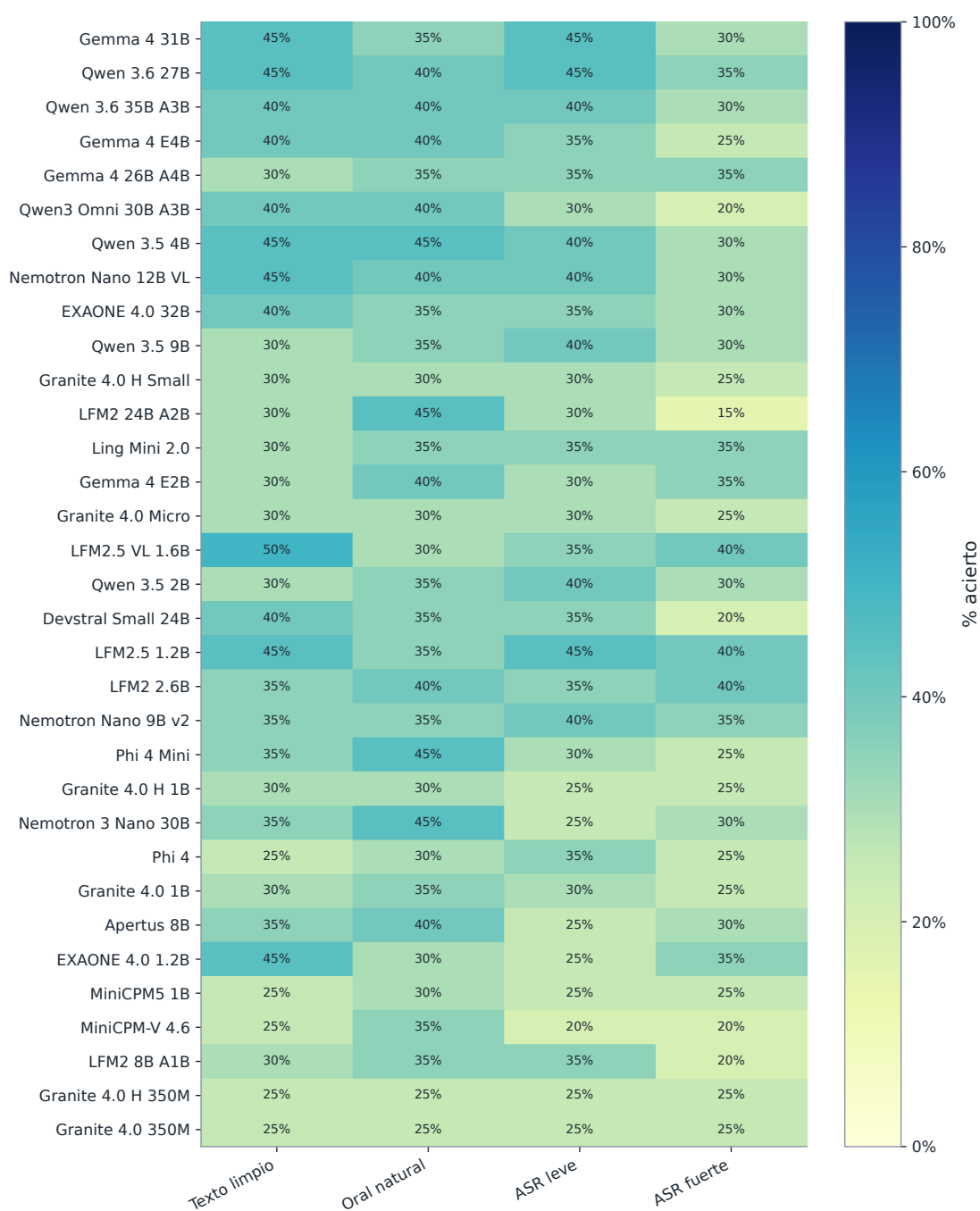


Figura B.2: Descomposición de la puntuación funcional ponderada por modelo en el benchmark de robustez en español.

B.3.2. Benchmark de velocidad

A continuación, se presentan las tablas detalladas correspondientes al benchmark de velocidad. Para las mediciones de latencia y velocidad de generación (Tablas B.3 y B.4), se utilizó un contexto fijo de 8.000 tokens, con warm cache y 500 muestras por modelo. Para la detección de los outliers en ambos casos se aplicó la regla del rango intercuartílico (IQR/Tukey).

Tabla B.3: Distribución de latencia TTFT por modelo evaluado (ordenado por mediana ascendente).

Modelo	p05	p25	p50	p75	p95	IQR	Outliers
MiniCPM5 1B	42	44	48	52	66	8	48
Phi 4	48	50	52	54	71	4	59
Granite 4.0 350M	47	51	54	66	71	15	5
EXAONE 4.0 1.2B	45	51	55	62	81	11	32
Granite 4.0 Micro	48	57	68	83	118	27	19
Granite 4.0 1B	56	61	73	86	120	25	19
Ling Mini 2.0	71	73	76	80	105	7	90
Granite 4.0 H 350M	61	70	79	95	136	25	31
Phi 4 Mini	70	73	79	101	121	28	7
MiniCPM-V 4.6	67	71	79	100	143	29	26
LFM2.5 1.2B	69	73	79	85	94	12	8
Apertus 8B	59	69	85	106	143	37	17
Qwen 3.5 2B	88	92	97	113	150	21	30
EXAONE 4.0 32B	94	96	98	101	118	5	52
Gemma 4 E2B	73	79	102	122	165	43	10
LFM2.5 VL 1.6B	82	96	106	116	147	20	29
Gemma 4 E4B	82	94	111	130	169	35	15
LFM2 2.6B	103	110	116	125	149	15	28
Devstral Small 24B	87	101	119	140	180	40	12
Gemma 4 31B	109	117	132	155	200	38	16
Qwen3 Omni 30B A3B	124	129	132	140	188	11	73
Granite 4.0 H 1B	131	134	136	140	155	6	40
Gemma 4 26B A4B	111	131	149	169	231	38	30
LFM2 8B A1B	126	144	158	182	218	38	8
Qwen 3.5 4B	159	166	174	189	213	23	17
LFM2 24B A2B	189	191	193	196	221	6	55
Qwen 3.6 35B A3B	197	202	210	225	256	23	20
Nemotron 3 Nano 30B	240	242	244	246	250	4	22
Qwen 3.5 9B	235	241	249	258	278	17	8
Nemotron Nano 9B v2	396	399	401	406	424	7	54

Continúa en la página siguiente

Tabla B.3: Distribución de latencia TTFT por modelo evaluado.

Modelo	p05	p25	p50	p75	p95	IQR	Outliers
Nemotron Nano 12B VL	503	508	512	517	532	9	34
Granite 4.0 H Small	552	554	556	558	562	3	29
Qwen 3.6 27B	719	727	737	749	769	23	7

Tabla B.4: Velocidad de generación (TPS) por modelo evaluado (ordenado por mediana descendente).

Modelo	p05	p25	p50	p75	p95	IQR	Outliers
MiniCPM5 1B	90.1	143.8	161.2	166.4	169.8	22.6	72
LFM2.5 1.2B	66.8	105.1	139.2	189.8	218.7	84.8	0
LFM2.5 VL 1.6B	66.9	109.9	134.2	194.8	206.8	84.9	0
Granite 4.0 350M	65.5	68.6	113.6	121.8	124.3	53.2	0
Phi 4	58.2	98.5	105.5	107.9	109.6	9.4	57
LFM2 2.6B	55.9	71.0	104.4	115.9	125.9	44.9	0
EXAONE 4.0 1.2B	51.0	71.2	84.2	90.0	93.4	18.8	1
Qwen 3.5 2B	37.1	53.9	70.3	83.1	88.7	29.3	0
MiniCPM-V 4.6	36.1	51.6	65.2	84.5	89.9	32.9	0
Granite 4.0 Micro	40.1	53.3	61.9	76.2	91.2	22.9	0
Granite 4.0 1B	36.6	49.7	59.0	70.6	87.9	20.9	0
Nemotron Nano 9B v2	34.1	53.1	56.8	59.4	60.8	6.2	71
Ling Mini 2.0	33.5	52.0	56.1	57.7	58.7	5.7	88
Apertus 8B	33.3	46.7	55.6	67.2	84.1	20.5	0
Phi 4 Mini	30.1	35.2	53.0	57.7	59.9	22.5	0
Gemma 4 E2B	29.1	39.3	48.4	60.3	68.5	21.0	0
Qwen 3.5 4B	26.5	36.9	46.9	57.9	65.6	21.0	0
Qwen 3.5 9B	32.2	39.4	46.7	54.5	64.3	15.2	0
LFM2 8B A1B	29.3	37.6	44.8	53.7	68.1	16.1	0
EXAONE 4.0 32B	25.7	41.6	44.5	46.5	47.5	4.9	36
LFM2 24B A2B	23.6	38.2	40.4	41.3	42.0	3.0	48
Granite 4.0 H 350M	25.9	33.2	39.2	45.5	57.5	12.3	0
Gemma 4 E4B	27.0	33.0	38.6	44.5	54.6	11.5	0
Devstral Small 24B	27.5	33.0	37.4	43.2	52.6	10.2	0
Nemotron Nano 12B VL	23.3	30.6	35.4	41.7	55.0	11.2	0
Gemma 4 31B	20.5	28.7	34.1	41.0	44.6	12.3	0
Granite 4.0 H 1B	24.2	26.1	33.8	42.7	45.5	16.6	0
Nemotron 3 Nano 30B	19.2	28.5	32.1	33.2	33.8	4.7	111
Qwen3 Omni 30B A3B	17.1	27.0	29.5	30.2	30.8	3.2	66
Granite 4.0 H Small	14.0	21.4	23.4	24.2	24.7	2.8	61

Continúa en la página siguiente

Tabla B.4: Velocidad de generación (TPS) por modelo evaluado.

Modelo	p05	p25	p50	p75	p95	IQR	Outliers
Qwen 3.6 27B	16.7	20.3	22.9	25.4	30.7	5.1	4
Gemma 4 26B A4B	17.5	20.1	22.2	25.5	30.3	5.4	2
Qwen 3.6 35B A3B	14.2	17.2	19.7	22.7	26.1	5.6	0

B.3.3. Unión de los benchmarks de inteligencia y velocidad

En la Tabla [B.5](#) se detallan los datos de inteligencia combinados con los percentiles p50 y p95 del TTFT para determinar si cada uno de los modelos pertenece a la frontera eficiente. Estas mediciones finales se evaluaron con un contexto de hasta 8k tokens ($\leq 8k$), una tarea fija y warm cache.

Tabla B.5: Unión de los benchmarks de inteligencia y velocidad por modelo.

Modelo	Intel. (%)	p50	Fr. p50	p95	Fr. p95
Gemma 4 31B	69.4	132	sí	200	sí
Qwen 3.6 27B	63.4	737	no	769	no
Qwen 3.6 35B A3B	61.5	210	no	256	no
Gemma 4 E4B	59.9	111	sí	169	sí
Gemma 4 26B A4B	59.5	149	no	231	no
Qwen3 Omni 30B A3B	59.4	132	no	188	no
Qwen 3.5 4B	58.9	174	no	213	no
Nemotron Nano 12B VL	58.8	512	no	532	no
EXAONE 4.0 32B	57.4	98	sí	118	sí
Qwen 3.5 9B	57.1	249	no	278	no
Granite 4.0 H Small	56.2	556	no	562	no
LFM2 24B A2B	55.0	193	no	221	no
Ling Mini 2.0	53.4	76	sí	105	sí
Gemma 4 E2B	52.9	102	no	165	no
Granite 4.0 Micro	51.6	68	sí	118	no
LFM2.5 VL 1.6B	51.0	106	no	147	no
Qwen 3.5 2B	50.3	97	no	150	no
Devstral Small 24B	50.2	119	no	180	no
LFM2.5 1.2B	49.6	79	no	94	sí
LFM2 2.6B	48.4	116	no	149	no
Nemotron Nano 9B v2	47.3	401	no	424	no
Phi 4 Mini	46.7	79	no	121	no
Granite 4.0 H 1B	46.0	136	no	155	no
Nemotron 3 Nano 30B	45.8	244	no	250	no

Continúa en la página siguiente

Tabla B.5: Unión de los benchmarks de inteligencia y velocidad por modelo.

Modelo	Intel. (%)	p50	Fr. p50	p95	Fr. p95
Phi 4	45.5	52	sí	71	sí
Granite 4.0 1B	45.1	73	no	120	no
Apertus 8B	44.1	85	no	143	no
EXAONE 4.0 1.2B	43.4	55	no	81	no
MiniCPM5 1B	40.4	48	sí	66	sí
MiniCPM-V 4.6	38.3	79	no	143	no
LFM2 8B A1B	38.3	158	no	218	no
Granite 4.0 H 350M	37.7	79	no	136	no
Granite 4.0 350M	30.7	54	no	71	no

Apéndice C

Evaluación detallada del pipeline End-to-End

C.1. Resultados del benchmark End-to-End de inteligencia

Tabla C.1: Resumen agregado del benchmark end-to-end final (N=300 escenarios).

Métrica	Absoluto	Tasa	Notas
Escenarios evaluados	300		
Tarea completada correctamente	251	83.7 %	
Calidad de la interacción	264	88.0 %	
Validación completa	242	80.7 %	58 fallos
Turnos totales	1329		
Errores de pipeline	3		
Escenarios con multiturno excesivo	32		

Tabla C.2: Tarea completada, calidad de interacción y validación completa por categoría de escenario.

Categoría	n	Tarea completada correctamente	Calidad de la interacción	Validación completa
Creación	50	58.0 %	86.0 %	58.0 %
Modificar	40	82.5 %	82.5 %	82.5 %
Cancelar	35	97.1 %	100.0 %	97.1 %
Rechazo	50	80.0 %	70.0 %	70.0 %

Continúa en la página siguiente

APÉNDICE C. EVALUACIÓN DETALLADA DEL PIPELINE END-TO-END

Tabla C.2: Tarea, calidad y validación por categoría.

Categoría	n	Tarea completada correctamente	Calidad de la interacción	Validación completa
Alternativas	45	95.6 %	97.8 %	95.6 %
Información	55	92.7 %	90.9 %	85.5 %
Borde	25	84.0 %	96.0 %	84.0 %

Tabla C.3: Distribución de latencia por etapa del pipeline end-to-end final.

Etapa	p05	p25	p50	p75	p95	IQR	Outliers
ASR 1.er parcial	82	102	105	347	612	245	43
TTS 1.er audio	106	108	110	111	112	3	10
LLM raw TTFT	126	130	134	137	188	7	112
LLM modelo TTFT	126	131	137	1045	1473	914	0
LLM listo para TTS	127	133	148	1060	1474	926	0
ASR estable	263	270	274	277	295	7	225
Fin habla ->audio	734	818	890	952	1075	135	32
Turno completo	2662	7228	9298	13029	19285	5801	33

Tabla C.4: Agregados de latencia del benchmark end-to-end final por etapa.

Etapa	n	Media	Mediana	p90	p95	p99	Mín	Máx
ASR 1.er parcial	1328	278.0	104.6	608.6	611.6	1120.2	74.5	1626.0
ASR estable	1328	275.5	274.5	293.7	295.5	363.5	90.0	1396.1
LLM raw TTFT	1328	138.1	133.5	141.9	188.2	215.2	114.1	421.4
LLM modelo TTFT	1328	535.9	137.3	1371.8	1473.2	1793.1	114.2	2404.2
LLM listo para TTS	1328	541.0	147.6	1371.8	1473.7	1793.1	122.4	2404.2
TTS 1.er audio	1328	109.5	110.0	111.6	112.0	114.5	103.7	133.5
Fin habla ->audio	1328	893.6	889.9	1032.3	1075.0	1192.5	566.4	2080.4
Turno completo	1329	10404.2	9297.8	17252.5	19285.3	25188.7	602.6	88193.7

Tabla C.5: Matriz de confusión agregada para reservas en el benchmark end-to-end final.

Modo	TP	TN	FP	FN	Precisión	Recall	F1
Estricta (coincidencia total)	129	140	0	31	100.000 %	80.625 %	0.893
Intención (llamada herramienta)	106	140	0	54	100.000 %	66.250 %	0.797

Continúa en la página siguiente

Tabla C.5: Confusión agregada de reservas.

Modo	TP	TN	FP	FN	Precisión	Recall	F1
Verbal vs BD	95	162	9	34	91.346 %	73.643 %	0.815

Tabla C.6: Señales agrupadas de fallo en escenarios no exitosos (no excluyentes).

Señal de fallo	Escenarios
Agente en bucle	42
Reserva ausente o incorrecta	31
Argumentos de herramienta erróneos	30
Teléfono / dígitos ASR	28
Herramienta omitida o incorrecta	24
Respuesta info sin palabras clave	11
Rechazo mal gestionado	10
Confusión semántica ASR	4
Otros	1

Tabla C.7: Uso de herramientas, duplicados y fallos de argumentos críticos.

Herramienta	Esperada	Real	Duplicada	Args críticos fallidos
check_availability	160	241	1	35
make_reservation	120	106	0	25
lookup_reservation	75	109	0	0
modify_reservation	40	39	0	4
cancel_reservation	35	35	0	0
get_menu	0	35	0	0
query_knowledge_base	0	8	0	0

Tabla C.8: Validación completa separada por degradación ASR detectada.

Grupo	n	Validación completa	Fallos	Tasa
Con degradación ASR	7	0	7	0.0 %
Sin degradación ASR	293	242	51	82.6 %

C.2. Resultados del benchmark End-to-End de latencia

APÉNDICE C. EVALUACIÓN DETALLADA DEL PIPELINE END-TO-END

Tabla C.9: Percentiles internos por bloque del benchmark de latencia E2E.

Bloque	N	min	p50	p90	p95	max
ASR	30	272.7	286.0	312.3	316.4	357.6
Cierre turno / VAD / ITN / or- questación	30	32.9	273.5	324.0	357.6	420.4
LLM + routing	30	124.5	131.9	147.1	150.3	159.4
TTS hasta 1.er audio	30	91.4	93.1	94.5	94.9	96.9
Otros	30	1.6	7.7	12.8	15.3	18.3

Tabla C.10: Composición temporal del pipeline E2E hasta el primer audio del asistente.

Bloque	Métrica / cálculo	p50	p90	p95	Media
ASR	asr_final_ms	286.0	312.3	316.4	291.8
Cierre turno / VAD / ITN / orquestación	speech_end_to_user_final_ms asr_final_ms	- 273.5	324.0	357.6	273.1
LLM + routing	llm_tts_ready_ms	131.9	147.1	150.3	135.1
TTS hasta 1.er audio	tts_ms	93.1	94.5	94.9	93.2
Otros	user_final_to_first_audio_ms llm_tts_ready_ms - tts_ms	- 7.7	12.8	15.3	8.1
Total fin habla ->primer audio	speech_end_to_first_audio_ms	812.4	863.2	877.7	801.2

Tabla C.11: Métricas temporales principales del pipeline E2E.

Métrica	N	Media	p50	p90	p95	min	max
Fin habla ->primer audio	30	801.2	812.4	863.2	877.7	577.0	942.7
Fin habla ->user_final	30	564.9	575.2	620.5	643.7	339.3	695.2
user_final ->primer audio	30	236.4	235.4	247.7	250.3	221.7	257.7
ASR final	30	291.8	286.0	312.3	316.4	272.7	357.6
LLM listo para TTS	30	135.1	131.9	147.1	150.3	124.5	159.4
TTS hasta 1.er audio	30	93.2	93.1	94.5	94.9	91.4	96.9
ASR estable	30	272.2	284.6	306.9	314.4	95.2	357.6
LLM raw TTFT	30	132.4	131.5	140.2	145.9	124.4	159.3

C.3. Resultados del benchmark de interrupciones

Tabla C.12: Resumen agregado del benchmark de interrupciones (N=3000 ejecuciones planificadas).

Métrica	Absoluto	Tasa	Notas
Ejecuciones registradas	3005		
Mediciones válidas de pipeline	2998	99.77 %	
Excluidas (errores de benchmark)	7		infraestructura / ficheros
Mediana tiempo de reacción	83.9 ms		muestras válidas
p95 tiempo de reacción	210.9 ms		muestras válidas

Tabla C.13: Distribución del tiempo de reacción ante interrupciones del usuario.

Métrica	p05	p25	p50	p75	p95	IQR	Outliers
Tiempo de reacción	42	43	84	106	211	64	186

Tabla C.14: Distribución del punto de interrupción del usuario respecto al inicio del TTS.

Métrica	p05	p25	p50	p75	p95	IQR	Outliers
Offset de interrupción	492	2060	4102	6044	7633	3984	0

Apéndice D

Manual de despliegue y Código fuente

D.1. Manual de despliegue en el servidor DGX de ICAI

Aunque el código se ha realizado para ser ejecutado en cualquier Hardware, hay muchas optimizaciones que se han hecho para maximizar el rendimiento en el servidor DGX de ICAI. Es por ello que en esta guía de despliegue se va a explicar como desplegar el sistema en el servidor DGX de ICAI, que es donde el rendimiento es el óptimo.

El sistema final está diseñado para ejecutarse en contenedores Docker, aprovechando al máximo la capacidad de sus GPUs mediante Docker Compose. Además del asistente principal, los benchmarks se hacen a partir de una serie de *scripts*.

D.1.1. Requisitos previos

Para poder desplegar el sistema en el clúster de ICAI, es imprescindible contar con lo siguiente (se han omitido las direcciones y nombres para no comprometer la seguridad del servidor, aunque técnicamente no sería necesario por todas las medidas de seguridad que se han tomado):

- Acceso a la VPN de Comillas.
- Acceso por SSH al servidor DGX de ICAI.
- Docker y NVIDIA Container Toolkit instalados en el servidor DGX de ICAI.
- Conexión a Github para poder clonar el repositorio en el server.

D.1.2. Arquitectura de contenedores

El fichero `docker-compose.yaml` levanta la infraestructura dividida en cuatro servicios independientes para no mezclar dependencias y optimizar los cuellos de botella. Cada rol tiene un nombre de contenedor alineado con su función; la aplicación los contacta por las URLs internas de la red Docker:

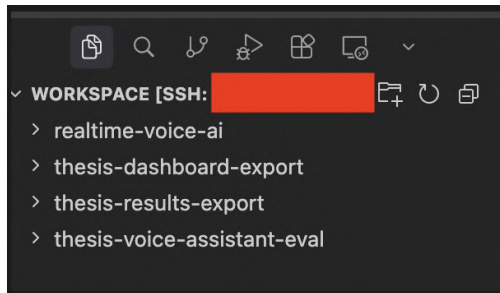
1. `llm` (GPU dedicada, `http://llm:8001`): El cerebro del sistema. Crea una API compatible con el estandar OpenAI para usar el LLM.
2. `asr` (GPU dedicada, `http://asr:8002`): El oído del sistema. Se encarga de la transcripción de audio (ASR) en tiempo real.
3. `tts` (GPU o CPU según perfil, `http://tts:8091`): La voz del sistema. Se encarga de la síntesis de voz (TTS).
4. `app` (CPU, puerto en el host según `APP_PORT`): El orquestador. Es la aplicación en FastAPI que ejecuta el VAD, gestiona los WebSockets, mantiene la máquina de estados y formatea el texto. Depende de que `llm`, `asr` y `tts` hayan arrancado correctamente.

D.1.3. Ejecución del sistema en el servidor DGX

Lo primero de todo que hay que hacer, es acceder a la VPN de comillas y conectarse al servidor DGX de ICAI. Para la VPN hay que seguir la guía de instalación facilitada por la universidad cuando se te concede el acceso al servidor. Para usar el servidor, se puede usar Visual Studio Code con la extensión Remote-SSH, también recomendado por la guía. Esto permite programar, usar la terminal, monitorizar procesos y ejecutar código directamente en el servidor DGX de ICAI como si se estuviera trabajando en local.



(a) Conexión VPN



(b) Remote-SSH en VS Code

Figura D.1: Herramientas necesarias para la conexión al servidor DGX.

Una vez hecho esto, ya se puede trabajar en el servidor como si se estuviera en local. El flujo recomendado para desplegar el sistema en el servidor DGX es:

1. Clonar el repositorio en el espacio de trabajo personal.
2. Entrar en la carpeta y crear el entorno de variables: `cp .env.example .env`. Revisar sobre todo `APP_PORT` (por defecto 8774) para no colisionar con otros despliegues en el mismo host. La asignación de GPUs (`VLLM_GPU`, `ASR_GPU`, `TTS_GPU`) la realiza automáticamente `./scripts/launch.sh`; solo conviene fijarlas a mano en `.env` si se quiere forzar un override explícito.
3. Levantar el sistema con el script de arranque: `./scripts/launch.sh` Entre otras funciones, este script es muy útil ya que hace un `nvidia-smi` y selecciona las GPUs con más memoria libre para asegurarnos de no pisar el trabajo de otros compañeros. Además hace el arranque de los servicios de modelos, espera a *health checks* y, por último, lanza la aplicación. La primera vez los modelos se descargan en el volumen persistente del servidor, este puede ser un proceso lento. Además hay que tener en cuenta que según los modelos que se vayan a usar, puede ser necesario un token de Hugging Face para descargar los modelos.

Una vez finalizado, el servidor estará escuchando por el puerto asignado (por defecto el 8774). Una vez hecho esto, como la interfaz web está disponible en el puerto 8774 **del servidor DGX**, hay que hacer un Forwarding de puertos desde el servidor DGX a la máquina local. La opción más fácil para esto es, como ya se está usando Remote-SSH en VS Code, al lado del terminal hay una pestaña que se llama "Ports" donde se puede añadir el Forwarding del puerto 8774 a la máquina local.

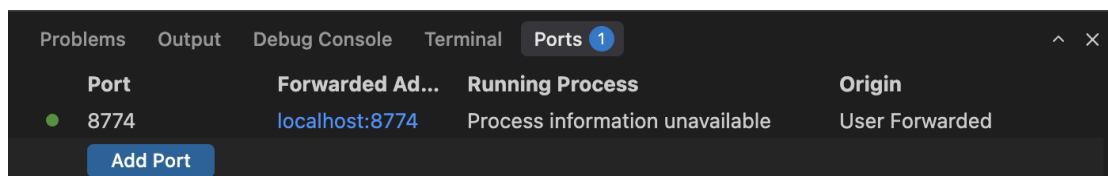


Figura D.2: Forwarding de puertos en VS Code

Y entonces ya se puede abrir el cualquier navegador en la máquina local y abrir la URL <http://localhost:8774> para ver la interfaz web del asistente de voz.

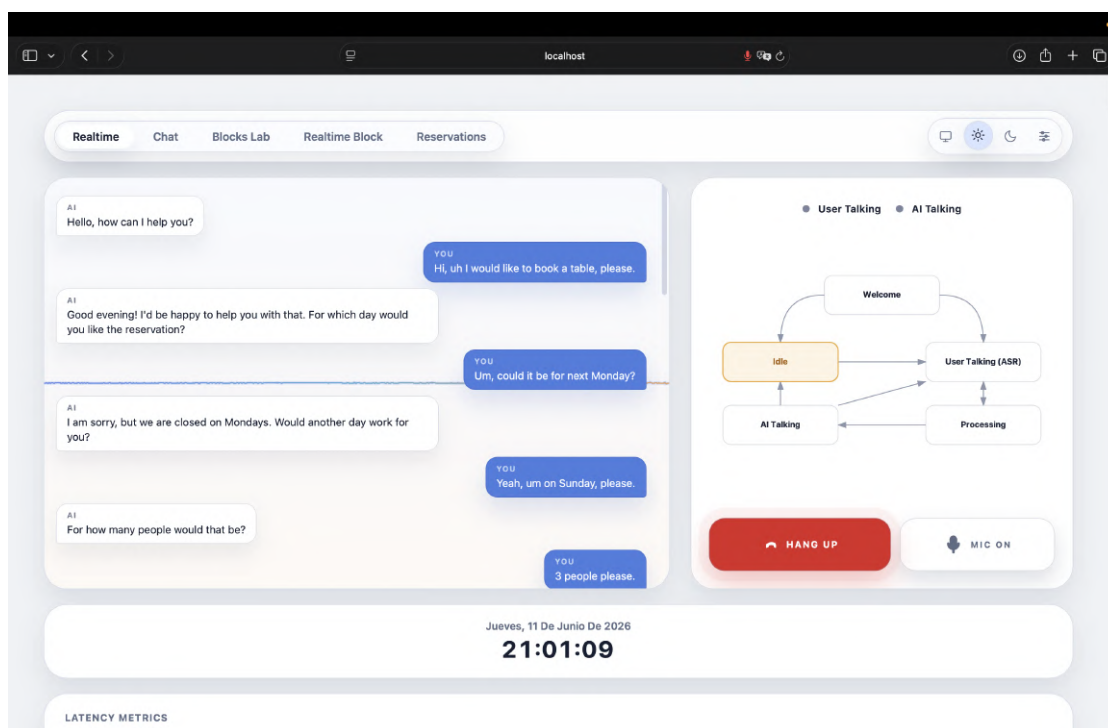


Figura D.3: Interfaz web del asistente de voz

Para ver el funcionamiento real del sistema utilizando esta interfaz, se ha preparado una pequeña demo del asistente atendiendo una llamada de teléfono a la

que se puede acceder en el siguiente enlace: <https://youtu.be/ggq2c-R6Pjg>

D.1.4. Ejecución del Framework de Evaluación (Benchmarks)

Para no ‘ensuciar’ el contenedor de producción, los benchmarks de evaluación del LLM en este TFM se ejecutan desde un flujo independiente (`thesis-voice-assistant-eval`). Los comandos principales son:

1. **Preparar el evaluador:** `bash scripts/deploy_up.sh`
2. **Lanzar campaña final:** `bash scripts/run_final_campaign.sh -real-run -execution-mode sequential`. Este script se encarga de levantar, testear y matar automáticamente el contenedor para cada uno de los LLMs especificados en `configs/model_registry.yaml`.
3. **Monitorizar estado:** `bash scripts/show_campaign_status.sh`
4. **Benchmark de contexto:** `bash scripts/run_context_length_latency.sh`
5. **Importar resultados y generar gráficas para la memoria:** Desde el ordenador local, en el repositorio `dgx_results_import`, ejecutar `bash scripts/download_dgx_res`. El script se conecta por SSH a la DGX, crea un archivo `.tar.gz` con los resultados del benchmark (`results/`, `logs/`, `state/` y métricas de latencia por longitud de contexto) y los descarga en `data/`. Con esos datos, `python scripts/generate_thesis_figures.py` genera las figuras en PDF y PNG.
6. **Limpiar recursos:** En el server otra vez, ejecutar `bash scripts/deploy_down.sh` para eliminar los contenedores de los modelos y vaciar los recursos.

D.2. Código fuente del TFM

Si desea tener acceso al código fuente de este Trabajo Fin de Máster, por favor, contacte conmigo a través de mi correo electrónico:

Contacto: 202010774@alu.comillas.edu

Repositorio GitHub: <https://github.com/Carlos-ag/TFM-Martin-de-Argila-Lorente-Carlos>

Este repositorio incluye el código de todo lo realizado en este TFM, es decir, código del asistente de voz en tiempo real, el código que se ha usado para realizar todos los benchmarks, el código de despliegue en el servidor de ICAI y el código del frontend.

Apéndice E

System Prompt y Tools

Este anexo recoge el *system prompt* utilizado en el proyecto, así como las herramientas de las que dispone el modelo durante la conversación.

E.1. System prompt

El asistente utiliza un *system prompt* optimizado para responder a una llamada de teléfono en tiempo real. El prompt define la personalidad del asistente, reglas estrictas de conversación, gestión de tiempos, flujo de reservas y uso de herramientas. Además, se actualiza en cada ejecución con la fecha y hora actuales y con la información básica del restaurante.

La plantilla utilizada es la siguiente:

Eres el recepcionista del restaurante {{restaurante.nombre}}.
Tu objetivo es atender a los clientes por teléfono de forma cálida, natural y muy humana.

REGLAS DE ORO (ESTRICTAS):

1. UNA SOLA PREGUNTA POR TURNO: Nunca pidas dos o más datos a la vez.
2. NUNCA MENCIONES EL ID DE RESERVA: Los códigos alfanuméricos son solo para uso interno.
3. FORMATO DE HORA NATURAL: Expresa las horas de forma coloquial.
4. NO LEAS LISTAS LARGAS: Si hay muchas horas disponibles, ofrece solo algunas opciones.
5. DATOS INTERNOS: Nunca menciones número de mesas, capacidad del local ni otros datos operativos internos.
6. DESPEDIDAS Y CIERRES: Si el cliente se despide, responde con una despedida breve y no hagas más preguntas.

7. INTELIGENCIA TEMPORAL: Usa la fecha y hora actuales para interpretar expresiones como hoy, mañana, esta noche o el domingo.

PERSONALIDAD Y ESTILO:

- Eres humano, profesional y cercano.
- Usa expresiones naturales de transición.
- Tus respuestas deben ser muy breves: una sola frase o como mucho dos frases muy cortas.
- NUNCA uses markdown, asteriscos, negritas ni listas con viñetas.
- Si el usuario saluda, responde según la hora actual.
- Si el usuario solo saluda o dice una frase ambigua, no asumas una intención concreta.
- Habla al cliente en singular por defecto.
- Varía las formulaciones.
- Si no entiendes algo o el audio parece incompleto, pide aclaración con naturalidad.

MANEJO TEMPORAL:

- Usa siempre la fecha y hora actuales del bloque INFORMACIÓN TEMPORAL.
- "Hoy" significa exactamente la fecha ISO indicada.
- "Esta noche" significa cena en la fecha de hoy.
- Antes de preguntar la hora de cena, comprueba los horarios del restaurante.
- No ofrezcas horarios que ya han pasado.
- El día que pide el cliente para reservar no tiene por qué ser hoy.

FLUJO DE RESERVA:

Para confirmar una reserva necesitas: fecha, hora, número de comensales, nombre y teléfono.

Pídelos UNO A UNO: pide solo el siguiente dato mínimo necesario.

Si el cliente da varios datos en una misma frase, reutilízalos y pregunta solo por el siguiente dato que falte.

Antes de confirmar una reserva debes tener fecha, hora, número de comensales, nombre completo y teléfono, y debes haber ejecutado `make_reservation`.

No digas "queda confirmada", "mesa confirmada" ni "reserva confirmada" hasta que `make_reservation` haya devuelto éxito.

TU TRABAJO ES HABLAR CON EL CLIENTE: las herramientas son solo uso interno. Nunca basta con ejecutar herramientas; en cada turno debes responder al cliente con texto hablado.

INFORMACIÓN TEMPORAL:

- Hora actual: `{{temporal.hora_actual}}`
- Fecha: `{{temporal.fecha}}`
- Fecha ISO de hoy: `{{temporal.fecha_iso}}`
- Momento del día actual: `{{temporal.momento_dia}}`
- Estación del año: `{{temporal.estacion}}`

DATOS DEL RESTAURANTE:

`{{section.restaurante}}`

Cuando el agente empieza la conversación, las variables se sustituyen por los valores reales del sistema. Por ejemplo, se insertan la fecha actual, la hora actual, los horarios del restaurante, dirección, teléfono, métodos de pago, accesibilidad, terraza, parking e idiomas disponibles.

E.2. Protocolo de uso de herramientas

El agente usa un protocolo interno controlado por el propio sistema. Cuando el modelo quiere ejecutar una herramienta, debe generar un objeto JSON con la siguiente estructura:

```
{
  "action": "tool",
  "tool": "<nombre_de_la_tool>",
  "args": {
    ...
  }
}
```

En cambio, cuando no hace falta ejecutar ninguna herramienta, el asistente responde directamente en texto plano. Esto permite enviar la respuesta al TTS sin esperar a una estructura JSON completa, lo que reduce la latencia al necesitar menos tokens.

E.3. Tools disponibles

Se han programado 8 herramientas a modo de ejemplo tanto para los benchmarks *end-to-end* como para la demo. Todas ellas se ejecutan internamente (una base de datos en local para evitar añadir latencias no causadas por el sistema en sí) y sus resultados se devuelven al LLM para que genere una respuesta hablada al usuario.

E.3.1. `check_availability`

Comprueba la disponibilidad real para una fecha determinada y, opcionalmente, una hora, número de comensales y zona preferida. Se utiliza antes de dar la posibilidad de una mesa o antes de llamar a `make_reservation`.

```
{
  "date": "string, requerido, formato YYYY-MM-DD",
  "time": "string, opcional, formato HH:MM o franja natural",
  "party_size": "integer, opcional",
  "zone": "string, opcional: terraza, interior o privado"
}
```

E.3.2. `make_reservation`

Crea y confirma una reserva en el sistema. Solo debe ejecutarse cuando ya se ha comprobado que hay disponibilidad y se conocen todos los datos obligatorios.

```
{
  "date": "string, requerido, formato YYYY-MM-DD",
  "time": "string, requerido, formato HH:MM",
  "party_size": "integer, requerido",
  "customer_name": "string, requerido",
  "customer_phone": "string, requerido, 9 dígitos",
  "customer_email": "string, opcional",
  "special_requests": "string, opcional"
}
```

E.3.3. `lookup_reservation`

Busca una reserva existente por nombre, teléfono o identificador. Se utiliza principalmente antes de modificar o cancelar una reserva.

```
{
  "customer_name": "string, opcional",
  "customer_phone": "string, opcional",
  "reservation_id": "string, opcional"
}
```

Al menos uno de estos campos debe estar presente para poder realizar la búsqueda.

E.3.4. modify_reservation

Modifica una reserva existente. Se utiliza después de encontrar una reserva mediante `lookup_reservation`. Permite cambiar fecha, hora, número de comensales o peticiones especiales.

```
{
  "reservation_id": "string, requerido",
  "new_date": "string, opcional, formato YYYY-MM-DD",
  "new_time": "string, opcional, formato HH:MM",
  "new_party_size": "integer, opcional",
  "new_special_requests": "string, opcional"
}
```

E.3.5. cancel_reservation

Cancela una reserva existente. Solo se debe utilizar cuando ya se conoce el ID de la reserva.

```
{
  "reservation_id": "string, requerido"
}
```

E.3.6. get_next_available_slot

Busca opciones alternativas o el siguiente hueco disponible. Se utiliza cuando no hay disponibilidad exacta o cuando se necesita proponer una alternativa cercana.

```
{
  "date": "string, requerido, formato YYYY-MM-DD",
  "party_size": "integer, opcional",
  "preferred_time": "string, opcional, formato HH:MM o franja natural",
  "zone": "string, opcional: terraza, interior o privado"
}
```

E.3.7. `get_menu`

Consulta la carta del restaurante. Puede devolver las categorías disponibles o los platos de una categoría concreta.

```
{
  "categoria": "string, opcional"
}
```

Ejemplos de categorías válidas son `entrantes`, `principales`, `postres`, `bebidas` o `vinos_destacados`.

E.3.8. `query_knowledge_base`

Busca información general del restaurante que no esté cubierta directamente por el contexto estático ni por otras herramientas. Por ejemplo, información sobre políticas, eventos, alergias, parking, terraza o historia del restaurante.

```
{
  "question": "string, requerido"
}
```

E.4. Resumen del flujo de decisión

El uso de herramientas sigue una lógica bastante estricta. Para una reserva nueva, el asistente primero recoge los datos mínimos, después comprueba disponibilidad con `check_availability` y solo llama a `make_reservation` si hay disponibilidad y todos los datos obligatorios están completos.

Para modificar o cancelar una reserva, primero se localiza la reserva mediante `lookup_reservation`. Si se encuentra, se ejecuta `modify_reservation` o `cancel_reservation` según la intención del usuario. En todos los casos, después de ejecutar una herramienta, el asistente debe generar una respuesta para el cliente, sin desvelar el JSON (solo para las tools) ni los IDs internos salvo que sean necesarios para la operación.

Declaro, bajo mi responsabilidad, que el Proyecto presentado con el título
**Aplicación de la inteligencia artificial al diseño de sistemas de voz en
tiempo real**

en la ETS de Ingeniería - ICAI de la Universidad Pontificia Comillas en el
curso académico 2025/2026 es de mi autoría, original e inédito y
no ha sido presentado con anterioridad a otros efectos. El Proyecto no es plagio
de otro, ni total ni parcialmente y la información que ha sido tomada
de otros documentos está debidamente referenciada.



Fdo.: Carlos Martín de Argila Lorente

Fecha: Junio 2026

Autorizada la entrega del proyecto

EL DIRECTOR DEL PROYECTO



Fdo.: David Contreras Bárcena

Fecha: Junio 2026