

Inteligencia artificial en la investigación científica: desafíos éticos y regulatorios en el contexto mexicano

Artificial intelligence in scientific research: ethical and regulatory challenges in the mexican context

Cortés-Moreno GY^{1*} 

1. Directora de Investigación, Comisión Nacional de Arbitraje Médico, Ciudad de México, México.

RESUMEN

La adopción acelerada de la inteligencia artificial (IA) ha transformado los procesos de análisis, modelación y producción de conocimiento en la investigación científica, particularmente en el ámbito de la salud. Aunque estas tecnologías ofrecen mejoras sustanciales en eficiencia, precisión y capacidad predictiva, su desarrollo ha superado la evolución de los marcos éticos y regulatorios, especialmente en países de ingresos medios como México. El objetivo de este artículo es examinar los principales desafíos éticos y regulatorios asociados al uso de IA en investigación científica, con énfasis en el contexto mexicano y en áreas prioritarias como salud pública, análisis epidemiológico, microbioma, diagnóstico asistido y gestión de riesgos clínicos.

Se realizó una revisión narrativa de la literatura científica, de documentos regulatorios nacionales e internacionales, así como de estudios recientes sobre la adopción de la inteligencia artificial en laboratorios y ensayos clínicos. Los hallazgos muestran avances importantes, pero también evidencian brechas en capacitación, gobernanza de datos, representatividad algorítmica, explicabilidad, reproducibilidad y responsabilidad jurídica. Persisten riesgos éticos derivados de sesgos en los datos, la opacidad de modelos complejos, las limitaciones del consentimiento informado y la fabricación algorítmica. En el ámbito regulatorio, México carece de estándares específicos para la validación, certificación, auditorías de equidad y supervisión del ciclo de vida de los modelos, en contraste con marcos como el Reglamento de Inteligencia Artificial de la Unión Europea, el de la Organización Mundial de la Salud (OMS), el de la Organización para la Cooperación y el Desarrollo Económicos (OCDE) y el de la Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura (por sus siglas en inglés, UNESCO).

Se concluye que México enfrenta la necesidad urgente de construir un marco integral de gobernanza algorítmica que garantice un uso ético, seguro y transparente de la IA en investigación científica. El fortalecimiento institucional y la formación especializada son elementos clave para avanzar hacia una adopción responsable y alineada con estándares internacionales.

Palabras clave: Inteligencia Artificial; ética en investigación; gobernanza algorítmica; regulación de IA; sesgos algorítmicos.

Autor(a) de Correspondencia:

Cortés-Moreno GY.
Directora de Investigación en la Comisión Nacional de Arbitraje Médico.
Dirección: Av. Marina Nacional No.60, Piso 14, Alcaldía Miguel Hidalgo, Col. Tacuba, Ciudad de México. C.P. 11410.
Correo electrónico: gabriela.cortesm@conamed.gpb.mx

Citar como:

Cortés-Moreno GY.
Inteligencia Artificial en la Investigación Científica: Desafíos Éticos y Regulatorios en el Contexto Mexicano.
Rev CONAMED. 2025;30 (Supl. 1): 274-282.

Fecha de recepción:

18 de noviembre 2025

Fecha de aceptación:

04 de diciembre 2025

ABSTRACT

The accelerated adoption of artificial intelligence (AI) has transformed processes of analysis, modelling and knowledge production within scientific research, particularly in the health sciences. Although these technologies offer substantial improvements in efficiency, accuracy and predictive capacity, their development has outpaced the evolution of ethical and regulatory frameworks, especially in middle-income countries such as Mexico. This article aims to examine the main ethical and regulatory challenges associated with the use of AI in scientific research, with emphasis on the Mexican context and on priority areas including public health, epidemiological analysis, microbiome research, assisted diagnosis and clinical risk management.

A narrative review of the scientific literature, as well as national and international regulatory documents and recent studies on the adoption of artificial intelligence in laboratories and clinical trials, was conducted. The findings show significant progress but also reveal gaps in training, data governance, algorithmic representativeness, explainability, reproducibility, and legal accountability. Ethical risks persist, stemming from data biases, the opacity of complex models, limitations in informed consent, and algorithmic fabrication. From a regulatory perspective, Mexico lacks specific standards for model validation, certification, equity audits, and life-cycle oversight, in contrast to frameworks such as the European Union Artificial Intelligence Act, those of the World Health Organization (WHO), the Organisation for Economic Co-operation and Development (OECD), and the United Nations Educational, Scientific and Cultural Organization (UNESCO).

It is concluded that Mexico faces an urgent need to establish a comprehensive framework for algorithmic governance that ensures the ethical, safe and transparent use of AI in scientific research. Institutional strengthening and specialised training are essential to advance towards responsible adoption aligned with international standards.

Keywords: Artificial intelligence; research ethics; algorithmic governance; AI regulation; algorithmic bias.

INTRODUCCIÓN

La adopción acelerada de la inteligencia artificial (IA) en la investigación científica ha transformado profundamente los métodos tradicionales de generación y análisis de datos. Desde la automatización de procesos experimentales hasta la modelación predictiva y el descubrimiento asistido por algoritmos, la IA se ha consolidado como un componente esencial en múltiples disciplinas, particularmente en las ciencias de la salud. Sin embargo, este avance ha ocurrido a un ritmo que supera la capacidad de regulación, supervisión y evaluación ética en diversos sistemas de ciencia y tecnología.^{1,2}

En México, el desarrollo normativo relacionado con la IA en investigación científica permanece en una fase inicial. Aunque existe regulación en materia de protección de datos personales, bioética e investigación en seres humanos, aún no se cuenta con un marco específico que establezca estándares para el uso de algoritmos, modelos predictivos o sistemas autónomos en investigación.³ Esta ausencia de regulación genera incertidumbre jurídica y expone a las personas participantes a riesgos no previstos adecuadamente.

La IA también plantea desafíos críticos como transparencia algorítmica, gobernanza de datos, adecuación del consentimiento informado, sesgos y discriminación, explicabilidad, reproducibilidad científica, seguridad y responsabilidad jurídica. Estos elementos inciden directamente en la integridad de la investigación y en la protección de derechos de las personas involucradas.⁴

De manera paralela, organismos internacionales han avanzado en directrices éticas y regulatorias sobre IA. Entre ellos destacan la UNESCO, la Organización Mundial de la Salud (OMS), la Organización para la Cooperación y el Desarrollo Económicos (OCDE) y la Unión Europea, que han establecido principios orientadores para el uso responsable de estas tecnologías en ciencia y salud⁵⁻⁷. Analizar estos marcos resulta fundamental para identificar oportunidades de fortalecimiento del contexto normativo nacional.

Por ello, este artículo tiene como propósito examinar los principales desafíos éticos y regulatorios asociados al uso de IA en la investigación científica, con énfasis en el ámbito de

la salud y en el contexto mexicano. A través de una revisión narrativa, se identifican brechas, riesgos y oportunidades que permitan orientar la construcción de un marco de gobernanza robusto, responsable y alineado con estándares internacionales.

Antecedentes y Panorama

La inteligencia artificial (IA) tiene sus raíces en los desarrollos teóricos de mediados del siglo XX. Los trabajos de Alan Turing plantearon las primeras interrogantes sobre la posibilidad de que una máquina pudiera razonar, sentando así las bases del pensamiento computacional.⁸ Posteriormente, las conferencias de Dartmouth en 1956 consolidaron el término "artificial intelligence" (AI) como un campo de estudio formal, integrando matemáticas, lógica, estadística y ciencias de la computación.⁹

Durante décadas, los avances en IA estuvieron limitados por el poder computacional y la disponibilidad de datos. Sin embargo, a partir de la década de 2010, la combinación de big data, poder de cómputo acelerado y la evolución del aprendizaje profundo (deep learning) impulsaron un cambio radical en su capacidad de procesamiento y en su potencial para la investigación científica.^{10,11}

En el ámbito biomédico, la IA ha demostrado ser una herramienta potente para el análisis de imágenes médicas, la predicción de riesgos, el descubrimiento de fármacos, el análisis genómico y la optimización de sistemas de salud. Estudios recientes han mostrado su capacidad para igualar o incluso superar el desempeño clínico humano en tareas específicas, especialmente en radiología, dermatología y patología digital.¹²⁻¹³

A nivel global, la adopción de IA en la investigación científica ha sido más acelerada en países que cuentan con infraestructura robusta, amplia disponibilidad de datos y marcos regulatorios desarrollados, como Estados Unidos, Canadá, Reino Unido y la Unión Europea.¹⁴ De manera paralela,

organismos internacionales como la Organización Mundial de la Salud (OMS) y la UNESCO han publicado lineamientos éticos y de gobernanza para orientar su uso responsable en salud y ciencia.^{5,6}

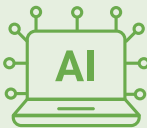
En contraste, América Latina presenta un desarrollo desigual y fragmentado. En México, aunque existen avances significativos en campos como salud pública, análisis epidemiológico, genómica, diagnóstico asistido y automatización de procesos,¹⁵ estas iniciativas no se encuentran plenamente reguladas ni estandarizadas. La falta de infraestructura homogénea, la limitada disponibilidad de datos interoperables y la ausencia de políticas nacionales de IA dificultan su adopción sistemática en investigación científica.

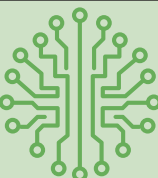
La aparición reciente del Centro Público de Formación en Inteligencia Artificial, impulsado por el Gobierno de México, representa un hito para fortalecer las capacidades nacionales en esta materia y puede contribuir a disminuir las brechas existentes en desarrollo, capacitación y gobernanza algorítmica.¹⁶

Tipos de Inteligencia Artificial y sus Aplicaciones en la Investigación Científica

La inteligencia artificial (IA) comprende un conjunto amplio de técnicas y modelos computacionales diseñados para realizar funciones que tradicionalmente requieren capacidades humanas, como analizar información, aprender de datos, reconocer patrones, tomar decisiones o generar contenido. En el ámbito de la investigación científica, estos sistemas han adquirido un papel central debido a su capacidad para procesar grandes volúmenes de datos con velocidad, precisión y eficiencia.

En la *Tabla 1*, se describen los principales subtipos de IA relevantes para la investigación, junto con sus aplicaciones más frecuentes en ciencia y salud.

Tipo de IA / Subtipo	Descripción breve	Aplicaciones científicas	Referencias
 Aprendizaje Automático	Algoritmos que identifican patrones en datos.	• Predicción de riesgos • Diagnóstico • Segmentación • Genómica • Fenotipado	17. Esteva A, Robicquet A, Ramsundar B, et al. A guide to deep learning in healthcare. Nat Med. 2019;25:24–29. 18. Obermeyer Z, Emanuel EJ. Predicting the future — Big data, machine learning, and clinical medicine. Science. 2016;353(6301):446–52.

 Aprendizaje supervisado	Entrenado con datos etiquetados.	• Diagnóstico asistido • Clasificación • Triage algorítmico	1. Topol EJ. High-performance medicine: the convergence of human and artificial intelligence. Lancet. 2019;393:507–22.
 Aprendizaje no supervisado	Identifica patrones sin etiquetas.	• Clustering de pacientes • Perfiles de riesgo • Genómica	19. Deo RC. Machine learning in medicine. Circ Res. 2015;116:109–26.
 Aprendizaje por refuerzo	Aprende por recompensas/penalizaciones.	• Optimización experimental • Robótica de laboratorio	20. Silver D, Huang A, Maddison CJ, et al. Mastering the game of Go with deep reinforcement learning. Nature. 2016;529:484–9.
 Aprendizaje Profundo	Redes neuronales profundas.	• Diagnóstico por imágenes • Histopatología digital • Predicción clínica • Reconocimiento de voz	21. LeCun Y, Bengio Y, Hinton G. Deep learning. Nature. 2015;521:436–44. 22. Rajpurkar P, Chen E, Banerjee O, Topol EJ. AI in health and medicine. Nat Med. 2022 Jan;28:31–8. doi: 10.1038/s41591-021-01614-0.
 Procesamiento de Lenguaje Natural (NLP)	Interpretación y generación de lenguaje humano.	• Lectura automatizada • Minería de expedientes • Análisis de textos • Asistentes virtuales	23. Devlin J, Chang MW, Lee K, Toutanova K. BERT: Pre-training of deep bidirectional transformers. In: Proceedings of NAACL-HLT; 2019. 24. Bommasani R, Hudson DA, Adeli E, et al. On the opportunities and risks of foundation models. arXiv:2108.07258. 2021.
 Visión Computacional	Análisis de imágenes y video.	• Microscopía digital • Biomarcadores • Rehabilitación • Control de calidad	25. Ronneberger O, Fischer P, Brox T. U-Net: Convolutional Networks for Biomedical Image Segmentation. In: MICCAI; 2015. p. 234–41. 26. Litjens G, Kooi T, Bejnordi BE, et al. A survey of deep learning in medical image analysis. Med Image Anal. 2017;42:60–88.
 Modelos Generativos (IA generativa)	Crean datos nuevos (imágenes/moléculas/texto).	• Imágenes sintéticas • Simulación • Diseño de fármacos • Generación de hipótesis	27. Karras T, Laine S, Aila T. A Style-Based Generator Architecture for Generative Adversarial Networks. In: CVPR; 2019. p. 4401–10. 28. Sanchez-Lengeling B, Aspuru-Guzik A. Inverse molecular design using machine learning: Generative models for matter engineering. Science. 2018;361:360–5.

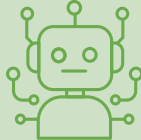
 Sistemas Expertos	Basados en reglas lógicas.	• Apoyo a comités • Decisiones éticas • Seguridad del paciente	29. Shortliffe EH. Artificial intelligence in medicine: Weighing the accomplishments, hype, and promise. Yearb Med Inform. 2019;28(1):257–262. doi:10.1055/s-0039-1677891.
 IA Autónoma y Robótica Inteligente	IA integrada en robots con diferentes niveles de autonomía.	• Laboratorios autónomos • Robots quirúrgicos • Cribado automatizado	30. King RD, Rowland J, Oliver SG, et al. The automation of science. Science. 2009;324(5923):85–89. 31. Yang GZ, Cambias J, Cleary K, et al. Medical robotics—Regulatory, ethical, and legal considerations for increasing levels of autonomy. Sci Robot. 2017;2(4):eaam8638.

Tabla 1. Panorama de las tecnologías de inteligencia artificial aplicadas a la investigación científica en salud.
Nota: Los artículos citados en este cuadro sirven como referentes conceptuales y metodológicos para comprender el uso de distintas técnicas de inteligencia artificial en la investigación científica y en el campo de la salud.

Aplicaciones Transversales de la Inteligencia Artificial en el Contexto Científico Mexicano

En México, la adopción de inteligencia artificial (IA) en la investigación científica y en los servicios de salud avanza mediante iniciativas diversas, aunque aún de forma fragmentada y con niveles heterogéneos de madurez tecnológica. En salud pública, los modelos predictivos se han empleado para anticipar brotes epidemiológicos, estimar cargas de enfermedad y optimizar la asignación de recursos sanitarios, especialmente durante emergencias recientes. Sin embargo, la evidencia más actual revela que la integración de IA en laboratorios clínicos y de investigación sigue siendo limitada: un estudio nacional con 362 profesionales reportó que solo 0.5% se consideran expertos en IA, 4.7% han recibido formación formal y cerca del 60% presenta baja familiaridad con aplicaciones específicas, pese a que 78% reconoce que estas herramientas ya intervienen en procesos como análisis de datos, detección de errores o verificación de resultados.³² Esta disparidad muestra que el potencial de la IA en el país no se ha traducido plenamente en prácticas consolidadas, en parte por la falta de capacitación técnica, la percepción de riesgo laboral y la ausencia de lineamientos normativos claros.

Aun en este escenario, México ha comenzado a emplear la IA de manera transversal en distintos campos científicos, lo cual refleja un ecosistema en transición. En el estudio del

microbiana humano, ambiental y alimentario, algoritmos de “Aprendizaje Automático y Aprendizaje Profundo” permiten caracterizar comunidades microbianas, identificar patrones taxonómicos y explorar asociaciones entre perfiles microbianos y enfermedades crónicas. Estas aproximaciones han fortalecido proyectos de investigación biomédica, nutricional y genómica desarrollados en instituciones académicas del país,³³ impulsando líneas de trabajo que requieren análisis complejos y grandes volúmenes de datos.

De manera paralela, el análisis de riesgos clínicos se ha consolidado como un campo en rápida expansión donde la IA desempeña un papel cada vez más sistemático y sofisticado. Una revisión de 142 estudios publicados entre 2013 y 2024 identificó tres áreas principales en las que los algoritmos se utilizan de forma consistente en ensayos clínicos: seguridad, eficacia y riesgos operativos. En seguridad, los modelos predicen toxicidad, severidad y probabilidad de eventos adversos con desempeños que alcanzan un área bajo la curva ROC de hasta 96%; en eficacia, se emplean métodos clásicos y de aprendizaje profundo para anticipar respuesta terapéutica, desenlaces clínicos y efectos del tratamiento; y en riesgos operativos, los modelos permiten anticipar la probabilidad de éxito por fase, estimar aprobaciones regulatorias, detectar riesgos de terminación anticipada y evaluar la calidad metodológica de los protocolos. Estas aplicaciones se nutren de fuentes tan diversas como “clinicaltrials

gov", bases regulatorias internacionales, modelos sintéticos de pacientes y representaciones basadas en *"Modelos de Lenguaje de Gran Escala"*, lo que ha permitido automatizar tareas de vigilancia, priorización de riesgos y monitoreo continuo.³⁴ En conjunto, estos avances muestran que la IA ya forma parte de un ecosistema más amplio de gestión integral del riesgo clínico y de investigación.

A pesar de estos progresos, persisten barreras significativas que limitan la consolidación de la IA en el país: la ausencia de un marco regulatorio específico, la heterogeneidad de la infraestructura institucional y las diferencias en la calidad y disponibilidad de datos. Estos retos afectan la escalabilidad, estandarización y uso seguro de las herramientas algorítmicas, y subrayan la necesidad de articular esfuerzos interinstitucionales para construir un ecosistema nacional de IA robusto, transparente y alineado con estándares internacionales.

Desafíos Éticos

El uso de inteligencia artificial (IA) en la investigación científica ofrece beneficios sustanciales, pero también plantea dilemas éticos complejos que influyen en la calidad, legitimidad y equidad de los resultados. En contextos como el mexicano (caracterizado por capacidades tecnológicas heterogéneas, disparidades institucionales y ausencia de lineamientos específicos) estos desafíos se

intensifican y requieren una evaluación rigurosa antes, durante y después de la implementación de modelos algorítmicos.³⁵

Entre los riesgos más relevantes destacan la presencia de sesgos algorítmicos derivados de bases de datos no representativas; la opacidad o falta de explicabilidad de modelos complejos, que dificulta la validación científica; los riesgos asociados al manejo de datos sensibles, especialmente cuando se desconocen los alcances del uso secundario; y la persistente insuficiencia del consentimiento informado tradicional para explicar implicaciones técnicas a personas no expertas.³⁶

Asimismo, la integridad científica enfrenta amenazas ante el uso de herramientas generativas capaces de producir datos sintéticos, imágenes modificadas o textos automatizados sin controles adecuados de trazabilidad. A esto se suma la ambigüedad en materia de responsabilidad ética y jurídica cuando decisiones o errores son atribuibles a algoritmos, así como la posibilidad de profundizar inequidades al emplear modelos entrenados con datos sesgados o no representativos.

La *Tabla 2* sintetiza estos riesgos y sus dimensiones principales, subrayando áreas críticas para garantizar una implementación responsable, transparente y alineada con estándares internacionales de ética en investigación.

Desafío ético	Fundamentación	Riesgos principales
Sesgos algorítmicos y falta de representatividad	Los modelos pueden aprender y amplificar sesgos presentes en los datos, especialmente en contextos con bases no representativas.	• Predicciones injustas • Discriminación hacia grupos vulnerables • Resultados no generalizables • Distorsión científica por datos incompletos
Opacidad y falta de explicabilidad	Muchos modelos funcionan como "cajas negras", dificultando comprender cómo generan sus predicciones.	• Imposibilidad de validar científicamente • Problemas de reproducibilidad • Toma de decisiones sin evidencia interpretable • Pérdida de confianza de participantes y comunidades
Privacidad y manejo de datos sensibles	La IA requiere grandes volúmenes de datos clínicos, genómicos o personales sin regulaciones específicas para usos secundarios o movilidad de datos.	• Riesgo de reidentificación • Filtraciones o uso indebido de información • Anonimización insuficiente • Vulneración de derechos de privacidad

Consentimiento informado insuficiente	Los modelos actuales de consentimiento no contemplan riesgos algorítmicos, usos secundarios de datos ni fallas del modelo.	• Consentimientos incompletos o poco comprensibles • Falta de transparencia sobre riesgos técnicos • Dificultad para explicar reutilización de datos
Integridad científica y fabricación algorítmica	La IA generativa puede producir datos falsos, imágenes manipuladas o manuscritos automatizados sin validación.	• Fabricación o manipulación de datos • Falta de trazabilidad • Riesgo de mala práctica científica • Pérdida de rigor y credibilidad
Responsabilidad ética y atribución de decisiones	No existe claridad normativa sobre quién es responsable ante daños derivados de decisiones algorítmicas.	• Vacíos de responsabilidad jurídica • Riesgo institucional • Incertidumbre para CEI y personal investigador
Impacto en la equidad y justicia social	La IA puede amplificar desigualdades existentes cuando algunos grupos quedan subrepresentados o excluidos.	• Brechas entre instituciones • Discriminación hacia grupos rurales o indígenas • Dependencia tecnológica desigual • Resultados sesgados
Autonomía y relación profesional-paciente	La IA puede modificar dinámicas de confianza y desplazar la supervisión humana en investigaciones clínicas.	• Delegación excesiva a sistemas automáticos • Reducción de supervisión profesional • Riesgo de decisiones no validadas • Deterioro de la relación profesional-paciente

Tabla 2. Desafíos éticos de la Inteligencia Artificial en la Investigación Científica.

Desafíos Regulatorios

La integración de IA en la investigación científica avanza con rapidez, mientras que la capacidad regulatoria en México evoluciona de forma más lenta, lo que produce un escenario de vacíos normativos, fragmentación institucional y ausencia de estándares unificados para evaluar, validar y supervisar sistemas algorítmicos. Si bien existe un andamiaje normativo en protección de datos personales, bioética, investigación en salud y dispositivos médicos, estas normas no contemplan los riesgos específicos asociados al uso de IA, tales como trazabilidad algorítmica, gobernanza del ciclo de vida de datos, auditorías de equidad o certificación tecnológica.^{3,15}

Las competencias regulatorias también se encuentran dispersas: el Instituto Nacional de Transparencia, Acceso a la Información y Protección de Datos, (INAI) supervisa datos personales; la Comisión Nacional de Bioética (CONBIOÉTICA) establece lineamientos éticos; La Comisión Federal para la Protección contra Riesgos Sanitarios

(COFEPRIS) regula tecnologías en salud; y la Secretaría de Ciencia, Humanidades, Innovación y Tecnología (SECIHTI) coordina políticas científicas. Sin una entidad articuladora que integre estos esfuerzos, persisten inconsistencias en criterios, mecanismos de supervisión y alcances regulatorios. Otro reto importante radica en los vacíos respecto al uso secundario, movilidad, anonimización avanzada y gobernanza de datos, que son elementos esenciales para entrenar modelos de IA, para los cuales no existe legislación específica. Asimismo, la desigualdad en infraestructura computacional, acceso a datos interoperables y capacidades técnicas limita la posibilidad de implementar modelos de manera segura y estandarizada en todo el país.

México tampoco cuenta aún con estándares nacionales obligatorios para validación, certificación o explicabilidad de modelos, en contraste con los marcos desarrollados por la Unión Europea (AI Act), la OMS o la OCDE. Esta falta de alineación dificulta la comparabilidad internacional, limita la adopción segura y reduce la competitividad científica del país.^{6,7,14}

En conjunto, estos elementos muestran la urgencia de avanzar hacia un marco regulatorio integral, coherente y adaptado al contexto nacional, que permita un uso responsable, seguro y éticamente robusto de la IA en la investigación científica.

CONCLUSIÓN

La inteligencia artificial ofrece oportunidades relevantes para fortalecer la calidad, rapidez y alcance de la investigación científica; sin embargo, su adopción en México se desarrolla en un contexto de brechas regulatorias, infraestructura desigual y ausencia de lineamientos específicos que garanticen transparencia, equidad y protección de las personas participantes. Los retos éticos se entrelazan con vacíos normativos en validación, certificación y responsabilidad jurídica, limitando una implementación segura y responsable.

Resulta indispensable avanzar hacia un marco regulatorio integral, alineado con estándares internacionales y adaptado al contexto nacional, que contemple criterios de validación algorítmica, auditorías de equidad y explicabilidad, reglas claras para el manejo de datos sensibles, lineamientos para tecnologías generativas y el fortalecimiento de las capacidades de los Comités de Ética en Investigación. La reciente creación del Centro Público de Formación en Inteligencia Artificial representa un paso estratégico para el desarrollo de talento y capacidades nacionales que apoyen esta transformación.

Consolidar una estrategia nacional en inteligencia artificial para la investigación permitirá impulsar un uso ético y seguro de estas tecnologías, además de posicionar a México como un referente regional en innovación científica con enfoque de derechos humanos y justicia algorítmica.

REFERENCIAS

1. Topol EJ. High-performance medicine: the convergence of human and artificial intelligence. *Lancet*. 2019;393(10175):507–22. doi: 10.1016/S0140-6736(18)30117-8
2. Esteva A, Robicquet A, Ramsundar B, Kuleshov V, DePristo M, Chou K, et al. A guide to deep learning in healthcare. *Nat Med*. 2019;25(1):24–29. doi: 10.1038/s41591-018-0316-z

3. México. Congreso de la Unión. Ley Federal de Protección de Datos Personales en Posesión de los Particulares. *Diario Oficial de la Federación*; 2010 jul 5
4. Morley J, Floridi L. The limits of machine learning in healthcare: what clinicians should know. *Sci Eng Ethics*. 2020;26:1333–53. doi: 10.1007/s11948-019-00151-x
5. UNESCO. Recommendation on the Ethics of Artificial Intelligence [Internet]. Paris: UNESCO; 2021.[cited 2025 Dec 3] Disponible en: <https://unesdoc.unesco.org>
6. World Health Organization. Ethics and governance of artificial intelligence for health: WHO guidance [Internet]. Geneva: World Health Organization; 2021. [Cited 2025 Dec 3]. Disponible en: <https://www.who.int/publications/i/item/9789240029200>
7. OECD. OECD Principles on Artificial Intelligence. Paris: Organisation for Economic Co-operation and Development; 2019. Disponible en: <https://oecd.ai>
8. Turing AM. Computing machinery and intelligence. *Mind*. 1950 Oct;59(236):433–60. doi: 10.1093/mind/LIX.236.433
9. McCarthy J, Minsky ML, Rochester N, Shannon CE. A proposal for the Dartmouth Summer Research Project on Artificial Intelligence. Hanover (NH): Dartmouth College; 1955. Available from: <https://ojs.aaai.org/aimagazine/index.php/aimagazine/article/view/1904>
10. LeCun Y, Bengio Y, Hinton G. Deep learning. *Nature*. 2015;521(7553):436–44. doi: 10.1038/nature14539
11. Jordan MI, Mitchell TM. Machine learning: Trends, perspectives, and prospects. *Science*. 2015;349(6245):255–60. doi: 10.1126/science.aaa8415
12. Esteva A, Kuprel B, Novoa R, et al. Dermatologist-level classification of skin cancer with deep neural networks. *Nature*. 2017;542:115–8. doi: 10.1038/nature21056
13. Rajpurkar P, Irvin J, Zhu K, et al. CheXNet: Radiologist-level pneumonia detection on chest X-rays with deep learning. *arXiv*. 2017. doi: 10.48550/arXiv.1711.05225
14. European Union. Artificial Intelligence Act. Brussels: EU; 2021
15. Secretaría de Economía; Centro de Investigación y Docencia Económicas (CIDE). Agenda Nacional Mexicana de Inteligencia Artificial 2030. México: Secretaría de Economía; 2021. Available from: https://wp.oecd.ai/app/uploads/2022/01/Mexico_Agenda_Nacional_Mexicana_de_IA_2030.pdf

16. Gobierno de México. Centro Público de Formación en Inteligencia Artificial. Presidencia de México; 2024. Disponible en: <https://www.gob.mx/presidencia/prensa/presidenta-claudia-sheinbaum-encabeza-presentacion-del-nuevo-centro-publico-de-formacion-en-inteligencia-artificial>
17. Rajkomar A, Dean J, Kohane I. Machine learning in medicine. *N Engl J Med*. 2019;380:1347–58. doi: 10.1056/NEJMr1814259
18. Obermeyer Z, Emanuel EJ. Predicting the future — Big data and machine learning in medicine. *Science*. 2016;353(6301):446–52. doi: 10.1126/science.aad9837
19. Deo RC. Machine learning in medicine. *Circulation*. 2015 Nov 17;132(20):1920–30. doi: 10.1161/CIRCULATIONAHA.115.001593
20. Silver D, Huang A, Maddison CJ, Guez A, Sifre L, van den Driessche G, et al. Mastering the game of Go with deep neural networks and tree search. *Nature*. 2016 Jan 28;529(7587):484–9. doi: 10.1038/nature16961.26
21. LeCun Y, Bengio Y, Hinton G. Deep learning. *Nature*. 2015;521(7553):436–44. doi: 10.1038/nature14539
22. Rajpurkar P, Chen E, Banerjee O, Topol EJ. AI in health and medicine. *Nat Med*. 2022 Jan;28:31–8. doi: 10.1038/s41591-021-01614-0
23. Devlin J, Chang MW, Lee K, Toutanova K. BERT: Pre-training of deep bidirectional transformers for language understanding. In: *Proceedings of NAACL-HLT*; 2019. doi: 10.48550/arXiv.1810.04805
24. Bommasani R, Hudson DA, Adeli E, Altman R, Arora S, von Arx S, et al. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258*. 2022. Available from: <https://arxiv.org/abs/2108.07258>
25. Ronneberger O, Fischer P, Brox T. U-Net: Convolutional networks for biomedical image segmentation. In: *MICCAI*; 2015. p. 234–41
26. Litjens G, Kooi T, Bejnordi BE, et al. A survey on deep learning in medical image analysis. *Med Image Anal*. 2017;42:60–88. doi: 10.1016/j.media.2017.07.005
27. Karras T, Laine S, Aila T. A Style-Based Generator Architecture for Generative Adversarial Networks. In: *CVPR*; 2019. p. 4401–10
28. Sanchez-Lengeling B, Aspuru-Guzik A. Inverse molecular design using machine learning. *Science*. 2018;361:360–5. doi: 10.1126/science.aat2663
29. Shortliffe EH, Sepúlveda MJ. Artificial intelligence in medicine: Weighing accomplishments and promise. *Yearb Med Inform*. 2019;28(1):257–62. doi: 10.1055/s-0039-1677891
30. King RD, Rowland J, Oliver SG, et al. The automation of science. *Science*. 2009;324(5923):85–9. doi: 10.1126/science.1165620
31. Yang GZ, Cambias J, Cleary K, et al. Medical robotics: Regulatory, ethical, and legal considerations. *Sci Robot*. 2017;2(4):eaam8638. doi: 10.1126/scirobotics.aam8638
32. Sánchez-González JM, Morán-Moguel MC, Rivera-Cisneros AE, Ramírez-Barba EJ, Sierra-Amor RI, Portillo-Gallo JH, et al. Conceptualization of the use of artificial intelligence by clinical or research laboratory professionals: challenges for its implementation in Mexico. *EJIFCC*. 2025 Oct 1;36(3):366–77
33. Hernández Medina R, Kutuzova S, Nielsen KN, Johansen J, Hansen LH, Nielsen M, et al. Machine learning and deep learning applications in microbiome research. *ISME Commun*. 2022;2:98. doi: 10.1038/s43705-022-00182-y
34. Teodoro D, Naderi N, Yazdani A, Zhang B, Bornet A. A scoping review of artificial intelligence applications in clinical trial risk assessment. *npj Digit Med*. 2025;8:486. doi: 10.1038/s41746-025-01234-x
35. Resnik DB, Hosseini M. The ethics of using artificial intelligence in scientific research: new guidance needed for a new tool. *AI Ethics*. 2024 May;5(2):1499–521. doi: 10.1007/s43681-024-00493-8
36. Hanna, M. G., Pantanowitz, L., Jackson, B., Palmer, O., Visweswaran, S., Pantanowitz, J., & Rashidi, H. (2025). Ethical and Bias Considerations in Artificial Intelligence/Machine Learning. *Modern Pathology*, 38, Article 100686

Copyright © 2025 Comisión Nacional de Arbitraje Médico.
Todos los derechos reservados

Cortés Moreno GY. ORCID: 0000-0002-4506-8223

Conflicto de intereses:

La autora declara que no existen conflictos de interés personales, comerciales, financieros ni de otra índole que puedan influir en el contenido, resultados o interpretación del presente artículo.

Financiamiento: Este trabajo no recibió apoyo financiero de ninguna fuente pública, privada ni institucional.