

El cerebro del agente: qué es un LLM

Datos de esta lección

Fase	F1 — Cimientos, seguridad y arquitectura del sistema
Módulo	M1.2 — Modelos mentales: radiografía de un agente telefónico
Puntos de profesionalidad al superar el test	15

Objetivo práctico

Al terminar esta lección podrás explicar con tus propias palabras qué es un LLM (modelo de lenguaje de gran tamaño), qué puede y qué no puede hacer de forma fiable, y por qué es la pieza que decide qué responde tu agente telefónico.

Resultado que obtendrás

Un modelo mental claro y correcto de qué es un LLM, que usarás en cada lección posterior en la que aparezcan palabras como «modelo», «prompt» o «temperatura».

Dónde encaja dentro del sistema final

En el sistema TELÉFONO → AGENTE → CEREBRO → VOZ → AUTOMATIZACIÓN → CALENDARIO, el LLM es exactamente el CEREBRO: recibe lo que ha dicho la persona que llama (ya convertido en texto) y decide qué responder, cuándo pedir un dato, y cuándo activar una automatización.

Conocimientos previos

- Módulo 1.1 completo (preparación del entorno de trabajo en Windows 11 con NVDA).

Herramientas necesarias

- Un navegador con NVDA activo, en modo de exploración (lección «Configura el navegador para leer con comodidad con NVDA»).

Posible coste

Esta lección es conceptual y no requiere ningún gasto ni ninguna cuenta.

Modelo mental

Piensa en el cerebro y el manual de instrucciones que utiliza un telefonista para decidir qué responder. El telefonista no improvisa sin criterio: tiene un manual con el tono de la empresa, la información que puede dar y la que no, y una enorme experiencia acumulada de conversaciones anteriores parecidas. Un LLM funciona de forma comparable: no «busca en una base de datos» una respuesta exacta, sino que genera la respuesta más adecuada según patrones aprendidos a partir de una enorme cantidad de texto, y según las instrucciones concretas que se le dan en cada conversación.

Explicación

LLM son las siglas de *Large Language Model* (modelo de lenguaje de gran tamaño). Es un programa entrenado con grandes cantidades de texto que aprende a predecir, palabra a palabra, cuál es la continuación más adecuada de un texto dado. Ejemplos de LLM que probablemente ya conoces son ChatGPT (de OpenAI), Claude (de Anthropic) y Gemini (de Google).

Un LLM no «sabe» hechos como los sabe una base de datos: no consulta una tabla con la respuesta exacta. Genera texto de forma probabilística, lo cual tiene dos consecuencias muy importantes para un agente telefónico profesional: primero, un LLM puede «inventar» información que suena convincente pero es incorrecta (a esto se le llama *alucinación*); segundo, cuanto más clara, concreta y limitada sea la instrucción que le des (el prompt), más fiable será su comportamiento. Por eso el Prompt Maestro que construirás en el módulo 1.3 incluye explícitamente la prohibición de inventar datos que el agente no tiene.

Vocabulario nuevo

LLM

Piénsalo así: El cerebro y el manual de instrucciones del telefonista.

Modelo de lenguaje de gran tamaño (*Large Language Model*): programa entrenado para generar texto coherente y adecuado al contexto, a partir de una instrucción o conversación.

Ejemplo: ChatGPT, Claude y Gemini son interfaces de conversación construidas sobre distintos LLM.

Alucinación

Piénsalo así: Cuando el telefonista, para no quedar mal, inventa una respuesta en lugar de decir «no lo sé».

Situación en la que un LLM genera información incorrecta o inventada, presentándola con la misma seguridad que una información correcta.

Ejemplo: Un agente mal configurado podría inventar un precio o un horario que no existe en la información real del negocio.

Preparación

No se requiere preparación técnica. Basta con leer con atención esta lección.

Procedimiento paso a paso

1. **Contexto:** Vas a comprobar que puedes localizar por tu cuenta, sólo con el teclado, la definición de «LLM» dentro del Glosario acumulativo de esta plataforma (una página real de esta web, no un ejemplo).

Acción de teclado: Abre el enlace «Glosario» de la navegación principal (o ve a la dirección /glosario.php). Una vez dentro, pulsa **H** repetidamente para recorrer los encabezados hasta llegar al encabezado de nivel 2 que dice «L».

Respuesta esperada de NVDA: NVDA debería anunciar un texto similar a «L, encabezado de nivel 2».

Qué significa: Has llegado al bloque de términos que empiezan por «L», donde está agrupado «LLM».

Acción siguiente: Sigue pulsando **H** una vez más hasta el encabezado de nivel 3 «LLM» y lee su definición completa avanzando línea a línea con la flecha abajo.

2. **Contexto:** Ahora localizas el segundo término de esta lección, «Alucinación», sin recorrer todo el glosario letra por letra.

Acción de teclado: Pulsa **Ctrl + F** para abrir la búsqueda del navegador, escribe «Alucinación» y pulsa **Enter**.

Respuesta esperada de NVDA: El navegador resalta la primera aparición del término en la página; al cerrar la búsqueda y mover el foco hacia esa zona, NVDA debería anunciar el encabezado de nivel 3 «Alucinación».

Qué significa: Confirmas que el término está documentado en el glosario y puedes leer su definición completa sin depender de la navegación letra a letra.

Acción siguiente: Cierra la búsqueda con `Escape` y lee la definición y el ejemplo de «Alucinación» con la flecha abajo.

Errores frecuentes y diagnóstico

Asumir que el LLM «sabe» automáticamente los datos concretos de un negocio (horarios, precios, servicios) sin que nadie se los haya dado.

Cómo localizarlo: Si el agente responde con información plausible pero que nadie ha introducido en su configuración, es una alucinación, no un fallo de conexión.

Solución: Este problema se resuelve dando al LLM, dentro del prompt o mediante herramientas de consulta, exactamente los datos reales que debe usar, y prohibiéndole explícitamente inventar lo que no conoce.

Comprobación final

Explica en voz alta, con tus propias palabras y sin mirar esta lección, qué es un LLM usando la analogía del telefonista y su manual de instrucciones. Si consigues explicarlo sin usar la palabra «inteligencia artificial» como comodín, la comprobación es positiva.

Qué acabas de conseguir

Tienes el modelo mental correcto del «cerebro» del sistema, necesario para entender por qué el Prompt Maestro (módulo 1.3) es la pieza más importante para controlar el comportamiento del agente.

Ejercicio práctico

Piensa en tres preguntas que un cliente podría hacerle a un agente telefónico de fontanería (por ejemplo, sobre precios, horarios y zonas de cobertura) y anota, para cada una, si un LLM sin ninguna información adicional podría responder con datos

reales o tendría que inventarlos.

Resumen de lo imprescindible

- LLM significa modelo de lenguaje de gran tamaño: genera texto, no consulta una base de datos exacta.
- Un LLM puede alucinar (inventar información) si no se le da la información correcta ni se le limita explícitamente.
- El LLM es el «cerebro» del sistema: decide qué responder a partir del prompt y de la conversación.
- Cuanto más clara y concreta sea la instrucción (prompt), más fiable es el comportamiento del LLM.

No necesitas aprender todavía

Todavía no necesitas saber cómo se entrena un LLM ni qué es una red neuronal. Este curso no requiere ese nivel técnico: se centra en usar el LLM correctamente a través de un prompt, no en construirlo.

Estoy preparado para continuar si puedo...

- Explicar con mis palabras qué es un LLM usando la analogía del telefonista.
- Explicar qué es una alucinación y por qué ocurre.
- Entender por qué darle al LLM información real y límites claros reduce el riesgo de que invente datos.

Test de comprobación

Para registrar el resultado y obtener Puntos de profesionalidad, realiza este test desde la versión web de la lección. Las opciones se muestran aquí sin señalar la respuesta correcta.

1. Un agente telefónico responde con seguridad que el negocio abre los domingos, pero esa información nunca se le proporcionó y el negocio en realidad cierra los domingos. ¿Qué está ocurriendo?
 - a. Un fallo de conexión telefónica.
 - b. Una alucinación del LLM: ha generado una respuesta plausible sin tener el dato real.
 - c. Un problema del servicio de voz (TTS).
2. ¿Por qué se dice que un LLM «genera» texto en lugar de «consultarlo» en una

base de datos?

- a. Porque predice, palabra a palabra, la continuación más adecuada según patrones aprendidos, no porque busque una respuesta exacta guardada.
 - b. Porque los LLM no pueden acceder nunca a ningún dato externo bajo ninguna circunstancia.
 - c. Porque «generar» y «consultar» son la misma operación con distinto nombre.
3. Dos negocios usan el mismo LLM para su agente telefónico. El negocio A le proporciona en el prompt sus horarios y precios reales y le prohíbe explícitamente inventar datos; el negocio B usa un prompt genérico sin esa información. ¿Qué negocio tiene más riesgo de sufrir alucinaciones sobre horarios y precios?
- a. El negocio B, porque su LLM tendrá que generar una respuesta plausible sin disponer del dato real ni de una instrucción que se lo impida.
 - b. El negocio A, porque darle más información al LLM aumenta siempre el riesgo de que invente datos.
 - c. Ambos tienen exactamente el mismo riesgo, porque el riesgo de alucinación no depende del prompt.

Fuentes

- [How to work with large language models — OpenAI \(consultado el 2026-08-30\)](#)
- [Intro to Claude — Anthropic \(consultado el 2026-08-30\)](#)

Última actualización de este contenido: **2026-02-10**.