

Fecha: 15-01-2026
Medio: El Mercurio
Supl.: El Mercurio - Innovacion
Tipo: Noticia general
Título: Ramón López de Mántaras: "Hay gente que usa la IA como si fuera su médico de cabecera. Esta imprudencia me sorprende"

Pág.: 4
Cm2: 385,5
VPE: \$ 5.064.267

Tiraje: 126.654
Lectoría: 320.543
Favorabilidad: ☐ No Definida

EXPERTO CON MÁS DE 50 AÑOS DE TRAYECTORIA EN EL CAMPO DE LA INTELIGENCIA ARTIFICIAL

Ramón López de Mántaras: "Hay gente que usa la IA como si fuera su médico de cabecera. Esta imprudencia me sorprende"

Crítico de las exageraciones en torno a esta tecnología, lamenta la forma en que ha entrado en la vida cotidiana. Espera que pronto la humanidad reconozca "que fuimos imprudentes al desplegar estos sistemas sin estudiar bien sus repercusiones". **FERNANDA GUAJARDO S.**

Ramón López de Mántaras lleva más de medio siglo investigando inteligencia artificial, desde mucho antes de que el concepto se volviera parte de la conversación cotidiana. Profesor emérito del Consejo Superior de Investigaciones Científicas (CSIC) español y fundador del Instituto de Investigación en Inteligencia Artificial (IIIA) de Barcelona, ha sido testigo directo de los ciclos de entusiasmo, estancamiento y resurgimiento de esta tecnología. Esa perspectiva de largo plazo es la que marca su mirada crítica sobre el momento actual, justo cuando la IA generativa se ha instalado en la vida diaria de millones de personas.

"Más que el avance tecnológico en sí mismo, lo que puede ser más sorprendente es el enorme *hype*, esta moda exagerada de la inteligencia artificial", afirma. Hasta hace pocos años, recuerda, la IA operaba de forma casi invisible para los usuarios. Estaba en buscadores o sistemas de recomendación, pero no se percibía como algo cercano. "Con la irrupción de los grandes modelos de lenguaje, como ChatGPT, esto ha llegado ahora muy directamente a las casas de cada uno", explica.

Desde el punto de vista técnico, López de Mántaras reconoce que los resultados logrados por estos sistemas son llamativos, pero insiste en que la fascinación pública ha ido mucho más rápido que la comprensión real de cómo funcionan. "Se basan en algo relativamente simple: predecir cuál es la siguiente palabra o el siguiente píxel. No comprenden de verdad el lenguaje en el sentido en que

lo comprendemos nosotros", señala.

Esa distancia entre capacidades reales y expectativas ha generado, a su juicio, un fenómeno preocupante: la antropomorfización de la tecnología. "Cómo tantísima gente se ha quedado fascinada, cómo se proyecta en estos sistemas la idea de que son realmente inteligentes, que tienen emociones, sentimientos y que saben de todo", dice. Una tendencia que, advierte, puede terminar "a costa de deshumanizarnos a nosotros mismos".

El investigador confiesa que le sorprende el uso de estas herramientas en ámbitos sensibles. "Hay gente que los utiliza como psicoterapeutas, como médicos, que les entrega resultados de análisis esperando que el sistema les diga qué tienen, como si fuera su médico de cabecera. Esta imprudencia me sorprende", afirma.

En su presentación de hoy en Congreso Futuro 2026, López de Mántaras abordará uno de los aspectos que considera más invisibilizados del debate: el trabajo humano que sostiene a la inteligencia artificial. Citando estimaciones del Banco Mundial, señala que entre 250 y 400 millones de personas,

principalmente en países del hemisferio sur, trabajan corrigiendo errores, refinando respuestas y entrenando modelos en condiciones extremadamente precarias. "Son personas que trabajan de manera invisible y en condiciones terribles. Son las que posibilitan que creamos que estos sistemas son tan inteligentes", sostiene.

Esa realidad, agrega, también desmonta la idea de una inteligencia completamente autónoma. "No son tan artificiales, porque hay centenares de millones de humanos afinando y mejorando manualmente estos sistemas", explica.

Para López de Mántaras, la responsabilidad frente a los errores de la IA es clara. "La responsabilidad siempre es de humanos, nunca de la inteligencia artificial. No es un agente moral, es *software*", afirma. Y apunta directamente a las grandes tecnológicas: "El responsable último es el CEO. Se cometió un error enorme al desplegar ChatGPT sin antes haber hecho pruebas serias de validación, corrección de sesgos y evaluación de riesgos".

Más allá de los fallos técnicos o regulatorios, el científico advierte por

un impacto más profundo: la pérdida gradual de habilidades humanas básicas. "Es absolutamente ridículo e innecesario pedirle a ChatGPT que te escriba un correo de tres líneas", dice. "Escribir es pensar. Cuando dejamos de hacerlo, perdemos una gimnasia mental fundamental".

Lo mismo ocurre, plantea, con la lectura reemplazada por resúmenes automáticos. "¿Alguien puede creer seriamente que leer un resumen de diez páginas es lo mismo que leer un libro de 300? Todo eso se pierde", afirma.

A este escenario se suma otro riesgo estructural: entrenar nuevos modelos con contenidos generados por la propia IA. "Esto puede llevar a un colapso del sistema. *Garbage in, garbage out*. Si metes basura, obtienes basura", advierte. Según explica, el aumento exponencial de recursos ya no se traduce en mejoras equivalentes, lo que sugiere que estos modelos podrían estar alcanzando un límite.

Cuando se le pregunta cómo le gustaría que se evaluara esta etapa dentro de 20 años, su respuesta es directa: "Me gustaría que dijéramos que nos equivocamos. Que fuimos imprudentes al desplegar estos sistemas sin estudiar bien sus repercusiones, incluso en la salud mental". Y concluye: "Espero que mucho antes de 20 años nos demos cuenta de que lo estamos haciendo mal".



El experto español expondrá hoy en Congreso Futuro 2026.

RICHARD SALGADO