



DOI: <https://doi.org/10.5554/22562087.e1178>

De modelos y monocultivos: riesgos epistémicos de la inteligencia artificial en la investigación médica

Of models and monocultures: epistemic risks of artificial intelligence in medical research

Jorge Iván Alvarado-Sánchez^{a-c} 

^a Departamento de Cuidado Intensivo, Fundación Santa Fe de Bogotá. Bogotá, Colombia.

^b Departamento de Ciencias Fisiológicas, Facultad de Medicina, Universidad Nacional de Colombia. Bogotá, Colombia.

^c Editor en Jefe, Colombian Journal of Anesthesiology. Bogotá, Colombia.

Correspondencia: Sociedad Colombiana de Anestesiología y Reanimación (S.C.A.R.E.), Cra 15a # 120 – 74. Bogotá, Colombia.

E-mail: jjalvarados@unal.edu.co

Cómo citar este artículo: Alvarado-Sánchez JI. Of Models and Monocultures: Epistemic Risks of Artificial Intelligence in Medical Research. Colombian Journal of Anesthesiology. 2026;54:e1178.

“El verdadero peligro no es que los ordenadores empiecen a pensar como los humanos, sino que los humanos empiecen a pensar como los ordenadores.”

—Sydney J. Harris

La inteligencia artificial (IA) está transformando la investigación médica. Desde la automatización de revisiones bibliográficas hasta la detección de patrones en grandes conjuntos de datos, la IA promete eficiencia y acceso a una escala sin precedentes (1). Los modelos de lenguaje y los asistentes de redacción también apoyan a investigadores en países de ingresos bajos y medios (LMICs) al mejorar la claridad y la traducción, lo que potencialmente amplía la participación en la publicación global (2,3). Sin embargo, bajo esta promesa yace un peligro más silencioso: la IA puede reforzar desigualdades históricas en quién produce, valida y difunde el conocimiento médico.

Inequidad epistémica en medicina

La inequidad epistémica se refiere al reconocimiento desigual de qué conocimientos se consideran autoritativos. En medicina, esto es relevante porque las perspectivas excluidas de los sistemas dominantes—a menudo provenientes de regiones, lenguas o tradiciones subrepresentadas—aportan ideas fundamentales sobre la

atención al paciente, las prácticas culturales y las realidades de los sistemas de salud que los estándares globales por sí solos no pueden captar. Ignorar tal diversidad reduce la evidencia y limita la innovación (4).

El riesgo de monocultivos del conocimiento

En el centro se encuentra el riesgo de los monocultivos del conocimiento: la dominancia de un único marco epistémico que eleva la precisión computacional como el árbitro supremo de la verdad (5). Este cambio privilegia lo que es medible sobre lo que es significativo, marginando la sabiduría clínica tácita y la comprensión contextual de la enfermedad (6). Los grandes modelos de lenguaje (LLMs) ilustran este riesgo. Bender et al. han demostrado que los LLMs entrenados con corpus como Wikipedia o Common Crawl amplifican las hegemonías culturales y lingüísticas presentes en sus datos, consolidando visiones del mundo dominantes (7). Alvero y colaboradores, al analizar más de 25.000 ensayos generados por IA, encontraron que los resultados se asemejaban más a la escritura de grupos socialmente privilegiados (8). Si la IA homogeneiza la expresión hacia normas dominantes, también puede homogeneizar la manera en que se plantean las preguntas de investigación y se abordan los problemas. Para la investigación médica, esto significa que las perspectivas alternativas pueden ser silenciadas antes incluso de ser consideradas (9).

Manifestaciones de los monocultivos

Los monocultivos se manifiestan al menos en tres ámbitos:

Datos clínicos. Los modelos entrenados predominantemente con conjuntos de datos de países de altos ingresos se exportan como estándares globales, a pesar de las diferencias drásticas en las cargas de enfermedad y las rutas de atención en otros lugares (10–12).

Literatura científica. Las herramientas de descubrimiento impulsadas por IA, como Elicit, Connected Papers y Litmaps, dependen de corpus como Semantic Scholar y CrossRef. Los estudios muestran que más del 90% de los documentos indexados en Web of Science y Scopus están en inglés, mientras que el español y el portugués representan apenas un 1% en conjunto (13). El propio Elicit reconoce su dependencia exclusiva de Semantic Scholar, omitiendo estructuralmente los estudios no escritos en inglés (14). Cuando estos sesgos se integran en los flujos de trabajo de IA, se automatizan y amplifican.

Entornos informativos cotidianos. Los algoritmos de las redes sociales entregan contenido alineado con las preferencias del usuario, creando “burbujas de filtro”. De manera similar, las herramientas médicas basadas en IA pueden priorizar estudios muy citados o en inglés, reforzando perspectivas dominantes y filtrando alternativas relevantes a nivel regional.

Lo qué hace diferente a la IA

Los escépticos podrían argumentar que estas inequidades ya existen. Lo que distingue a la IA es cómo transforma su escala y funcionamiento. Primero, la IA acelera la exclusión al automatizarla: los sesgos que antes aparecían esporádicamente en decisiones editoriales o de financiación ahora operan de manera continua e invisible. Segundo, la IA incrementa la opacidad: mien-

tras que las decisiones humanas pueden ser escrutadas, los filtros algorítmicos actúan como cajas negras. Tercero, la IA universaliza patrones provenientes de conjuntos de datos dominantes, transformando normas locales en estándares globales. En conjunto, estas características magnifican las inequidades, institucionalizándolas de formas más rápidas, difíciles de detectar y más complejas de cuestionar.

Hacia la pluralidad epistémica

Si no se controla, la IA corre el riesgo de estrechar, en lugar de diversificar, el conocimiento médico. Para contrarrestar esto, son esenciales tres principios:

Pluralidad. Las herramientas deben incorporar bases de datos no anglófonas (por ejemplo, SciELO, LILACS, AJOL), literatura gris y revistas con enfoque regional. Incluir múltiples tradiciones de evidencia garantiza que las preguntas de investigación no se formulen exclusivamente desde los sistemas dominantes.

Transparencia. Los desarrolladores deberían revelar las fuentes, los criterios de clasificación y las exclusiones, de manera similar a las etiquetas nutricionales en los alimentos. Sin esto, los usuarios no pueden saber qué evidencia ha sido eliminada silenciosamente.

Humildad en la interpretación. Más allá de la transparencia y la diversidad, la humildad es crucial. Los modelos deberían mostrar no solo los estudios incluidos, sino también los excluidos; los sistemas clínicos deberían aclarar el alcance de sus datos de entrenamiento; los editores deberían tratar los resultados derivados de IA como provisionales, no definitivos. La humildad recuerda a los investigadores que la IA es un complemento, no un sustituto, del juicio humano. Tal humildad podría servir como contrapeso a los monocultivos del conocimiento descritos anteriormente, asegurando que la automatización no se convierta en una autoridad incuestionable.

CONCLUSIÓN

La inteligencia artificial no solo refleja las inequidades existentes; las transforma. Al incorporar sesgos en infraestructuras algorítmicas que operan a gran escala, la IA corre el riesgo de crear monocultivos del conocimiento que privilegian lo universal sobre lo particular, lo medible sobre lo significativo. La novedad no reside en la existencia de la inequidad, sino en cómo la IA cambia su funcionamiento, alcance y detectabilidad.

Si se utiliza de manera responsable, la IA podría ampliar la diversidad epistémica en la medicina. Pero esto requiere decisiones de diseño conscientes guiadas por la pluralidad, la transparencia y la humildad. De lo contrario, las herramientas destinadas a iluminar la complejidad pueden terminar oscureciendo el conocimiento que más necesitamos.

Uso de la inteligencia artificial generativa

El autor utilizó ChatGPT (OpenAI, GPT-5, versión de septiembre de 2025) para mejorar la fluidez y legibilidad del texto en inglés.

Conflictos de interés

JIAS es el Editor en Jefe de la Colombian Journal of Anesthesiology.

REFERENCIAS

1. Wang H, Fu T, Du Y, Gao W, Huang K, Liu Z, et al. Scientific discovery in the age of artificial intelligence. *Nature*. 2023;620(7972):47–60. <https://doi.org/10.1038/s41586-023-06221-2>
2. Kapoor MC. Navigating the Impact of Artificial Intelligence on Medical Writing. *Ann Card Anaesth*. 2025;28(2):105–6. https://doi.org/10.4103/aca.aca.14_25
3. Gallifant J, Afshar M, Ameen S, Aphinyanaphongs Y, Chen S, Cacciamani G, et al. The TRIPOD-LLM reporting guideline for studies using large language models. *Nat Med*.

- 2025;31(1):60-9. <https://doi.org/10.1038/s41591-024-03425-5>
4. Alvarado-Sánchez JI. The invisible weight of knowledge: epistemic inequity in global critical care research. *Intensive Care Med.* 2025;51:1724-6. <https://doi.org/10.1007/s00134-025-08039-0>
 5. Messeri L, Crockett MJ. Artificial intelligence and illusions of understanding in scientific research. *Nature.* 2024;627(8002):49-58. <https://doi.org/10.1038/s41586-024-07146-0>
 6. Pratt B, de Vries J. Where is knowledge from the global South? An account of epistemic justice for a global bioethics. *J Med Ethics.* 2023;49(5):325-34. <https://doi.org/10.1136/jme-2022-108291>
 7. Bender EM, Gebru T, McMillan-Major A, Shmitchell S. On the Dangers of Stochastic Parrots. In: *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency.* New York, NY, USA: ACM; 2021. p. 610-23. <https://doi.org/10.1145/3442188.3445922>
 8. Alvero AJ, Lee J, Regla-Vargas A, Kizilcec RF, Joachims T, Antonio AL. Large language models, social demography, and hegemony: comparing authorship in human and synthetic text. *J Big Data.* 2024;11(1):138. <https://doi.org/10.1186/s40537-024-00986-7>
 9. Santos B de S. *Epistemologies of the South Justice Against Epistemicide.* 2014.
 10. Obermeyer Z, Powers B, Vogeli C, Mullainathan S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science (1979).* 2019;366(6464):447-53. <https://doi.org/10.1126/science.aax2342>
 11. Durán JM, Jongsma KR. Who is afraid of black box algorithms? On the epistemological and ethical basis of trust in medical AI. *J Med Ethics.* 2021;47:329-35. <https://doi.org/10.1136/medethics-2020-106820>
 12. Obermeyer Z, Powers B, Vogeli C, Mullainathan S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science (1979).* 2019;366(6464):447-53. <https://doi.org/10.1126/science.aax2342>
 13. Vera-Baceta MA, Thelwall M, Kousha K. Web of Science and Scopus language coverage. *Scientometrics.* 2019;121(3):1803-13. <https://doi.org/10.1007/s11192-019-03264-z>
 14. ELICIT. Information and advice from the Elicit team (Internet). (citado 2025 Sep 22). Disponible en: <https://support.elicit.com/en/articles/553025>