

Aplicación de modelos de lenguaje de gran escala en capacitación personalizada para entrevistas técnicas

Itzel Cabrera, Diego Flores, Bella Martinez-Seis, Obdulia Pichardo-Lagunas

Instituto Politécnico Nacional,
Unidad Profesional Interdisciplinaria en Ingeniería y Tecnologías Avanzadas,
México

{bcmartinez,opichardola}@ipn.mx

Resumen. Prepararse para una entrevista técnica es un desafío que requiere no solo conocimientos teóricos, sino también práctica efectiva y personalizada. El presente trabajo propone integrar Modelos de Lenguaje de Gran Escala (LLMs, por sus siglas en inglés) en el proceso de capacitación para entrevistas técnicas. El prototipo propuesto proporciona experiencias de aprendizaje personalizadas basadas en los conocimientos seleccionados del candidato. El sistema web ofrece una interfaz flexible, interactiva y atractiva. Analizamos el diseño y la implementación del sistema en un entorno de entrevista técnica en el área de Desarrollo Web. Especialistas del área evaluaron el rendimiento del sistema en tres aspectos: preguntas generadas, respuestas generadas y evaluación de las respuestas.

Palabras clave: Modelos de lenguaje de gran escala, entrevistas técnicas, aplicaciones AI.

Leveraging Large Language Models for Personalized Technical Interview Preparation

Abstract. Preparing for a technical interview is a challenging task that requires not only theoretical knowledge but also effective and personalized practice. This study proposes the integration of Large Language Models (LLMs) into the training process for technical interviews. The proposed prototype provides personalized learning experiences based on the candidate's selected knowledge areas. The web-based system offers a flexible, interactive, and engaging interface. We analyze the system's design and implementation in the context of technical interviews in the field of Web Development. Experts in the field evaluated the system's performance across three key aspects: the quality of generated questions, the quality of generated answers, and the evaluation of those answers.

Keywords: Large language models, technical interviews, AI application.

1. Introducción

El proceso de reclutamiento técnico representa un desafío tanto para las empresas como para los candidatos, especialmente en la evaluación de conocimientos específicos y habilidades prácticas. Con la evolución de la Inteligencia Artificial, los Modelos de Lenguaje de Gran Escala (LLMs, por sus siglas en inglés) han emergido como una herramienta que podría mejorar la eficiencia y precisión de estas evaluaciones.

Las LLMs ofrecen una solución más versátil y efectiva para la capacitación en entrevistas técnicas en comparación con otras tecnologías de Inteligencia Artificial (IA). Por ejemplo, los chatbots convencionales pueden responder preguntas pero con respuestas predefinidas. Los sistemas de evaluación automatizada basados en redes neuronales, requieren conjuntos de datos especializados y entrenamiento costoso; mientras otros sistemas expertos requieren conjuntos de reglas predefinidas que limitan su capacidad de adaptación a distintos perfiles de candidatos. En contraste, los LLMs pueden ajustar la dificultad de las preguntas, generar explicaciones según el nivel del usuario y ofrecer escenarios específicos según el área de especialización.

Por otro lado, en el caso particular de México, durante el 2020 hubo una aceleración en la digitalización de las empresas y es por ello que, las vacantes en el sector tecnológico aumentaron un 57%; colocando a los desarrolladores de front-end y back-end como dos de las profesiones más solicitadas [6]. Con relación a lo expuesto, se observa que el área de Desarrollo de Web es una de las áreas tecnológicas más importantes en el panorama socioeconómico.

En este artículo, exploramos cómo los LLMs pueden aplicarse en entrevistas técnicas personalizadas en el área de Desarrollo Web, ajustando preguntas y respuestas según el nivel de conocimiento del practicante. La presente propuesta abarca desde la generación de preguntas hasta la evaluación de ellas.

Para potenciar la LLM se hace uso del proceso conocido como prompt engineering para refinar y mejorar la entrada (prompt) a la Inteligencia Artificial Generativa (GenAI). En este desarrollo, se hace uso de los LLMs en tres procesos: (1) generación de preguntas correspondientes al área y nivel del candidato en torno al Desarrollo Web, (2) generación de respuestas esperadas para las preguntas y (3) evaluación de las respuestas del candidato a través del asistente de voz.

El desarrollo final corresponde al asistente inteligente de voz para entrevistas técnicas que abarca la extracción de datos en el currículum vitae del usuario, la generación de preguntas personalizadas a su perfil técnico, la transformación de las respuestas orales del usuario a texto y la evaluación de las preguntas.

El resto del artículo se compone de: Sección 2 explora el estado del arte de los LLMs enfocado principalmente a reclutamiento de personal, Sección 3 describe la propuesta del sistema, en esta se describirán los módulos que la componen así como la arquitectura usada en su implementación, la Sección 4 muestra la evaluación y resultados del LLM y la Sección 5 presenta las conclusiones.

2. LLMs para el reclutamiento de personal

En esta sección, presentaremos brevemente los antecedentes de los LLMs junto con una descripción general de los estudios relacionados centrados en el uso de LLMs enfocados al proceso de reclutamiento de personal como la generación de currículos, emparejamiento de vacantes y entrevistas entre reclutador y candidatos.

Los Modelos de Lenguaje Extendidos o Modelos de Lenguaje de Gran Escala (LLM) son una categoría de modelos funcionales entrenados sobre enormes cantidades de datos que los hacen capaces de generar lenguaje natural, entre otros tipos de contenidos, para realizar una amplia gama de tareas [4]. Los LLMs estiman la probabilidad de la generación de un símbolo o secuencia de símbolos dentro de una secuencia más extensa de símbolos, un LLM toma mayor potencial cuando se introduce el concepto de transformador (Transformer en inglés) que comprende de un codificador y decodificador.

Los LLMs se encuentran inmersos en varios procesos enfocados a la educación, capacitación y aprendizaje [11]. El uso de ellos para el aprendizaje de otros idiomas es muy concurrido, por ejemplo [15] integra GPT API para personalizar cursos de inglés. También se está usando para integrarla en metaversos [1] sobre todo involucrados en educación. Es evidente que cualquier proceso que involucre el intercambio de preguntas y respuestas podría ser apoyado por los LLMs sin embargo, aún hay aspectos por mejorar [5] como la identificación de respuestas o comunicación persuasiva en entrevistas periodísticas [7].

En el área de reclutamiento laboral existen trabajos que hacen uso de los LLMs como parte del proceso principal. Aguinalde et al. [2] las usaron para obtener las responsabilidades y tareas requeridas en ofertas para pasantes de ingeniería aeroespacial y posteriormente, consolidando la entrada para entrenar algoritmos de clasificación. DeBaer et al. [3] utilizan LLMs para generación de ejemplos de conversaciones para entrevistas de trabajo.

MockLLM [13] realiza generación de entrevistas simuladas y evaluación bilateral, con el objeto de encontrar las mejores relaciones persona-trabajo. Enfocado hacia la generación de preguntas para reclutadores de TI, Pasaribu et al. [10] dividen las entrevistas en conductual y técnico. Para el primero, usó fine-tuning en el modelo Longformer y para el segundo few-shot learning en el LLM GPT-4. Sin embargo, dichos trabajos se enfocan en las reclutadoras y no en los candidatos, como se propone en este trabajo.

Enfocado a los candidatos existen propuestas como la creación de currículos [14]. La Universidad Politécnica de Bucharest [12] desarrolló un simulador para entrevistas de trabajo combinando realidad virtual, reconocimiento de emociones y chatbots, este último a través del framework Pandarobots. ITEM [8] es una propuesta de entrenamiento para entrevistas con Realidad Virtual usando GenAI, donde el LLM de GPT genera preguntas para evaluar las habilidades de comunicación y liderazgo del estudiante. La inmersión en el metaverso genera un enfoque transformador, sin embargo, deja a un lado las entrevistas técnicas que son retomadas en esta propuesta.

3. Entrenamiento para entrevistas técnicas

Las entrevistas laborales técnicas se pueden clasificar en tres: estructuradas, no estructuradas y semiestructuradas. El trabajo fue desarrollado para entrevistas estructuradas cuyas preguntas se enfocan en las aptitudes del candidato evitando opiniones sesgadas, además de ser de corta duración ya que las preguntas se concretan antes de comenzar. La estructura de la aplicación está compuesta por cuatro módulos:

1. **Tzin**¹. Captura de campos clave de CV.
2. **Waach**². Generación de preguntas técnico-conceptuales personalizados.
3. **Marli**³. Procesamiento de voz a texto.
4. **Chari**⁴. Evaluación de las respuestas.

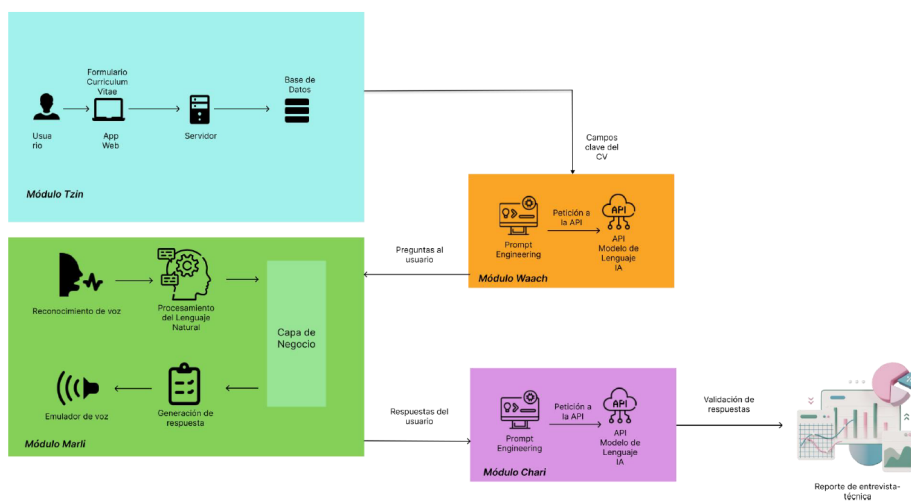


Fig. 1. Diagrama a bloques del prototipo de entrenamiento para entrevistas técnicas.

A continuación, se detalla cada módulo y se describe cómo se comunican. El uso de LLMs se presenta en los módulos Waach (Generación de preguntas) y Chari (Evaluación de respuestas). Posteriormente se describe la arquitectura usada para la implementación.

¹ Del Náhuatl clásico, significa *amado, respetado, querido*.

² Del maya yucateco, significa *soltar*.

³ Unión de dos nombres.

⁴ Diminutivo de Charango, instrumento musical.

3.1. Módulo Tzin

El módulo Tzin permite obtener los campos claves de un CV por competencias y nivel, el cual permite el registro de conocimiento y aptitudes a partir de categorías (con respecto a lo profesional) y a partir de cada una se pueden ingresar campos con su respectivo nivel de conocimiento.

3.2. Módulo Waach

El Módulo Waach obtiene las preguntas técnico-conceptuales a partir de los campos clave proporcionados por el módulo Tzin. Se evaluaron cualitativamente Gemini y GPT4 como se describe en la Sección 4.1. Para ello, se enlaza a la LLM cuyo diseño y evaluación de *prompt* se detalla en la Sección 4.2. La LLM utilizada es GPT4-o mini.

3.3. Módulo Marli

El módulo Marli es el encargado de procesamiento de voz. Marli requiere como entrada las preguntas generadas por Waach. El usuario responde verbalmente a las preguntas y se realiza la transformación a texto con la API SpeechRecognition de NPM debido a que es multilingüístico y no tiene costo ni límite.

3.4. Módulo Chari

El módulo Chari tiene como objetivo mostrar la evaluación de las respuestas (correcta o incorrecta); con el fin de determinar el desempeño del usuario en dicha entrevista. Para ello se hace uso nuevamente de la LLM GPT4-o mini. Finalmente, esta evaluación se guarda y se muestra en la aplicación web; de esta forma, se busca que el usuario visualice su progreso dadas las entrevistas que haya realizado.

3.5. Arquitectura e implementación

Para el proyecto se cuenta con una base de datos relacional donde el usuario tiene el histórico de sus currículums vitae, cada uno con una o varias categorías con su respectivo nivel. Además, se tiene un registro de las entrevistas donde se almacenan las preguntas y evaluación. El manejador de base de datos usado es SQL Server 2019, debido a su compatibilidad con Spring JPA y el entorno distribuido.

Para la arquitectura, se consideró la de microservicios (ver Figura 2) debido a que la propuesta involucra diferentes módulos que deben integrarse a una aplicación web. De esta forma, la aplicación cuenta con los siguientes principios: resiliencia, alta escalabilidad y disponibilidad.

La plataforma distribuida para la transmisión de datos utilizada es Kafka debido a su capacidad para manejar flujos de datos en tiempo real y de alta

velocidad entre módulos como la generación de preguntas, el procesamiento de respuestas y la evaluación de desempeño. Kafka permite almacenar mensajes durante un tiempo configurable, lo que facilita reprocesar datos históricos y escalar de manera eficiente a medida que aumenta el volumen de usuarios y entrevistas. Además, su modelo de consumo independiente (donde los consumidores rastrean su propio progreso) asegura flexibilidad para manejar tareas asincrónicas complejas, como la interacción con modelos de lenguaje y el procesamiento de entrevistas, características críticas para la aplicación.

La comunicación se lleva a cabo de la siguiente forma según la Figura 2:

1. El usuario realiza la conexión hacia el servidor de aplicaciones donde se encuentre el aplicativo bajo el protocolo HTTPS.
2. El aplicativo front-end realiza la solicitud a un recurso que contenga algún microservicio dentro del back-end, este será interceptado por el Gateway.
3. El Gateway es el filtro, para autenticar todas las peticiones que quieran acceder a las diferentes instancias de los microservicios y denegar aquellas que no estén autenticadas.
4. El Gateway conoce los microservicios que se hayan descubierto/registrado.
5. 6. 7. El Gateway realiza la petición y el balanceo de carga hacia la instancia, junto con el endpoint del recurso al cual quiere acceder, así mismo con el body request o query params que llegara a ocupar la API.
8. Publica en el bus de datos los campos del CV.
9. Se avalúa la entrevista con ayuda del LLM.

4. Evaluación y resultados

En este capítulo, se presenta el análisis cualitativo entre LLMs y los resultados obtenidos a partir de la evaluación del modelo de lenguaje de gran escala seleccionado. La evaluación se divide en varias secciones dentro del contexto de entrevistas: en primer lugar, se explora el diseño y evaluación de prompts, elementos fundamentales que guían la interacción con el modelo; en segundo lugar, se evalúa por expertos en el área, la capacidad del modelo para generar preguntas y respuestas relevantes en el contexto de entrevistas técnicas; y finalmente, se proporciona una visión general sobre el rendimiento del modelo.

4.1. Análisis cualitativo de LLMs

Los LLM son una categoría de modelos funcionales que permiten comprender y generar lenguaje natural. Se analizaron las LLM más competitivas en el año 2024: GPT, Gemini y BERT [9]. Para los dos primeros se hizo un análisis cualitativo en tres rubros: (1) generación de preguntas, (2) alcance de contenidos y (3) calificar respuestas. Siendo relevantes los dos primeros para el módulo Waach y el último para el módulo Marli. El análisis cualitativo comprende la prueba de prompts relacionados con el área y la evaluación de la respuesta por parte de expertos con respecto a la completitud, correctitud y objetividad. En

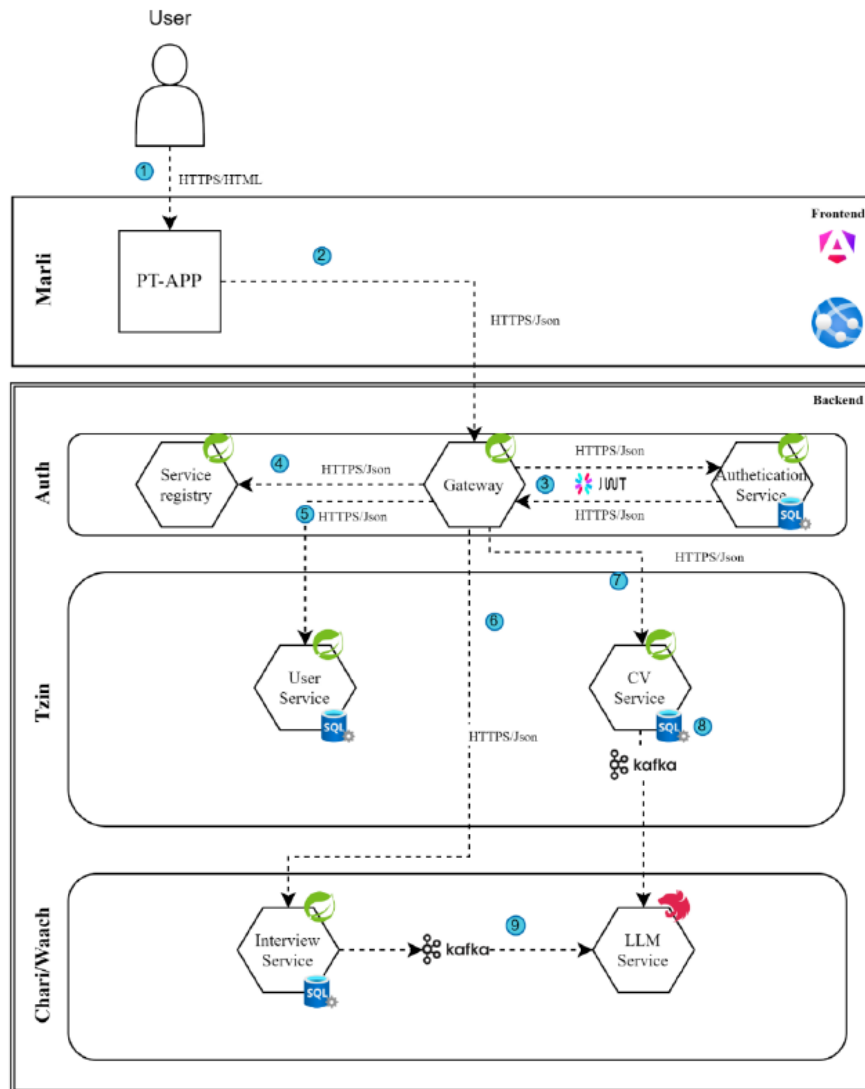


Fig. 2. Diagrama de arquitectura del sistema de entrenamiento para entrevistas técnicas.

la evaluación de **generación de preguntas**, GPT tuvo un mejor rendimiento que Gemini, ya que siguió directrices específicas y generó preguntas conceptuales en el formato indicado, mientras que Gemini se enfocó en preguntas prácticas fuera del alcance requerido. En la **evaluación de alcance de contenidos**, ambos modelos tuvieron un rendimiento similar, pero Gemini ofreció respuestas más contextuales y con enlaces de respaldo, lo que puede reducir alucinaciones

y validar resultados. En la **calificación de respuestas**, ambos modelos son comparables, pero Gemini proporciona contexto adicional en sus evaluaciones, mientras que GPT es más estricto con errores tipográficos, ortográficos y gramaticales. En conclusión, tanto GPT como Gemini son opciones competitivas para generar preguntas y evaluar respuestas. Sin embargo, debido a que se requieren respuestas concretas por parte del LLM, se optó por GPT.

4.2. Diseño y evaluación de *promts*

De acuerdo con la metodología GPEI: metodología para un Prompt Engineering eficiente, en el diseño del prompt de la aplicación se siguieron los siguientes pasos:

1. Definición de metas:
 - Generar preguntas-respuestas de acuerdo a los campos clave del CV.
 - Evaluar respuestas del usuario, teniendo como referencia las respuestas correctas que el LLM haya generado
2. Elementos a incluir en el prompt:
 - Incluir detalles.
 - Pedir al modelo que tome un rol.
 - Proveer ejemplos.
 - Especificar la extensión de la respuesta.
 - Pedir al modelo que responda usando algún texto de referencia.
3. Etapa de iteración consiste en probar una versión del prompt y si la respuesta no cumple con los criterios de evaluación, entonces el prompt se modifica de forma inmediata.

La evolución del *prompt* para generar las **preguntas con sus respuestas** fue la siguiente: El prompt inicial contenía los elementos enumerados con anterioridad considerando los campos del CV del usuario. En la segunda iteración se solicitó expresamente el número de preguntas y se dieron ejemplos. En la tercera, se aumentó el contexto y se agregó el enfoque de las preguntas y se solicitó la respuesta correcta. En la cuarta se limitó la extensión de la respuesta. La quinta se clarificó que se requiere en idioma español.

El *prompt* inicial para **calificar respuestas** indicaba qué rol tenía, qué realizarías, descripción de qué se le dará, en qué formato, cuál es el objetivo, cuál es el formato de respuesta y se intercalan definiciones para evitar ambigüedades. En la segunda iteración se aclaró que la evaluación la hará desde su perspectiva considerando correcta si la respuesta tiene noción del tema. En la tercera, se agregó la respuesta correcta para tenerlo como referencia en la evaluación. La cuarta, se aclara que el enfoque sea el contenido. Finalmente, en la quinta, se especifica el idioma.

4.3. Evaluación de Generación de Preguntas-Respuestas por la LLM

En la evaluación participaron 3 expertos que se describen a continuación:

1. Primer experto. 12 años de trayectoria, cargos como desarrollador Java, líder técnico y arquitecto de soluciones, desempeñando tareas de planeación y gestión de personal, así mismo validando el conocimiento y capacidades técnicas del equipo que tuviera a su cargo.
2. Segundo experto. 6 años de trayectoria, desempeñándose como desarrollador Java y líder de proyectos, participe de procesos de reclutamiento para labores dentro del Desarrollo Web.
3. Tercer experto. 26 años de experiencia, gerente de gestión de proyectos y diferentes gerencias, pasante de maestría en ciberseguridad, participante en subáreas del desarrollo web. Ha diseñado y validado de habilidades técnicas en procesos de reclutamiento.

Se realiza el análisis de los resultados obtenidos con el modelo GPT-4-o mini. Se evaluaron los resultados proporcionados por la LLM en 3 aspectos:

1. Evaluación de preguntas generadas. Se evalúa la correspondencia con la categoría y nivel. Se espera que la pregunta generada corresponda al tópico y nivel ingresado, los niveles considerados fueron: Trainee, Junior y Senior.
2. Evaluación de respuestas generadas. Se valida la respuesta generada por la LLM es correcta con respecto a la pregunta generada.
3. Evaluación de calificación de respuestas. Se valida que la calificación asignada por la LLM de las respuestas del usuario sea correcta.

Se proporcionó a los expertos las respuestas generadas por el LLM. Ellos etiquetaron como correcto o incorrecto según el aspecto que estén evaluando. La participación de tres expertos permite discernir si lo arrojado por el LLM es correcto, ya que podrían los tres coincidir en que sea correcto o no, pero podría darse el caso de tener diferencia de opinión, en ese caso si dos de tres coinciden es la evaluación que se compara con el LLM.

Evaluación de preguntas generadas. Para el nivel Trainee se tuvo una coincidencia del 88.8 % de la evaluación manual de los expertos con respecto al nivel. Para el nivel de Junior la coincidencia fue también de 88.8 % y para el nivel de Senior fue del 100 %. Se puede concluir que las preguntas que GPT 4-o mini proporciona son las adecuadas al nivel correspondiente del usuario. Esto sugiere que el Prompt para la generación de preguntas-respuestas permite al modelo generar cuestionamientos coherentes y relevantes para el contexto técnico.

Evaluación de respuestas generadas. Los expertos evaluaron las respuestas generadas en los tres niveles, de tal forma que el 92.59 % de las respuestas era considerada correcta por los expertos. Se puede concluir que las respuestas que GPT4-o mini proporciona son las adecuadas tanto en contenido como en extensión. Este resultado confirma que el modelo no solo interpreta correctamente las preguntas, sino que también genera respuestas claras, completas y bien estructuradas. La capacidad del modelo para mantener consistencia en la calidad de las respuestas respalda la idea de que el diseño del Prompt para la generación de preguntas-respuestas es robusto y eficiente.

Evaluación de calificaciones generadas. El 66 % de las calificaciones son correctas según el promedio de los expertos. Aunque este porcentaje es significativo, también evidencia un área de mejora en el diseño del Prompt final para la evaluación de respuestas y en la configuración del modelo para evaluar respuestas de manera más precisa. Este resultado sugiere la necesidad de realizar ajustes al Prompt final para la evaluación de respuestas para mejorar la precisión de las calificaciones, así como considerar estrategias complementarias, como la incorporación de ejemplos adicionales o criterios más específicos de evaluación.

4.4. Análisis de evaluación.

Los tres expertos consultados coinciden en que el modelo realiza las tareas asignadas de manera coherente y eficiente en la mayoría de los casos. Sin embargo, mientras que los resultados de los dos primeros rubros reflejan un alto nivel de desempeño, el tercer rubro indica que existe margen para optimizar el modelo en la tarea de calificación. Esto resalta la importancia de iterar en el diseño del prompt para evaluar las respuestas y de explorar ajustes adicionales para maximizar la efectividad del modelo en todos los aspectos evaluados. Este análisis subraya el potencial del modelo para aplicaciones prácticas, mientras que también enfatiza la necesidad de una mejora continua para alcanzar una mayor precisión y consistencia en sus resultados.

5. Conclusión

La propuesta implementó una arquitectura de microservicios que integra Kafka para la comunicación en tiempo real y el uso de LLMs para dar solución al entrenamiento personalizado en entrevistas técnicas en el área de Desarrollo Web.

Las tareas principales del proyecto son la generación de preguntas y evaluación de respuestas. La generación de preguntas se realiza a partir de los campos del CV proporcionados por los usuarios, lo que asegura que las entrevistas sean personalizadas y se adapten a las competencias de cada candidato. Este enfoque hace que las entrevistas sean más relevantes y efectivas en la evaluación de habilidades técnicas. Al mismo tiempo, las respuestas correctas generadas por el modelo sirven como referencia para la evaluación posterior de las respuestas del usuario, asegurando que la evaluación sea coherente y alineada con el conocimiento técnico esperado.

Durante el análisis, un desafío importante fue la valoración y comparación de los diferentes LLM existentes en el mercado; esto debido a la constante actualización de modelos. Tanto GPT como Gemini fueron opciones competitivas para generar preguntas y evaluar respuestas. Se presenta una evaluación manual dada por expertos en el área según sus conocimientos, lo cual genera una validación que depende de los etiquetadores. Sin embargo, no se requirió un banco de datos para el desarrollo del proyecto.

Referencias

1. Abdullakutty, F., Qayyum, A., Qadir, J.: Trustworthy AI for educational metaverses. *Authorea Preprints*, (2024)
2. Aguinalde, P., Shin, J., Carroll, B. F., Crippen, K. J.: Leveraging large language models to automatically investigate core tasks within undergraduate engineering work-integrated learning experiences. In: 2024 IEEE Frontiers in Education Conference (FIE). pp. 1–8. IEEE (2024)
3. De Baer, J., Doğruöz, A. S., Demeester, T., Develder, C.: Single-vs. dual-prompt dialogue generation with LLMs for job interviews in human resources. *arXiv preprint arXiv:2502.18650*, (2025)
4. IBM: Modelos de Lenguaje Grande (LLM) (2025), <https://www.ibm.com/mx-es/topics/large-language-models>, consultado el 25 de marzo de 2025
5. Kamerlin, S. C. L., Ratcliff, W. C.: Using AI to prepare for academic interviews—don't trade authenticity for polish (2025)
6. LinkedIn Business: Jobs on the Rise - Tendencias en Contratación (2025), <https://business.linkedin.com/es-mx/talent-solutions/resources/talent-acquisition/jobs-on-the-rise-cont-fact>, consultado el 25 de marzo de 2025
7. Lu, M., Cho, H. J., Shi, W., May, J., Spangher, A.: Newsinterview: a dataset and a playground to evaluate LLMs' ground gap via informational interviews. *arXiv preprint arXiv:2411.13779*, (2024)
8. Nofal, A. B., Ali, H., Hadi, M., Ahmad, A., Qayyum, A., Johri, A., Al-Fuqaha, A., Qadir, J.: Ai-enhanced interview simulation in the metaverse: Transforming professional skills training through VR and generative conversational AI. *Computers and Education: Artificial Intelligence*, vol. 8, pp. 100347 (2025)
9. Ozdemir, S.: Quick start guide to large language models: strategies and best practices for using ChatGPT and other LLMs. Addison-Wesley Professional (2023)
10. Pasaribu, D. K. H., Dewandaru, A., Saptawati, G. A. P.: Development of LLM-based system for IT talent interview. In: 2024 IEEE International Conference on Data and Software Engineering (ICoDSE). pp. 108–113. IEEE (2024)
11. Sallam, M.: ChatGPT utility in healthcare education, research, and practice: systematic review on the promising perspectives and valid concerns. In: *Healthcare*. vol. 11, pp. 887. MDPI (2023)
12. Stanica, I., Dascalu, M.-I., Bodea, C. N., Moldoveanu, A. D. B.: VR job interview simulator: where virtual reality meets artificial intelligence for education. In: 2018 Zooming innovation in consumer technologies conference (ZINC). pp. 9–12. IEEE (2018)
13. Sun, H., Lin, H., Yan, H., Zhu, C., Song, Y., Gao, X., Shang, S., Yan, R.: Facilitating multi-role and multi-behavior collaboration of large language models for online job seeking and recruiting. *arXiv preprint arXiv:2405.18113*, (2024)
14. Sunico, R. J., Pachchigar, S., Kumar, V., Shah, I., Wang, J., Song, I.: Resume building application based on LLM (large language model). In: 2023 International Conference on Computing, Communication, and Intelligent Systems (ICCCIS). pp. 486–492. IEEE (2023)
15. Yan, H.: Let's chat: Integrating large language models into blended learning of English for specific purposes. In: 2023 International Symposium on Educational Technology (ISET). pp. 187–192. IEEE (2023)