

From Dr. Google to Dr. AI: medical misinformation, clinical trust, and ethical responsibility in contemporary medicine

Del Dr. Google al Dr. IA: desinformación médica, confianza clínica y responsabilidad ética en la medicina contemporánea

Álvaro Barrios-Núñez^{1a}, Lorena Cudris-Torres^{2b*}, Lilian Danielle Bolaño Acosta^{3c},
Paola García-Roncallo^{4d}, Andrés Fernando Ramírez-Giraldo^{5d},
Deisi Johanna Duque-Torres^{6e}, Margaret Carolina Arzuaga Mendoza^{7f}

SUMMARY

Generative artificial intelligence has rapidly entered health information seeking among patients, caregivers, students, and professionals. Its ability to produce immediate, personalized, and persuasive answers offers opportunities for health literacy, medical education, and clinical communication. However, it also creates an emerging form of medical misinformation: content that is technically plausible and linguistically convincing yet potentially erroneous, incomplete, biased, or fabricated. This critical reflection article analyzes the transition from “Dr. Google” to “Dr. AI” and discusses its implications for patient safety, the physician-patient relationship, equity, medical education, and ethical governance. It argues that primary care and professional training should incorporate digital curation, critical verification,

and clinical communication competencies regarding artificial intelligence, avoiding both technophobic rejection and uncritical adoption of systems that remain insufficiently validated.

Keywords: Artificial intelligence, health misinformation, medical ethics, health literacy, primary care, medical education

RESUMEN

La inteligencia artificial generativa se ha incorporado con rapidez a la búsqueda de información sanitaria por parte de pacientes, cuidadores, estudiantes y profesionales. Su capacidad para producir respuestas inmediatas, personalizadas y persuasivas representa una oportunidad para la alfabetización en salud, la educación médica y la comunicación clínica. Sin embargo, también inaugura una forma emergente de desinformación médica: contenidos técnicamente verosímiles, lingüísticamente convincentes y potencialmente erróneos, incompletos, sesgados o fabricados. Este artículo de reflexión crítica analiza la transición del “Dr. Google” al “Dr. IA” y discute sus implicaciones para la seguridad del

ORCID:

¹ [0000-0003-4153-8950](https://orcid.org/0000-0003-4153-8950)

² [0000-0002-3120-4757](https://orcid.org/0000-0002-3120-4757)

³ [0000-0003-4286-4301](https://orcid.org/0000-0003-4286-4301)

⁴ [0000-0002-9223-2129](https://orcid.org/0000-0002-9223-2129)

⁵ [0000-0002-4709-9891](https://orcid.org/0000-0002-4709-9891)

⁶ [0000-0002-6076-3267](https://orcid.org/0000-0002-6076-3267)

⁷ [0000-0001-6507-4260](https://orcid.org/0000-0001-6507-4260)

***Corresponding author:** Lorena Cudris-Torres, Senior Lecturer 3, Department of Social Sciences, Universidad de la Costa, Barranquilla, Colombia. E-mail: lcudris3@cuc.edu.co

Recibido: 10 de junio 2026

Aceptado: 29 de julio 2026

- a. Clínica General del Norte, Barranquilla, Colombia.
- b. Centro de Investigación en Salud, Bienestar y Autocuidado “SABIA”, Universidad de la Costa, Barranquilla, Colombia. c
- c. Universidad Nacional Abierta y a Distancia, Valledupar, Colombia
- d. Universidad de la Costa, Barranquilla, Colombia.
- e. Fundación Universitaria del Área Andina, Universidad Tecnológica de Pereira, Pereira, Colombia.
- f. Universidad Popular del Cesar, Valledupar, Colombia

paciente, la relación médico-paciente, la equidad, la educación médica y la gobernanza ética. Se propone que la atención primaria y la formación profesional incorporen competencias de curaduría digital, verificación crítica y comunicación clínica sobre inteligencia artificial, evitando tanto el rechazo tecnofóbico como la adopción acrítica de sistemas aun insuficientemente validados.

Palabras clave: *Inteligencia artificial, desinformación en salud, ética médica, alfabetización en salud, atención primaria, educación médica*

INTRODUCTION

For more than two decades, the metaphor of “Dr. Google” has captured a recurring tension in medical practice: patients arrive at consultations with online information that may be useful, fragmented, alarmist, or false. The rise of generative artificial intelligence (AI) is transforming this situation. Patients no longer merely consult websites; they interact with systems that respond in natural language, organize symptoms, explain diagnoses, simplify scientific articles, suggest questions for medical visits, and offer recommendations in a confident, seemingly expert tone. This shift marks a qualitative change: seeking health information is no longer a dispersed browsing experience, but a personalized conversation mediated by an algorithm.

Recent literature acknowledges that large language models and multimodal systems can support medical education, documentation, patient communication, the interpretation of complex information, and some clinical support tasks (1-5). However, risks are also noted, including hallucinations, biases, privacy, opacity, cognitive dependence, digital inequities, and the mass production of misinformation (2, 6-10). In healthcare, these risks take on a unique ethical dimension: an incorrect response is not only an informational problem but also a potential threat to patient safety.

This article argues that generative AI should be understood as a new informational actor in medicine. It does not replace the professional, but

it does influence how patients and doctors build trust, interpret evidence, and make decisions. The critical question is not only whether AI can provide accurate answers, but under what conditions an AI-generated response might influence health behaviors, delay consultations, reinforce false beliefs, encourage self-medication, weaken treatment adherence, or exacerbate inequalities.

From this perspective, the article critically examines the relationship among generative AI, medical misinformation, and clinical trust, drawing on recent literature in public health, bioethics, medical education, health regulation, and primary care. It is not intended as a systematic review, but rather as a critical narrative reflection aimed at identifying dilemmas, risks, and courses of action for contemporary medicine that is humanistic, technologically competent, and ethically responsible.

From “Google” to “AI”: a mutation of the health information ecosystem

Digital inquiries about symptoms, diagnoses, and treatments predate generative AI. However, traditional search engines present visible, heterogeneous, and comparable results: institutional websites, blogs, videos, forums, advertisements, or scientific documents. Users can identify the source of these results, though not always with sufficient skill. Generative chatbots reconfigure this experience: they integrate information into a single, fluid, and conversational response, often without clearly indicating the sources’ origin, hierarchy, or quality.

This change increases the persuasive power of the message; a response written with a clinical structure, technical vocabulary, and empathetic tone can generate an illusion of authority. Hence, some authors warn that generative AI not only produces information, but also apparent trust (11-13). The difference between a dubious website and a convincing algorithmic response is crucial: the latter may seem more personalized, more neutral, and more “scientific”, although it is not necessarily more correct.

The literature on AI in medicine has shown potential benefits in education, research support, answering patient questions, and processing clinical information (1, 14-18). However, the very attributes that make these systems attractive—speed, constant availability, natural language processing, and simplification capabilities—can exacerbate harm when the information generated is false, incomplete, or decontextualized. Consequently, the medical and ethical question must shift from fascination with the technical capabilities to an evaluation of their conditions of use, limitations, and responsibilities.

Algorithmic medical misinformation: a distinct threat from traditional misinformation

Health misinformation is not new; the COVID-19 pandemic demonstrated the speed with which rumors, conspiracy theories, false cures, and anti-vaccine content can circulate on social media and affect collective behavior. What is new in the era of generative AI is the possibility of producing misinformation at scale, at low cost, with high personalization, and with an appearance of legitimacy. De Angelis et al. termed this phenomenon an “AI-driven infodemic,” noting the risk that large language models could amplify public health threats through automated, persuasive content (6).

Empirical evidence is beginning to confirm this concern. An analysis published in *The British Medical Journal* evaluated safeguards, risk mitigation, and transparency of language models against instructions aimed at producing health misinformation; its findings suggest that barriers are possible but inconsistent and dependent on the model and system design (7). More recently, a study in *The Lancet Digital Health* showed that models can be susceptible to false medical information when it is presented in formats that simulate clinical legitimacy, such as discharge notes or professional language (8).

Algorithmic medical misinformation has its own characteristics. First, it can be generated instantly and at large volumes. Second, it can

adapt to the user’s language level, fear, age, culture, or question. Third, it can combine true data with erroneous inferences, making it more difficult to identify. Fourth, it can omit critical warnings, such as the need for in-person evaluation when warning signs are present. Fifth, it can be deliberately used for commercial, political, or ideological purposes, including promoting unproven products, discrediting vaccines, or distorting scientific evidence.

In this scenario, the harm is not limited to factual error; AI can reinforce confidence in a flawed explanation, validate risky decisions, or provide a false sense of control. In topics such as cancer, mental health, pregnancy, cardiovascular disease, self-medication, vaccination, or acute pain, incorrect answers can delay care, promote dangerous interventions, or increase anxiety. Therefore, AI-generated misinformation should be considered a patient safety issue, not just a communication challenge.

Hallucinations and verisimilitude: when clinical language simulates evidence

One of the most discussed risks of generative AI is hallucination: the production of plausible but false, unverifiable, or nonexistent content. In healthcare, this phenomenon can manifest as fabricated bibliographic citations, incorrect dosages, improbable diagnoses presented as probable, inappropriate interpretation of symptoms, invention of clinical guidelines, exaggeration of therapeutic benefits, or minimization of risks. Recent reviews of large language models in medicine identify accuracy, bias, privacy, transparency, and the presence of convincing but inaccurate content as central concerns (33,34). The problem is not only that the model can be wrong, but that it can be wrong despite excellent writing.

The literature on AI and decision-making warns that manipulated, false, or unsupervised results can lead to misunderstandings of reality and affect decisions based on seemingly reliable information (9). In medicine, textual plausibility

can be especially dangerous because many patients lack the criteria to distinguish between a valid clinical explanation and a convincing but incorrect formulation. Even professionals in training may interpret a fluent response as a sign of competence.

Large-scale language models do not “understand” disease the way a clinician does when facing a patient. They operate using statistical patterns derived from large datasets, and while they may exhibit remarkable performance in specific tests or tasks, that performance does not equate to contextual clinical judgment. Clinical practice requires integrating history, physical examination, trajectory, uncertainty, communication, patient preferences, available resources, and ethical implications. AI can assist, but it cannot single-handedly assume the responsibility for making decisions. Therefore, hallucination should be interpreted as an epistemological and moral risk: epistemological because it blurs the boundary between validated knowledge and plausible text; moral because it shifts the burden of verification onto users who are not always able to evaluate the response. Patient safety demands that medicine not delegate to the patient the responsibility for detecting errors that even some clinical systems fail to make transparent adequately.

Clinical trust in the era of the algorithmic intermediary

The doctor-patient relationship is built on trust, technical expertise, clear communication, and an acknowledgment of vulnerability. Generative AI introduces a third element to this relationship. Patients can consult the chatbot before their appointment, during the diagnostic process, or after receiving medical advice. In some cases, the tool can facilitate understanding and prepare relevant questions. In others, it can erode trust, stimulate unfounded doubts, or lead patients to confront their doctor with an unvalidated, algorithmically generated “second opinion.”

This issue should not be approached by a corporate defense of medical authority; patients have the right to be informed, to ask questions,

and to participate in their decisions actively. Autonomy requires access to understandable information. However, autonomy is weakened when the information that fuels it is opaque, false, or manipulated. Autonomy based on misinformation is not empowerment, but rather exposure to harm.

At this point, the medical response cannot be to prohibit or ridicule the use of AI; such an attitude can deepen the distance between professionals and patients. The alternative is to incorporate the conversation about AI into the consultation: ask if the patient has sought information from chatbots, review the findings together, explain which elements are correct, which are insufficient, and why a particular recommendation requires clinical evaluation. Trust is not protected by denying the technology, but by practicing critical clinical education.

This shift implies an emerging communicative competence: the ability of healthcare professionals to engage in dialogue with patients who arrive informed, overinformed, or misinformed by AI. Contemporary medicine requires clinical digital literacy: not only knowing how to use tools but also understanding their limitations, biases, and risks to translate this knowledge into reliable guidance.

Health literacy, the digital divide, and new inequalities

The risks of generative AI are not evenly distributed. People with lower health literacy, lower digital literacy, limited access to care, or prior experiences of institutional distrust may be more vulnerable to accepting an automated response as valid. Primary care has been identified as a strategic space for addressing this vulnerability due to its proximity, continuity, and community function (1).

Health literacy involves the ability to access, understand, evaluate, and use information to make appropriate decisions. When information becomes conversational, personalized, and seemingly

expert, literacy must be broadened: users need to know that an AI response is not equivalent to a diagnosis, that it may omit warning signs, that it does not always cite verifiable sources, and that it may respond differently depending on how the question is phrased. This new literacy cannot fall solely on the individual; educational systems, health services, scientific societies, and public policies must promote it.

Inequality deepens when groups with greater educational capital use AI as a critical complement, while groups with less access to care use it as a substitute for consultation. In Latin American contexts, where geographic, economic, and administrative barriers to accessing services persist, AI could become an apparent solution to structural gaps. The danger is that the promise of digital accessibility ends up normalizing secondary care for vulnerable populations.

Ethical discussion must avoid two extremes. The first is denying the benefits of tools that can translate technical language, facilitate health education, and support self-care. The second is assuming that technological availability equates to health equity. Equity demands that AI strengthen, not replace, the public responsibility to ensure safe, humane, and evidence-based care.

Medical education: digital curation, critical thinking, and prevention of cognitive dependency

Medical education is one of the fields where generative AI is generating the most enthusiasm; it can support personalized learning, clinical simulation, question generation, explanation of complex concepts, and immediate feedback (3,20-25). However, the literature also warns of risks of dependency, loss of human skills, biases, privacy issues, misinformation, and weakening of critical thinking if used without supervision (3,20,23,26).

The educational response should not be to exclude AI from training, but rather to integrate it critically. Future professionals must learn to formulate appropriate questions, verify answers, detect omissions, compare them with

clinical guidelines, recognize uncertainty, and communicate limitations to patients. In other words, the competency is not “using ChatGPT,” but rather practicing digital curation in environments saturated with automated information.

Digital medical curation involves selecting, comparing, contextualizing, and translating information; it is an epistemological and ethical skill. Students who accept a generated answer without verifying it may reproduce errors; professionals who use it to save time without supervision may compromise patient safety; researchers who incorporate fabricated text or references may damage scientific integrity. On this point, international editorial recommendations are clear: AI tools should not be considered authors, and their use must be transparently declared when they participate in the production of manuscripts (27).

Medical education, therefore, must incorporate modules on responsible AI, algorithmic bias, data protection, clinical validation, communicating uncertainty, and the ethics of automation. It should also promote a culture of double-checking: no AI response that could affect a clinical decision should be accepted without verification against validated sources, professional judgment, and the patient’s specific conditions.

Ethical governance and regulation: from enthusiastic innovation to verifiable accountability

The development of AI in healthcare requires a strong governance framework. The World Health Organization (WHO) has emphasized principles such as protecting autonomy, promoting well-being and safety, transparency, explainability, accountability, inclusion, equity, and sustainability (4,5). These principles are especially relevant for generative systems because their responses can vary, their sources are not always traceable, and their performance can change with updates.

Healthcare regulation faces a challenge: many general-purpose chatbots are not presented as

medical devices but are used by the public for medical purposes. This gray area makes oversight difficult. A system may avoid declaring clinical intent and still answer questions about symptoms, medications, or diagnoses. The boundary between general information and health advice becomes porous when the user asks personal questions and the system responds in an individualized tone.

International regulatory experience is beginning to recognize these challenges. The Food and Drug Administration (FDA) has advanced action plans for AI and machine learning software as medical devices, emphasizing the entire product lifecycle and the need for risk-proportional oversight (28). In Europe, the AI Regulation classifies certain systems used for medical purposes as high-risk, requiring risk management, data quality, clear information, and human oversight (29). Analysis of its implications for digital medical products shows that regulated devices, diagnostics, and clinical support tools will need to articulate requirements for safety, data, transparency, and post-market surveillance (32). While these frameworks do not solve all the problems of open generative AI, they do provide guidance: health innovation must be verifiable, auditable, and accountable.

Bioethics offers additional criteria: beneficence requires that AI provide real clinical value; non-maleficence requires preventing foreseeable errors; autonomy requires comprehensible and transparent information; justice requires avoiding biases and gaps; and accountability requires identifying who is responsible for harm. The UNESCO Recommendation on AI Ethics reinforces this approach by placing human rights, dignity, transparency, equity, and human oversight as guiding principles (30). In turn, WHO regulatory considerations insist that AI systems for health require assessment, risk management, monitoring, and governance throughout their entire lifecycle (31). The contemporary challenge is that AI distributes responsibilities among developers, institutions, practitioners, regulators, and users. Precisely for this reason, governance cannot be reduced to terms and conditions of use.

Primary care as the first line of health re-education against AI

Primary care occupies a privileged place in this discussion; it is the level of the system where common questions are answered, early risks are detected, individuals and families are supported longitudinally, and community trust is built. Therefore, it can serve as a central setting for re-educating on the responsible use of AI in healthcare (1).

This re-education should not be approached as an isolated campaign, but rather as a daily clinical practice. Primary care teams can teach simple criteria: do not use AI when warning signs are present; do not change medications based on a chatbot's recommendation; do not replace consultations with automated responses; verify information with institutional sources; be wary of content that promises quick cures; ask the healthcare professional when an answer contradicts clinical indications; and understand that AI can support the preparation of questions but not confirm personal diagnoses.

Primary care can also guide patients to reliable sources, identify local misinformation patterns, and communicate risks in clear language. In communities with low digital literacy, this educational role is as important as vaccination, prenatal care, or the management of chronic diseases. AI literacy should be an integral part of health promotion. The challenge is that healthcare professionals themselves need training. Doctors, nurses, or psychologists cannot be expected to provide guidance on AI if they lack the basic skills to understand how it works, its risks, and its responsible uses. Community re-education must be accompanied by ongoing professional development.

Latin America: Opportunity, Technological Dependence, and Health Sovereignty

In Latin America, the debate on generative AI in healthcare must be situated within a context of

inequality, fragmented systems, digital divides, and technological dependence. The region can benefit from tools that improve information translation, educational support, administrative management, and health communication. However, it can also be exposed to models trained with data, languages, epidemiological realities, and cultural values that are not sufficiently representative of its populations.

Digital health sovereignty implies that countries are not simply consumers of external technologies, but rather actors capable of evaluating, adapting, regulating, and producing solutions tailored to their needs. This requires investment in research, data infrastructure, ethics committees, regulation, professional training, and social participation. It also demands that universities and scientific societies assume an active role in the critical evaluation of AI and in the production of local knowledge.

Generative AI can be an ally in democratizing health information. Still, it can also reproduce epistemic coloniality if it imposes responses based on decontextualized evidence, renders local knowledge invisible, or generalizes inappropriate recommendations. In public health, cultural relevance is not an embellishment, but a condition for effectiveness.

Critical proposal: minimum principles for prudent integration

Based on the reviewed literature, ten minimum principles are proposed for addressing generative AI in healthcare from a clinical, ethical, and educational perspective:

- Primacy of clinical judgment: no AI-generated response should replace professional assessment when there is a diagnostic or therapeutic risk.
- Transparency: patients, professionals, and institutions must know when they are interacting with AI systems and what their limitations are.

- Verification: all information with potential clinical impact must be verified against guidelines, institutional sources, or reliable scientific literature.
- Non-substitution of care: AI can guide questions, but not replace consultation, physical examination, or follow-up.
- Data protection: sensitive patient data should not be entered into unauthorized systems or systems without privacy guarantees.
- Equity: digital solutions should be designed to reduce, not widen, access and literacy gaps.
- Human oversight: clinical users must have responsible professionals, traceability, and corrective mechanisms.
- Continuing medical education: AI training should include ethics, bias, communication, patient safety, and digital curation.
- Institutional responsibility: Hospitals, universities, and healthcare services must create explicit policies for the use of AI.
- Ongoing evaluation: Systems must be audited in real-world scenarios, considering performance, harm, bias, and social acceptability.

These principles are not intended to stifle innovation, but rather to prevent medicine from confusing novelty with progress. Technology is truly effective in healthcare when it improves outcomes, respects rights, reduces harm, and strengthens public trust.

CONCLUSIONS

Generative AI represents one of the most rapid transformations in the health information ecosystem. Its ability to produce understandable, personalized, and readily available answers can support health literacy, medical education, and

clinical communication. However, this same ability can generate medical misinformation disguised as rigorous information, exacerbate inequities, erode clinical trust, and shift decision-making toward opaque systems.

The shift from “Dr. Google” to “Dr. AI” is not merely a technological evolution; it is an ethical mutation of contemporary medicine. Authority is no longer contested solely between doctor and patient, but also between clinical knowledge, scientific evidence, digital platforms, and generative models capable of simulating certainty. The response should not be technophobia or uncritical adoption, but rather informed prudence.

Medicine needs to take on a renewed educational role: teaching patients, students, and professionals how to engage with AI critically. Primary care, due to its close ties with the community, can lead to this healthcare re-education. Universities must train professionals capable of using AI without relinquishing critical thinking. Scientific journals must demand transparency and rigor. Regulators must close gray areas in which general-purpose systems end up functioning as healthcare advisors.

Ultimately, generative AI will be as beneficial or risky as the ethical, clinical, and social conditions of its integration. The crucial question is not whether medicine will use AI, but whether it can do so without sacrificing the trust, responsibility, and human dignity that underpin the medical act.

REFERENCES

1. Coronado-Vázquez V, Allande-Cussó R, Caparrós-González RA, Gómez-Salgado J. Inteligencia artificial y desinformación en salud: la necesidad de reeducación desde la atención primaria. *Aten Primaria*. 2026;58:103460.
2. Nunes R, Nunes SB. Inteligencia artificial: un puente hacia el futuro de la medicina. *Rev Bioet*. 2025;33:e4115ES.
3. Chávez-Martínez O, Ragacini LA. Educación médica e inteligencia artificial: perspectivas y desafíos éticos. *Rev Med Inst Mex Seguro Soc*. 2025;63(5):e6736.
4. World Health Organization. Ethics and governance of artificial intelligence for health: guidance on large multi-modal models. Geneva: World Health Organization; 2024.
5. World Health Organization. Ethics and governance of artificial intelligence for health: WHO guidance. Geneva: World Health Organization; 2021.
6. De Angelis L, Baglivo F, Arzilli G, Privitera GP, Ferragina P, Tozzi AE, et al. ChatGPT and the rise of large language models: the new AI-driven infodemic threat in public health. *Front Public Health*. 2023;11:1166120.
7. Menz BD, Kuderer NM, Bacchi S, Modi ND, Latif U, Jayanth S, et al. Current safeguards, risk mitigation, and transparency measures for large language models against the generation of health misinformation: repeated cross-sectional analysis. *The British Medical Journal*. 2024;384:e078538.
8. Omar M, Sorin V, Wieler LH, Charney AW, Kovatch P, Horowitz CR, et al. Mapping the susceptibility of large language models to medical misinformation across clinical notes and social media: a cross-sectional benchmarking analysis. *Lancet Digit Health*. 2026;8(1):100949.
9. Medina Grajales ME, Uyasán Giraldo RA, Valenzuela Gutiérrez RR. Riesgos de la inteligencia artificial generativa en la producción de la información para la toma de decisiones. Bogotá: Universidad EAN; 2024.
10. Huo B, Boyle A, Marfo N, Tangamornsuksan W, Steen JP, McKechnie T, et al. Large language models for chatbot health advice studies: a systematic review. *JAMA Netw Open*. 2025;8:e2457879.
11. Pal A, Wangmo T, Bharadia T, Ahmed-Richards M, Bhandari MB, Kachhadiya R, et al. Generative AI/LLMs for plain language medical information for patients, caregivers and general public: opportunities, risks and ethics. *Patient Prefer Adherence*. 2025;19:2227-2249.
12. Sallam M. ChatGPT utility in healthcare education, research, and practice: systematic review on the promising perspectives and valid concerns. *Healthcare (Basel)*. 2023;11(6):887.
13. Ayers JW, Poliak A, Dredze M, Leas EC, Zhu Z, Kelley JB, et al. Comparing physician and artificial intelligence chatbot responses to patient questions posted to a public social media forum. *JAMA Intern Med*. 2023;183(6):589-596.
14. Templin T, Perez MW, Sylvia S, Leek J, Sinnott-Armstrong N. Addressing 6 challenges in

- generative AI for digital health: a scoping review. *PLOS Digit Health*. 2024;3(5):e0000503.
15. Ning Y, Teixayavong S, Shang Y, Savulescu J, Nagaraj V, Miao D, et al. Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist. *Lancet Digit Health*. 2024;6:e848-e856.
 16. Lee P, Bubeck S, Petro J. Benefits, limits, and risks of GPT-4 as an AI chatbot for medicine. *N Engl J Med*. 2023;388(13):1233-1239.
 17. Haug CJ, Drazen JM. Artificial intelligence and machine learning in clinical medicine. *N Engl J Med*. 2023;388(13):1201-1208.
 18. Rajpurkar P, Lungren MP. The current and future state of AI interpretation of medical images. *N Engl J Med*. 2023;388(21):1981-1990.
 19. Gordon M, Daniel M, Ajiboye A, Uraiby H, Xu NY, Bartlett R, et al. A scoping review of artificial intelligence in medical education: BEME Guide No. 84. *Med Teach*. 2024;46(4):446-470.
 20. Saroha S. Artificial intelligence in medical education: promise, pitfalls, and practical pathways. *Adv Med Educ Pract*. 2025;16:1039-1046.
 21. Civaner MM, Uncu Y, Bulut F, Chalil EG, Tatli A. Artificial intelligence in medical education: a cross-sectional needs assessment. *BMC Med Educ*. 2022;22(1):772.
 22. Azer SA, Guerrero APS. The challenges imposed by artificial intelligence: are we ready in medical education? *BMC Med Educ*. 2023;23(1):680.
 23. Ghorashi N, Ismail A, Ghosh P, Al Haddad M, Patel R, Rahman S, et al. AI-powered chatbots in medical education: potential applications and implications. *Cureus*. 2023;15(8):e43271.
 24. Boscardin CK, Gin B, Golde PB, Hellevig A, Choi J, O'Sullivan PS, et al. ChatGPT and generative artificial intelligence for medical education: potential impact and opportunity. *Acad Med*. 2024;99(1):22-27.
 25. Knopp MI, Warm EJ, Weber D, Durning SJ, Hemmer PA, Fazio SB, et al. AI-enabled medical education: threads of change, promising futures, and risky realities across four potential future worlds. *JMIR Med Educ*. 2023;9:e50373.
 26. Mondal H. Ethical engagement with artificial intelligence in medical education. *Adv Physiol Educ*. 2025;49(1):163-165.
 27. International Committee of Medical Journal Editors. Use of AI by authors. ICMJE; 2026. Available from: <https://www.icmje.org/recommendations/browse/artificial-intelligence/ai-use-by-authors.html>
 28. US Food and Drug Administration. Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) Action Plan. Silver Spring: FDA; 2021.
 29. European Commission. Artificial intelligence in healthcare. Brussels: European Commission; 2024. Available from: https://health.ec.europa.eu/ehealth-digital-health-and-care/artificial-intelligence-healthcare_en
 30. UNESCO. Recommendation on the ethics of artificial intelligence. Paris: UNESCO; 2022.
 31. World Health Organization. Regulatory considerations on artificial intelligence for health. Geneva: World Health Organization; 2023.
 32. Aboy M, Minssen T, Vayena E. Navigating the EU AI Act: implications for regulated digital medical products. *NPJ Digit Med*. 2024; 7:237.
 33. Haltaufderheide J, Ranisch R. The ethics of ChatGPT in medicine and healthcare: a systematic review on large language models (LLMs). *NPJ Digit Med*. 2024;7:183.
 34. Omiye JA, Gui H, Rezaei SJ, Zou J, Daneshjou R. Large language models in medicine: the potential and pitfalls: a narrative review. *Ann Intern Med*. 2024;177(2):210-220.