



IALAB

AI academy
by doinGlobal

AI

BREAKING

NEWS

AGOSTO 2025 | EPISODIO 2

NOVEDADES EN IA DE LA SEMANA



Gemini incorpora memoria automatica y modo temporal

Google incorporó una configuración nueva a su asistente Gemini que permite que el sistema retenga automáticamente detalles y preferencias clave compartidas durante conversaciones anteriores. Esta funcionalidad, llamada Personal Context, se encuentra activa para los usuarios del modelo 2.5 Pro en algunos países y estará disponible próximamente para la versión 2.5 Flash.

Cuando el ajuste está activado, Gemini usa el historial de chats para brindar respuestas más naturales y relevantes. Esta memoria automática extiende una función previa, que requería que el usuario pida explícitamente que Gemini recordara algo.

En paralelo, también introdujo un modo de privacidad denominado Temporary Chats, concebido como una sesión de chat temporal que no se guarda ni se incorpora al aprendizaje del modelo. Estas conversaciones desaparecen después de hasta 72 horas y no influyen en interacciones posteriores.

Además, se agregaron controles adicionales para regular el tratamiento de los archivos o fotos enviados a Gemini, y el antiguo ajuste Gemini Apps Activity será renombrado como Keep Activity, conservando la posibilidad de desactivarlo por completo



OpenAI amplía opciones en GPT-5 con nuevos modos de uso y reactivación de modelos anteriores

OpenAI habilitó un selector que permite elegir entre los modos Auto, Fast y Thinking al usar GPT-5. Según confirmó el CEO Sam Altman, el modo "Thinking" brinda mayor profundidad de razonamiento para tareas complejas, mientras que "Fast" prioriza la rapidez y "Auto" delega al sistema la selección del modo más adecuado según el diálogo.

El modo Thinking ahora permite enviar hasta 3 000 mensajes por semana, lo que representa un incremento considerable frente al límite anterior, y ofrece una ventana de contexto de 196 000 tokens, lo que facilita el procesamiento de conversaciones o documentos extensos en una sola interacción.

Además, Pay-subscriptores pueden reactivar GPT-4o desde el selector, y también acceder a modelos heredados como o3 y o4-mini mediante una opción de configuración dentro de ChatGPT.

Tras su lanzamiento, GPT-5 ha recibido una oleada de críticas por parte de usuarios que comparan su desempeño con el de GPT-4o. Muchos destacan que sus respuestas son más cortas, menos detalladas y con un tono más frío o distante, describiéndolo como “plano”, “sin creatividad” o incluso “lobotomizado”, en contraste con el estilo más personal y empático de su predecesor.

Desde UBA IALAB testeamos GPT-5 en 82 pruebas en comparación con sus modelos anteriores

Con la presentación oficial de ChatGPT-5, desde UBA IALAB decidimos retomar una línea de trabajo iniciada en marzo de 2023; evaluar si los modelos de lenguaje más avanzados logran resolver errores ya identificados en versiones anteriores. Para ello, aplicamos nuevamente un conjunto de 82 pruebas en las que GPT-3.5 y GPT-4 habían fallado, y comparamos los resultados con GPT-4o, GPT-o1 y ahora GPT-5.

Los datos muestran que, al menos en este conjunto de casos, no se observan mejoras significativas ni en precisión ni en comportamiento sesgado. En términos de respuestas correctas, GPT-3.5 no resolvió ningún caso (0%). GPT-o1 fue el que obtuvo mejor rendimiento, con un 68% de aciertos, seguido por GPT-4o (60%) y GPT-5 (61%). En relación con el sesgo, GPT-3.5 volvió a ser el peor con un 40%, mientras que GPT-4o alcanzó un 34% y tanto GPT-o1 como GPT-5 mostraron un 27%.

La prueba se diseñó incluyendo únicamente ejemplos en los que los modelos evaluados en 2023 habían fallado. En enero de 2025, ese mismo conjunto de 82 pruebas fue reutilizado para evaluar a GPT-4o y GPT-o1. Con la llegada de GPT-5, se aplicó de forma idéntica, sin modificaciones en los enunciados ni en el sistema de evaluación, lo que permite comparar los resultados de manera directa entre versiones.

Para ver el libro con todas las pruebas que hicimos en 2023 hace click [acá](https://lnkd.in/d_2zueCK):
https://lnkd.in/d_2zueCK



Claude incorpora memoria y amplía su capacidad de contexto a un millón de tokens

Anthropic habilitó una nueva función para Claude que permite buscar y recuperar información de conversaciones previas, facilitando retomar proyectos sin volver a empezar. Los usuarios pueden activar esta capacidad, solicitar específicamente que Claude revise otros chats, y hacer avanzar el diálogo con contexto relevante.

La herramienta ya está disponible para suscriptores de los planes Max, Team y Enterprise, con acceso planificado para otros niveles próximamente. Claude no almacena recuerdos de forma automática; el buscado de conversaciones anteriores se realiza solo cuando el usuario lo solicita desde la configuración.

Además, Claude 4 Sonnet ahora admite hasta un millón de tokens de contexto en la API, lo que representa un aumento de cinco veces respecto del límite anterior. Esta capacidad permite procesar más de 75.000 líneas de código o 2.500 páginas de texto. Esto facilita el análisis técnico de repositorios, el examen de conjuntos documentales grandes o la comparación de múltiples fuentes dentro de una misma interacción

xAI avanza con un selector de modelos y lanza Grok Imagine

xAI habilitó un selector que permite a los usuarios elegir explícitamente qué versión de Grok utilizar en cada interacción. La interfaz ofrece cuatro alternativas: Grok 3 para respuestas rápidas, Grok 4 para razonamiento más profundo, Grok 4 Heavy para tareas complejas en paralelo, y un modo automático que elige el modelo más adecuado según el tipo de consulta.

Por otro lado presentó Grok Imagine, una nueva herramienta que permite transformar texto o imágenes en clips animados de seis a quince segundos con sonido. Disponibles para usuarios SuperGrok y PremiumPlus en iOS, y ahora también accesible para Android.

Se ofrecen cuatro modos de estilo y animación: Normal, Fun, Custom y Spicy. La activación del modo Spicy refleja una apuesta deliberada de xAI por permitir imágenes con menos restricciones, una diferenciación clara frente a sus competidores. Esta estrategia ha generado críticas por los riesgos de deepfakes no consensuados, fugas de contenido sensible y el desafío regulatorio del nuevo panorama legal sobre pornografía generada por inteligencia artificial.



Mistral lanza Medium 3.1 con mejoras en rendimiento, tono conversacional y búsquedas web

Mistral AI publicó una nueva versión de su modelo base, Mistral Medium 3.1, ya disponible como modelo predeterminado en su asistente Le Chat y accesible mediante API.

En cuanto al rendimiento, la versión 3.1 incorpora mejoras que impactan directamente en la capacidad de razonamiento, en la generación de código y en la comprensión de entradas más complejas. Estos cambios están orientados a optimizar la respuesta del modelo en tareas profesionales y técnicas sin recurrir a modelos de mayor costo o carga computacional.

Una novedad destacada de esta versión es el ajuste del tono conversacional. El modelo ha sido calibrado para ofrecer respuestas más coherentes, menos abruptas y con un estilo más apropiado al contexto del intercambio.

Otro cambio significativo se encuentra en la optimización de las búsquedas web integradas. El modelo ahora interpreta de manera más precisa las consultas que requieren apoyo externo, genera resultados más pertinentes y presenta respuestas con información mejor alineada a la actualidad.

La actualización ya está activa en todos los entornos habilitados, sin necesidad de configuración adicional. En el caso de integraciones mediante API, la nueva versión puede invocarse directamente con el endpoint correspondiente.



IALAB



academy
by doinGlobal



*También podés escuchar el resumen de audio
de las novedades haciendo [click aquí](#).*