

Estudio | Uso de las IAs para escribir reduce capacidad del cerebro

Un reciente estudio del [MIT](#) (Massachusetts Institute of Technology) ha arrojado luz sobre el «costo cognitivo» asociado al uso de **Modelos de Lenguaje Grandes (LLMs)**, como **ChatGPT**, en el contexto educativo, específicamente en la tarea de escribir ensayos. La investigación, aún en fase de pre-publicación y pendiente de revisión por pares, sugiere que el uso de estas herramientas podría tener un impacto medible en las habilidades de aprendizaje de los usuarios.

El estudio, que busca comprender cómo la adopción masiva de los LLMs afecta a humanos y empresas, se centró en evaluar la carga cognitiva de los participantes. Para ello, se asignaron **54 participantes** a tres grupos iniciales durante **tres sesiones**: un **grupo de LLM** (usando la herramienta), un **grupo de motor de búsqueda** y un **grupo «solo cerebro»** (sin herramientas). En una cuarta sesión, 18 de estos participantes cambiaron de grupo: el grupo LLM pasó a no usar herramientas (referido como **LLM-a-Cerebro**), y el grupo «solo cerebro» comenzó a usar LLM (**Cerebro-a-LLM**).

Para obtener una comprensión profunda de las activaciones neuronales durante la escritura, los investigadores registraron la **actividad cerebral** de los **participantes** mediante **electroencefalografía (EEG)**. Además, se realizó un **análisis de procesamiento del lenguaje natural (NLP)** y se entrevistó a cada participante después de cada sesión. La calificación de los ensayos se llevó a cabo con la ayuda de **profesores humanos** y un **juez de IA** especialmente diseñado.

Hallazgos más relevantes:

- **Diferencias en los patrones de conectividad neural:** El análisis de EEG reveló que los grupos de LLM, motor de búsqueda y «solo cerebro» presentaban **patrones de conectividad neural significativamente diferentes**, lo que refleja distintas estrategias cognitivas.
- **Reducción sistemática de la conectividad cerebral con el apoyo externo:** La conectividad cerebral disminuyó sistemáticamente con la cantidad de apoyo externo. El grupo «solo cerebro» exhibió las redes más fuertes y de

mayor alcance, el **grupo de motor de búsqueda** mostró un compromiso intermedio, y la **asistencia del LLM** provocó el acoplamiento general más débil.

▪ **Impacto en el cambio de herramientas (Sesión 4):**

- Los participantes del grupo **LLM-a-Cerebro** mostraron una **conectividad neural más débil** y una menor activación de las redes alfa y beta.
- Los participantes del grupo **Cerebro-a-LLM** demostraron una **mayor recuperación de la memoria** y una reactivación de nodos occipito-parietales y prefrontales extendidos, lo que probablemente respalda el procesamiento visual, similar a lo que se observa con frecuencia en el grupo del motor de búsqueda.

▪ **Baja «propiedad» y capacidad de citación en el grupo LLM:** Los participantes del grupo LLM reportaron una **baja «propiedad» de sus ensayos** en las entrevistas. Además, su **capacidad para citar** de los ensayos que habían escrito minutos antes fue notablemente inferior. El grupo del motor de búsqueda mostró una fuerte «propiedad», aunque menor que el grupo «solo cerebro».

▪ **Homogeneidad lingüística dentro de los grupos:** Se encontró una consistencia en el reconocimiento de entidades nombradas (NERs), n-gramas y la ontología de temas dentro de cada grupo.

▪ **Peor desempeño del grupo LLM a lo largo del tiempo:** A lo largo de las cuatro sesiones, que se llevaron a cabo durante cuatro meses, los participantes del grupo LLM mostraron un **peor desempeño** que sus contrapartes del grupo «solo cerebro» en todos los niveles: neural, lingüístico y de puntuación.

Consideraciones y advertencias importantes:

El estudio se presenta como una **guía preliminar** para fomentar una mejor comprensión de los impactos cognitivos y prácticos de la IA en los entornos de aprendizaje. Los autores enfatizan que el documento es una **pre-impresión** y, como tal, aún no ha sido revisado por pares. Por lo tanto, todas las conclusiones deben tratarse con precaución y considerarse preliminares.

Los investigadores también señalan **varias limitaciones** y futuras vías de investigación necesarias:

- **Número limitado de participantes** de un área geográfica y académica específica. Se necesita una muestra más grande y diversa en futuros estudios.

- El estudio se realizó con **ChatGPT**; los resultados no pueden generalizarse directamente a otros modelos de LLM sin más investigación.
- Futuros estudios podrían incluir el uso de LLMs con **otras modalidades** además del texto, como el audio.
- No se dividió la tarea de escritura de ensayos en **subtareas** (generación de ideas, redacción, etc.), lo que podría ofrecer un análisis más profundo.
- El análisis de EEG actual se centró en los patrones de conectividad, pero no examinó los **cambios en la potencia espectral**, lo que podría proporcionar información adicional. La **fMRI** podría ser el siguiente paso para una localización más precisa de las contribuciones corticales profundas.
- Los hallazgos son **dependientes del contexto** (escritura de ensayos en un entorno educativo) y **no pueden generalizarse** a otras tareas.
- Futuros estudios deberían explorar los **impactos longitudinales** del uso de herramientas en la retención de la memoria, la creatividad y la fluidez en la escritura.

Es crucial evitar términos sensacionalistas como «estúpido», «tonto», «cerebro dañado», «pasividad», «recorte», «colapso», «escáneres cerebrales», «los LLM te hacen dejar de pensar», «impacto negativo» o «hallazgos aterradores» al referirse a este estudio, ya que los autores no utilizan dicha terminología y esto desvirtuaría el rigor de la investigación.

El estudio no recibió financiación externa. Los investigadores planean futuros estudios, incluyendo uno sobre la «codificación de la vibra».

Con información de VF