

¿La inteligencia artificial va a reemplazar a los médicos?

Will artificial intelligence replace doctors?

Sebastián Sguiglia Schütz^a 

Resumen

En este artículo, se plantea una reflexión crítica sobre dos frases que se repiten con cierta frecuencia en el ámbito de la salud, ya sea con miedo o con entusiasmo: “la inteligencia artificial (IA) va a reemplazar a los médicos” y “los médicos que utilizan la IA van a reemplazar a los que no la utilizan”. El autor señala las limitaciones actuales de la IA y el bombo publicitario que la rodea, y proponiendo una visión alternativa, en la cual la industria tecnológica y sus productos son tratados con el mismo nivel de escepticismo, definido como la disposición a examinar críticamente las afirmaciones, creencias o conocimientos, evitando aceptarlos como verdaderos sin pruebas o argumentos adecuados. A su vez, retoma la postura de algunos autores que definen a los grandes modelos de lenguaje como ChatGPT como máquinas capaces de generar desinformación y comenta riesgos derivados del uso indiscriminado de estas tecnologías.

Abstract

This article offers a critical reflection on two statements that are frequently repeated in the healthcare field, whether out of fear or enthusiasm: “Artificial intelligence (AI) is going to replace doctors” and “Doctors who use AI are going to replace those who do not.” The author highlights the current limitations of AI and the hype surrounding it, proposing an alternative perspective in which the tech industry and its products are treated with the same level of skepticism, defined as the willingness to critically examine claims, beliefs, or knowledge, avoiding accepting them as true without adequate evidence or arguments. In turn, he echoes the stance of some authors who define large language models like ChatGPT as machines capable of generating misinformation and discusses risks arising from the indiscriminate use of these technologies.

Palabras clave: Inteligencia Artificial Generativa, Inteligencia Artificial, Alfabetización Digital, Medicina Basada en la Evidencia.
Key words: Generative Artificial Intelligence, Artificial Intelligence, Computer Literacy, Evidence-Based Medicine.

Sguiglia Schütz S. ¿La inteligencia artificial va a reemplazar a los profesionales médicos?. Evid Actual Pract Ambul. 2026;29(2):e007204. Available from: <https://doi.org/10.51987/evidencia.v28i3.7204>

^a Servicio de Medicina Familiar y Comunitaria, Hospital Italiano de Buenos Aires; Secretaría de Investigación, Universidad Hospital Italiano de Buenos Aires, Argentina. sebastian.sguiglia@hospitalitaliano.org.ar <https://orcid.org/0000-0002-5218-5241>



Introducción

Un fantasma recorre los centros de salud, hospitales, jornadas y congresos: la idea de que la inteligencia artificial (IA) reemplazará a los profesionales de la salud. Un concepto relacionado, pero menos radical, puede expresarse de la siguiente manera: “la IA no va a reemplazar a los médicos, pero los médicos que usan la IA van a reemplazar a los que no la utilizan”^{1,2}. En esta editorial, vamos a defender la postura de que estas frases son principalmente eslóganes propuestos por la industria tecnológica, que busca persuadir a los profesionales médicos para que adoptemos ciertos productos. Como profesionales, deberíamos conocer estas herramientas, sus ventajas y limitaciones, evitando tanto el miedo como el exceso de confianza que surgen del bombo publicitario.

La inteligencia artificial es un término amplio

En primer lugar, cabe destacar que ambas frases mencionadas en la introducción son un poco vagas, ya que no aclaran a qué tipo de IA se están refiriendo. La IA puede definirse de manera general como un dispositivo capaz de realizar tareas que normalmente requieren (o requerían hasta un momento histórico determinado) la acción de una persona dotada de inteligencia^{3,4}. Bajo esta definición, objetos como las calculadoras digitales o los correctores automáticos de ortografía incorporados a los procesadores de texto son formas de IA⁴.

Las herramientas computarizadas de apoyo a la toma de decisiones (CDSS, por sus iniciales en inglés) también son variantes de IA de “la vieja escuela”, que aplican un conjunto de reglas fijas para producir determinados resultados. Un ejemplo son los

sistemas que detectan automáticamente posibles interacciones entre fármacos prescritos a un mismo paciente en la historia clínica electrónica y emiten una alerta a los usuarios⁵.

Otras herramientas de IA sirven para realizar diagnósticos o predicciones acerca del mundo; por ejemplo, los sistemas utilizados para analizar de forma automática radiografías de tórax y detectar la presencia de hallazgos anormales, como los neumotórax⁶. En general, los creadores de estos programas no indican explícitamente a las herramientas cómo reconocer estos hallazgos, sino que les proveen un gran número de radiografías, de las cuales algunas están clasificadas (etiquetadas) como normales, otras como neumotórax y otras pueden presentar otros hallazgos, como derrames pleurales o consolidaciones. Los sistemas identifican patrones en estas imágenes que luego pueden utilizar para clasificar nuevas imágenes. Este proceso se conoce como aprendizaje automático o *machine learning*, específicamente en su variante supervisada, en la que los sistemas aprenden a partir de ejemplos etiquetados. También existen técnicas de aprendizaje automático no supervisado mediante las cuales los programas pueden identificar patrones en conjuntos de datos no etiquetados, por ejemplo, subtipos de pacientes que se comportan de manera similar.

Cuando hablamos de la IA en la actualidad, solemos referirnos a la IA generativa, un conjunto de herramientas diseñadas para producir contenidos nuevos como textos, audio o imágenes. Los grandes modelos de lenguaje (del inglés *large languages models*: LLM) como ChatGPT®, Gemini® o Claude®, entrenados mediante aprendizaje automático con un enorme volumen de textos tomados de la Internet y luego refinados para que sus respuestas se ajusten a determinados formatos, pertenecen a esta categoría⁷. La función principal de estos modelos es predecir, dada una serie de palabras de un texto, cuál o cuáles son las siguientes palabras más probables. Esta capacidad puede utilizarse para generar textos nuevos con gran rapidez y verosimilitud en respuesta a instrucciones (denominadas *prompts*) de los usuarios.

Los grandes modelos de lenguaje, máquinas de generar “chamuyo”

Los LLM tienen muchas aplicaciones útiles, como generar borradores, traducir, corregir textos o facilitar el aprendizaje simulando interacciones con un tutor o asistente⁷⁻⁹. Más allá de sus evidentes ventajas, una limitación fundamental de estas herramientas es que la generación de textos se basa en la probabilidad, es decir en la predicción de la siguiente palabra más probable, de acuerdo a los patrones que el modelo “aprendió” en su base de datos de entrenamiento. Por lo tanto, no hay garantías de que el texto generado sea preciso o verdadero. Sin embargo, en líneas generales, si asumimos que los textos de la base de datos contienen más verdades que falsedades, la probabilidad de que el texto generado

sea verdadero es más alta que tirar una moneda¹⁰. La generación de respuestas que no son fácticamente correctas se conoce como alucinaciones, una metáfora criticada por antropomorfizar a los LLM y por plantear este proceso como un evento excepcional, en lugar de un fenómeno frecuente y una limitación propia de este tipo de herramientas¹⁰. Existen estrategias para mitigar las alucinaciones, entre ellas mejorar el diseño del *prompt* y conectar al LLM con una base de datos específica, pero hasta la actualidad este riesgo no puede eliminarse por completo^{7,10-12}.

Por este motivo, los LLM fueron caracterizados por algunos autores como máquinas de producir textos fluidos y convincentes, pero que en última instancia no tienen una relación directa con la verdad, una forma de discurso típica de (malos) políticos en campaña, estudiantes en problemas o estafadores, que en inglés se conoce como *bullshit* y podríamos traducir al español rioplatense como “chamuyo”^{10,13}. A esta limitación se suman otros riesgos, algunos comunes a otras formas de IA, como las posibilidades de amplificar sesgos racistas, clasistas o machistas presentes en los textos con los que fueron entrenados, facilitar la producción de desinformación y los discursos de odio, comprometer la privacidad y seguridad de los datos, vulnerar los derechos de propiedad intelectual y generar un impacto ambiental y económico negativo^{11,14,15}.

Los médicos, ¿vamos a ser reemplazados?

Teniendo en cuenta lo anterior, es evidente que la IA en general y la generativa, en particular, son herramientas fallibles que no pueden reemplazar por completo a los profesionales médicos. Además, la conexión emocional, el lenguaje no verbal, los fines humanísticos de la profesión que contrastan con los intereses comerciales de muchas empresas tecnológicas y la capacidad de rendir cuentas de manera individual por los actos y decisiones tomadas son aspectos altamente valorados por los usuarios del sistema de salud que no pueden ser transmitidos con facilidad a una computadora⁴.

¿Pero es posible que los médicos que usan las herramientas de IA reemplacen a los que no las usan? Creo que, nuevamente, es importante ser específicos para dar una respuesta más acertada. Por ejemplo, una pregunta de intervención debería formularse en formato PICO, con una población, una intervención, un comparador y uno o más desenlaces definidos¹⁶. Cuando los profesionales de salud que se desempeñan en un contexto determinado (P) utilizan una herramienta determinada de IA (I), en comparación con cuando no la utilizan, ¿logran en forma consistente una mejora en un desenlace determinado, por ejemplo, menor riesgo de errores? Una vez planteada la pregunta, deberíamos intentar responderla aplicando el método científico, idealmente mediante investigaciones como ensayos clínicos o, en su defecto, estudios cuasiexperimentales bien diseñados, prestando

atención a posibles amenazas a la validez como la selección de controles inadecuados y la medición de desenlaces subrogados que no tengan significancia clínica real.

Si encontramos que una determinada herramienta de IA es efectiva para resolver un problema determinado y además es segura, costo-efectiva, equitativa, etc., entonces deberíamos ir más allá de la utilización individual y promover en forma activa desde las sociedades científicas la adopción a gran escala de esta tecnología. En lugar de optar por posturas extremas, según las cuales una tecnología determinada o las personas que la usan van a reemplazar a quienes no lo hacen, podríamos decir que tal o cual tecnología parece funcionar o no en determinado contexto y que sería bueno probarla. En definitiva, deberíamos tomar las afirmaciones de la industria tecnológica de manera tan crítica (y no sesgada) como lo hacemos con la industria farmacéutica.

Reflexión final

En el contexto actual, vemos que existe una fuerte presión social y un lobby de las empresas

tecnológicas para utilizar la IA cada vez con mayor frecuencia y en un número diverso de situaciones. Además, en el mercado proliferan una gran variedad de marcas similares entre sí cuyas diferencias probablemente sean imperceptibles para la mayoría de los usuarios y, por otro lado, hay evidencia de que el uso irrestricto de la IA puede afectar en forma negativa el aprendizaje, causar daños en la salud mental en personas susceptibles y conducir a errores por exceso de confianza, motivo por el cual no deberíamos tomar decisiones importantes bajo su influencia^{17,18}. Debemos reflexionar sobre la necesidad de regular estas tecnologías para que se adapten a las necesidades de la sociedad en lugar de adoptarlas acríticamente según las necesidades de la industria u otros intereses y, sobre todo, abordarlas con sobriedad y una sana dosis de escepticismo¹⁴.

Recibido el 05/02/2026, aceptado el 30/03/2026 y
publicado el 20/04/2026



Licencia Creative Commons Atribución-NoComercial-CompartirIgual 4.0 Internacional
(CC BY-NC-SA 4.0)

Declaración de uso de inteligencia artificial

No fueron utilizadas herramientas de inteligencia artificial en la preparación de este manuscrito.

Fuentes de financiamiento / Conflicto de interés de los autores

Este artículo no contó con fuentes de financiamiento específicas. El autor declaró no poseer conflictos de interés.

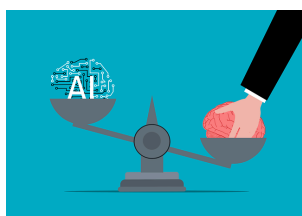
Declaración de revisión editorial

Artículo no evaluado por pares externos. Por sus características, este artículo fue evaluado por el Comité Editorial de Evidencia

Referencias

1. Sezgin E. Artificial intelligence in healthcare: Complementing, not replacing, doctors and healthcare providers. Digit Health [Internet]. 2023 [cited 24 ene 2026];9. DOI: <https://doi.org/10.1177/20552076231186520> PMID: [37426593](https://pubmed.ncbi.nlm.nih.gov/37426593/)
2. Akhtar R. AI Won't Replace You. A Human Using AI Will. Forbes [Internet]. 2025 [cited 24 ene 2026]; Available from: <https://www.forbes.com/sites/reeceakhtar/2025/03/27/ai-wont-replace-you-a-human-using-ai-will/>
3. National Library of Medicine. Artificial Intelligence [Internet]. 1986 [cited 27 ene 2026]. Available from: <https://www.ncbi.nlm.nih.gov/mesh/68001185>
4. Mundzic A, Bogdanffy R, Sundemo D, et al. Exploring Primary Care Patients' Perspectives on Artificial Intelligence: Systematic Literature Review and Qualitative Meta-Synthesis. JMIR AI [Internet]. 2025 [cited 24 ene 2026];4:e72211. DOI: <https://doi.org/10.2196/72211> PMID: [41259713](https://pubmed.ncbi.nlm.nih.gov/41259713/)
5. Luna DR, Rizzato Ledo DA, Otero CM, et al. User-centered design improves the usability of drug-drug interaction alerts: Experimental comparison of interfaces. J Biomed Inform [Internet]. 2017 [cited 24 ene 2026];66:204-13. DOI: <https://doi.org/10.1016/j.jbi.2017.01.009> PMID: [28108211](https://pubmed.ncbi.nlm.nih.gov/28108211/)
6. Mosquera C, Diaz FN, Binder F, et al. Chest x-ray automated triage: A semiologic approach designed for clinical implementation, exploiting different types of labels through a combination of four Deep Learning architectures. Comput Methods Programs Biomed [Internet]. 2021 [cited 24 ene 2026];206:106130. DOI: <https://doi.org/10.1016/j.cmpb.2021.106130> PMID: [34023576](https://pubmed.ncbi.nlm.nih.gov/34023576/)

7. Shanmugam D, Agrawal M, Movva R, et al. Generative Artificial Intelligence in Medicine. Annu Rev Biomed Data Sci [Internet]. 2025 [cited 24 ene 2026];8:199-226. DOI: <https://doi.org/10.1146/annurev-biodatasci-103123-095332> PMID: [40101213](https://pubmed.ncbi.nlm.nih.gov/40101213/)
8. Toyos IS, Ortiz J, Ratti MFG, et al. Inteligencia artificial en la escritura científica: oportunidades y desafíos en ciencias de la salud. Evid Actual Práct Ambul [Internet]. 2024 [cited 24 ene 2026];27:e007138. DOI: <https://doi.org/10.51987/evidencia.v28i1.7138>
9. Abdunour REE, Gin B, Boscardin CK. Educational Strategies for Clinical Supervision of Artificial Intelligence Use. N Engl J Med [Internet]. 2025;393:786-97. DOI: <https://doi.org/10.1056/nejmra2503232> PMID: [40834302](https://pubmed.ncbi.nlm.nih.gov/40834302/)
10. Hicks MT, Humphries J, Slater J. ChatGPT is bullshit. Ethics Inf Technol [Internet]. 2024 [cited 24 ene 2026];26:38. Available from: <https://link.springer.com/article/10.1007/s10676-024-09775-5>
11. Weidinger L, Uesato J, Rauh M, et al. Taxonomy of Risks posed by Language Models. En: Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency [Internet]. New York, NY, USA: Association for Computing Machinery; 2022 [cited 24 ene 2026]. p. 214-29. DOI: <https://doi.org/10.1145/3531146.3533088>
12. Dilmegami C, Daldal A. AI Hallucination: Compare top LLMs like GPT-5.2 in 2026 [Internet]. AIMultiple. 2026 [cited 27 ene 2026]. Available from: <https://research.aimultiple.com/ai-hallucination/>
13. Narayanan A. ChatGPT is a bullshit generator. But it can still be amazingly useful [Internet]. 2025 [cited 27 ene 2026]. Available from: <https://www.normaltech.ai/p/chatgpt-is-a-bullshit-generator-but>
14. AI as Normal Technology [Internet]. Knight First Amendment Institute. [cited 22 ene 2026]. Available from: <http://knightcolumbia.org/content/ai-as-normal-technology>
15. Hartzog W, Silbey JM. How AI Destroys Institutions [Internet]. Rochester, NY: Social Science Research Network; 2025 [cited 27 ene 2026]. Available from: <https://papers.ssrn.com/abstract=5870623>
16. Volij C, Carrara C, Kopitowski K. Preguntas para resolver dudas clínicas. En: Medicina Familiar y Práctica ambulatoria. 4.a ed. Editorial Médica Panamericana; 2025. p. 36-9.
17. Kosmyrna N, Hauptmann E, Yuan YT, et al. Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task [Internet]. arXiv; 2025 [cited 27 ene 2026]. DOI: <https://doi.org/10.48550/arXiv.2506.08872>
18. Looi JC, Allison S, Bastiampillai T, et al. Illusions of intelligence, connection and reality: Perils of large-language AI models for people with severe mental illness. Australas Psychiatry [Internet]. 2025;34(1):5-7. DOI: <https://doi.org/10.1177/10398562251380544> PMID: [40952393](https://pubmed.ncbi.nlm.nih.gov/40952393/)



[Imagen de portada](#) de [Mohamed Hassan](#) bajo licencia de contenidos de [Pixabay](#)