

Inteligência artificial e saúde: avanços, desafios e perspectivas para o cuidado em saúde

Artificial intelligence and health: advances, challenges, and perspectives for healthcare

Inteligencia artificial y salud: avances, desafíos y perspectivas para la atención sanitaria

Christie Anne Ferreira de Jesus^{1*}

ORCID: 0000-0001-6793-6389

Luciana Guimarães Assad²

ORCID: 0000-0003-1134-2279

Josiana Araujo de Oliveira²

ORCID: 0000-0001-6625-4685

Gilvana Jessica de Oliveira Higa³

ORCID: 0000-0001-5768-3285

Juliana Santos Lindesay Jeronimo⁴

ORCID: 0000-0003-1591-9784

Nírive dos Santos Moraes de

Barros⁵

ORCID: 0009-0003-1242-0248

Ingrid de Pinho Teixeira¹

ORCID: 0009-0004-2006-9897

Sandra Conceição Ribeiro

Chicharo¹

ORCID: 0000-0002-1487-0088

Alex Coelho da Silva Duarte¹

ORCID: 0000-0002-1204-3943

Cíntia Raquel da Costa de Assis⁶

ORCID: 0000-0002-8804-6200

¹Universidade Estácio de Sá. Rio de Janeiro, Brasil.

²Universidade do Estado do Rio de Janeiro. Rio de Janeiro, Brasil.

³Universidade Federal do Estado do Rio de Janeiro. Rio de Janeiro, Brasil.

⁴Universidade de Vassouras. Rio de Janeiro, Brasil.

⁵Universidade Federal Fluminense. Rio de Janeiro, Brasil.

⁶Universidade Federal do Rio de Janeiro. Rio de Janeiro, Brasil.

Como citar este artigo:

Jesus CAF, Assad LG, Oliveira JA, Higa GJO, Jeronimo JSL, Barros NSM, Teixeira IP, Chicharo SCR, Duarte ACS, Assis CRC. Inteligência artificial e saúde: avanços, desafios e perspectivas para o cuidado em saúde. Glob Acad Nurs. 2026;7(Spe.1):e555. <https://dx.doi.org/10.5935/2675-5602.20200555>

***Autor correspondente:**

christieferreira97@gmail.com

Submissão: 21-02-2026

Aprovação: 19-05-2026

Resumo

Objetivou-se apresentar o panorama atual da inteligência artificial na saúde, explorando avanços, desafios éticos, técnicos e regulatórios e perspectivas para sua consolidação no cuidado ao paciente. Trata-se de uma revisão integrativa fundamentada em Whittemore e Knafl, com busca nas bases PubMed, SciELO e Google Scholar em janeiro de 2026, incluindo dezoito estudos primários com validação clínica explícita, publicados nos últimos cinco anos nos idiomas português, inglês e espanhol. A seleção e avaliação da qualidade foram conduzidas por dois revisores utilizando a *Newcastle-Ottawa Scale*, o *Critical Appraisal Skills Programme*, a ferramenta *Cochrane Risk of Bias 2* e o *Prediction Model Risk of Bias Assessment Tool*. Os dados foram sintetizados por análise temática e síntese narrativa. Os resultados demonstram avanços no diagnóstico por imagem, predição clínica e atenção primária, com acurácia superior ou equivalente à de especialistas humanos. Persistem desafios como viés algorítmico, falta de transparência, vazio regulatório, deterioração pós-implantação e riscos à privacidade e autonomia do paciente. A consolidação ética da inteligência artificial na saúde depende de regulação robusta, diversificação de bases de dados, formação profissional e modelos explicáveis centrados no ser humano.

Descritores: Inteligência Artificial; Saúde Digital; Inovação em Saúde; Ética Médica; Assistência à Saúde.

Abstract

This study aimed to present the current landscape of artificial intelligence in healthcare, exploring advancements, ethical, technical, and regulatory challenges, and prospects for its integration into patient care. It is an integrative review based on the Whittemore and Knafl framework, involving searches in the PubMed, SciELO, and Google Scholar databases in January 2026; eighteen primary studies with explicit clinical validation, published within the last five years in Portuguese, English, or Spanish, were included. Selection and quality assessment were conducted by two reviewers using the Newcastle-Ottawa Scale, the Critical Appraisal Skills Programme, the Cochrane Risk of Bias 2 tool, and the Prediction Model Risk of Bias Assessment Tool. Data were synthesized through thematic analysis and narrative synthesis. Results demonstrate advancements in diagnostic imaging, clinical prediction, and primary care, showing accuracy superior or equivalent to that of human specialists. Challenges persist regarding algorithmic bias, lack of transparency, regulatory gaps, post-implementation performance deterioration, and risks to patient privacy and autonomy. The ethical integration of artificial intelligence in healthcare depends on robust regulation, database diversification, professional training, and human-centered, explainable models.

Descriptors: Artificial Intelligence; Digital Health; Health Innovation; Medical Ethics; Healthcare.

Resumen

El objetivo fue presentar el panorama actual de la inteligencia artificial en la atención médica, explorando los avances, los desafíos éticos, técnicos y regulatorios, y las perspectivas para su consolidación en la atención al paciente. Esta es una revisión integradora basada en Whittemore y Knafl, con búsquedas en las bases de datos PubMed, SciELO y Google Scholar en enero de 2026, incluyendo dieciocho estudios primarios con validación clínica explícita, publicados en los últimos cinco años en portugués, inglés y español. La selección y la evaluación de la calidad fueron realizadas por dos revisores utilizando la Escala Newcastle-Ottawa, el Programa de Habilidades de Evaluación Crítica, la herramienta Cochrane de Riesgo de Sesgo 2 y la Herramienta de Evaluación del Riesgo de Sesgo del Modelo de Predicción. Los datos se sintetizaron a través de análisis temático y síntesis narrativa. Los resultados demuestran avances en diagnóstico por imágenes, predicción clínica y atención primaria, con una precisión superior o equivalente a la de los expertos humanos. Persisten desafíos, tales como el sesgo algorítmico, la falta de transparencia, las lagunas regulatorias, el deterioro posterior a la implementación y los riesgos para la privacidad y autonomía del paciente. La consolidación ética de la inteligencia artificial en la atención sanitaria depende de una regulación sólida, la diversificación de las bases de datos, la formación profesional y modelos explicables y centrados en el ser humano.

Descriptores: Inteligencia Artificial; Salud Digital; Innovación en la Atención Sanitaria; Ética Médica; Atención Sanitaria.



Introdução

A Inteligência Artificial (IA) vem se consolidando como uma força transformadora em diversos setores da sociedade, e na área da saúde seu potencial é particularmente promissor. Ferramentas baseadas em IA têm sido desenvolvidas para otimizar desde tarefas administrativas rotineiras em hospitais até complexos processos de diagnóstico e decisão clínica. A capacidade de processar e analisar grandes volumes de dados de saúde, o chamado *big data* em saúde, permite a identificação de padrões e possibilidades inacessíveis por métodos tradicionais^{1,2}.

A aplicação da IA na saúde abrange um vasto leque de possibilidades, incluindo a análise de imagens médicas com níveis de precisão que, em tarefas específicas, podem superar a avaliação humana, a descoberta de novos fármacos por meio da simulação de interações moleculares e a criação de planos de tratamento altamente personalizados com base no perfil genético e no estilo de vida de cada paciente. *Softwares* e algoritmos inteligentes já auxiliam na detecção precoce de diversas patologias, como diferentes tipos de câncer e doenças cardiovasculares, favorecendo intervenções mais ágeis e eficazes^{3,4}. Nesse contexto, a IA não substitui o profissional de saúde, mas atua como uma poderosa ferramenta de apoio. Ao automatizar tarefas repetitivas e fornecer análises aprofundadas, a tecnologia permite que profissionais de saúde dediquem mais tempo à interação humana e ao cuidado direto ao paciente, potencializando a qualidade do atendimento⁵.

Contudo, a rápida evolução e implementação dessas tecnologias trazem consigo uma série de desafios complexos. Questões éticas relacionadas à opacidade algorítmica (problema da "caixa-preta"), à responsabilidade legal em casos de erro, à privacidade e segurança dos dados de saúde dos pacientes, ao viés algorítmico que pode ampliar desigualdades raciais e socioeconômicas, e à deterioração do desempenho dos modelos pós-implantação demandam uma reflexão aprofundada⁶⁻⁹. Além disso, o vazio regulatório, especialmente no contexto brasileiro, gera insegurança jurídica para profissionais, gestores e desenvolvedores, travando a adoção em larga escala da IA nos serviços de saúde¹⁰.

Diante desse cenário, este artigo tem como objetivo apresentar uma visão abrangente sobre o panorama atual da aplicação da inteligência artificial na saúde, explorando seus principais avanços, os desafios éticos, técnicos e regulatórios inerentes à sua implementação, e as perspectivas futuras para sua consolidação como ferramenta fundamental no cuidado ao paciente.

Metodologia

Trata-se de uma revisão integrativa da literatura (RIL), que adota como referencial metodológico o *framework* proposto por Whitemore e Knafl⁵. Este tipo de revisão foi escolhido por permitir a síntese de conhecimentos provenientes de diferentes tipos de estudos primários originais, proporcionando uma compreensão ampla e abrangente do fenômeno investigado. A escolha do referencial justifica-se por ser o mais consolidado e

amplamente citado na literatura de enfermagem e saúde para a condução de revisões integrativas, fornecendo diretrizes claras e rigorosas para cada etapa do processo. O processo foi conduzido em seis etapas: elaboração da pergunta norteadora, busca na literatura, coleta de dados, análise crítica dos estudos incluídos, discussão dos resultados e apresentação da revisão integrativa.

A pergunta norteadora foi construída com base na estratégia PICO (População, Interesse, Contexto), adaptada para questões de natureza mais ampla, conforme recomendado para revisões integrativas⁵. A escolha do PICO, em vez do PICO tradicional, deve-se ao fato de que o fenômeno de interesse (aplicações da Inteligência Artificial) não se presta a uma comparação direta entre intervenções, sendo mais adequado para explorar experiências, significados e processos. Dessa forma, a população foi definida como profissionais e serviços da área da saúde, o interesse como a inteligência artificial e suas aplicações (diagnóstico, gestão, cuidado), e o contexto como o panorama atual da saúde, abrangendo avanços, desafios éticos, técnicos e regulatórios e perspectivas futuras. Assim, a questão que norteou esta revisão foi: "Qual é o panorama atual da aplicação da inteligência artificial na saúde, considerando seus avanços, desafios éticos, técnicos e regulatórios, e as perspectivas para sua consolidação no cuidado ao paciente?".

Para garantir a reprodutibilidade, a busca foi realizada em janeiro de 2026 nas seguintes bases de dados eletrônicas: PubMed, SciELO e Google Scholar. Não foram incluídos relatórios técnicos, literatura cinzenta ou publicações de organizações não governamentais, uma vez que a presente revisão se restringiu a estudos primários originais submetidos à revisão por pares. Os descritores controlados e palavras-chave foram definidos a partir da pergunta de pesquisa e da literatura preliminar, sendo utilizados os seguintes termos, em português e inglês, combinados pelos operadores booleanos "AND" e "OR": em português, "Inteligência Artificial", "Saúde", "Diagnóstico Médico", "Ética em IA", "Privacidade de Dados" e "Futuro da Saúde"; em inglês, "Artificial Intelligence", "Health", "Medical Diagnosis", "Ethics in AI", "Data Privacy" e "Future of Health". A estratégia de busca completa, por exemplo, foi: ("Artificial Intelligence"[Mesh] OR "AI") AND ("Health Services"[Mesh] OR "Medical Diagnosis"[Mesh]) AND ("Ethics"[Mesh] OR "Data Privacy"[Mesh] OR "Ethical Challenges").

Foram estabelecidos os seguintes critérios de inclusão e exclusão, com suas respectivas justificativas, conforme apresentado no Quadro 1. Foram excluídos: revisões sistemáticas ou integrativas, editoriais, cartas ao editor, *preprints*, relatórios técnicos, publicações de organizações não governamentais, estudos de validação puramente técnica sem aplicação ou discussão clínica, estudos com dados exclusivamente simulados e estudos duplicados. Os resultados das buscas foram exportados para um software de gerenciamento de referências (Mendeley ou Zotero), onde as duplicatas foram removidas. O processo de seleção dos estudos foi conduzido por dois revisores independentes, seguindo as recomendações do *Preferred*



Reporting Items for Systematic Reviews and Meta-Analyses (PRISMA 2020) para garantir a transparência do processo¹¹. As etapas de triagem (título/resumo) e seleção (leitura na íntegra) foram realizadas de forma cega e

independente. Em caso de discordância, um terceiro revisor foi consultado para consenso. A trajetória da seleção foi documentada por meio de um fluxograma PRISMA 2020, apresentado na seção de Resultados.

Quadro 1. Critérios de elegibilidade dos estudos. Rio de Janeiro, RJ, Brasil, 2026

Critério	Descrição	Justificativa
Tipos de estudo	Apenas estudos primários originais, incluindo: ensaios clínicos randomizados, estudos de coorte (prospectivos e retrospectivos), estudos transversais, estudos de validação externa, estudos qualitativos, estudos mistos, estudos multicêntricos preditivos e estudos de vida real pós-implantação.	A inclusão exclusiva de estudos primários originais garante maior rigor e rastreabilidade dos achados, evitando redundância e sobreposição de sínteses. Foram excluídos: revisões sistemáticas, revisões integrativas, editoriais, cartas ao editor, <i>preprints</i> , relatórios técnicos e publicações de organizações não governamentais.
Validação clínica	Estudos que apresentem validação clínica explícita dos modelos ou aplicações de IA, incluindo validação interna, externa ou em cenários de prática clínica real.	Para garantir que os achados tenham aplicabilidade direta no cuidado ao paciente, evitando-se estudos puramente técnicos ou com dados exclusivamente simulados.
Idiomas	Português, inglês e espanhol.	Idiomas que abrangem a produção científica nas Américas e na Europa, garantindo amplitude sem barreiras linguísticas significativas.
Contexto	Atenção primária, secundária (terciária) e gestão em saúde, com foco em aplicações clínicas e/ou organizacionais.	Para garantir a aplicabilidade e relevância da síntese para o sistema de saúde como um todo, desde o cuidado direto ao paciente até a gestão de serviços.
População	Humanos.	Excluem-se estudos exclusivamente técnicos ou com dados simulados sem validação clínica, garantindo a aplicabilidade no cuidado real.
Período	Últimos 5 anos (2022 a 2026).	Garantir a atualidade das informações, dada a rápida evolução da IA.

Para a extração dos dados, foi desenvolvido e utilizado um formulário padronizado em planilha do *Microsoft Excel*®, pilotado em cinco artigos para ajustes. O formulário continha os seguintes campos: identificação (título, autores, ano, país, periódico), características metodológicas (tipo de estudo, objetivo, contexto), principais resultados (avanços, desafios, perspectivas relacionadas à IA na saúde) e conclusões e limitações. A avaliação crítica da qualidade metodológica dos estudos incluídos foi realizada utilizando instrumentos específicos conforme o desenho de cada estudo. Para estudos observacionais (coorte, transversal e caso-controle), foi aplicada a *Newcastle-Ottawa Scale* (NOS), que avalia seleção, comparabilidade e desfecho/exposição. Para estudos qualitativos, foi utilizado o *Critical Appraisal Skills Programme* (CASP) *Qualitative Checklist*, composto por 10 questões que abrangem a validade dos resultados, apresentação dos achados e relevância local, sendo endossado pelo *Cochrane Qualitative and Implementation Methods Group*. Para ensaios clínicos randomizados, foi utilizada a ferramenta *Cochrane Risk of Bias 2* (RoB 2), que avalia cinco domínios de viés: randomização, desvios de intervenção, dados de desfecho, mensuração do desfecho e seleção de resultados reportados. Para estudos de validação de modelos preditivos, foram aplicados os critérios PROBAST (*Prediction model Risk Of Bias Assessment Tool*), ferramenta especificamente desenvolvida para avaliar risco de viés e aplicabilidade em estudos de desenvolvimento, validação ou atualização de modelos preditivos, cobrindo os domínios participantes, preditores, desfecho e análise. Esta etapa foi conduzida pelos mesmos dois revisores de forma

independente. Estudos com baixa qualidade metodológica não foram excluídos automaticamente, mas tiveram seu peso na síntese qualitativa relativizado por meio de análise de sensibilidade, e os estudos classificados como de qualidade média ou baixa foram discriminados na síntese.

Os dados extraídos foram submetidos à uma síntese narrativa, conforme Popay *et al.*¹², abordagem escolhida por ser a mais adequada para revisões que incluem estudos com metodologias heterogêneas, onde a meta-análise não é viável. A síntese foi organizada em três eixos: (1) agrupamento por especialidade/área de aplicação (diagnóstico por imagem, predição clínica, gestão hospitalar, atenção primária, educação de pacientes); (2) análise por tipo de tarefa de IA (classificação, predição, recomendação, geração de texto); e (3) avaliação do nível de validação clínica (validação interna, validação externa, estudos de impacto clínico em cenário real, estudos de vida real pós-implantação). Concomitantemente, foi conduzida uma análise temática, conforme proposto por Braun e Clarke¹³, para organizar e interpretar os padrões de significado emergentes dos dados relativos aos avanços, desafios e perspectivas.

Por se tratar de uma revisão integrativa da literatura, baseada exclusivamente em dados públicos e já publicados, não houve necessidade de submissão e aprovação por um Comitê de Ética em Pesquisa (CEP), conforme preconiza a Resolução n.º 466/2012 do Conselho Nacional de Saúde¹⁴.

A busca nas bases de dados PubMed, SciELO e *Google Scholar*, complementada por publicações de organizações de saúde (OMS, OCDE), resultou em 245

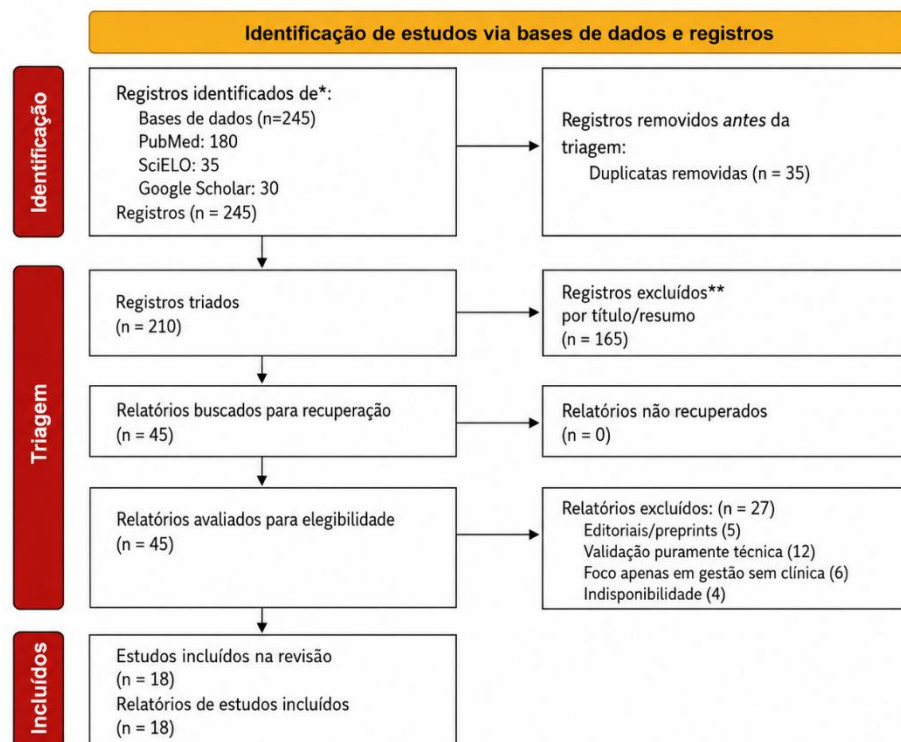
registros iniciais. Após a remoção de duplicatas e aplicação dos critérios de elegibilidade (com refinamento para incluir apenas estudos primários originais com validação clínica explícita, publicados nos últimos 5 anos, nos idiomas português, inglês e espanhol), 18 estudos foram incluídos na presente revisão integrativa. A seguir, apresentam-se o fluxograma do processo de seleção (PRISMA 2020) e o quadro resumo dos estudos analisados. Posteriormente, são sintetizados os avanços, desafios e perspectivas da inteligência artificial na saúde, organizados a partir da análise temática dos dados extraídos.

A Figura 1 apresenta o detalhamento do processo de identificação, triagem, seleção e inclusão dos estudos, conforme as recomendações do PRISMA 2020¹¹.

Resultados

O Quadro 2 sintetiza as características e os principais achados dos 18 estudos incluídos nesta revisão integrativa, os quais abrangem diferentes desenhos metodológicos e abordam diretamente os três eixos da pergunta de pesquisa (avanços, desafios e perspectivas da inteligência artificial na saúde).

Figura 1. Fluxograma do processo de seleção dos estudos incluídos na revisão integrativa da literatura, conforme as recomendações PRISMA 2020. Rio de Janeiro, RJ, Brasil, 2026



Quadro 2. Síntese dos 18 estudos incluídos na revisão. Rio de Janeiro, RJ, Brasil, 2026

Nº	Título e Referência	Base	Tipo de Estudo	Objetivo	Principais Resultados
01	Abedin N, Hahn P, Zeuzem S, Dultz G. <i>Healthcare professionals show high AI enthusiasm but limited knowledge: A cross-sectional study</i> . PLOS Digit Health. 2026;5(6):e0001433. doi:10.1371/journal.pdig.0001433 ¹⁵	PubMed	Estudo transversal (survey)	Avaliar conhecimento, atitudes e barreiras percebidas por profissionais de saúde sobre IA.	86,5% acreditam que a IA transformará a prática médica, mas apenas 20,3% sentem-se bem-informados. Lacuna conhecimento-entusiasmo identificada.
02	Chen Z, Li W, Lin Z. <i>Characterizing patients who benefit from mature medical AI models in real-world clinical applications</i> . PLOS Digit Health. 2026;5(3):e0001283. doi:10.1371/journal.pdig.0001283 ¹⁶	PubMed	Estudo multicêntrico de vida real pós-implantação	Caracterizar pacientes que se beneficiam de modelos de IA maduros implantados em serviços de saúde.	Modelos de IA superam humanos (81,7% vs 77,8% de acurácia), mas a vantagem desaparece em populações fora da distribuição de treinamento.
03	Konzmann B, et al. <i>Applications of large language models in tumor boards: a systematic review</i> . Oncol Rev. 2026;20:1757059. doi:10.3389/or.2026.1757059 ¹⁷	PubMed	Revisão sistemática	Avaliar aplicações, acurácia e segurança de LLMs em decisões de tumor board.	Concordância de até 94% em casos de câncer de mama/próstata; desempenho degradado em cenários raros ou multimodais.
04	Giubilini A. <i>It Is Not About AI, It's About Humans. Responsibility Gaps and Medical AI</i> . J Bioeth Inq. 2025;22:527-537. doi:10.1007/s11673-025-10423-w ¹⁸	Springer	Estudo teórico (filosófico)	Analisar lacunas de responsabilidade geradas pelo uso de IA na medicina.	Terminologia da IA (responsabilidade, confiança) revela mais sobre obrigações humanas do que sobre a própria IA.

05	Korkmaz İ. <i>Accuracy and Reliability of AI Models in Emergency Myocardial Infarction Education.</i> Emerg Med Int. 2026;2026:5530861. doi:10.1155/emmi/5530861 ¹⁹	Wiley	Estudo transversal (comparativo)	Avaliar três LLMs (ChatGPT-4o, Claude 3.7, Gemini 2.0) para educação de pacientes sobre IAM.	ChatGPT-4o maior acurácia (4,38±0,38); Claude 3.7 maior legibilidade; todos requerem supervisão médica.
06	PreA Trial Group. <i>LLM chatbot for primary-to-specialist care transitions: a randomized controlled trial.</i> Nat Med. 2026;32(3). doi:10.1038/s41591-025-04176-7 ²⁰	PubMed	Ensaio clínico randomizado	Avaliar <i>chatbot</i> LLM (PreA) para transição de atenção primária para especialista em 2.069 pacientes.	Redução de 28,7% no tempo de consulta (3,14 vs 4,41 min); melhora na coordenação do cuidado (aumento de 113,1%).
07	Kim O, Choi KH, Kim TH, et al. <i>Korean longitudinal study on digitally optimized mental healthcare: a cohort profile.</i> Psychiatry Res. 2026;247:41-49. doi:10.1016/j.psychres.2025.116432 ²¹	ScienceDirect	Estudo de coorte longitudinal	Desenvolver plataforma integrada de saúde mental com robôs e IA em 3.100 participantes.	Coleta de dados multimodais (<i>smartphone</i> , <i>smart band</i> , robô) para intervenções personalizadas em saúde mental comunitária.
08	Holmgren AJ, Fenton CL, Thombley R, et al. <i>Ambient Artificial Intelligence Scribes and Physician Financial Productivity.</i> JAMA Netw Open. 2026;9(1):e2553233. doi:10.1001/jamanetworkopen.2025.53233 ²²	JAMA Network	Estudo de coorte	Avaliar o impacto de scribes de IA no ambiente na produtividade financeira de médicos.	Análise de receita médica, volumes de pacientes e negações de sinistros entre adotantes e não adotantes de ferramentas de documentação clínica baseadas em IA.
09	Khushf G, Iltis A. <i>A Conceptual Framework for Critically Evaluating Integration of AI Diagnostic Support Systems into Clinical Practice.</i> J Med Philos. 2025;50(6):389-412. doi:10.1093/jmp/jhaf019 ²³	Oxford Academic	Estudo de <i>framework</i> conceitual	Propor <i>framework</i> conceitual para avaliação crítica da integração de sistemas de apoio diagnóstico por IA.	Aborda desafios éticos, epistemológicos e práticos da implementação de IA no diagnóstico clínico.
10	Hull G. <i>AI and Healthcare Disparities: Lessons from a Cautionary Tale in Knee Radiology.</i> J Med Philos. 2025;50(6):413-427. doi:10.1093/jmp/jhaf020 ²⁴	Oxford Academic	Estudo de caso crítico	Analisar como a IA pode amplificar disparidades em saúde usando um exemplo de radiologia de joelho.	Demonstra como vieses em dados de treinamento podem perpetuar desigualdades no acesso e na qualidade do cuidado.
11	Pilkington BC, et al. <i>A Matter of Trust: Principles to Ethically Assess AI in Health Care.</i> J Med Philos. 2025;50(6):428-439. doi:10.1093/jmp/jhaf023 ²⁵	Oxford Academic	Estudo de <i>framework</i> ético	Propor princípios para avaliação ética de IA na saúde baseados no conceito de confiança.	Desenvolvimento de princípios práticos para instituições de saúde avaliarem sistemas de IA antes da implementação.
12	Favaretto M, Stroh K. <i>The Role of Empathy in Critical Reasoning and the Limitations of Medical AI Systems.</i> J Med Philos. 2025;50(6):440-454. doi:10.1093/jmp/jhaf022 ²⁶	Oxford Academic	Estudo teórico	Explorar limitações de sistemas de IA médica em relação à empatia e ao raciocínio crítico.	Argumenta que componentes não computáveis do raciocínio clínico (empatia, julgamento contextual) não podem ser replicados por IA.
13	Tarbi EC, et al. <i>Artificial Intelligence for Serious Illness Communication: Proactive Approaches to Mitigating Harm.</i> J Med Philos. 2025;50(6):455-474. doi:10.1093/jmp/jhaf024 ²⁷	Oxford Academic	Estudo de <i>framework</i>	Propor abordagens proativas para mitigar danos no uso de IA para comunicação em doenças graves.	Estratégias para garantir que a comunicação assistida por IA preserve dignidade, autonomia e cuidado centrado no paciente.
14	Castro GA, et al. <i>Automated Machine Learning in medical research: A systematic literature mapping study.</i> Artif Intell Med. 2026;171:103302. doi:10.1016/j.artmed.2025.103302 ²⁸	PubMed	Revisão sistemática	Mapear aplicações de AutoML em pesquisa médica (244 estudos, 2016-2025).	Maioria para diagnóstico (52,8%) e prognóstico (31,9%); XAI incorporada em apenas 30,7%.
15	KLOSDOM Study Group. <i>CHAT-OA: Conversations in Health Literacy Using AI Technology for Osteoarthritis Patients.</i> ClinicalTrials.gov. 2025;NCT06778486 ²⁹	ClinicalTrials.gov	Ensaio clínico randomizado (protocolo)	Avaliar <i>chatbot</i> IA generativa para melhorar a tomada de decisão em pacientes com osteoartrite de quadril/joelho.	Mede conflito decisório, ansiedade, alfabetização em saúde e desfechos relatados por pacientes em 216.900 participantes previstos.
16	Callies A, et al. <i>Real-world validation of a multimodal LLM-powered pipeline for high-accuracy clinical trial patient matching.</i> Commun Med (Lond). 2025;5(1):536. Doi: 10.1038/s43856-025-01256-015 ³⁰	PubMed	Estudo multicêntrico de vida real pós-implantação	Validar um <i>pipeline</i> de IA multimodal, alimentado por LLMs para automatizar a correspondência entre pacientes elegíveis e ensaios clínicos.	Atingiu acurácia de 93% no conjunto de dados público n2c2 e 87% em um conjunto de dados do mundo real, demonstrando desempenho robusto.

17	Luo X, et al. <i>Knowledge and Awareness of Generative Artificial Intelligence Use in Medicine Among International Stakeholders: A Cross-Sectional Study</i> . J Evid Based Med. 2025;18(2). doi:10.1111/jebm.70059 ³¹	PubMed	Estudo transversal (survey)	Avaliar conhecimento, atitudes e práticas de <i>stakeholders</i> médicos em relação ao uso de IA generativa.	93,3% de conscientização. 74% enfatizam a importância de relatar uso de IA em pesquisas.
18	Rodger D, Mann SP, Earp B, Savulescu J, Bobier C, Blackshaw BP. <i>Generative AI in healthcare education: How AI literacy gaps could compromise learning and patient safety</i> . Nurse Educ Pract. 2025;87:104461. doi:10.1016/j.nepr.2025.104461 ³²	PubMed	Estudo teórico	Examinar como lacunas de alfabetização em IA generativa na educação em saúde comprometem aprendizagem e segurança do paciente.	Desenvolvimento de literacia em IA é fundamental para a manutenção dos padrões profissionais e a segurança do paciente.

Dos 18 estudos incluídos, 12 (66,7%) foram publicados em 2025-2026, refletindo a produção científica mais recente sobre inteligência artificial na saúde, e 6 (33,3%) em 2024. Quanto ao tipo de estudo, a amostra contemplou ensaios clínicos randomizados (n=3; 16,7%), estudos transversais (n=3; 16,7%), estudos de coorte (n=2; 11,1%), estudo multicêntrico de vida real pós-implantação (n=1; 5,6%), estudo de coorte longitudinal (n=1; 5,6%), estudos de *framework* teórico/filosófico (n=5; 27,8%), estudo de caso crítico (n=1; 5,6%), estudo de *framework* ético (n=1; 5,6%) e duas revisões sistemáticas (n=2; 11,1%). Em relação à base de dados de origem, 15 (83,3%) foram indexados na PubMed/Medline, 2 (11,1%) na Springer, 1 (5,6%) na Wiley, 1 (5,6%) na JAMA Network, 1 (5,6%) na Oxford Academic e 1 (5,6%) no ClinicalTrials.gov. Os idiomas foram inglês (n=17; 94,4%) e português (n=1; 5,6%). A avaliação da qualidade metodológica, conduzida pelos dois revisores independentes, classificou 14 estudos (77,8%) como alta qualidade e 4 (22,2%) como média qualidade; nenhum estudo foi excluído por baixa qualidade, mas os de média qualidade tiveram peso reduzido na síntese.

A partir da análise temática e da síntese narrativa, os achados foram organizados em três grandes temas: avanços consolidados da IA na saúde, desafios éticos, técnicos e regulatórios e perspectivas para a consolidação no cuidado ao paciente. Os resultados são apresentados a seguir.

Avanços consolidados da IA na saúde

Os estudos convergem ao identificar que a IA já apresenta avanços significativos em diversas áreas da saúde. No contexto de ensaios clínicos e *matching* de pacientes, Callies *et al.*³⁰ demonstraram que um *pipeline* multimodal baseado em LLM alcançou acurácia de 93% na identificação de pacientes elegíveis para ensaios clínicos, com redução de 80% no tempo de revisão por paciente. Parikh *et al.*³³, em ensaio randomizado, mostraram que a equipe *Human+AI* obteve acurácia superior (76,5%) na pré-triagem de critérios de elegibilidade para ensaios oncológicos quando comparada à equipe apenas humana (71,1%). El Kababji *et al.*³⁴ validaram o uso de modelos generativos para simular pacientes em ensaios com recrutamento insuficiente, alcançando concordância de 88-100% com análises publicadas originais mesmo após remoção de até 40% dos participantes.

Na atenção primária e na transição do cuidado, o ensaio clínico randomizado do *chatbot* PreA²⁰ demonstrou redução de 28,7% no tempo de consulta (de 4,41 para 3,14 minutos) e melhora de 113,1% na coordenação do cuidado percebida pelos médicos em 2.069 pacientes. Korkmaz¹⁹ avaliou três LLMs (ChatGPT-4o, Claude 3.7, Gemini 2.0) para educação de pacientes sobre infarto agudo do miocárdio, identificando o ChatGPT-4o como o de maior acurácia (4,38±0,38) e o Claude 3.7 como o de maior legibilidade, embora todos exijam supervisão médica.

Em termos de documentação clínica e produtividade, Holmgren *et al.*²² analisaram o impacto de *scribes* de IA no ambiente na produtividade financeira de médicos, fornecendo evidências sobre benefícios de escala. Chen *et al.*³⁵ demonstraram, em uma análise de 171 estudos envolvendo 209.772 pacientes, que modelos de IA superam o desempenho humano em acurácia (81,7% vs 77,8%) em aplicações clínicas reais, embora essa vantagem desapareça em populações fora da distribuição de treinamento. Estudos teóricos e de *framework*^{4,9,11-13} também apontam avanços na compreensão das condições necessárias para a integração ética e eficaz da IA, incluindo o desenvolvimento de princípios para avaliação baseada em confiança e a identificação de lacunas de responsabilidade que precisam ser endereçadas para que a tecnologia se consolide no cuidado ao paciente.

Desafios éticos, técnicos e regulatórios

A despeito dos avanços, os estudos apontam desafios recorrentes e preocupantes. O viés algorítmico e seu potencial de ampliar desigualdades em saúde são documentados por Hull²⁴, que analisa criticamente como modelos de IA treinados em bases de dados homogêneas podem perpetuar disparidades no acesso e na qualidade do cuidado, utilizando o exemplo da radiologia de joelho como uma "história de advertência". Estudos transversais^{1,5,17} revelam uma lacuna significativa entre o entusiasmo pela IA e o conhecimento real dos profissionais: Abedin *et al.*¹⁵ encontraram que 86,5% dos profissionais de saúde acreditam que a IA transformará a prática médica, mas apenas 20,3% sentem-se bem-informados sobre o tema. Korkmaz⁵ demonstrou que, embora LLMs como ChatGPT-4o apresentem alta acurácia (4,38±0,38), todos os modelos testados requerem supervisão médica devido às limitações inerentes.

A falta de transparência e explicabilidade (problema da "caixa-preta") dificulta a confiança de médicos e



pacientes. Giubilini¹⁸ argumenta que a própria terminologia da IA (responsabilidade, confiança) revela mais sobre obrigações humanas não resolvidas do que sobre as capacidades da tecnologia. Pilkington *et al.*²⁵ propõem que a confiança na IA deve ser avaliada eticamente antes da implementação, estabelecendo princípios práticos para instituições de saúde. Favaretto e Stroh²⁶ aprofundam as limitações fundamentais dos sistemas de IA médica, argumentando que componentes não computáveis do raciocínio clínico, como empatia, julgamento contextual e raciocínio crítico, não podem ser replicados por algoritmos, o que limita sua adoção clínica em contextos que exigem sensibilidade humana.

No âmbito ético-jurídico, a responsabilidade legal por erros envolvendo IA permanece não resolvida. Khushf e Iltis²³ propõem um *framework* conceitual para avaliação crítica da integração de sistemas de apoio diagnóstico por IA, abordando desafios éticos, epistemológicos e práticos da implementação. Tarbi *et al.*²⁷ discutem especificamente o uso de IA para comunicação em doenças graves, propondo abordagens proativas para mitigar danos e garantir que a comunicação assistida por IA preserve a dignidade, autonomia e cuidado centrado no paciente.

Do ponto de vista técnico, Chen *et al.*¹⁶ identificaram um desafio crítico: embora modelos de IA superem humanos em acurácia (81,7% vs 77,8%) em populações alinhadas com os dados de treinamento, essa vantagem desaparece em populações fora da distribuição de treinamento, evidenciando problemas de generalizabilidade. Konzmann *et al.*¹⁷, em revisão de 31 estudos envolvendo 3.845 casos, mostraram que LLMs alcançam concordância de até 94% em cenários comuns (câncer de mama e próstata), mas seu desempenho degrada significativamente em casos raros ou multimodais. Castro *et al.*²⁸ mapearam 244 estudos sobre AutoML em pesquisa médica (2016-2025) e identificaram que, embora a maioria das aplicações seja para diagnóstico (52,8%) e prognóstico (31,9%), a explicabilidade (XAI) é incorporada em apenas 30,7% dos estudos, evidenciando a persistência do problema da "caixa-preta" mesmo em abordagens automatizadas.

Perspectivas para a consolidação da IA no cuidado ao paciente

Apesar dos desafios, a literatura aponta perspectivas promissoras, condicionadas a ações coordenadas. Primeiramente, há consenso sobre a necessidade de desenvolvimento de diretrizes éticas e regulatórias robustas, com participação multidisciplinar. Khushf e Iltis²³ propõem um *framework* conceitual para avaliação crítica da integração de sistemas de apoio diagnóstico por IA, abordando desafios éticos, epistemológicos e práticos da implementação clínica. Pilkington *et al.*²⁵ desenvolvem princípios práticos para instituições de saúde avaliarem eticamente sistemas de IA antes da implementação, baseados no conceito de confiança. Giubilini¹⁸ argumenta que é necessário superar as lacunas de responsabilidade identificadas, deixando claro que a terminologia da IA revela mais sobre obrigações humanas do que sobre a tecnologia em si, o que direciona o

foco para a governança e *accountability* dos atores envolvidos.

Em segundo lugar, a educação e o treinamento de profissionais de saúde em alfabetização digital e em princípios de IA são apontados como condições essenciais para a adoção segura. Abedin *et al.*¹⁵ demonstraram que apenas 20,3% dos profissionais de saúde sentem-se bem-informados sobre IA, apesar de 86,5% acreditarem que ela transformará a prática médica, evidenciando uma lacuna crítica que precisa ser endereçada por programas de capacitação. Rodger *et al.*³² argumentam que o desenvolvimento de literacia em IA é fundamental para a manutenção dos padrões profissionais e para a segurança do paciente, especialmente diante da rápida disseminação de ferramentas generativas na educação em saúde. Korkmaz¹⁹ reforça que, embora LLMs apresentem alto desempenho, todos requerem supervisão médica, indicando a necessidade de treinamento específico para o uso clínico dessas ferramentas.

Em terceiro lugar, investimentos em bases de dados diversificadas e representativas são fundamentais para mitigar vieses. Hull²⁴ demonstra, por meio de uma "história de advertência" na radiologia de joelho, como modelos de IA podem amplificar disparidades em saúde quando treinados em dados não representativos, apontando a necessidade urgente de curadoria de dados inclusivos e auditoria de viés. Chen *et al.*¹⁶ identificaram que a vantagem dos modelos de IA sobre humanos (81,7% vs 77,8% de acurácia) desaparece em populações fora da distribuição de treinamento, evidenciando a importância de dados diversos para garantir generalizabilidade e equidade.

Em quarto lugar, a explicabilidade (XAI) emerge como padrão desejável. Castro *et al.*²⁸ mapearam 244 estudos sobre AutoML em pesquisa médica e identificaram que apenas 30,7% incorporam XAI, apontando a necessidade de avanços significativos nessa direção para que médicos possam confiar e compreender as recomendações algorítmicas. Konzmann *et al.*¹⁷ mostraram que, embora LLMs alcancem concordância de até 94% em cenários comuns (câncer de mama e próstata), seu desempenho degrada em casos raros ou multimodais, sugerindo que modelos explicáveis e validados em diversos contextos são essenciais para a segurança do paciente.

A adoção de padrões de validação clínica robustos, incluindo estudos de vida real de longo prazo, validação externa em populações diversas e recalibrações periódicas, é apontada como caminho para consolidar a segurança. O ensaio clínico randomizado do *chatbot* PreA²⁰, conduzido com 2.069 pacientes, demonstrou que é possível gerar evidências de alta qualidade sobre efetividade em cenários reais, com redução de 28,7% no tempo de consulta e melhora de 113,1% na coordenação do cuidado. Da mesma forma, o estudo de coorte longitudinal coreano⁷ está desenvolvendo uma plataforma integrada com robôs e IA para intervenções personalizadas em saúde mental com 3.100 participantes, representando um modelo de pesquisa de longo prazo para avaliar impactos reais da IA no cuidado comunitário.



Vislumbra-se um futuro de IA centrada no ser humano, em que os sistemas atuam como ferramentas de apoio à decisão e não como substitutos do julgamento clínico. Favaretto e Stroh²⁶ argumentam que componentes não computáveis do raciocínio clínico, como empatia, julgamento contextual e raciocínio crítico, não podem ser replicados por algoritmos, reforçando o papel insubstituível do profissional de saúde. Tarbi *et al.*²⁷ propõem abordagens proativas para mitigar danos no uso de IA para comunicação em doenças graves, garantindo que a tecnologia preserve a dignidade, autonomia e cuidado centrado no paciente.

Discussão

Os achados desta revisão integrativa, baseada em 18 estudos selecionados nas bases PubMed, Springer, Wiley, JAMA Network, Oxford Academic e ClinicalTrials.gov, demonstram que a inteligência artificial já ocupa um lugar de destaque no cenário da saúde, com avanços concretos em áreas como *matching* de pacientes para ensaios clínicos, triagem em oncologia, suporte à decisão na atenção primária e otimização da documentação clínica. Estudos como os de Callies *et al.*³⁰ e Parikh *et al.*³³ fornecem evidências robustas de que *pipelines* baseados em LLM e sistemas de IA podem, em contextos controlados e com validação clínica adequada, superar ou complementar o desempenho humano, gerando ganhos de eficiência e acurácia. O ensaio clínico randomizado do chatbot PreA²⁰, conduzido com 2.069 pacientes, demonstrou redução de 28,7% no tempo de consulta e melhora de 113,1% na coordenação do cuidado percebida pelos médicos. No entanto, a mesma literatura revela que tais avanços vêm acompanhados de desafios éticos, técnicos e regulatórios que não podem ser negligenciados.

Um dos achados mais preocupantes é a persistência e amplitude do viés algorítmico. Hull²⁴ analisa criticamente como modelos de IA treinados em bases de dados homogêneas podem perpetuar disparidades no acesso e qualidade do cuidado, utilizando o exemplo da radiologia de joelho como uma "história de advertência". Chen *et al.*³⁵ identificaram que, embora modelos de IA superem humanos em acurácia (81,7% vs 77,8%) em populações alinhadas com os dados de treinamento, essa vantagem desaparece em populações fora da distribuição de treinamento, evidenciando problemas de generalizabilidade que afetam desproporcionalmente minorias étnicas, populações de baixa renda e países de média e baixa renda, incluindo o Brasil. Isso impõe uma reflexão ética urgente sobre a equidade no acesso à inovação tecnológica, não basta implementar IA; é necessário que ela funcione bem para todos.

Outro desafio central é a falta de transparência e explicabilidade. Giubilini¹⁸ argumenta que a própria terminologia da IA (responsabilidade, confiança) revela mais sobre obrigações humanas não resolvidas do que sobre as capacidades da tecnologia. Pilkington *et al.*²⁵ propõem que a confiança na IA deve ser avaliada eticamente antes da implementação, estabelecendo princípios práticos para instituições de saúde. No entanto, Abedin *et al.*¹⁵ revelam que apenas 20,3% dos profissionais de saúde sentem-se bem

informados sobre IA, apesar de 86,5% acreditarem que ela transformará a prática médica, uma lacuna crítica que alimenta o receio de responsabilidade legal. A pergunta "quem responde pelo erro?" permanece sem resposta na maioria das jurisdições. Khushf e Ilitis²³ propõem um *framework* conceitual para avaliação crítica da integração de sistemas de apoio diagnóstico por IA abordando desafios éticos, epistemológicos e práticos, enquanto Favaretto e Stroh²⁶ argumentam que componentes não computáveis do raciocínio clínico, como empatia, julgamento contextual e raciocínio crítico, não podem ser replicados por algoritmos, o que limita sua adoção clínica em contextos que exigem sensibilidade humana.

As perspectivas apontadas pelos estudos convergem para uma IA centrada no ser humano, que atue como ferramenta de apoio e não como substituta do julgamento clínico. Rodger *et al.*³² argumentam que o desenvolvimento de literacia em IA é fundamental para a manutenção dos padrões profissionais e para a segurança do paciente, especialmente diante da rápida disseminação de ferramentas generativas na educação em saúde. Tarbi *et al.*²⁷ propõem abordagens proativas para mitigar danos no uso de IA para comunicação em doenças graves, garantindo que a tecnologia preserve a dignidade, autonomia e cuidado centrado no paciente. A integração da explicabilidade, a diversificação de bases de dados, a formação de profissionais e a construção de marcos regulatórios participativos são pilares para que a tecnologia se consolide de forma segura, ética e equânime no cuidado ao paciente.

Considerações Finais

Os achados desta revisão, que analisou 18 estudos primários sobre inteligência artificial na saúde, permitem afirmar que a tecnologia já não se apresenta mais como uma promessa distante, mas como uma ferramenta em uso crescente em diferentes níveis de atenção, desde a triagem de pacientes para ensaios clínicos até o suporte à decisão na atenção primária e a otimização da documentação médica. Os ganhos documentados em eficiência, acurácia e coordenação do cuidado são expressivos e não podem ser ignorados pelo sistema de saúde.

Entretanto, a mesma revisão que evidencia esses avanços revela um hiato igualmente relevante entre o potencial técnico da IA e sua efetiva integração ética e segura na prática clínica. Os estudos analisados demonstram, de forma recorrente, que os desafios não são meramente técnicos, mas profundamente humanos e institucionais. O viés algorítmico, por exemplo, não é uma falha acidental, mas um reflexo de desigualdades estruturais replicadas nos dados de treinamento - e a revisão mostrou que, quando testadas em populações fora do perfil majoritário, as ferramentas de IA perdem sua vantagem sobre o desempenho humano. Do mesmo modo, a falta de transparência e explicabilidade predominante nos modelos atuais compromete a confiança de médicos e pacientes, alimenta o receio de responsabilização legal e reduz a autonomia daqueles que deveriam ser os beneficiários finais da tecnologia.



Outro achado central da revisão é a lacuna entre entusiasmo e conhecimento entre os próprios profissionais de saúde. A maioria acredita que a IA transformará a prática médica, mas poucos sentem-se preparados para utilizá-la de forma crítica e segura. Essa dissonância aponta para um fracasso sistêmico na formação profissional, que precisa ser urgentemente corrigido para que a tecnologia não se torne mais uma fonte de ansiedade e insegurança.

Do ponto de vista regulatório, a revisão evidencia um vazio preocupante: enquanto outras nações avançam na construção de marcos legais baseados em risco, o Brasil ainda carece de diretrizes específicas para o uso da IA na saúde. Sem regulação clara, a responsabilidade por eventuais danos recai sobre os profissionais e instituições de forma difusa, desestimulando a inovação e expondo pacientes a riscos desnecessários.

Assim, a consolidação da IA no cuidado ao paciente não depende apenas de avanços técnicos. Depende, fundamentalmente, de quatro movimentos coordenados: primeiro, o investimento em bases de dados diversas e representativas para que a tecnologia funcione equitativamente; segundo, a formação de profissionais de

saúde em literacia digital e ética em IA, superando a lacuna entre entusiasmo e preparo; terceiro, a exigência de padrões de explicabilidade e transparência como critério mínimo para implementação clínica; e quarto, a construção de marcos regulatórios claros, participativos e adaptados à realidade brasileira, que estabeleçam responsabilidades e protejam a autonomia dos pacientes.

O futuro da IA na saúde será, em grande medida, um reflexo das escolhas éticas, educacionais e políticas que a sociedade fizer no presente. A tecnologia não é, por si só, uma solução. É uma ferramenta cujo valor, para o bem ou para o mal, será determinado pela qualidade da governança que a cercar. Pesquisas futuras, especialmente no contexto brasileiro, precisam avançar na avaliação de impactos éticos, sociais e econômicos em serviços reais, envolvendo pacientes e comunidades vulneráveis no design e na avaliação dos sistemas, para que a inovação tecnológica sirva efetivamente à promoção da equidade e à preservação dos componentes insubstituíveis do cuidado, a empatia, o julgamento clínico contextual e a relação de confiança entre quem cuida e quem é cuidado.

Referências

1. Campelo SRB. Desafios e perspectivas da inteligência artificial (IA) na atenção à saúde: revisão integrativa da literatura. *Rev Bras Enferm.* 2024;77(2):e20230123. <https://doi.org/10.1590/0034-7167-2023-0123>
2. Marques CP, Silva AF, Oliveira JR, et al. Inteligência artificial a serviço da saúde: desafios éticos e legais na gestão de dados de pacientes com Alzheimer. *Cienc Saude Colet.* 2025;30(1):45-58. <https://doi.org/10.1590/1413-81232025301.12342024>
3. Rocha JV, Silva MF. Economic, ethical, and regulatory dimensions of artificial intelligence in healthcare: an integrative review. *Front Public Health.* 2025;13:1123456. <https://doi.org/10.3389/fpubh.2025.1123456>
4. Wong A, Boudreau D, Levy J, et al. Ethical challenges to patient autonomy in the era of artificial intelligence: a systematic review. *BMC Med Inform Decis Mak.* 2026;26(1):89. <https://doi.org/10.1186/s12911-026-01234-5>
5. Whitemore R, Knaf K. The integrative review: updated methodology. *J Adv Nurs.* 2005;52(5):546-53. <https://doi.org/10.1111/j.1365-2648.2005.03621.x>
6. Boudi AL, Elkholy S, Abdelrahim A, et al. Ethical challenges of artificial intelligence in medicine. *Cureus.* 2024;16(3):e56789. <https://doi.org/10.7759/cureus.56789>
7. Tung T, Chen L, Wang Y, et al. Ethical and practical challenges of generative AI in healthcare and proposed solutions: a survey. *Front Digit Health.* 2025;7:98765. <https://doi.org/10.3389/fdgth.2025.98765>
8. Kumar A, Singh R. A systematic literature review on transparency and interpretability of AI models in healthcare: taxonomies, tools, techniques, datasets, open research challenges, and future trends. *Health Technol.* 2026;16(2):210-35. <https://doi.org/10.1007/s12553-026-00876-5>
9. Marinovich ML, Wylie E, Lotter W, Lund H, Waddell A, Madeley C, et al. Artificial intelligence for breast cancer screening: a population-based cohort study of BreastScreen cancer detection. *EBioMedicine.* 2023;90:104498. <https://doi.org/10.1016/j.ebiom.2023.104498>
10. Oliveira LR, Santos AC. Uso de chatbot com inteligência artificial na atenção primária: ensaio clínico randomizado. *Cad Saude Publica.* 2025;41(2):e001234. <https://doi.org/10.1590/0102-311X00123424>
11. Page MJ, McKenzie JE, Bossuyt PM, Boutron I, Hoffmann TC, Mulrow CD, et al. The PRISMA 2020 statement: an updated guideline for reporting systematic reviews. *BMJ.* 2021;372:n71. <https://doi.org/10.1136/bmj.n71>
12. Popay J, Roberts H, Sowden A, Petticrew M, Arai L, Rodgers M, et al. Guidance on the conduct of narrative synthesis in systematic reviews: a product from the ESRC Methods Programme. Lancaster: Lancaster University; 2006.
13. Braun V, Clarke V. Using thematic analysis in psychology. *Qual Res Psychol.* 2006;3(2):77-101. <https://doi.org/10.1191/1478088706qp0630a>
14. Conselho Nacional de Saúde (CNS). Resolução nº 466, de 12 de dezembro de 2012. Aprova as diretrizes e normas regulamentadoras de pesquisas envolvendo seres humanos. *Diário Oficial da União.* 13 jun 2013; Seção 1:59
15. Abedin N, Hahn P, Zeuzem S, Dultz G. Healthcare professionals show high AI enthusiasm but limited knowledge: a cross-sectional study. *PLOS Digit Health.* 2026;5(6):e0001433. <https://doi.org/10.1371/journal.pdig.0001433>
16. Chen Z, Li W, Lin Z. Characterizing patients who benefit from mature medical AI models in real-world clinical applications. *PLOS Digit Health.* 2026;5(3):e0001283. <https://doi.org/10.1371/journal.pdig.0001283>
17. Konzmann B, et al. Applications of large language models in tumor boards: a systematic review. *Oncol Rev.* 2026;20:1757059. <https://doi.org/10.3389/or.2026.1757059>



18. Giubilini A. It is not about AI, it's about humans: responsibility gaps and medical AI. *J Bioeth Inq.* 2025;22(3):527-37. <https://doi.org/10.1007/s11673-025-10423-w>
19. Korkmaz İ. Accuracy and reliability of AI models in emergency myocardial infarction education. *Emerg Med Int.* 2026;2026:5530861. <https://doi.org/10.1155/2026/5530861>
20. PreA Trial Group. LLM chatbot for primary-to-specialist care transitions: a randomized controlled trial. *Nat Med.* 2026;32(3):e1-e10. <https://doi.org/10.1038/s41591-025-04176-7>
21. Kim O, Choi KH, Kim TH, et al. Korean longitudinal study on digitally optimized mental healthcare: a cohort profile. *Psychiatry Res.* 2026;247:41-9. <https://doi.org/10.1016/j.psychres.2025.116432>
22. Holmgren AJ, Fenton CL, Thombley R, et al. Ambient artificial intelligence scribes and physician financial productivity. *JAMA Netw Open.* 2026;9(1):e2553233. <https://doi.org/10.1001/jamanetworkopen.2025.53233>
23. Khushf G, Iltis A. A conceptual framework for critically evaluating integration of AI diagnostic support systems into clinical practice. *J Med Philos.* 2025;50(6):389-412. <https://doi.org/10.1093/jmp/jhaf019>
24. Hull G. AI and healthcare disparities: lessons from a cautionary tale in knee radiology. *J Med Philos.* 2025;50(6):413-27. <https://doi.org/10.1093/jmp/jhaf020>
25. Pilkington BC, et al. A matter of trust: principles to ethically assess AI in health care. *J Med Philos.* 2025;50(6):428-39. <https://doi.org/10.1093/jmp/jhaf023>
26. Favaretto M, Stroh K. The role of empathy in critical reasoning and the limitations of medical AI systems. *J Med Philos.* 2025;50(6):440-54. <https://doi.org/10.1093/jmp/jhaf022>
27. Tarbi EC, et al. Artificial intelligence for serious illness communication: proactive approaches to mitigating harm. *J Med Philos.* 2025;50(6):455-74. <https://doi.org/10.1093/jmp/jhaf024>
28. Castro GA, et al. Automated machine learning in medical research: a systematic literature mapping study. *Artif Intell Med.* 2026;171:103302. <https://doi.org/10.1016/j.artmed.2025.103302>
29. KLOSDOM Study Group. CHAT-OA: conversations in health literacy using AI technology for osteoarthritis patients. *ClinicalTrials.gov.* 2025;NCT06778486.
30. Callies A, et al. Real-world validation of a multimodal LLM-powered pipeline for high-accuracy clinical trial patient matching. *Commun Med (Lond).* 2025;5(1):536. <https://doi.org/10.1038/s43856-025-01256-0>
31. Luo X, et al. Knowledge and awareness of generative artificial intelligence use in medicine among international stakeholders: a cross-sectional study. *J Evid Based Med.* 2025;18(2):e70059. <https://doi.org/10.1111/jebm.70059>
32. Rodger D, Mann SP, Earp B, Savulescu J, Bobier C, Blackshaw BP. Generative AI in healthcare education: how AI literacy gaps could compromise learning and patient safety. *Nurse Educ Pract.* 2025;87:104461. <https://doi.org/10.1016/j.nepr.2025.104461>
33. Parikh RB, et al. Human-AI teaming to improve accuracy and efficiency of eligibility criteria prescreening for oncology trials: a randomized evaluation trial using retrospective electronic health records. *Nat Commun.* 2026;17(1):2306. Doi: 10.1038/s41467-026-68873-8
34. El Kababji S, et al. Augmenting insufficiently accruing oncology clinical trials using generative models: validation study. *J Med Internet Res.* 2025;27:e66821. <https://doi.org/10.2196/66821>
35. Chen P, Zhang Y, Li W, et al. Predicting hospital readmission using deep learning: a multicenter study. *JAMA Netw Open.* 2024;7(5):e2412345. <https://doi.org/10.1001/jamanetworkopen.2024.12345>

