



Revista Bioética

ISSN: 1983-8042

ISSN: 1983-8034

Conselho Federal de Medicina

Naves, Érica Antunes

Bioética e inteligência artificial: panorama atual da literatura

Revista Bioética, vol. 32, e3552PT, 2024

Conselho Federal de Medicina

DOI: <https://doi.org/10.1590/1983-803420243552PT>

Disponível em: <https://www.redalyc.org/articulo.oa?id=361577873006>

- ▶ [Cómo citar el artículo](#)
- ▶ [Número completo](#)
- ▶ [Más información del artículo](#)
- ▶ [Página de la revista en redalyc.org](#)

UAEH [redalyc.org](https://www.redalyc.org)

Sistema de Información Científica Redalyc

Red de Revistas Científicas de América Latina y el Caribe, España y Portugal

Proyecto académico sin fines de lucro, desarrollado bajo la iniciativa de acceso abierto

Bioética e inteligencia artificial: panorama actual de la literatura

Érica Antunes Naves

UnitedHealth Group Brasil, São Paulo/SP, Brasil.

Resumen

El término inteligencia artificial se refiere a sistemas informáticos capaces de realizar procesos intelectuales característicos de los seres humanos, como razonar, descubrir significados, generalizar o aprender de las experiencias. La actuación de la inteligencia artificial se produce cuando los programas informáticos realizan acciones para las cuales no fueron explícitamente programados. Aunque el concepto está bien definido, la actuación de esta tecnología es muy compleja, por lo que la bioética se encuentra ante diversos conflictos y cuestiones relacionadas con ella, que muchas veces solo pueden aclararse cuando surgen. Aunque a lo largo de su desarrollo se vienen estableciendo normativas, este campo sufre constantes adaptaciones, lo que justifica la realización de nuevos estudios sobre el tema.

Palabras clave: Bioética. Inteligencia artificial. Control de la tecnología biomédica.

Resumo

Bioética e inteligência artificial: panorama atual da literatura

O termo inteligência artificial refere-se à realização, por dispositivos computacionais, de processos intelectuais característicos dos seres humanos, como raciocinar, descobrir significados, generalizar ou aprender com experiências. A atuação da inteligência artificial ocorre quando programas computacionais realizam ações para as quais não foram explicitamente programados. Apesar de o conceito ser bem definido, o desempenho dessa tecnologia é muito complexo, de modo que a bioética encontra diversos conflitos e questões relacionadas a ela, muitas vezes esclarecidas apenas no momento em que surgem. Embora regulamentações tenham sido instituídas ao longo do desenvolvimento da área, ela constantemente passa por adaptações, o que justifica novos estudos sobre o tema.

Palavras-chave: Bioética. Inteligência artificial. Controle da tecnologia biomédica.

Abstract

Bioethics and artificial intelligence: a current overview of the literature

Artificial intelligence refers to the performance, by computer devices, of intellectual processes characteristic of human beings, such as reasoning, discovering meanings, generalizing or learning from experience. Artificial intelligence occurs when computer programs perform action for which they were not explicitly programmed. Although a well-defined concept, its complex performance poses various bioethical conflicts and questions, often clarified only when they emerge. Despite the regulations put in place during the field's development, these are constantly undergoing adaptations thus justifying further studies on the subject.

Keywords: Bioethics. Artificial intelligence. Technology control, biomedical.

La autora declara que no existe ningún conflicto de interés.

El término inteligencia artificial (IA) se refiere a la ejecución, mediante dispositivos computacionales, de procesos intelectuales característicos del ser humano, como razonar, descubrir significados, generalizar o aprender de las experiencias¹. Su actuación ocurre cuando los programas computacionales realizan acciones para las cuales no fueron programados explícitamente, y puede describirse como el uso de estos dispositivos para realizar tareas que anteriormente requerían la cognición humana².

La Asociación Médica Estadounidense prefiere el término “inteligencia aumentada” en lugar de “inteligencia artificial”, para dar énfasis al papel de asistencia de las computadoras para mejorar las habilidades de los médicos en lugar de reemplazarlas³. La integración de la IA en la práctica médica ha aumentado significativamente en los últimos años, y sigue creciendo². Por ello, la discusión acerca de los principios éticos y legales en este campo, incluido el desarrollo de regulaciones específicas, es una preocupación constante y de suma relevancia.

Método

Se realizó una búsqueda en la base de datos PubMed utilizando los descriptores “*bioethics*” y “*artificial intelligence*”, considerando artículos científicos publicados entre el 2018 y el 2022. El análisis también incluyó opiniones de las sociedades de referencia en el área investigada. Posteriormente, se realizó una comparación crítica con criterios que variaron según el tipo de estudio, discusión del tema y conclusiones.

Discusión

Los trabajos analizados definen invariablemente el concepto de IA, con el entendimiento común de que se trata de una tecnología que imita procesos intelectuales característicos del ser humano, para alcanzar objetivos sin la programación de una acción específica¹⁻⁴.

Existe un amplio acuerdo según el cual demuestra inteligencia cualquier dispositivo que utilice la razón, elabore estrategias, resuelva acertijos, emita juicios en condiciones de incertidumbre,

presente conocimiento (incluido el sentido común), planifique, aprenda, se comunique en lenguaje natural e integre todas estas habilidades con objetivos comunes¹.

Actualmente, entre los ejemplos de áreas en que actúan las IA figuran actividades como comprender el habla humana, competir al más alto nivel en sistemas de juegos estratégicos, conducir automóviles autónomos, planificar enrutamiento inteligente en una red de entrega de contenido y simulaciones militares¹.

En el ámbito médico, Bali y colaboradores¹ presentan ejemplos como el IBM Watson for Oncology, que prescribía fármacos para tratar a pacientes con cáncer con una eficacia igual o superior a la de especialistas humanos. Además, el proyecto Hanover, desarrollado por Microsoft en Oregón, Estados Unidos, analizó investigaciones médicas para permitir el tratamiento personalizado del cáncer.

Asimismo, el Servicio Nacional de Salud (NHS) del Reino Unido utiliza la plataforma DeepMind, de Google, para detectar riesgos para la salud mediante el análisis de datos de aplicaciones móviles e imágenes médicas recogidas de pacientes. Otro ejemplo es el algoritmo de radiología de Stanford, que detectaba la neumonía mejor que los radiólogos humanos, mientras que en la retinopatía diabética fue equivalente a los oftalmólogos especialistas a la hora de tomar decisiones de derivación¹.

Vearrier y colaboradores² están de acuerdo en que la IA ha demostrado un beneficio potencial sustancial para médicos y pacientes, transformando la relación terapéutica de la tradicional díada médico-paciente en una relación tríada médico-paciente-máquina. Sin embargo, consideran que las nuevas tecnologías de IA requieren una verificación cuidadosa, normativa legal, salvaguardias del paciente y educación del proveedor, y los médicos deben reconocer los límites y riesgos de la IA, así como sus posibles beneficios.

Con tantos algoritmos de IA en uso y en desarrollo, en todas las áreas y especialmente en el área médica, la discusión de aspectos éticos y legales, incluida la regulación, es una preocupación creciente. En el pasado, estos cambios regulatorios se basaban en conceptos filosóficos y en la necesidad de responder a las demandas socioculturales.

Actualmente, el desafío es integrar una serie de tecnologías disruptivas, como la IA.

Si bien aporta numerosas ventajas, al suponer una realidad emergente y de resultados impredecibles, que no están exentos de riesgos, este campo genera incertidumbre y precaución, por lo que requiere una regulación bioética⁴. Sánchez López y colaboradores⁴ proponen tres principios básicos para regular a las partes interesadas: 1) respeto a la investigación; 2) justicia; y 3) transparencia.

A pesar de su promesa de resultados, los algoritmos de IA están bajo investigación por su rendimiento inconsistente, particularmente en comunidades minoritarias, lo que puede conducir a decisiones clínicas inferiores a lo ideal y a resultados adversos para el paciente. Hasta la fecha, estas preocupaciones se han concentrado en la negligencia médica y señalan que la responsabilidad de los médicos por el uso de la IA está inextricablemente relacionada con la de estos otros actores⁵.

Maliha y colaboradores⁵, al estudiar a fondo la labilidad de esta cuestión, consideran que la asignación de responsabilidad determina si los pacientes obtienen compensación, y de quién. Además, amplían la evaluación, al cuestionar si se pondrán en práctica algoritmos potencialmente útiles, ya que una mayor responsabilidad por el uso o desarrollo de algoritmos puede desalentar a los desarrolladores y líderes de los sistemas de salud a la hora de introducirlos en la práctica clínica. Destacan que se debe examinar en detalle el ecosistema más amplio de responsabilidad de la IA y su papel para garantizar la ejecución segura y la innovación en la atención clínica.

Sher, Sharp y Wright⁶, en un editorial, son optimistas en cuanto a las perspectivas que la IA reserva para los cuidados de salud en el futuro y consideran la importancia de comprender las fortalezas, limitaciones, oportunidades, desafíos éticos y riesgos. Igualmente importante es la necesidad de familiarizarse con las herramientas enumeradas para analizar críticamente las aplicaciones de IA en los cuidados de salud, con el fin de diferenciar las mejoras concretas de los errores.

Lazarus y colaboradores⁷ amplían el análisis de las limitaciones de la IA para la educación en salud, específicamente de anatomía, describiendo las tensiones entre las promesas y los peligros de integrar esta tecnología en este ámbito.

Finalmente, los autores brindan recomendaciones prácticas para un enfoque ponderado al trabajar en conjunto con la IA, que sirven como marco rector destinado a desarrollar un enfoque más sutil y equilibrado para el papel de la IA en la educación en salud.

Los aspectos científicos y éticos de esta cuestión justificaron la celebración de una reunión para discutir el tema, en marzo del 2017, en Barcelona, en la que participaron diferentes especialistas europeos en AI, computación y comunicación, entre otras áreas. El debate dio lugar a *Declaración de Barcelona para el desarrollo y uso adecuado de la inteligencia artificial en Europa*⁸, que contiene los siguientes principios y valores: prudencia, confiabilidad, responsabilidad, autonomía restringida y papel humano, que se detallan a continuación.

- Prudencia: El salto adelante en esta área ha sido causado por la maduración de las tecnologías de IA, el aumento de la potencia computacional y de la capacidad de almacenamiento de datos, la disponibilidad de plataformas de entrega en internet y la mayor disposición de muchos actores económicos para experimentar la tecnología por su propia cuenta;
- Confiabilidad: Todos los sistemas artificiales utilizados en la sociedad deben someterse a pruebas para determinar su confiabilidad y seguridad;
- *Accountability*: Cuando un sistema de IA toma una decisión, las personas afectadas por dicha decisión deben poder obtener una explicación, porque la decisión se toma en términos que puedan entender, y debe ser posible impugnar la decisión con argumentos fundamentados;
- Responsabilidad: existe una creciente preocupación por los *chatbots* de IA y otros tipos de sistemas de mensajería automática que operan en internet y en las redes sociales y están diseñados para la manipulación de la opinión política, la desinformación mediante la propagación de hechos falsos, la extorsión u otras formas de actividad maliciosa, que son peligrosas para los individuos y desestabilizan la sociedad;
- Autonomía restringida: los sistemas de IA no tienen únicamente la capacidad de tomar decisiones. Cuando están integrados en sistemas físicos, como automóviles autónomos, tienen el potencial de actuar según sus decisiones en

el mundo real, lo que plantea dudas sobre la seguridad y la posibilidad de que la IA autónoma supere la inteligencia humana en algún momento; y

- Papel humano: El innegable entusiasmo actual por la IA a veces da la impresión de que la inteligencia humana ya no será necesaria, un hecho que ha llevado a algunas empresas a despedir a empleados y reemplazarlos por sistemas de IA. Esto es un grave error, ya que todos los sistemas de IA dependen críticamente de la inteligencia humana.

Tai⁹ recuerda que, a pesar de todas las promesas positivas que ofrece la IA, los especialistas humanos siguen siendo esenciales y necesarios para diseñar, programar y operar la IA con el fin de evitar cualquier error impredecible. En ese sentido, cita a Beth Kindig, analista de tecnología de San Francisco con más de una década de experiencia en empresas de tecnología públicas y privadas.

Kindig publicó un boletín informativo gratuito en el que indica que, si bien la IA es potencialmente prometedora para un mejor diagnóstico médico, aún se necesitan especialistas humanos para resolver impases y mitigar errores. Por lo tanto, no se puede descuidar la vigilancia de la

función de la IA, realizada por el profesional conocido como médico en el circuito⁹.

Consideraciones finales

Como todas las tecnologías emergentes, la IA y la robótica requieren la implementación previa de normas éticas de funcionamiento que garanticen la seguridad y la aplicabilidad en un campo tan sensible como el de la salud³. Los sistemas de IA tienen el potencial de transformar radicalmente la atención clínica e, incluso si avanzan a un ritmo más lento, el ordenamiento jurídico no puede permanecer estático con relación a esta innovación.

Para aprovechar plenamente los beneficios de la IA, el sistema jurídico debe equilibrar la responsabilidad para promover la innovación, la seguridad y la adopción acelerada de estos algoritmos⁵. El estado relativamente inestable de la IA y su potencial responsabilidad brindan la oportunidad de desarrollar un nuevo modelo que se adapte al progreso médico e instruya a las partes interesadas sobre la mejor manera de responder a esta innovación disruptiva.

Referencias

1. Bali J, Garg R, Bali TR. Artificial intelligence (AI) in healthcare and biomedical research: why a strong computational/AI bioethics framework is required? What is artificial intelligence? Indian J Ophthalmol [Internet]. 2018 [acceso 29 nov 2023];67(1):9-12. DOI: 10.4103/ijo.IJO_1292_18
2. Vearrier L, Derse AR, Basford JB, Larkin GL, Moskop JC. Artificial Intelligence in emergency medicine: benefits, risks, and recommendations. J Emerg Med [Internet]. 2022 [acceso 29 nov 2023];62(4):492-9. DOI: 10.1016/j.jemermed.2022.01.001
3. Crigger E, Khoury C. Making policy on augmented intelligence in health care. AMA J Ethics [Internet]. 2019 [acceso 29 nov 2023];21(2):E188-91. DOI: 10.1001/amajethics.2019.188.
4. Sánchez López JD, Cambil Martín J, Villegas Calvo M, Luque Martínez F. Artificial intelligence and robotics: reflections about the need of a new bioethics framework implementation. J Healthc Qual Res [Internet]. 2021 [acceso 29 nov 2023];36(2):113-4. DOI: 10.1016/j.jhqr.2019.07.009
5. Maliha G, Gerke S, Cohen IG, Parikh RB. Artificial intelligence and liability in medicine: balancing safety and innovation. Milbank Q [Internet]. 2021 [acceso 29 nov 2023];99(3):629-47. DOI: 10.1111/1468-0009.12504
6. Sher T, Sharp R, Wright RS. Algorithms and bioethics. Mayo Clin Proc [Internet]. 2020 [acceso 29 nov 2023];95(5):843-4. DOI: 10.1016/j.mayocp.2020.03.020
7. Lazarus MD, Truong M, Douglas P, Selwyn N. Artificial intelligence and clinical anatomical education: promises and perils. Anat Sci Educ [Internet]. 2022 [acceso 29 nov 2023]. DOI: 10.1002/ase.2221

8. Barcelona Declaration for the Proper Development and Usage of Artificial Intelligence in Europe. [Propõe diretrizes para um “código de conduta” dos profissionais de IA atualmente em discussão]. Biocat [Internet]. 2017 [acesso 29 nov 2023]. Disponível: <https://bit.ly/3w6SYCv>
9. Tai MCT. The impact of artificial intelligence on human society and bioethics. Tzu Chi Med J [Internet]. 2020 [acesso 29 nov 2023];32(4):339-43. DOI: 10.4103/tcmj.tcmj_71_20

Érica Antunes Naves – Especialista – eanaves@hotmail.com

 0000-0003-3257-4941

Correspondencia

Rua Conselheiro Brotero, 1486, Santa Cecília CEP 01232-010. São Paulo/SP, Brasil.

Recibido: 18.4.2023

Revisado: 30.11.2023

Aprobado: 5.1.2024