

INNOVACIÓN e INTELIGENCIA ARTIFICIAL EN LA INVESTIGACIÓN CIENTÍFICA

Diego Alejandro Fernández Cando, Ibeth Patricia Párraga
Solórzano, Franklin Santiago Sosa Cifuentes, Delly Maribel San
Martín Torres, Carlos Alejandro Segarra Ramos & Cynthia
Shakira Enríquez Fierro



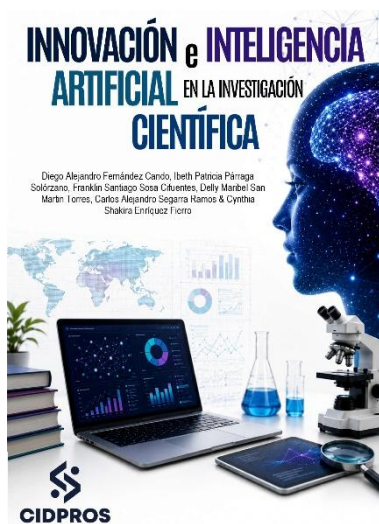
CIDPROS

Centro de innovación y desarrollo profesional

Innovación e inteligencia artificial en la investigación científica

Diego Alejandro Fernández Cando, Ibeth Patricia
Párraga Solórzano, Franklin Santiago Sosa
Cifuentes, Delly Maribel San Martín Torres, Carlos
Alejandro Segarra Ramos & Cynthia Shakira
Enríquez Fierro





Datos bibliográficos

ISBN	978-9907-9562-7-6
Título del libro	Innovación e inteligencia artificial en la investigación científica
Autores	Diego Alejandro Fernández Cando Ibeth Patricia Párraga Solórzano Franklin Santiago Sosa Cifuentes Delly Maribel San Martín Torres Carlos Alejandro Segarra Ramos Cynthia Shakira Enríquez Fierro
Editorial	CIDPROS EDITORIAL
Materia	001.4 - Investigación
Público objetivo	Profesional / académico
Año	2026
Número de edición	1
Tamaño	2.78Mb
Soporte	Libro digital descargable
Formato	PDF (.pdf)
Idioma	Español
DOI	https://doi.org/10.67166/dms1bc09
	Hecho en Ecuador / Made in Ecuador

Dr. Diego Alejandro Fernández Cando MSc.

Unidad Educativa Particular San Francisco Javier, La Escuela Javeriana de Loja.

fcalex1711@gmail.com

<https://orcid.org/0009-0007-2425-0169>

Loja, Loja, Ecuador

Semblanza

Dr. Diego Alejandro Fernández Cando, MSc., educador, investigador, escritor y líder académico ecuatoriano, originario de Loja, Ecuador, es Máster en Enseñanza del Inglés como Lengua Extranjera por el Centro Panamericano de Estudios Superiores de México y Licenciado en Ciencias de la Educación, mención Idioma Inglés, por la Universidad Nacional de Loja. Su sólida formación se complementa con múltiples diplomados internacionales en TEFL, TEYL, TESOL, TEAL y TEOL, además de certificaciones profesionales y técnicas avaladas por el Ministerio del Trabajo del Ecuador, SETEC y SENESCYT, consolidando un perfil integral orientado a la excelencia educativa, la investigación científica y la gestión institucional.



Actualmente se desempeña como docente investigador del Instituto Tecnológico Internacional Los Andes, Vicerrector de la Unidad Educativa Particular San Francisco Javier de Loja y coordinador académico de Easy English School of Languages, institución con nueve sucursales que promueve el acceso al aprendizaje del inglés y genera significativo impacto social y laboral. Su trayectoria incluye destacados cargos de liderazgo nacional e internacional, como Presidente de la Asociación Nacional de Profesores de Lenguas Extranjeras del Ecuador y Presidente de APIZSE, desde donde impulsa el desarrollo profesional docente, la innovación pedagógica y la dignificación del maestro de idiomas. Autor de más de 20 libros y de más de 35 artículos científicos, ha consolidado una carrera caracterizada por la producción intelectual, la formación continua y la transformación educativa.

Su excelencia ha sido reconocida con múltiples distinciones nacionales e internacionales, entre ellas varios Doctorados Honoris Causa, premios en investigación científica, liderazgo académico, excelencia educativa y reconocimientos globales como Embajador para la Paz 2026, Primer Lugar en Excelencia Educativa en Educa Latinoamérica 2026, y prestigiosos galardones otorgados por instituciones de América, Europa, Asia y organismos multilaterales. Dr. Fernández Cando representa una nueva generación de educadores humanistas cuya labor integra ciencia, innovación, liderazgo ético y compromiso social, proyectando desde Ecuador una visión transformadora de la educación como motor de desarrollo, paz y progreso global.

MSc. Ibeth Patricia Párraga Solórzano

Escuela Superior Politécnica Agropecuaria de Manabí “Manuel Félix López”

ipparraga@espam.edu.ec

<https://orcid.org/0009-0001-2732-8810>

Manta, Manabí, Ecuador

Semblanza

Ibeth Patricia Párraga Solórzano es una profesional ecuatoriana nacida en la ciudad de Manta, provincia de Manabí, cuya trayectoria ha estado marcada por el compromiso con la gestión pública, el desarrollo organizacional, la investigación y la docencia universitaria. Es Magíster en Administración de Empresas con mención en Recursos Humanos y cursa el Doctorado en Dirección de Proyectos en la Universidad de Investigación e Innovación de México (UIIX), consolidando una línea académica orientada a la innovación, la planificación estratégica y la mejora continua de las organizaciones. Cuenta con más de quince años de experiencia profesional en los sectores público y privado. Inició su carrera en el sistema financiero, fortaleciendo competencias en administración, planificación y gestión institucional. Posteriormente, asumió responsabilidades de alta dirección en entidades públicas y paralelamente, ha desarrollado una amplia trayectoria como asesora independiente.



En el ámbito académico, se desempeña como docente de la Escuela Superior Politécnica Agropecuaria de Manabí Manuel Félix López (ESPAM MFL), donde contribuye a la formación de profesionales en áreas relacionadas con la administración, la gestión de la innovación organizacional y la productividad. Su producción científica se centra en asociatividad productiva, desarrollo rural, gestión organizacional y dirección de proyectos, promoviendo propuestas orientadas al fortalecimiento de capacidades locales y al desarrollo territorial sostenible.

Como parte de su actividad investigativa, ha presentado ponencias científicas en escenarios académicos internacionales, entre ellos eventos desarrollados en la República del Perú, compartiendo los resultados de investigaciones vinculadas con sus estudios doctorales.

Convencida de que la educación, la investigación y la planificación estratégica constituyen herramientas fundamentales para el desarrollo de la sociedad, orienta su labor hacia la formación de profesionales íntegros, el fortalecimiento institucional y la generación de conocimiento con impacto social.

Mgs. Franklin Santiago Sosa Cifuentes

Investigador Independiente

franklinsosac_@hotmail.com

<https://orcid.org/0009-0009-3959-3206>

Riobamba, Chimborazo, Ecuador

Semblanza

Franklin Santiago Sosa Cifuentes es un Psicólogo Educativo, de nacionalidad ecuatoriana nacido el 09 de junio de 1994 en la ciudad de Riobamba, cuya vida ha estado marcado por su perseverancia, en un entorno urbano que fortaleció su sensibilidad intra e interpersonal enfocando en contribuir positivamente a las personas y la comunidad para el servicio del bien común. Desde temprana edad desarrolló una profunda empatía, respeto y tolerancia por el comportamiento humano, las emociones y los procesos mentales de los seres humanos.



Cuenta con una sólida formación académica: es Licenciado en Psicología Educativa, Orientador Vocacional y Familiar desde 2017, además obtuvo el título de Magíster en Intervención Psicológica en el Desarrollo y la Educación, el cual le permite generar espacios de apoyo psicopedagógico en las distintas etapas del desarrollo humano, así como estoy plenamente capacitado en las distintas especialidades de problemas de aprendizaje. Además, es especializado en las áreas de actividades de capacitación y administración educativa.

A lo largo de su trayectoria, ha destacado tanto en el ejercicio libre de la profesión dado que, cuenta con el centro de capacitaciones virtuales conocido como REPSICO-ECUADOR, el cual se desempeña como propietario y presidente, permitiendo ejecutar capacitaciones a nivel nacional e internacional, brindado ponencias en la Universidad de los Hemisferios y al Instituto Superior Tecnológico José Ortega y Gasset. Finalmente, fue miembro de la Federación de Estudiantes Universitarios del Ecuador para la Directiva de la Asociación Estudiantil de la Carrera de Psicología Educativa en la Universidad Nacional de Chimborazo de la ciudad de Riobamba.

Mgs. Delly Maribel San Martin Torres

Universidad Técnica de Machala
dsanmartin@utmachala.edu.ec
<https://orcid.org/0000-0002-4680-4042>
Machala, El Oro, Ecuador

Semblanza

Delly Maribel San Martin Torres es una ingeniera química ecuatoriana, nacida el 22 de octubre de 1973 en Machala. Su trayectoria profesional se ha desarrollado principalmente en la docencia universitaria, dentro del área de Química, con experiencia en la enseñanza de Química General, Química Inorgánica, Tecnología inorgánica, Química Analítica Cualitativa y Cuantitativa, así como en fundamentos relacionados con la estructura de los materiales. Actualmente, fortalece su perfil académico mediante la realización de un Doctorado en Ciencias Aplicadas y está próxima a culminar una Maestría en Internacionalización de la Educación superior, procesos formativos que consolidan su interés por la investigación, la innovación educativa y la proyección internacional de la educación superior.



A lo largo de su ejercicio profesional, ha participado en los ejes sustantivos de la educación superior: docencia, investigación y vinculación con la sociedad. Su trabajo integra enfoques pedagógicos innovadores y propuestas relacionadas con la sostenibilidad, la valorización de residuos agroindustriales, el estudio de materiales funcionales y el aprovechamiento responsable de los recursos naturales. Su interés académico se vincula con el fortalecimiento de la educación superior y la generación de soluciones científicas aplicadas a problemáticas ambientales y sociales.

Su perfil se distingue por el compromiso con la excelencia educativa, la responsabilidad social universitaria y la formación de profesionales capaces de analizar, transformar y aportar al desarrollo de su entorno. Desde su práctica docente, promueve una visión de la Química como una ciencia esencial para comprender la materia, optimizar procesos y contribuir a la construcción de una sociedad más sostenible.

MD. Carlos Alejandro Segarra Ramos

Investigador Independiente

carlosegarra1@hotmail.com

<https://orcid.org/0009-0004-0602-2363>

Machala, El Oro, Ecuador

Semblanza

Carlos Alejandro Segarra Ramos es un médico general ecuatoriano, nacido el 1 de marzo de 2001 en la ciudad de Machala, provincia de El Oro. Ha residido toda su vida en esta ciudad, un entorno que ha contribuido a fortalecer su compromiso con el bienestar de la comunidad y a desarrollar una visión de la medicina centrada en la atención integral, la empatía y el servicio a las personas.

Desde sus primeros años mostró interés por las ciencias de la salud y por el cuidado de quienes más lo necesitan, lo que orientó su vocación hacia el ejercicio de la medicina con responsabilidad, ética y sentido humanitario. Actualmente se desempeña como médico rural, función que ejerce desde enero de 2026, experiencia que le ha permitido conocer de cerca las necesidades sanitarias de la población y reafirmar su compromiso con la promoción, prevención y recuperación de la salud.

En el ámbito académico, obtuvo el título de Médico General y continúa fortaleciendo su formación profesional mediante la actualización constante de sus conocimientos. En la actualidad cursa una Maestría en Salud Pública y realiza un Diplomado en Cuidado del Paciente Crítico y Urgencias, con el propósito de ampliar sus competencias clínicas y contribuir al fortalecimiento de los servicios de salud mediante una atención basada en la evidencia, la calidad y la mejora continua.



PhD. Cynthia Shakira Enríquez Fierro

Universidad Internacional del Ecuador – UIDE

cyenriquezfi@uide.edu.ec

<https://orcid.org/0009-0002-5389-9892>

Quito – Pichincha - Ecuador

Semblanza

Cynthia Shakira Enríquez Fierro es docente universitaria, investigadora y estratega en comunicación, marketing y publicidad, con experiencia académica y profesional en diversas universidades e instituciones de educación superior. Su trayectoria se enfoca en áreas como comunicación digital, branding, reputación corporativa, comportamiento del consumidor, storytelling, publicidad estratégica y análisis de medios digitales. Ha participado en proyectos académicos, científicos y de formación vinculados a investigación, innovación y desarrollo de estrategias comunicacionales aplicadas a contextos contemporáneos. Su trabajo integra enfoques críticos, tecnológicos y creativos, especialmente en temas relacionados con transformación digital, eWOM, activismo de marca, medios digitales y comunicación estratégica, aportando tanto en docencia como en dirección y acompañamiento de proyectos académicos y científicos.





El contenido y las ideas expuestas en esta obra se encuentran protegidos por la normativa vigente en materia de propiedad intelectual y constituyen derechos exclusivos de su(s) autor(es).

Todos los derechos reservados © 2026

Sinopsis

Innovación e Inteligencia Artificial en la Investigación Científica analiza la transformación de la producción de conocimiento mediante sistemas inteligentes capaces de apoyar la búsqueda documental, la formulación de problemas, el diseño metodológico, el análisis de datos, la experimentación y la comunicación académica. La obra distingue entre aprendizaje automático, aprendizaje profundo, inteligencia artificial generativa, procesamiento del lenguaje natural, visión artificial, minería de textos, sensores inteligentes, gemelos digitales y laboratorios automatizados, explicando sus aplicaciones y límites en diferentes campos científicos.

A lo largo de seis capítulos se examinan la innovación científica, la investigación interdisciplinaria, las tecnologías emergentes, la gestión inteligente del conocimiento, la calidad metodológica, la integridad académica y la gobernanza institucional. El libro aborda también los riesgos asociados con alucinaciones, sesgos algorítmicos, opacidad, fuga de datos, dependencia tecnológica, vulneración de la privacidad y dificultades para garantizar transparencia, explicabilidad y reproducibilidad.

Desde una perspectiva crítica, se sostiene que la inteligencia artificial no sustituye el razonamiento científico ni la responsabilidad del investigador. Su contribución depende de la calidad de los datos, la pertinencia metodológica, la supervisión humana y la existencia de políticas responsables. La obra propone comprender estas tecnologías como infraestructuras de apoyo para una ciencia más colaborativa, abierta, rigurosa y socialmente pertinente.

Palabras clave: inteligencia artificial; investigación científica; innovación tecnológica; integridad académica; gobernanza científica.

Synopsis

Innovation and Artificial Intelligence in Scientific Research analyzes the transformation of knowledge production through intelligent systems capable of supporting literature searches, problem formulation, methodological design, data analysis, experimentation, and academic communication. The work distinguishes among machine learning, deep learning, generative artificial intelligence, natural language processing, computer vision, text mining, intelligent sensors, digital twins, and automated laboratories, explaining their applications and limitations across different scientific fields.

Across six chapters, the book examines scientific innovation, interdisciplinary research, emerging technologies, intelligent knowledge management, methodological quality, academic integrity, and institutional governance. It also addresses the risks associated with hallucinations, algorithmic bias, opacity, data leakage, technological dependence, privacy violations, and difficulties in ensuring transparency, explainability, and reproducibility.

From a critical perspective, the book argues that artificial intelligence does not replace scientific reasoning or the researcher's responsibility. Its contribution depends on data quality, methodological relevance, human oversight, and the existence of responsible policies. The work proposes understanding these technologies as support infrastructures for a more collaborative, open, rigorous, and socially relevant form of science.

Keywords: artificial intelligence; scientific research; technological innovation; academic integrity; scientific governance.

Índice de contenido

Sinopsis.....	11
Synopsis.....	12
Introducción.....	16
Capítulo 1. Innovación científica impulsada por la inteligencia artificial.....	19
1.1 Transformación de los modelos tradicionales de investigación científica.....	20
1.2 Inteligencia artificial como motor de innovación en la producción de conocimiento	23
1.3 Formulación de problemas, preguntas e hipótesis asistida por inteligencia artificial.....	26
1.4 Diseño de metodologías innovadoras mediante sistemas inteligentes.....	29
1.5 Descubrimiento de patrones, relaciones y tendencias en datos científicos.....	32
1.6 Experimentación, simulación y validación de resultados mediante inteligencia artificial.....	36
1.7 Evaluación del impacto científico, tecnológico y social de la innovación basada en inteligencia artificial	39
Capítulo 2. Investigación interdisciplinaria con inteligencia artificial.....	43
2.1 Fundamentos y alcance de la investigación interdisciplinaria asistida por inteligencia artificial	44
2.2 Integración de datos, métodos y lenguajes entre disciplinas científicas.....	47
2.3 Convergencia entre ciencias naturales, sociales, humanas y tecnológicas	51
2.4 Modelos colaborativos para equipos científicos interdisciplinarios	55
2.5 Aplicaciones interdisciplinarias en salud, educación, ingeniería y ciencias ambientales	58
2.6 Barreras epistemológicas, metodológicas e institucionales de la interdisciplinariedad	61
Capítulo 3. Tecnologías emergentes para la investigación científica.....	65
3.1 Aprendizaje automático y aprendizaje profundo en el análisis científico	66
3.2 Inteligencia artificial generativa aplicada a la producción académica.....	69

3.3	Procesamiento del lenguaje natural y minería de textos científicos	72
3.4	Visión artificial para el análisis de imágenes, muestras y fenómenos experimentales.....	75
3.5	Internet de las cosas, sensores inteligentes y recopilación automatizada de datos	78
3.6	Computación en la nube, ciencia de datos y procesamiento de alto rendimiento	82
3.7	Gemelos digitales, robótica científica y laboratorios automatizados.....	85
Capítulo 4. Gestión del conocimiento mediante inteligencia artificial		88
4.1	Fundamentos de la gestión inteligente del conocimiento científico	89
4.2	Búsqueda y recuperación automatizada de literatura académica.....	92
4.3	Clasificación, síntesis y análisis semántico de información científica	94
4.4	Sistemas de recomendación para la identificación de fuentes y tendencias de investigación.....	97
4.5	Repositorios inteligentes, datos abiertos y preservación del conocimiento...	101
4.6	Grafos de conocimiento, redes de citación y visualización de dominios científicos	104
4.7	Transferencia, reutilización y circulación del conocimiento asistida por inteligencia artificial	107
Capítulo 5. Calidad, ética e integridad en la investigación con inteligencia artificial .		111
5.1	Calidad, confiabilidad y validación de resultados generados con inteligencia artificial.....	112
5.2	Alucinaciones, errores y límites de los sistemas generativos	115
5.3	Sesgos algorítmicos y desigualdades en la producción científica	118
5.4	Transparencia, explicabilidad y reproducibilidad de los procesos investigativos	121
5.5	Integridad académica, autoría y responsabilidad científica	124
5.6	Privacidad, protección de datos y seguridad de la información científica.....	127

Capítulo 6. Gobernanza y futuro de la investigación científica con inteligencia artificial	131
6.1 Gobernanza institucional de la inteligencia artificial en la investigación científica	132
6.2 Regulación y marcos normativos para el uso responsable de sistemas inteligentes.....	135
6.3 Propiedad intelectual y uso legítimo de datos y contenidos científicos.....	138
6.4 Ciencia abierta, innovación colaborativa y democratización del conocimiento	140
6.5 Financiamiento, infraestructura y desigualdades en el acceso a la inteligencia artificial.....	143
6.6 Nuevos paradigmas y perspectivas futuras de la investigación científica	147
Bibliografía.....	150

Introducción

La investigación científica atraviesa una transformación profunda impulsada por la expansión de la inteligencia artificial, el crecimiento de los datos digitales y el desarrollo de infraestructuras capaces de automatizar operaciones antes reservadas al trabajo humano. Estas tecnologías intervienen en la búsqueda de literatura, la formulación de problemas, el diseño metodológico, el procesamiento de información, la experimentación, la escritura académica y la difusión de resultados. Su incorporación modifica los tiempos, recursos y formas de producir conocimiento, pero también exige revisar los criterios con los que se determina su calidad, pertinencia y validez.

La inteligencia artificial no constituye una herramienta única ni homogénea. Comprende enfoques diferenciados, como aprendizaje automático, aprendizaje profundo, procesamiento del lenguaje natural, visión artificial, minería de textos, sistemas generativos, analítica de datos y automatización inteligente. Cada uno responde a funciones, capacidades y limitaciones particulares. Distinguir estas tecnologías resulta indispensable para evitar interpretaciones imprecisas y seleccionar soluciones coherentes con las preguntas de investigación, los tipos de datos disponibles y las exigencias epistemológicas de cada disciplina científica.

El potencial de estos sistemas se manifiesta en su capacidad para examinar grandes volúmenes de información, descubrir patrones, integrar fuentes heterogéneas y apoyar procesos experimentales complejos. Los algoritmos pueden reconocer relaciones difíciles de detectar mediante procedimientos exclusivamente manuales, mientras los laboratorios automatizados, sensores inteligentes y gemelos digitales permiten desarrollar ciclos continuos de observación, simulación y ajuste. Sin embargo, una mayor capacidad computacional no produce conocimiento de forma automática. Los resultados necesitan interpretación teórica, contrastación empírica y evaluación crítica para adquirir significado científico (Krenn et al., 2022).

El papel del investigador, lejos de desaparecer, adquiere nuevas responsabilidades. La automatización desplaza parte del trabajo operativo hacia los sistemas inteligentes, pero incrementa la necesidad de supervisar datos, parámetros, decisiones y resultados. Corresponde al equipo científico determinar qué tareas pueden delegarse, comprobar la exactitud de las respuestas generadas y reconocer los límites de cada modelo. Castillo (2023) destaca que la supervisión humana continúa siendo indispensable cuando estas

herramientas participan en la producción académica, debido a que no pueden asumir responsabilidad intelectual, ética ni jurídica por los contenidos elaborados.

La inteligencia artificial también favorece nuevas formas de colaboración interdisciplinaria. Problemas relacionados con salud, educación, ambiente, ingeniería, administración o ciencias sociales requieren integrar datos, métodos y lenguajes procedentes de campos distintos. Los sistemas inteligentes pueden facilitar conexiones entre conceptos, organizar información y representar relaciones mediante grafos de conocimiento. No obstante, la convergencia tecnológica no elimina las diferencias epistemológicas entre disciplinas. La investigación interdisciplinaria exige acuerdos sobre el objeto de estudio, los criterios de evidencia, la interpretación de resultados y la distribución de responsabilidades dentro de los equipos.

Junto con sus oportunidades, estas tecnologías introducen riesgos que afectan directamente la confiabilidad del conocimiento. Las alucinaciones generativas, los sesgos algorítmicos, la fuga de datos, el sobreajuste, la opacidad de los modelos y la falta de documentación pueden producir resultados plausibles, pero metodológicamente incorrectos. Kapoor y Narayanan (2023) advierten que determinados errores en la preparación y separación de datos generan evaluaciones artificialmente favorables. Por ello, la reproducibilidad requiere registrar fuentes, versiones, parámetros, transformaciones, código y criterios de selección empleados durante todo el proceso investigativo.

Las implicaciones éticas y jurídicas tampoco pueden separarse de la innovación científica. El tratamiento automatizado de información plantea interrogantes sobre privacidad, consentimiento, seguridad, propiedad intelectual, autoría y uso legítimo de contenidos. De igual manera, los modelos entrenados con información histórica pueden reproducir desigualdades sociales, culturales, lingüísticas y territoriales. La adopción responsable exige políticas institucionales, mecanismos de evaluación y estructuras de gobernanza que permitan identificar riesgos, establecer responsabilidades y garantizar que la eficiencia tecnológica no se alcance mediante prácticas incompatibles con los derechos de las personas.

La distribución desigual de infraestructura, financiamiento y competencias digitales constituye otro desafío decisivo. Las universidades y centros de investigación con mayores recursos pueden acceder a plataformas, bases de datos y capacidad

computacional avanzada, mientras instituciones de regiones menos favorecidas enfrentan mayores barreras. Esta desigualdad limita la participación científica y puede profundizar dependencias respecto de proveedores tecnológicos externos. La ciencia abierta ofrece posibilidades para ampliar el acceso, aunque necesita políticas de formación, cooperación e infraestructura compartida que permitan una apropiación crítica y no solamente el consumo pasivo de herramientas desarrolladas en otros contextos.

Esta obra examina la innovación científica impulsada por inteligencia artificial, la investigación interdisciplinaria, las tecnologías emergentes, la gestión inteligente del conocimiento, la calidad, la ética, la integridad y la gobernanza. Sus seis capítulos articulan aplicaciones, fundamentos metodológicos, riesgos y perspectivas futuras, manteniendo como eje la responsabilidad humana sobre las decisiones científicas. El propósito no es presentar la inteligencia artificial como una solución infalible, sino analizar las condiciones bajo las cuales puede contribuir a producir conocimiento confiable, transparente, reproducible, socialmente pertinente y científicamente justificable.



Capítulo

1.

Innovación científica impulsada por la inteligencia artificial

Diego Alejandro Fernández Cando



1.1 Transformación de los modelos tradicionales de investigación científica

La transformación de los modelos tradicionales de investigación científica no supone su eliminación, sino una reorganización de sus etapas, recursos y criterios de validación. Durante buena parte del siglo XX, la investigación se estructuró mediante secuencias relativamente lineales: formulación del problema, revisión teórica, diseño metodológico, recolección de datos, análisis y difusión. La incorporación de sistemas inteligentes también obliga a reconocer que toda observación científica supone procesos de selección y construcción. Desde el constructivismo operativo, Giordano (2022) explica que el conocimiento no constituye una reproducción directa de la realidad, sino el resultado de operaciones mediante las cuales se establecen distinciones y significados. En consecuencia, las clasificaciones algorítmicas tampoco son neutrales: dependen de datos, variables, categorías y objetivos definidos durante el diseño del sistema.

Este cambio modifica la relación entre investigador, método y evidencia. Castillo (2023) sostiene que la supervisión humana continúa siendo indispensable cuando la inteligencia artificial apoya tareas de escritura o producción científica. El investigador deja de actuar únicamente como ejecutor de procedimientos y asume funciones de selección, control, interpretación y evaluación crítica. La responsabilidad científica no disminuye con la automatización; se amplía, porque exige examinar la calidad de los datos, la pertinencia de los resultados y los límites de las inferencias producidas mediante herramientas computacionales.

Los modelos convencionales estaban condicionados por la capacidad humana para localizar, clasificar y analizar grandes volúmenes de información. La inteligencia artificial reduce parcialmente esa limitación mediante procesamiento automatizado, búsqueda avanzada y detección de relaciones difíciles de identificar manualmente. Sin embargo, una mayor capacidad computacional no garantiza una comprensión científica suficiente. Krenn et al. (2022) señalan que comprender implica relacionar resultados con explicaciones, conceptos y estructuras teóricas. Por ello, el rendimiento algorítmico debe integrarse con razonamiento epistemológico, conocimiento disciplinar y juicio crítico responsable del investigador.

La revisión de literatura constituye uno de los ámbitos donde la transformación resulta más visible. Polo y Casique (2025) proponen integrar grafos de conocimiento y redes neuronales para mejorar la recuperación de información documental. Estas tecnologías

permiten explorar conexiones entre conceptos, publicaciones, autores y campos disciplinares con mayor rapidez. De forma complementaria, la visualización de datos abiertos mediante grafos facilita la representación estructurada de relaciones científicas, la identificación de núcleos temáticos y el reconocimiento de tendencias que podrían permanecer dispersas en revisiones exclusivamente manuales (Ávila, 2024).

También cambia la manera de formular problemas científicos. En el enfoque tradicional, las brechas se identificaban principalmente mediante lectura especializada, experiencia disciplinar y comparación sistemática de antecedentes. Los sistemas inteligentes pueden apoyar este proceso mediante análisis de tendencias, agrupamiento temático y exploración de corpus extensos. No obstante, la baja frecuencia de un tema no determina por sí sola su relevancia científica. Una brecha debe conducir a preguntas delimitadas, justificadas y científicamente abordables, vinculadas con necesidades teóricas, metodológicas o sociales concretas (Arias y Artigas, 2022).

El diseño metodológico pasa igualmente de una lógica rígida a una arquitectura más adaptable. Falcón y Serpa (2021) explican que los métodos teóricos y empíricos cumplen funciones diferenciadas, aunque complementarias, dentro de la investigación. La inteligencia artificial puede apoyar la comparación de técnicas, la selección de variables y la simulación de escenarios, pero no sustituye la coherencia entre problema, objetivos, enfoque, muestra, instrumentos y evidencia requerida. La innovación adquiere valor cuando fortalece esa articulación y no cuando introduce herramientas sofisticadas sin justificación científica suficiente previa.

En el tratamiento de datos, el cambio es especialmente profundo. Los procedimientos estadísticos convencionales se complementan con aprendizaje automático, aprendizaje profundo y algoritmos de cómputo científico capaces de modelar relaciones complejas. Estas herramientas amplían la capacidad analítica, pero también introducen riesgos de opacidad, sobreajuste, pérdida de interpretabilidad y dependencia de configuraciones técnicas. La selección del modelo debe responder a la naturaleza de los datos y al propósito del estudio, no únicamente a la novedad de la herramienta utilizada (Hua et al., 2023; Joyce et al., 2023).

La automatización de laboratorios representa otra ruptura con la práctica experimental tradicional. Fushimi et al. (2025) describen un sistema autónomo destinado a apoyar investigaciones biotecnológicas mediante la coordinación de instrumentos, protocolos y

registros. Este tipo de infraestructura reduce la intervención manual continua y favorece ciclos automatizados de prueba, evaluación y ajuste. Sin embargo, su implementación exige trazabilidad, control institucional y delimitación de responsabilidades. Los laboratorios autónomos también plantean desafíos tecnológicos, normativos y organizacionales que deben analizarse junto con sus ventajas operativas (Tobias y Wahab, 2025).

Figura 1

Evolución de la investigación científica



Nota. La figura sintetiza la evolución de la investigación científica mediante procesos iterativos, capacidades computacionales, automatización y nuevas funciones del investigador. Elaboración propia.

La producción de conocimiento adquiere además una dimensión colaborativa entre investigadores y sistemas inteligentes. En lugar de considerar la inteligencia artificial como sustituto del trabajo científico, resulta más preciso comprenderla como una infraestructura de apoyo cognitivo y operativo. Su incorporación puede fortalecer procesos de búsqueda, análisis, organización documental y redacción académica, siempre que existan protocolos claros y formación adecuada. La calidad de esta colaboración depende de la supervisión, la capacidad crítica y la definición institucional de las tareas que pueden automatizarse (Cedeño et al., 2024).

La transformación tecnológica no elimina los fundamentos epistemológicos de la ciencia. Cabrera y Cepeda (2022) destacan que la epistemología orienta la producción del conocimiento al ofrecer criterios para analizar su origen, validez y alcance. Los sistemas

inteligentes pueden ampliar operaciones científicas, pero no reemplazan la discusión sobre causalidad, objetividad, explicación y justificación. La gnoseología y la epistemología continúan siendo necesarias para comprender cómo se construye, valida y comunica responsablemente el conocimiento producido mediante procedimientos asistidos por inteligencia artificial (García et al., 2023).

Los modelos emergentes también deben responder a exigencias más rigurosas de transparencia y reproducibilidad. La rapidez analítica pierde valor cuando los procedimientos, parámetros y decisiones no pueden examinarse o repetirse. La fuga de datos puede generar resultados artificialmente favorables y afectar la validez de estudios basados en aprendizaje automático (Kapoor y Narayanan, 2023). Por ello, documentar fuentes de datos, código, versiones, criterios de selección y configuraciones técnicas constituye una obligación metodológica central para garantizar investigaciones confiables, auditables y científicamente verificables (Semmelrock et al., 2025).

La transformación de los modelos tradicionales conduce hacia una investigación más conectada, automatizada e interdisciplinaria, pero también más exigente en términos de control crítico. Morin (1999) plantea que el conocimiento científico debe reconocer relaciones, incertidumbres y niveles de complejidad que no pueden reducirse a explicaciones aisladas. Desde esta perspectiva, el avance tecnológico no debe evaluarse solo por la velocidad de los procesos o la cantidad de datos analizados, sino por la calidad, pertinencia, responsabilidad y solidez del conocimiento que contribuye a producir.

1.2 Inteligencia artificial como motor de innovación en la producción de conocimiento

La inteligencia artificial se ha convertido en un motor de innovación porque amplía la capacidad de la investigación para procesar información, reconocer regularidades y generar nuevas rutas de análisis. Su aporte no se limita a acelerar tareas, sino que modifica la forma en que se construyen preguntas, se organizan evidencias y se relacionan hallazgos procedentes de fuentes diversas. Barradas (2023) sostiene que estas tecnologías transforman la actividad científica al introducir apoyos computacionales capaces de intervenir en múltiples fases de la producción de conocimiento.

La innovación científica impulsada por inteligencia artificial depende de su integración con marcos teóricos, criterios metodológicos y conocimientos disciplinares. Un sistema puede identificar correlaciones o proponer relaciones entre variables, pero la

interpretación de esos resultados exige fundamentos epistemológicos y capacidad crítica. La producción de conocimiento no surge automáticamente del procesamiento masivo de datos, sino de la articulación entre evidencia, explicación y justificación científica. Por ello, el valor de la inteligencia artificial está condicionado por la calidad de las preguntas y del razonamiento humano que orienta su uso (Krenn et al., 2022; Castelo, 2025).

En los procesos de descubrimiento, los algoritmos inteligentes permiten explorar espacios de información que resultarían difíciles de examinar mediante procedimientos exclusivamente manuales. Hua et al. (2023) destacan que los algoritmos aplicados al cómputo científico favorecen la resolución de problemas complejos y el análisis de grandes conjuntos de datos. Estas capacidades pueden acelerar la identificación de patrones, hipótesis plausibles y relaciones emergentes. Sin embargo, la innovación no radica únicamente en la velocidad, sino en la posibilidad de formular explicaciones más consistentes y abrir nuevas líneas de investigación.

La recuperación documental también adquiere una función estratégica dentro de la producción de conocimiento. La integración de redes neuronales y grafos de conocimiento permite establecer conexiones entre conceptos, documentos y campos científicos que podrían permanecer fragmentados en búsquedas tradicionales. Polo y Casique (2025) plantean que esta combinación metodológica fortalece la recuperación de información y mejora la organización del conocimiento disponible. De manera complementaria, la representación gráfica de datos de investigación facilita la identificación de vínculos temáticos, trayectorias conceptuales y áreas insuficientemente exploradas (Ávila, 2024).

La inteligencia artificial generativa introduce posibilidades adicionales para sintetizar, reformular y estructurar contenidos académicos. Puede apoyar la preparación de borradores, la comparación de argumentos y la organización inicial de información, siempre que sus resultados sean verificados. Su utilidad no elimina el riesgo de producir afirmaciones inexactas, referencias inexistentes o formulaciones carentes de fundamento. La supervisión humana resulta indispensable para evaluar la coherencia, la originalidad y la correspondencia entre el texto generado y la evidencia científica disponible (Castillo, 2023; Gao et al., 2023).

Tabla 1

Dimensiones de innovación de la inteligencia artificial en la producción de conocimiento

Dimensión	Aporte innovador	Condición metodológica	Riesgo principal
Exploración documental	Relaciona grandes volúmenes de literatura y datos	Delimitación precisa de criterios de búsqueda	Inclusión de información irrelevante
Generación de hipótesis	Sugiere asociaciones y patrones emergentes	Validación teórica y contrastación empírica	Confundir correlación con explicación
Análisis científico	Procesa estructuras complejas y múltiples variables	Selección justificada de modelos	Opacidad y sobreajuste
Automatización experimental	Ejecuta y ajusta procedimientos repetitivos	Protocolos trazables y supervisión técnica	Pérdida de control sobre decisiones
Escritura académica	Apoya síntesis, organización y reformulación	Verificación de contenido y autoría	Alucinaciones e integridad académica
Comunicación de resultados	Facilita visualización y adaptación de contenidos	Correspondencia con la evidencia original	Simplificación excesiva

Nota. La tabla sintetiza funciones innovadoras de la inteligencia artificial y las condiciones necesarias para que su incorporación contribuya a una producción de conocimiento válida, comprensible y responsable.

La innovación también se expresa en la automatización experimental. Los laboratorios autónomos integran instrumentos, software y sistemas de decisión para ejecutar procedimientos, registrar resultados y ajustar condiciones operativas. Fushimi et al. (2025) muestran cómo estas arquitecturas pueden apoyar la investigación biotecnológica mediante ciclos de experimentación más continuos. No obstante, la autonomía técnica debe acompañarse de trazabilidad, documentación y criterios de intervención humana. La capacidad de automatizar experimentos adquiere valor científico cuando las decisiones del sistema pueden ser examinadas, justificadas y reproducidas.

Otro aporte relevante consiste en favorecer formas de investigación interdisciplinaria. Los sistemas inteligentes pueden integrar datos, lenguajes y métodos procedentes de distintas áreas, facilitando el diálogo entre ciencias naturales, sociales, médicas y tecnológicas. Esta convergencia amplía las posibilidades de abordar problemas complejos que no

pueden comprenderse desde una sola disciplina. Sin embargo, la integración requiere compatibilidad conceptual y metodológica, porque reunir información heterogénea no garantiza una explicación coherente. La innovación depende de construir relaciones pertinentes entre conocimientos diversos y no solo de acumular datos (Morin, 1999).

La producción de conocimiento asistida por inteligencia artificial también exige revisar los criterios de autoría, transparencia e integridad académica. Ros y Pérez (2023) sostienen que herramientas como ChatGPT pueden apoyar la escritura científica, pero no deben reconocerse como autoras, porque no asumen responsabilidad sobre el contenido. Del mismo modo, el uso de sistemas generativos debe declararse cuando su intervención sea significativa. La innovación pierde legitimidad si oculta procedimientos, atribuye capacidades humanas a las herramientas o impide identificar quién tomó las decisiones intelectuales centrales del estudio.

La inteligencia artificial puede impulsar una ciencia más dinámica, conectada y capaz de explorar problemas complejos, pero su contribución depende de condiciones institucionales y metodológicas rigurosas. La innovación científica no consiste en adoptar herramientas por su novedad, sino en utilizarlas para fortalecer la calidad de las preguntas, la solidez de los análisis y la pertinencia de las explicaciones. Su incorporación debe acompañarse de supervisión, trazabilidad y evaluación crítica, especialmente cuando interviene en decisiones que afectan la validez del conocimiento producido (Cedeño et al., 2024; Contreras et al., 2024).

1.3 Formulación de problemas, preguntas e hipótesis asistida por inteligencia artificial

La formulación del problema constituye una fase decisiva porque orienta el alcance, la pertinencia y la viabilidad de toda investigación. La inteligencia artificial puede apoyar esta tarea mediante la exploración de literatura, la identificación de conceptos recurrentes y el reconocimiento de vacíos temáticos. Sin embargo, una ausencia documental no equivale automáticamente a un problema científico relevante. Arias y Artigas (2022) sostienen que las brechas deben vincularse con preguntas justificadas, delimitadas y susceptibles de abordaje mediante procedimientos metodológicos coherentes.

Los sistemas inteligentes permiten analizar grandes conjuntos de publicaciones para detectar tendencias, controversias, relaciones conceptuales y áreas poco exploradas. Esta capacidad resulta útil cuando el investigador trabaja con campos extensos o

interdisciplinarios, donde la revisión manual puede fragmentar la comprensión del problema. La inteligencia artificial contribuye a organizar antecedentes y comparar enfoques, pero la decisión final requiere conocimiento disciplinar. La pertinencia de un problema depende de su valor científico, social o metodológico, no solo de su frecuencia dentro de un corpus documental (Mata et al., 2024).

La asistencia algorítmica puede mejorar la delimitación del problema al proponer variables, poblaciones, contextos y relaciones potenciales. Cedeño et al. (2024) señalan que la inteligencia artificial ofrece apoyo en distintas actividades de investigación universitaria, incluida la organización de información relevante. No obstante, las sugerencias generadas deben examinarse para evitar formulaciones demasiado amplias, ambiguas o desconectadas de la evidencia disponible. Delimitar implica establecer con precisión qué se estudiará, desde qué perspectiva y bajo cuáles condiciones temporales, espaciales o conceptuales.

La formulación de preguntas también puede beneficiarse del procesamiento del lenguaje natural. Estas herramientas permiten reformular expresiones, identificar términos imprecisos y comparar distintas versiones de una misma interrogante. Garzón et al. (2025) destacan la relevancia de la comunicación entre humanos y sistemas de inteligencia artificial mediante procesamiento del lenguaje natural. Aun así, una pregunta bien redactada no siempre es científicamente adecuada. Debe expresar una relación investigable, responder al problema planteado y guardar correspondencia con los objetivos, el diseño y los datos posibles.

Las preguntas asistidas por inteligencia artificial requieren una revisión epistemológica que determine qué tipo de conocimiento pretenden producir. Cabrera y Cepeda (2022) explican que la epistemología orienta la construcción del conocimiento científico mediante criterios de validez, fundamento y alcance. Por ello, no basta con generar interrogantes gramaticalmente correctas o novedosas. El investigador debe valorar si la pregunta busca describir, explicar, comprender, comparar o predecir un fenómeno, y si el enfoque seleccionado resulta compatible con la naturaleza del objeto estudiado.

La inteligencia artificial también puede apoyar la generación inicial de hipótesis mediante el reconocimiento de asociaciones entre variables y patrones presentes en datos o publicaciones. Hua et al. (2023) muestran que los algoritmos inteligentes amplían las posibilidades del cómputo científico y del análisis de relaciones complejas. Sin embargo,

una asociación algorítmica no constituye por sí sola una hipótesis científica sólida. Esta debe derivarse de fundamentos teóricos, expresarse de manera contrastable y establecer una relación clara entre elementos observables dentro del estudio.

En investigaciones exploratorias, cualitativas o interpretativas, la asistencia inteligente no debe forzar la formulación de hipótesis cuando el diseño no las requiere. Bautista (2022) destaca que la investigación cualitativa responde a fundamentos epistemológicos y metodológicos propios, centrados en la comprensión de significados, experiencias y contextos. En estos casos, la inteligencia artificial puede apoyar la construcción de preguntas orientadoras, categorías preliminares o mapas conceptuales. Su utilidad depende de respetar la lógica del enfoque y evitar trasladar automáticamente estructuras propias de diseños explicativos o predictivos.

La calidad de las propuestas generadas está condicionada por la información proporcionada al sistema. Instrucciones vagas suelen producir problemas genéricos, preguntas superficiales o hipótesis difíciles de contrastar. Por esa razón, el investigador debe introducir contexto disciplinar, antecedentes, población, variables y propósito del estudio. Barradas (2023) plantea que la inteligencia artificial transforma la investigación al ampliar sus capacidades operativas. No obstante, dicha transformación exige interacción crítica, porque la herramienta responde a los datos y criterios disponibles, sin comprender plenamente las implicaciones científicas de cada decisión.

Uno de los principales riesgos consiste en aceptar formulaciones plausibles sin verificar su respaldo documental. Los modelos generativos pueden producir expresiones coherentes que aparentan rigor científico, aunque incluyan relaciones no demostradas o antecedentes inexistentes. Gao et al. (2023) evidencian que los textos generados por ChatGPT pueden resultar difíciles de distinguir de producciones reales. En la formulación investigativa, esta capacidad obliga a contrastar cada afirmación con fuentes válidas y a comprobar que el problema planteado corresponda con una brecha auténtica y documentada.

La supervisión humana debe abarcar también la evaluación ética de las preguntas y las hipótesis. Determinadas formulaciones pueden reproducir sesgos, utilizar categorías discriminatorias o proponer análisis que afecten la privacidad de personas y comunidades. La asistencia tecnológica no exime al investigador de examinar las consecuencias del estudio, la sensibilidad de los datos y la pertinencia de las variables seleccionadas. El uso

responsable exige transparencia, juicio crítico y respeto por los principios de integridad científica durante toda la construcción del proyecto (Contreras et al., 2024).

La inteligencia artificial puede fortalecer la formulación de problemas, preguntas e hipótesis cuando funciona como apoyo para explorar, comparar y depurar alternativas. Castillo (2023) advierte que la supervisión humana resulta esencial en el uso de herramientas generativas aplicadas a la actividad científica. La decisión final debe permanecer bajo responsabilidad del investigador, quien determina la relevancia del problema, la coherencia de las preguntas y la contrastabilidad de las hipótesis. La innovación metodológica surge de combinar capacidad computacional con razonamiento teórico, experiencia disciplinar y evaluación crítica (Callirgos et al., 2022).

1.4 Diseño de metodologías innovadoras mediante sistemas inteligentes

El diseño de metodologías innovadoras mediante sistemas inteligentes parte de una reconsideración de la relación entre problema, método, datos y decisiones analíticas. Estas tecnologías pueden apoyar la selección de técnicas, la organización de variables y la evaluación preliminar de alternativas metodológicas (Callirgos et al., 2022). Sin embargo, su intervención debe responder a una lógica científica previamente definida. Falcón y Serpa (2021) recuerdan que los métodos teóricos y empíricos cumplen funciones distintas dentro de la investigación, por lo que su combinación exige coherencia y no una automatización indiferenciada de procedimientos.

La inteligencia artificial puede facilitar la comparación entre diseños experimentales, descriptivos, correlacionales, predictivos o interpretativos, especialmente cuando existen múltiples opciones posibles. A partir de información contextual, los sistemas pueden sugerir rutas metodológicas, advertir incompatibilidades y simular escenarios de aplicación. Estas funciones resultan útiles para ampliar la deliberación del investigador, pero no sustituyen su juicio. La metodología debe corresponder con la naturaleza del problema, el tipo de conocimiento buscado y las condiciones reales de acceso a datos, participantes y recursos disponibles (Cedeño et al., 2024).

Bautista (2022) señala que la investigación cualitativa requiere decisiones coherentes con su fundamento epistemológico, su orientación comprensiva y su atención al contexto. Desde esta perspectiva, los sistemas inteligentes pueden apoyar la construcción de categorías preliminares, guías de entrevista, matrices de observación o rutas de codificación. No obstante, tales propuestas deben revisarse para evitar esquemas rígidos

que limiten la emergencia de significados. La innovación metodológica no consiste en imponer automatización sobre cualquier enfoque, sino en adaptar la herramienta a la lógica interna del diseño seleccionado.

En los estudios cuantitativos, los sistemas inteligentes pueden contribuir a estimar tamaños muestrales, comparar variables, detectar valores atípicos y seleccionar modelos de análisis. También permiten anticipar problemas de calidad, consistencia o distribución de los datos antes de ejecutar procedimientos definitivos. Estas ventajas amplían la capacidad de planificación, aunque pueden generar confianza excesiva en recomendaciones algorítmicas. La selección metodológica debe fundamentarse en supuestos verificables y no únicamente en indicadores de rendimiento computacional, especialmente cuando se utilizan algoritmos complejos o poco interpretables (Hua et al., 2023; Joyce et al., 2023).

La simulación constituye otra vía de innovación metodológica. Mediante modelos computacionales, el investigador puede explorar comportamientos posibles, evaluar condiciones iniciales y anticipar resultados bajo diferentes escenarios. Qian et al. (2022) explican que los gemelos digitales permiten representar objetos y procesos físicos mediante réplicas cibernéticas capaces de apoyar análisis y toma de decisiones. Aplicados a la investigación, estos entornos favorecen pruebas preliminares y reducción de riesgos, siempre que sus parámetros reproduzcan adecuadamente las condiciones del fenómeno estudiado y se documenten sus limitaciones.

Tabla 2
Comparación entre diseño metodológico convencional y diseño asistido por sistemas inteligentes

Criterio	Diseño metodológico convencional	Diseño asistido por sistemas inteligentes
Selección de métodos	Depende principalmente de revisión teórica y experiencia del investigador	Incorpora comparación automatizada de alternativas y simulación de escenarios
Delimitación de variables	Se realiza mediante antecedentes y razonamiento disciplinar	Puede complementarse con detección de patrones y asociaciones en grandes conjuntos de datos
Construcción de instrumentos	Se apoya en modelos previos, validación experta y pruebas piloto	Permite generar versiones preliminares y detectar inconsistencias semánticas o estructurales

Criterio	Diseño metodológico convencional	Diseño asistido por sistemas inteligentes
Planificación del análisis	Se define antes de procesar los datos	Puede ajustarse mediante evaluación anticipada de supuestos y rendimiento de modelos
Capacidad de adaptación	Los cambios suelen depender de revisión manual	Facilita ajustes iterativos durante el proceso, sujetos a control metodológico
Riesgo principal	Rigidez, lentitud o limitada exploración de alternativas	Dependencia algorítmica, opacidad y selección automática sin justificación científica
Responsabilidad	Recae directamente en el equipo investigador	Permanece en el equipo investigador, aunque intervengan sistemas automatizados

Nota. La comparación muestra que la asistencia inteligente amplía las opciones de diseño y evaluación metodológica, pero no reemplaza la responsabilidad científica ni la necesidad de justificar cada decisión según el problema y el enfoque adoptado.

La construcción de instrumentos también puede beneficiarse de sistemas capaces de analizar lenguaje, detectar ambigüedades y proponer estructuras preliminares. Garzón et al. (2025) destacan los avances del procesamiento del lenguaje natural en la interacción entre personas y sistemas inteligentes. En investigación, esta capacidad puede aplicarse a cuestionarios, guías de entrevista o protocolos de observación. Sin embargo, la claridad lingüística no garantiza validez de contenido. Cada instrumento debe someterse a revisión experta, pilotaje y evaluación de su correspondencia con las dimensiones que pretende medir o comprender.

Los sistemas inteligentes permiten asimismo adaptar metodologías durante la ejecución del estudio. En diseños secuenciales o iterativos, pueden identificar patrones emergentes, sugerir nuevas comparaciones y orientar ajustes en la recolección de datos. Esta flexibilidad resulta especialmente valiosa en investigaciones complejas, aunque introduce el riesgo de modificar decisiones después de observar resultados preliminares. Para preservar la integridad metodológica, los cambios deben registrarse, justificarse y distinguirse del plan original, evitando que la adaptación se convierta en una búsqueda oportunista de resultados favorables (Kapoor y Narayanan, 2023).

Fushimi et al. (2025) muestran que los sistemas autónomos de laboratorio pueden ejecutar protocolos, registrar resultados y ajustar condiciones experimentales con una intervención humana reducida. Este tipo de automatización permite metodologías más continuas, precisas y escalables. No obstante, la innovación depende de que cada operación sea

trazable y reproducible. La autonomía técnica debe acompañarse de protocolos de supervisión, criterios de interrupción y mecanismos de validación. Sin estos controles, la eficiencia experimental podría ocultar errores acumulativos, decisiones opacas o fallas difíciles de identificar posteriormente.

La reproducibilidad ocupa un lugar central en el diseño metodológico asistido por inteligencia artificial. Los procedimientos deben documentar fuentes de datos, versiones de software, parámetros, criterios de entrenamiento y decisiones de exclusión o transformación. Semmelrock et al. (2025) identifican barreras que dificultan la reproducción de investigaciones basadas en aprendizaje automático, entre ellas la falta de documentación y acceso a recursos. Diseñar metodologías innovadoras exige prever desde el inicio cómo podrán revisarse, repetirse y auditarse los procesos, en lugar de considerar la transparencia como una tarea posterior.

La explicabilidad también condiciona la elección de sistemas inteligentes dentro de una metodología. Zárate (2022) sostiene que la explicabilidad permite comprender los fundamentos de una decisión algorítmica y evaluar sus efectos. Cuando una herramienta influye en la clasificación de datos, la selección de casos o la generación de predicciones, el investigador debe conocer su funcionamiento y sus límites. No siempre será posible interpretar cada operación interna, pero sí justificar la herramienta, describir sus criterios y establecer controles que reduzcan la dependencia de resultados opacos.

El diseño metodológico innovador no se define por la cantidad de tecnologías incorporadas, sino por la calidad con que estas fortalecen la relación entre problema, evidencia y explicación. Contreras et al. (2024) advierten que el uso científico de inteligencia artificial exige responsabilidad, transparencia y supervisión. En consecuencia, cada herramienta debe seleccionarse por su pertinencia, validarse según el contexto y someterse a revisión crítica. La innovación metodológica surge cuando los sistemas inteligentes amplían las capacidades investigativas sin debilitar la coherencia, la trazabilidad ni la responsabilidad humana.

1.5 Descubrimiento de patrones, relaciones y tendencias en datos científicos

El descubrimiento de patrones en datos científicos constituye una de las contribuciones más relevantes de la inteligencia artificial a la investigación contemporánea. Los algoritmos pueden examinar grandes volúmenes de información, reconocer regularidades y señalar asociaciones que resultarían difíciles de identificar mediante análisis

exclusivamente manuales. Hua et al. (2023) explican que los avances en algoritmos inteligentes amplían las capacidades del cómputo científico y favorecen la resolución de problemas complejos. Su utilidad depende, no obstante, de la calidad, estructura y pertinencia de los datos analizados.

Un patrón representa una regularidad observable dentro de un conjunto de datos, pero no equivale necesariamente a una explicación científica. Hu (2021), al examinar la comprensión científica desde la propuesta de Hempel, muestra que explicar un fenómeno exige relacionarlo con estructuras argumentativas y conocimientos que permitan entender por qué ocurre. La inteligencia artificial puede identificar agrupamientos, secuencias, anomalías o relaciones estadísticas sin comprender por sí misma las causas que las producen. Por ello, los resultados algorítmicos deben interpretarse desde marcos teóricos y conocimientos disciplinares. La comprensión científica exige relacionar las regularidades detectadas con conceptos, mecanismos y argumentos capaces de explicar su significado dentro del fenómeno estudiado (Krenn et al., 2022).

La capacidad de reconocer relaciones complejas resulta especialmente valiosa cuando intervienen numerosas variables o fuentes heterogéneas. Los modelos inteligentes pueden comparar atributos, estimar dependencias y construir representaciones que muestran conexiones entre elementos dispersos. Ávila (2024) señala que los grafos de conocimiento permiten visualizar datos abiertos de investigación y representar vínculos mediante estructuras organizadas. Estas herramientas ayudan a descubrir redes conceptuales y trayectorias informativas, aunque las conexiones identificadas deben evaluarse para determinar si expresan relaciones sustantivas o simples coincidencias estadísticas.

El análisis de tendencias permite observar cambios, continuidades y desplazamientos dentro de series temporales o colecciones documentales. Mediante técnicas automatizadas, es posible examinar la evolución de temas, conceptos, indicadores y resultados científicos a lo largo del tiempo. La integración de grafos de conocimiento y redes neuronales fortalece la recuperación y organización de información procedente de corpus extensos (Polo y Casique, 2025). Sin embargo, una tendencia documental puede reflejar prioridades editoriales, disponibilidad de financiamiento o concentración geográfica, y no únicamente transformaciones reales del conocimiento.

En las ciencias de la salud, la identificación automatizada de patrones puede contribuir al reconocimiento temprano de situaciones críticas. Swaminathan et al. (2023) desarrollaron un sistema de procesamiento del lenguaje natural orientado a detectar mensajes asociados con crisis de salud mental. Este tipo de aplicación muestra cómo el análisis inteligente de textos puede apoyar intervenciones rápidas. No obstante, su funcionamiento requiere especial cautela debido a la sensibilidad de los datos, la posibilidad de falsos positivos y las consecuencias que una clasificación errónea puede generar para las personas involucradas.

Los modelos de lenguaje también pueden analizar relaciones presentes en documentos científicos y sintetizar evidencia distribuida entre múltiples fuentes. Tang et al. (2023) evaluaron su desempeño en la síntesis de evidencia médica, mostrando el potencial de estas herramientas para organizar información especializada. A pesar de ello, una síntesis automatizada puede omitir matices, confundir niveles de evidencia o formular conclusiones que exceden los resultados originales. La detección de tendencias debe acompañarse de revisión experta, comparación documental y verificación de la correspondencia entre afirmaciones y fuentes.

El reconocimiento de anomalías constituye otra función relevante dentro del análisis científico. Los sistemas inteligentes pueden señalar valores atípicos, cambios inesperados o combinaciones inusuales que ameritan una revisión específica. Estas señales pueden revelar errores de medición, problemas de calidad, eventos excepcionales o fenómenos no considerados inicialmente. Sin embargo, eliminar automáticamente los casos atípicos puede empobrecer el análisis y ocultar información valiosa. La decisión debe responder al diseño metodológico, al contexto de recolección y a una justificación científica claramente documentada (Mata et al., 2024).

En la investigación experimental, los patrones detectados pueden utilizarse para ajustar condiciones y orientar nuevas pruebas. Fushimi et al. (2025) describen un sistema de laboratorio autónomo capaz de apoyar procesos biotecnológicos mediante ciclos coordinados de experimentación. La combinación entre automatización y análisis inteligente permite reconocer resultados prometedores y modificar parámetros operativos con mayor rapidez. Aun así, cualquier ajuste debe quedar registrado y someterse a control humano, pues una modificación automática puede introducir sesgos, alterar la comparabilidad o dificultar la reproducción posterior del experimento.

Los gemelos digitales amplían estas posibilidades al representar procesos físicos mediante réplicas cibernéticas. Qian et al. (2022) explican que estas arquitecturas permiten modelar objetos, simular comportamientos y analizar escenarios futuros. En investigación, su aplicación facilita la identificación de tendencias bajo diferentes condiciones sin intervenir inmediatamente sobre el sistema real. No obstante, la validez de los patrones simulados depende de la calidad del modelo, de la actualización de los datos y de la correspondencia entre la representación digital y el fenómeno observado.

Figura 2

Descubrimiento de patrones en datos científicos



Nota. La figura representa las etapas para identificar, interpretar y analizar patrones y relaciones en datos científicos. Elaboración propia.

El descubrimiento algorítmico también enfrenta riesgos de sesgo. Un modelo puede reconocer patrones consistentes dentro de los datos disponibles y, al mismo tiempo, reproducir desigualdades o categorías problemáticas incorporadas en ellos. Omiye et al. (2023) evidencian que los grandes modelos de lenguaje pueden propagar enfoques médicos basados en categorías raciales. Este problema demuestra que una regularidad estadística no debe aceptarse automáticamente como válida o neutral. El análisis requiere examinar cómo se construyeron los datos y qué consecuencias genera su interpretación.

La reproducibilidad es esencial para distinguir patrones sólidos de resultados accidentales. Kapoor y Narayanan (2023) advierten que la fuga de datos puede producir rendimientos artificialmente elevados y afectar la validez de investigaciones basadas en aprendizaje automático. Por esta razón, los conjuntos de entrenamiento, validación y

prueba deben gestionarse con separación y criterios transparentes. También es necesario documentar transformaciones, parámetros y decisiones analíticas. Sin estas medidas, las tendencias detectadas podrían depender de errores metodológicos y no de relaciones científicas realmente generalizables.

El valor científico del descubrimiento de patrones, relaciones y tendencias no reside únicamente en la capacidad de procesar más información, sino en convertir esos hallazgos en conocimiento verificable y comprensible. Semmelrock et al. (2025) destacan que la reproducibilidad en investigaciones basadas en aprendizaje automático enfrenta barreras técnicas, documentales e institucionales. Superarlas exige combinar análisis computacional, interpretación disciplinar, transparencia y supervisión humana. La inteligencia artificial aporta señales y estructuras útiles, pero corresponde al investigador determinar su significado, alcance y pertinencia dentro del proceso científico.

1.6 Experimentación, simulación y validación de resultados mediante inteligencia artificial

La experimentación asistida por inteligencia artificial amplía la capacidad de diseñar, ejecutar y ajustar procedimientos científicos con mayor precisión y continuidad. Los sistemas inteligentes pueden coordinar instrumentos, registrar variables, detectar desviaciones y recomendar modificaciones durante el desarrollo de una prueba. Esta intervención no elimina la planificación metodológica, sino que incorpora nuevas formas de control y retroalimentación. Fushimi et al. (2025) muestran que los laboratorios autónomos pueden apoyar investigaciones biotecnológicas mediante la integración de dispositivos, software y decisiones operativas dentro de ciclos experimentales supervisados.

La automatización experimental resulta especialmente útil en tareas repetitivas, sensibles a variaciones mínimas o dependientes de numerosas combinaciones de parámetros. Al reducir intervenciones manuales, puede mejorar la consistencia de las mediciones y aumentar la cantidad de ensayos realizados. Sin embargo, una mayor productividad no garantiza por sí misma resultados válidos. Cada procedimiento debe conservar criterios de calibración, trazabilidad y control de calidad, además de mecanismos que permitan identificar errores técnicos o decisiones algorítmicas inadecuadas antes de que afecten toda la secuencia experimental (Tobias y Wahab, 2025).

La simulación permite estudiar fenómenos mediante representaciones computacionales capaces de reproducir comportamientos, condiciones y relaciones difíciles de observar directamente. Qian et al. (2022) explican que los gemelos digitales constituyen réplicas cibernéticas de objetos o procesos físicos que facilitan el análisis de escenarios. En investigación, estas herramientas permiten anticipar respuestas, comparar alternativas y reducir riesgos antes de intervenir sobre sistemas reales. Su validez depende de la calidad de los datos, la precisión de los parámetros y la correspondencia entre el modelo digital y el fenómeno representado.

Las aplicaciones de gemelos digitales abarcan sectores donde la experimentación directa puede resultar costosa, peligrosa o limitada por condiciones operativas. Singh et al. (2022) identifican usos en diferentes industrias, relacionados con monitoreo, predicción y optimización de procesos. Dentro de la investigación científica, estas arquitecturas pueden apoyar pruebas virtuales y evaluar modificaciones antes de aplicarlas físicamente. No obstante, los resultados simulados deben interpretarse como aproximaciones condicionadas por supuestos, simplificaciones y datos de entrada, no como reproducciones infalibles de la realidad observada.

La inteligencia artificial también puede intervenir en la validación mediante comparación de resultados, detección de anomalías y evaluación de rendimiento predictivo. Los algoritmos identifican inconsistencias, valores atípicos y diferencias entre observaciones reales y resultados esperados. Hua et al. (2023) señalan que los avances en cómputo científico inteligente favorecen el análisis de problemas complejos. Sin embargo, la validación no debe reducirse a métricas de precisión. Es necesario comprobar la coherencia teórica, la estabilidad del modelo, la calidad de los datos y la aplicabilidad de los resultados.

Tabla 3
Comparación entre experimentación convencional y experimentación asistida por inteligencia artificial

Criterio	Experimentación convencional	Experimentación asistida por inteligencia artificial
Ejecución de protocolos	Predomina la intervención manual y el seguimiento directo	Integra automatización, sensores y ajustes algorítmicos
Capacidad de repetición	Limitada por tiempo, recursos y disponibilidad del personal	Facilita ciclos continuos y mayor cantidad de ensayos

Criterio	Experimentación convencional	Experimentación asistida por inteligencia artificial
Ajuste de parámetros	Se realiza después de revisar resultados parciales	Puede efectuarse durante la ejecución mediante retroalimentación
Simulación previa	Depende de modelos simplificados o pruebas piloto	Permite escenarios complejos y réplicas digitales dinámicas
Detección de errores	Se basa en controles periódicos y revisión humana	Incorpora alertas automáticas y reconocimiento de anomalías
Validación	Prioriza contrastes estadísticos y replicación experimental	Añade evaluación predictiva, pruebas computacionales y análisis de estabilidad
Riesgo metodológico	Errores manuales, baja escalabilidad y variabilidad operativa	Opacidad, dependencia técnica y propagación automatizada de errores
Responsabilidad científica	Corresponde al equipo investigador	Permanece en el equipo, aunque intervengan sistemas autónomos

Nota. La comparación evidencia que la inteligencia artificial amplía la capacidad experimental y de simulación, pero exige controles adicionales para garantizar trazabilidad, interpretación y validación científica.

La validación de modelos requiere separar adecuadamente los datos utilizados para entrenamiento, ajuste y evaluación. Cuando esta división se realiza de forma incorrecta, el sistema puede aprender información que debería permanecer reservada para la prueba final. Kapoor y Narayanan (2023) advierten que la fuga de datos produce estimaciones artificialmente favorables y compromete la reproducibilidad. Por ello, la evaluación debe diseñarse antes del análisis, documentar cada transformación y evitar modificaciones orientadas exclusivamente a mejorar indicadores después de observar los resultados obtenidos.

La interpretabilidad constituye otro criterio relevante para validar resultados generados mediante sistemas inteligentes. Un modelo puede alcanzar un rendimiento elevado y, al mismo tiempo, ofrecer explicaciones insuficientes sobre sus decisiones. Joyce et al. (2023) relacionan la inteligencia artificial explicable con la necesidad de transparencia, interpretabilidad y comprensión. En contextos científicos, conocer qué variables influyen en una predicción permite evaluar su plausibilidad y detectar relaciones espurias. La validación debe considerar tanto la eficacia computacional como la posibilidad de justificar los resultados desde el conocimiento disciplinar.

La reproducibilidad experimental exige registrar versiones de software, parámetros, conjuntos de datos, secuencias de operación y criterios utilizados para aceptar o descartar

resultados. Semmelrock et al. (2025) identifican barreras técnicas y organizacionales que dificultan reproducir investigaciones basadas en aprendizaje automático. La automatización puede favorecer registros detallados, pero también generar procesos difíciles de reconstruir cuando la documentación es insuficiente. Diseñar flujos auditables desde el inicio permite determinar si los hallazgos dependen del fenómeno estudiado o de configuraciones particulares del sistema empleado.

La validación mediante inteligencia artificial no puede limitarse a confirmar el funcionamiento interno de un algoritmo. También debe examinar si sus resultados son transferibles a otros contextos, poblaciones o condiciones experimentales. Ciobanu et al. (2024) destacan la importancia de combinar reproducibilidad e interpretabilidad en aplicaciones médicas basadas en aprendizaje automático. Esta exigencia es aplicable a otros campos científicos, donde un modelo puede funcionar adecuadamente con datos controlados y perder rendimiento frente a variaciones reales, cambios instrumentales o poblaciones insuficientemente representadas.

La experimentación y la simulación asistidas por inteligencia artificial fortalecen la investigación cuando amplían la capacidad de prueba sin debilitar el rigor metodológico. Su implementación requiere supervisión humana, validación independiente y criterios claros para intervenir ante resultados inesperados. Contreras et al. (2024) advierten que el uso científico de estas tecnologías implica desafíos éticos relacionados con transparencia y responsabilidad. La innovación experimental adquiere legitimidad cuando los sistemas inteligentes complementan la decisión humana, documentan sus operaciones y permiten examinar críticamente cada resultado producido.

1.7 Evaluación del impacto científico, tecnológico y social de la innovación basada en inteligencia artificial

La evaluación debe distinguir entre producción de conocimiento, aplicación científica y desarrollo técnico. Lopardo (2024) diferencia la ciencia básica, la ciencia aplicada y la técnica como actividades relacionadas, pero no equivalentes. Esta precisión resulta necesaria porque una herramienta de inteligencia artificial puede representar un avance técnico sin generar conocimiento nuevo, o producir resultados científicos cuya aplicación práctica todavía sea limitada. Los indicadores de impacto deben responder, por tanto, a la naturaleza específica de la innovación examinada.

La evaluación del impacto científico de una innovación basada en inteligencia artificial debe examinar si realmente amplía la capacidad para formular preguntas, producir evidencia y generar explicaciones relevantes. Castelo (2025) advierte que la investigación tecnológica enfrenta dificultades cuando el desarrollo técnico no se vincula adecuadamente con necesidades científicas, productivas o sociales. Por ello, no basta con registrar mayor velocidad de procesamiento o incremento de publicaciones. El valor científico depende de la calidad metodológica, la originalidad de los hallazgos y su contribución al conocimiento acumulado. Barradas (2023) sostiene que la inteligencia artificial transforma la investigación cuando interviene de manera significativa en la producción científica, y no solo como recurso instrumental complementario.

Un primer criterio corresponde a la solidez y reproducibilidad de los resultados obtenidos. Las innovaciones basadas en aprendizaje automático pueden mostrar desempeños elevados y, al mismo tiempo, depender de errores en la preparación de datos, configuraciones poco transparentes o procedimientos difíciles de replicar. La evaluación debe considerar acceso a datos, documentación del código, descripción de parámetros y estabilidad de los resultados frente a nuevas pruebas. Sin estas condiciones, el impacto científico puede ser aparente y no representar un avance verificable (Kapoor y Narayanan, 2023; Semmelrock et al., 2025).

Ciobanu et al. (2024) destacan que la reproducibilidad y la interpretabilidad son elementos esenciales para valorar aplicaciones de aprendizaje automático en medicina. Este criterio puede extenderse a otras áreas científicas, porque un modelo difícil de comprender limita la posibilidad de revisar sus decisiones y explicar sus resultados. El impacto no debe medirse únicamente por precisión predictiva, sino también por la capacidad de justificar cómo se llegó a una conclusión. La explicabilidad fortalece la confianza, facilita la validación y permite identificar errores o relaciones espurias.

El impacto tecnológico se relaciona con la capacidad de una innovación para mejorar procesos, integrar sistemas y resolver problemas operativos de forma sostenible. Los laboratorios autónomos, por ejemplo, permiten coordinar instrumentos, automatizar protocolos y aumentar la continuidad experimental. Fushimi et al. (2025) muestran el potencial de estas arquitecturas en investigación biotecnológica. Sin embargo, su valoración debe incluir costos, mantenimiento, interoperabilidad, dependencia de proveedores y necesidad de personal especializado. Una solución tecnológicamente avanzada puede generar bajo impacto si no puede sostenerse institucionalmente.

La evaluación también debe considerar la transferencia de la innovación hacia contextos productivos, educativos, sanitarios o administrativos. Una herramienta puede ser científicamente prometedora, pero carecer de condiciones para su adopción práctica. Los gemelos digitales ilustran esta relación entre conocimiento y aplicación, al permitir representar procesos físicos, simular escenarios y apoyar decisiones antes de intervenir sobre sistemas reales (Qian et al., 2022). Su impacto tecnológico depende de la precisión de los modelos, la actualización de los datos y la capacidad organizacional para utilizar sus resultados.

El impacto social exige analizar quiénes se benefician de la innovación, quienes asumen sus riesgos y qué grupos pueden quedar excluidos. La inteligencia artificial puede ampliar el acceso a información, mejorar servicios y apoyar decisiones complejas, pero también reproducir desigualdades presentes en los datos. Kuppler et al. (2022) advierten que una predicción considerada justa no garantiza una decisión socialmente justa. Por ello, la evaluación debe integrar criterios de equidad, distribución de beneficios y consecuencias sobre poblaciones históricamente desfavorecidas.

Los sesgos algorítmicos constituyen un indicador central del impacto social. Omiye et al. (2023) evidencian que los grandes modelos de lenguaje pueden reproducir enfoques médicos basados en categorías raciales. Este hallazgo demuestra que una innovación puede alcanzar utilidad operativa mientras mantiene supuestos problemáticos o discriminatorios. Evaluar el impacto requiere revisar la composición de los datos, las categorías empleadas, las tasas de error entre grupos y los efectos de las decisiones automatizadas. La neutralidad tecnológica no debe asumirse sin evidencia suficiente.

La privacidad y la protección de datos forman parte de la valoración social y jurídica de estas innovaciones. Andrade et al. (2024) analizan la relación entre inteligencia artificial, tratamiento de datos personales y fundamentos legales. Una innovación puede ofrecer resultados valiosos y, al mismo tiempo, comprometer derechos mediante recolección excesiva, usos secundarios o falta de consentimiento. La evaluación debe incluir proporcionalidad, seguridad, anonimización y control sobre la información. El impacto positivo disminuye cuando la eficiencia tecnológica se obtiene mediante prácticas incompatibles con la protección de las personas.

Otro aspecto relevante es la distribución institucional de capacidades. La innovación basada en inteligencia artificial requiere infraestructura, financiamiento, acceso a datos y

formación especializada. Estas condiciones no se encuentran igualmente disponibles entre universidades, regiones o países. La ciencia abierta puede ampliar oportunidades, pero también reproducir ventajas acumulativas cuando ciertos actores concentran recursos y visibilidad (Ross et al., 2022). Por ello, la evaluación del impacto debe observar no solo resultados inmediatos, sino también desigualdades de acceso, dependencia tecnológica y posibilidades reales de apropiación local.

La evaluación integral del impacto científico, tecnológico y social exige combinar indicadores cuantitativos con análisis cualitativos y criterios éticos. Contreras et al. (2024) señalan que el uso de inteligencia artificial en investigación debe acompañarse de responsabilidad, transparencia y supervisión. Una innovación valiosa no es aquella que únicamente incrementa eficiencia, sino la que produce conocimiento confiable, fortalece capacidades, reduce riesgos y genera beneficios socialmente pertinentes. Medir estos efectos de manera articulada permite distinguir entre novedad tecnológica y transformación científica verdaderamente sostenible.



Capítulo

2.

Investigación interdisciplinaria con inteligencia artificial

Ibeth Patricia Párraga Solórzano

2.1 Fundamentos y alcance de la investigación interdisciplinaria asistida por inteligencia artificial

La investigación interdisciplinaria asistida por inteligencia artificial surge de la necesidad de abordar problemas cuya complejidad excede los límites de una sola disciplina. Su fundamento consiste en integrar conceptos, métodos, datos y formas de interpretación procedentes de campos distintos, sin reducirlos a una simple acumulación de perspectivas. Morin (1999) plantea que el conocimiento complejo exige reconocer relaciones, interdependencias e incertidumbres. Desde esta orientación, la inteligencia artificial funciona como mediadora capaz de conectar información heterogénea y apoyar procesos científicos más amplios.

La interdisciplinariedad se diferencia de la multidisciplinariedad porque no se limita a reunir especialistas que trabajan de forma paralela. Requiere interacción conceptual, adaptación metodológica y construcción de un lenguaje compartido entre áreas. Los sistemas inteligentes pueden facilitar este proceso mediante clasificación semántica, extracción de conceptos y reconocimiento de relaciones entre corpus especializados. Sin embargo, la integración no se produce automáticamente por disponer de más información. Debe existir un problema común que justifique la convergencia de conocimientos y oriente las decisiones analíticas (Duedra et al., 2020).

Krenn et al. (2022) sostienen que la inteligencia artificial puede contribuir a la comprensión científica cuando sus resultados se relacionan con explicaciones y estructuras conceptuales. Esta idea resulta central en la investigación interdisciplinaria, donde los algoritmos pueden detectar conexiones entre datos pertenecientes a dominios distintos. Su aporte no reside únicamente en procesar información, sino en hacer visibles relaciones que favorezcan nuevas interpretaciones. Aun así, corresponde a los investigadores determinar si esas conexiones poseen significado teórico o representan asociaciones superficiales generadas por proximidad estadística.

La recuperación de información constituye una base operativa para el trabajo interdisciplinario. Polo y Casique (2025) proponen integrar grafos de conocimiento y redes neuronales con el propósito de fortalecer la búsqueda documental. Estas herramientas permiten vincular términos, categorías y publicaciones procedentes de diferentes campos, reduciendo la fragmentación del conocimiento. Los grafos también facilitan la representación de trayectorias conceptuales y relaciones entre fuentes. Su uso

exige controlar la calidad de los metadatos, la ambigüedad terminológica y las diferencias semánticas que existen entre disciplinas.

La inteligencia artificial amplía el alcance interdisciplinario al procesar datos estructurados, textos, imágenes, señales y registros experimentales dentro de una misma arquitectura analítica. Hua et al. (2023) describen avances en algoritmos inteligentes aplicados al cómputo científico que permiten resolver problemas de elevada complejidad. Esta capacidad favorece estudios donde intervienen variables de naturaleza distinta, como ocurre en salud, educación, ingeniería o ciencias sociales. No obstante, integrar formatos heterogéneos requiere criterios de compatibilidad, normalización y validación para evitar interpretaciones metodológicamente inconsistentes.

Tabla 4
Síntesis de fundamentos de la investigación interdisciplinaria asistida por inteligencia artificial

Fundamento	Función de la inteligencia artificial	Exigencia científica
Problema complejo	Integra información procedente de varios campos	Delimitación compartida del objeto de estudio
Convergencia conceptual	Relaciona términos, categorías y modelos	Construcción de significados compatibles
Integración metodológica	Combina técnicas analíticas y tipos de datos	Coherencia entre métodos y preguntas
Mediación tecnológica	Automatiza búsqueda, clasificación y análisis	Supervisión humana y trazabilidad
Producción de explicaciones	Detecta relaciones y patrones emergentes	Interpretación teórica interdisciplinaria
Validación conjunta	Contrasta resultados desde distintas perspectivas	Criterios comunes de calidad y pertinencia
Alcance social	Apoya el estudio de problemas multidimensionales	Evaluación ética y contextual de efectos

Nota. La síntesis muestra que la inteligencia artificial facilita conexiones entre disciplinas, pero la integración científica depende de acuerdos conceptuales, metodológicos y éticos establecidos por los equipos de investigación.

La colaboración interdisciplinaria también transforma la distribución de responsabilidades dentro de los equipos. Especialistas en datos, expertos disciplinares, metodólogos y profesionales de áreas aplicadas deben coordinar decisiones sobre variables, modelos y criterios de interpretación. Cedeño et al. (2024) señalan que la inteligencia artificial puede fortalecer varias fases de la investigación universitaria, aunque su aprovechamiento exige capacidades institucionales. La innovación no depende

únicamente de disponer de herramientas, sino de crear condiciones para que los participantes comprendan sus aportes, límites y obligaciones dentro del proceso común.

Las diferencias epistemológicas pueden convertirse en una fuente de enriquecimiento o de conflicto. Bautista (2022) explica que los enfoques cualitativos responden a fundamentos orientados a comprender significados, experiencias y contextos, mientras otros diseños privilegian medición, predicción o explicación causal. La inteligencia artificial puede intervenir en ambos, pero no debe homogeneizar sus supuestos. Una investigación interdisciplinaria rigurosa reconoce estas diferencias y establece cómo se relacionarán los resultados, evitando que una perspectiva metodológica subordine a las demás sin una justificación explícita.

El alcance de estos estudios se extiende a problemas que combinan dimensiones técnicas, humanas y organizacionales. Los gemelos digitales, por ejemplo, permiten representar procesos físicos y analizar escenarios mediante réplicas cibernéticas (Qian et al., 2022). Su aplicación puede requerir conocimientos de ingeniería, informática, gestión y ciencias sociales para comprender tanto el funcionamiento técnico como sus efectos institucionales. La inteligencia artificial facilita la integración de estas dimensiones, aunque la interpretación final debe considerar las condiciones reales en las que la tecnología será utilizada.

La interdisciplinariedad asistida por inteligencia artificial también plantea desafíos de comunicación. Cada disciplina posee conceptos, unidades de análisis y criterios de evidencia particulares. Garzón et al. (2025) destacan la relevancia del procesamiento del lenguaje natural para mejorar la interacción entre personas y sistemas inteligentes. Estas capacidades pueden ayudar a traducir terminologías y organizar información especializada, pero no eliminan las diferencias de significado. Los equipos deben construir acuerdos sobre definiciones, variables y alcances para evitar que una aparente equivalencia lingüística produzca errores conceptuales.

Otro límite corresponde a los sesgos que pueden trasladarse entre disciplinas mediante modelos compartidos. Un sistema entrenado con categorías problemáticas puede reproducirlas en nuevos campos y otorgarles una apariencia de objetividad técnica. Omiye et al. (2023) demostraron que los grandes modelos de lenguaje pueden propagar enfoques médicos basados en categorías raciales. Este riesgo obliga a revisar los datos, las clasificaciones y las consecuencias de aplicar un modelo fuera de su contexto original.

La integración interdisciplinaria requiere validación contextual y responsabilidad ética conjunta.

La investigación interdisciplinaria asistida por inteligencia artificial alcanza su mayor valor cuando permite construir explicaciones más completas sobre problemas complejos, sin diluir la especificidad de cada disciplina. Contreras et al. (2024) advierten que el uso científico de estas tecnologías debe acompañarse de transparencia, supervisión y responsabilidad. Por ello, la colaboración debe sostenerse en acuerdos metodológicos, trazabilidad de decisiones y evaluación crítica de los resultados. La inteligencia artificial amplía las posibilidades de conexión, pero la coherencia científica continúa dependiendo del razonamiento humano.

2.2 Integración de datos, métodos y lenguajes entre disciplinas científicas

La integración de datos entre disciplinas científicas exige superar diferencias de formato, escala, procedencia y significado. Los conjuntos provenientes de salud, ingeniería, educación, ciencias sociales o humanidades no pueden combinarse únicamente por compartir una infraestructura digital. Cada campo produce información bajo criterios específicos de medición, clasificación y validación. La inteligencia artificial puede apoyar la normalización y el reconocimiento de correspondencias, pero la integración científica requiere examinar cómo fueron generados los datos y qué limitaciones acompañan su uso (Hua et al., 2023).

Uno de los principales desafíos consiste en asegurar compatibilidad semántica. Dos disciplinas pueden utilizar un mismo término con significados distintos o emplear expresiones diferentes para describir fenómenos semejantes. Garzón et al. (2025) destacan que el procesamiento del lenguaje natural favorece la comunicación entre seres humanos y sistemas inteligentes. Aplicado a la investigación interdisciplinaria, puede facilitar la identificación de equivalencias conceptuales, la clasificación temática y la traducción de vocabularios especializados. Sin embargo, estas correspondencias deben ser revisadas por expertos de las áreas involucradas.

Los grafos de conocimiento ofrecen una vía para representar relaciones entre conceptos, documentos, variables y entidades procedentes de campos distintos. Ávila (2024) muestra que estas estructuras permiten organizar y visualizar datos abiertos de investigación mediante ontologías y modelos vinculados. Su utilidad interdisciplinaria reside en hacer explícitas conexiones que podrían permanecer dispersas entre bases de datos o

publicaciones. No obstante, la construcción del grafo exige definir criterios comunes, resolver ambigüedades y justificar cada vínculo para evitar asociaciones técnicamente posibles, pero conceptualmente débiles.

Polo y Casique (2025) proponen combinar grafos de conocimiento y redes neuronales para fortalecer la recuperación de información documental. Esta integración puede ayudar a localizar evidencia relevante en corpus extensos y heterogéneos, especialmente cuando los términos de búsqueda varían entre disciplinas. La recuperación interdisciplinaria no debe limitarse a aumentar la cantidad de documentos encontrados. También debe valorar pertinencia, contexto y relación con el problema común, pues una búsqueda amplia puede introducir materiales incompatibles con el enfoque metodológico o con la unidad de análisis seleccionada.

La integración de métodos requiere reconocer que cada disciplina formula preguntas y valida resultados mediante tradiciones específicas. Falcón y Serpa (2021) distinguen la función de los métodos teóricos y empíricos dentro de la investigación. En un diseño interdisciplinario, la inteligencia artificial puede apoyar la combinación de análisis estadístico, procesamiento textual, simulación y clasificación automatizada. Sin embargo, la coexistencia de técnicas no garantiza coherencia. Cada método debe responder a una dimensión concreta del problema y articularse con los demás mediante una lógica explícita de complementariedad.

Los enfoques cualitativos plantean exigencias particulares dentro de esta integración. Bautista (2022) señala que la investigación cualitativa se orienta a comprender significados, experiencias y contextos desde fundamentos epistemológicos propios. La inteligencia artificial puede apoyar la codificación preliminar, la organización de testimonios y la detección de recurrencias discursivas. Sin embargo, no debe reducir la interpretación a frecuencias o patrones léxicos. La integración con métodos cuantitativos requiere conservar la profundidad contextual y evitar que los resultados narrativos se transformen únicamente en variables desprovistas de significado.

En los estudios cuantitativos, la combinación de bases procedentes de distintas disciplinas puede ampliar el poder explicativo y predictivo de los modelos. Los algoritmos inteligentes permiten analizar relaciones complejas y múltiples variables dentro de grandes conjuntos de datos (Hua et al., 2023). No obstante, la integración exige revisar escalas de medición, unidades, criterios de inclusión y calidad de los registros. Si estas

diferencias no se controlan, el modelo puede detectar patrones derivados de inconsistencias técnicas y no de relaciones reales entre los fenómenos examinados.

Los lenguajes científicos también reflejan distintas formas de construir conocimiento. Una ecuación, una categoría jurídica, una narrativa histórica o un indicador clínico no poseen la misma función epistemológica. Krenn et al. (2022) sostienen que la comprensión científica requiere relacionar resultados con explicaciones y estructuras conceptuales. La inteligencia artificial puede traducir, resumir o vincular expresiones especializadas, pero no garantiza equivalencia entre ellas. La integración rigurosa demanda interpretar qué representa cada lenguaje y cómo puede contribuir al problema interdisciplinario sin perder su sentido original.

La visualización constituye un recurso importante para reunir lenguajes y datos heterogéneos. Mediante grafos, mapas de relaciones y representaciones dinámicas, los equipos pueden observar conexiones entre variables, conceptos y fuentes. Estas herramientas facilitan la comunicación entre especialistas y permiten discutir hallazgos desde perspectivas diferentes. Sin embargo, una representación visual puede simplificar excesivamente relaciones complejas o sugerir cercanías que no poseen fundamento causal. La claridad gráfica debe acompañarse de criterios metodológicos y explicaciones sobre el alcance de cada vínculo representado (Ávila, 2024).

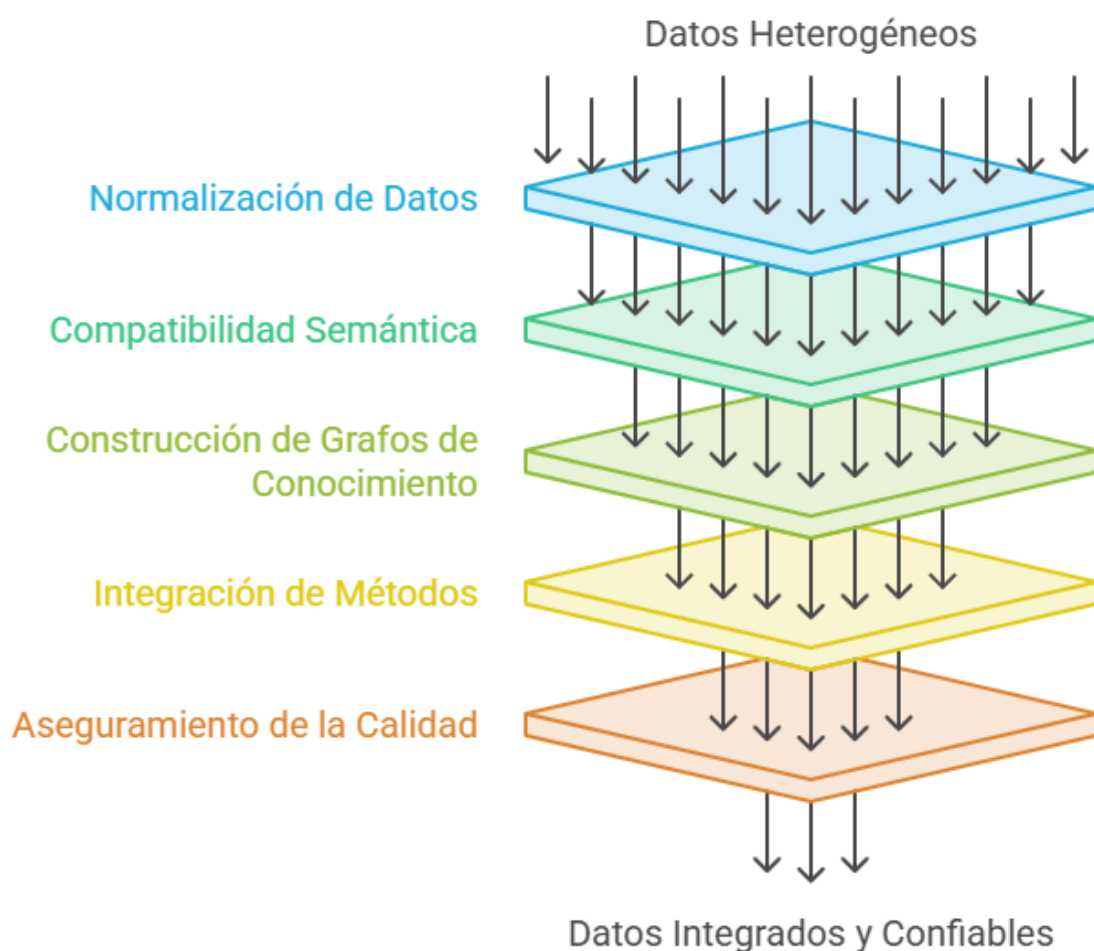
La integración interdisciplinaria también enfrenta problemas de trazabilidad. Cuando los datos atraviesan procesos de limpieza, transformación, codificación y análisis automatizado, puede resultar difícil reconstruir cómo se produjo un resultado. Semmelrock et al. (2025) identifican la documentación insuficiente como una barrera para la reproducibilidad en investigaciones basadas en aprendizaje automático. Por ello, cada modificación debe registrarse junto con sus criterios, responsables y efectos. La trazabilidad permite distinguir entre decisiones disciplinares, transformaciones técnicas y resultados generados por el sistema inteligente.

La calidad de la integración depende igualmente de la gobernanza de los datos. No todas las disciplinas manejan los mismos niveles de sensibilidad, acceso o restricción. En salud, educación o ciencias sociales, la combinación de bases puede incrementar riesgos de reidentificación y uso indebido. Andrade et al. (2024) examinan la relación entre inteligencia artificial y protección de datos personales desde una perspectiva legal. La integración científica debe incorporar principios de proporcionalidad, seguridad y control

de acceso, además de definir responsabilidades sobre almacenamiento, tratamiento y reutilización.

Figura 3

Proceso de integración de datos interdisciplinarios



Nota. La figura muestra el proceso de transformación de datos heterogéneos en información integrada y confiable mediante normalización, compatibilidad semántica, grafos de conocimiento y control de calidad. Elaboración propia.

Los sesgos pueden trasladarse de un campo a otro cuando se integran datos o modelos sin suficiente revisión contextual. Omiye et al. (2023) evidencian que los grandes modelos de lenguaje pueden reproducir enfoques médicos asociados con categorías raciales. Este problema muestra que una clasificación aceptada en determinado corpus puede resultar discriminatoria o inadecuada en otro contexto. La investigación interdisciplinaria requiere examinar el origen de las categorías, la representatividad de los datos y las consecuencias que tendría aplicar una misma lógica analítica en poblaciones diferentes.

La integración de datos, métodos y lenguajes alcanza valor científico cuando permite construir explicaciones más amplias sin borrar las diferencias entre disciplinas. Morin

(1999) plantea que el conocimiento complejo exige articular elementos diversos y reconocer sus interdependencias. La inteligencia artificial puede facilitar conexiones, organizar información y ampliar las posibilidades analíticas, pero no reemplaza los acuerdos epistemológicos y metodológicos del equipo. La coherencia final depende de una interpretación conjunta que preserve la especificidad de cada aporte y fortalezca la comprensión del problema compartido.

2.3 Convergencia entre ciencias naturales, sociales, humanas y tecnológicas

La convergencia entre ciencias naturales, sociales, humanas y tecnológicas responde a la necesidad de estudiar fenómenos que combinan dimensiones materiales, culturales, históricas, institucionales y computacionales. Ninguna disciplina puede explicar por sí sola problemas como la salud pública, el cambio tecnológico, la educación digital o la transformación productiva. Morin (1999) sostiene que el conocimiento complejo requiere vincular elementos diversos sin reducir sus diferencias. Desde esta perspectiva, la inteligencia artificial actúa como una infraestructura de conexión que facilita integrar datos, métodos y marcos interpretativos.

Las ciencias naturales aportan modelos de observación, experimentación y explicación de procesos físicos, biológicos o ambientales. Su convergencia con la inteligencia artificial amplía la capacidad de analizar grandes volúmenes de datos, reconocer regularidades y simular comportamientos complejos. Hua et al. (2023) destacan que los algoritmos inteligentes fortalecen el cómputo científico y la resolución de problemas de elevada complejidad. Sin embargo, los resultados computacionales deben mantenerse vinculados con mecanismos observables, condiciones experimentales y criterios de validación propios de cada campo científico.

Las ciencias sociales incorporan el análisis de instituciones, relaciones de poder, desigualdades y comportamientos colectivos. Su participación resulta indispensable cuando una innovación tecnológica afecta comunidades, modifica decisiones o redistribuye oportunidades. La inteligencia artificial puede procesar encuestas, registros administrativos y discursos, pero esos datos requieren interpretación histórica y contextual. Los sistemas algorítmicos no operan al margen de la sociedad, porque reflejan categorías, prioridades y sesgos presentes en sus datos y diseños (Kuppler et al., 2022; Kontiainen et al., 2022).

Las humanidades contribuyen con herramientas para examinar significados, valores, narrativas, lenguajes y experiencias humanas. Su convergencia con áreas tecnológicas permite cuestionar cómo se construyen las categorías empleadas por los sistemas inteligentes y qué concepciones de persona, conocimiento o verdad contienen. Zemelman (2021) resalta la importancia del pensamiento epistémico para problematizar la realidad sin quedar restringido por esquemas teóricos cerrados. Esta capacidad crítica evita que la integración interdisciplinaria se limite a adoptar modelos computacionales como representaciones neutrales e incuestionables.

Las disciplinas tecnológicas proporcionan arquitecturas, algoritmos, infraestructuras de datos y mecanismos de automatización. Su función dentro de la convergencia no consiste únicamente en suministrar herramientas, sino en participar en la construcción metodológica del problema. Los gemelos digitales, por ejemplo, integran modelos computacionales y representaciones de procesos físicos para simular escenarios y apoyar decisiones (Qian et al., 2022). Su desarrollo requiere conocimientos de ingeniería, informática y ciencias aplicadas, además de comprensión social sobre los contextos donde serán implementados y evaluados.

La convergencia alcanza mayor profundidad cuando existe un objeto común de investigación y no solo una yuxtaposición de aportes disciplinares. Duedra et al. (2020) proponen avanzar desde la compartimentalización hacia una ecología de saberes capaz de relacionar ciencia básica, aplicada, extensión y divulgación. Esta orientación exige construir preguntas compartidas, acordar conceptos y definir cómo contribuirá cada disciplina. La inteligencia artificial puede facilitar conexiones entre información heterogénea, pero no sustituye el diálogo necesario para establecer una estructura científica verdaderamente integrada.

Tabla 5
Contribuciones disciplinares a la investigación convergente asistida por inteligencia artificial

Campo disciplinar	Contribución principal	Aplicación de la inteligencia artificial	Riesgo de una integración insuficiente
Ciencias naturales	Explicación de procesos físicos, biológicos y ambientales	Modelado, predicción, clasificación y simulación	Confundir correlaciones con mecanismos causales

Campo disciplinar	Contribución principal	Aplicación de la inteligencia artificial	Riesgo de una integración insuficiente
Ciencias sociales	Análisis de instituciones, comportamientos y desigualdades	Procesamiento de datos sociales y detección de tendencias	Descontextualizar categorías y relaciones sociales
Humanidades	Interpretación de significados, valores y narrativas	Minería textual, análisis semántico y organización documental	Reducir experiencias humanas a patrones lingüísticos
Ciencias tecnológicas	Desarrollo de algoritmos, sistemas e infraestructuras	Automatización, integración y procesamiento computacional	Priorizar eficiencia sobre pertinencia científica
Metodología	Articulación entre preguntas, métodos y evidencia	Comparación de diseños y apoyo a decisiones	Combinar técnicas sin coherencia epistemológica
Ética y gobernanza	Evaluación de efectos, responsabilidades y derechos	Auditoría, trazabilidad y supervisión de sistemas	Naturalizar sesgos, opacidad o distribución desigual de riesgos

Nota. La tabla sintetiza los aportes de cada campo y muestra que la convergencia requiere complementariedad conceptual, metodológica y ética, no únicamente integración tecnológica.

La salud constituye un campo representativo de esta convergencia. Los sistemas inteligentes pueden analizar imágenes, registros clínicos, lenguaje y datos biomédicos, pero su aplicación exige considerar necesidades profesionales, interpretabilidad y consecuencias para los pacientes. Bienefeld et al. (2023) plantean que la inteligencia artificial explicable debe conectar los objetivos de los desarrolladores con las necesidades de los clínicos. Esta articulación demuestra que la eficacia técnica no es suficiente si el sistema no puede integrarse con prácticas médicas, responsabilidades institucionales y comprensión humana.

En educación, la convergencia reúne conocimientos pedagógicos, psicológicos, sociales, lingüísticos y tecnológicos. La inteligencia artificial puede apoyar análisis de información, organización de contenidos y procesos de investigación universitaria, pero sus resultados deben interpretarse según contextos institucionales y formativos. Cedeño et al. (2024) señalan que estas herramientas amplían oportunidades para la actividad investigativa académica. No obstante, una aplicación educativa rigurosa requiere comprender las prácticas docentes, las desigualdades de acceso y los efectos que la automatización puede producir sobre el aprendizaje y la autonomía.

La investigación ambiental y productiva también se beneficia de la convergencia entre modelado científico, ingeniería y análisis social. Los gemelos digitales permiten representar sistemas, evaluar cambios y anticipar comportamientos en sectores diversos (Singh et al., 2022). Sin embargo, una simulación tecnológicamente precisa puede resultar insuficiente si excluye impactos territoriales, económicos o comunitarios. La toma de decisiones requiere combinar evidencia física con análisis de costos, capacidades institucionales y consecuencias distributivas. La inteligencia artificial aporta escenarios, pero su valoración depende de criterios provenientes de varias disciplinas.

La integración de lenguajes científicos constituye uno de los principales retos de la convergencia. Un indicador cuantitativo, una categoría sociológica, una interpretación histórica y un modelo computacional no expresan el mismo tipo de evidencia. Krenn et al. (2022) explican que la comprensión científica supone relacionar resultados con explicaciones y estructuras conceptuales. Por ello, los equipos deben establecer equivalencias prudentes, reconocer incompatibilidades y evitar traducciones simplificadoras. La inteligencia artificial puede organizar términos y detectar conexiones, aunque no resuelve por sí sola las diferencias epistemológicas.

La convergencia también exige revisar cómo se distribuyen autoridad y responsabilidad dentro de los equipos. Los especialistas técnicos no deben decidir aisladamente qué categorías sociales utilizar, del mismo modo que los expertos disciplinares necesitan comprender los límites de los modelos computacionales. Contreras et al. (2024) destacan que el uso de inteligencia artificial en investigación plantea obligaciones relacionadas con transparencia, supervisión y responsabilidad. Estas exigencias deben asumirse colectivamente, mediante protocolos que definan quién selecciona datos, valida resultados, interpreta hallazgos y responde ante posibles consecuencias negativas.

Los sesgos evidencian la necesidad de integrar análisis técnico, social y humanístico. Omiye et al. (2023) muestran que los grandes modelos de lenguaje pueden reproducir enfoques médicos basados en categorías raciales. Un problema de este tipo no puede resolverse únicamente mediante ajustes algorítmicos, porque también involucra historia, lenguaje, desigualdad y prácticas institucionales. La convergencia permite examinar simultáneamente el funcionamiento del sistema, el origen de las categorías y sus efectos sobre grupos específicos, fortaleciendo una evaluación más completa de la innovación.

La convergencia entre ciencias naturales, sociales, humanas y tecnológicas amplía la capacidad de comprender problemas complejos y diseñar respuestas científicamente pertinentes. Su valor no radica en reunir más datos o disciplinas, sino en construir relaciones justificadas entre evidencias, conceptos y métodos. La inteligencia artificial facilita procesamiento, conexión y simulación, pero la coherencia final depende del trabajo interpretativo del equipo. Una investigación convergente sólida preserva la especificidad disciplinar, establece responsabilidades compartidas y somete las soluciones tecnológicas a evaluación científica, social y ética.

2.4 Modelos colaborativos para equipos científicos interdisciplinarios

Los modelos colaborativos para equipos científicos interdisciplinarios deben organizar la participación de especialistas con conocimientos, lenguajes y responsabilidades diferentes. La inteligencia artificial puede facilitar esta coordinación mediante sistemas de gestión documental, clasificación temática, análisis compartido y seguimiento de decisiones. Sin embargo, la colaboración no se produce únicamente por utilizar una plataforma común. Requiere objetivos definidos, acuerdos metodológicos y criterios explícitos para integrar aportes diversos dentro de una misma investigación (Cedeño et al., 2024).

Un modelo colaborativo eficaz comienza con la delimitación conjunta del problema. Cada disciplina debe identificar qué dimensión puede explicar, qué datos necesita y cuáles son los límites de su intervención. Morin (1999) plantea que el conocimiento complejo exige relacionar elementos heterogéneos sin reducirlos a una sola perspectiva. En consecuencia, la inteligencia artificial debe utilizarse como apoyo para conectar información y no como mecanismo que imponga uniformidad conceptual. La construcción inicial de acuerdos evita contradicciones posteriores entre métodos, variables y formas de interpretar resultados.

La distribución de funciones constituye otro componente central. Los equipos pueden incluir expertos disciplinares, metodólogos, analistas de datos, desarrolladores y responsables de ética o gobernanza. Cada participante debe conocer qué decisiones le corresponden y cuáles requieren deliberación colectiva. Contreras et al. (2024) destacan que el uso de inteligencia artificial en investigación exige transparencia, supervisión y responsabilidad. Estas obligaciones deben traducirse en protocolos que definan quién

selecciona los datos, valida los modelos, interpreta los hallazgos y responde ante posibles errores.

Los sistemas inteligentes pueden mejorar la comunicación al organizar documentos, resumir discusiones y detectar términos ambiguos entre disciplinas. Garzón et al. (2025) señalan que el procesamiento del lenguaje natural favorece la interacción entre humanos y sistemas de inteligencia artificial. Dentro de los equipos científicos, esta capacidad puede apoyar la construcción de glosarios compartidos y la identificación de conceptos equivalentes. No obstante, las traducciones automáticas deben revisarse, porque una semejanza lingüística no garantiza que dos disciplinas atribuyan el mismo significado a una categoría.

La colaboración también requiere mecanismos para integrar datos heterogéneos. Ávila (2024) muestra que los grafos de conocimiento permiten representar relaciones entre datos, conceptos y fuentes de investigación. Estas estructuras pueden funcionar como espacios comunes donde cada disciplina incorpora información bajo criterios previamente acordados. Su utilidad depende de la calidad de los metadatos, la claridad de las ontologías y la trazabilidad de cada vínculo. Sin estas condiciones, el sistema puede conectar elementos técnicamente compatibles, pero científicamente incongruentes.

Polo y Casique (2025) proponen combinar grafos de conocimiento y redes neuronales para mejorar la recuperación de información documental. Este enfoque puede fortalecer el trabajo colaborativo al permitir que los integrantes localicen evidencia relevante desde vocabularios distintos. La recuperación compartida reduce duplicaciones y facilita la comparación de antecedentes. Sin embargo, el equipo debe acordar criterios de inclusión, pertinencia y jerarquización de fuentes, porque una búsqueda automatizada amplia puede introducir documentos poco compatibles con el problema o con el enfoque metodológico adoptado.

Los modelos colaborativos deben preservar la especificidad epistemológica de cada disciplina. Bautista (2022) explica que la investigación cualitativa se orienta a comprender significados, experiencias y contextos mediante fundamentos propios. Sus resultados no deben subordinarse automáticamente a métricas cuantitativas ni traducirse en variables descontextualizadas. De igual manera, los análisis predictivos requieren criterios de validación que no siempre son aplicables a interpretaciones narrativas. La

colaboración rigurosa reconoce estas diferencias y define cómo se relacionarán los distintos tipos de evidencia.

La toma de decisiones puede organizarse mediante esquemas centralizados, distribuidos o híbridos. En un modelo centralizado, una coordinación principal integra los aportes; en uno distribuido, cada área conserva mayor autonomía; y en el modelo híbrido se combinan ambos mecanismos según la fase del estudio. La inteligencia artificial puede apoyar cualquiera de estas estructuras mediante alertas, registros y análisis comparativos. Su utilidad depende de que las decisiones relevantes permanezcan documentadas y puedan revisarse colectivamente (Semmelrock et al., 2025).

La evaluación de resultados debe realizarse desde múltiples perspectivas. Un modelo puede mostrar buen rendimiento computacional y, al mismo tiempo, producir interpretaciones débiles o consecuencias sociales problemáticas. Krenn et al. (2022) sostienen que la comprensión científica requiere relacionar resultados con explicaciones y estructuras conceptuales. Por ello, cada hallazgo debe ser examinado por especialistas técnicos y disciplinares. La validación conjunta permite detectar relaciones espurias, errores de contexto y conclusiones que exceden el alcance real de los datos analizados.

La equidad dentro del equipo también influye en la calidad de la colaboración. La concentración de recursos, infraestructura o conocimiento técnico puede otorgar mayor poder de decisión a ciertos participantes. Ross et al. (2022) advierten que la ciencia abierta puede reproducir ventajas acumulativas y amenazas a la equidad. En proyectos interdisciplinarios, resulta necesario garantizar acceso compartido a datos, herramientas y procesos de interpretación. La colaboración se debilita cuando algunos integrantes ejecutan tareas operativas sin participar en las decisiones intelectuales centrales.

Los modelos colaborativos alcanzan su mayor eficacia cuando combinan coordinación institucional, autonomía disciplinar y responsabilidad compartida. La inteligencia artificial puede organizar información, facilitar comunicación y apoyar análisis complejos, pero no reemplaza la deliberación científica. Cedeño et al. (2024) destacan que su incorporación en la investigación universitaria requiere capacidades y criterios de uso. Un equipo interdisciplinario sólido establece reglas claras, documenta sus decisiones y reconoce que la calidad del conocimiento depende tanto de la tecnología como de la cooperación entre sus integrantes.

2.5 Aplicaciones interdisciplinarias en salud, educación, ingeniería y ciencias ambientales

Las aplicaciones interdisciplinarias de inteligencia artificial adquieren especial relevancia cuando los problemas científicos combinan variables biológicas, sociales, técnicas y ambientales. Su utilidad no depende únicamente de procesar grandes volúmenes de información, sino de integrar conocimientos procedentes de campos distintos dentro de una misma estrategia analítica. Krenn et al. (2022) sostienen que la comprensión científica exige relacionar resultados, explicaciones y estructuras conceptuales. Por ello, cada aplicación debe articular capacidades computacionales con criterios disciplinares, metodológicos y contextuales claramente definidos.

En salud, la inteligencia artificial puede apoyar diagnóstico, clasificación de riesgos, síntesis de evidencia y análisis de registros clínicos. Bienefeld et al. (2023) destacan que la explicabilidad debe responder tanto a las necesidades de los profesionales sanitarios como a los objetivos de los desarrolladores. Esta articulación resulta indispensable porque una predicción técnicamente precisa puede ser poco útil si no puede interpretarse dentro de la práctica clínica. La aplicación interdisciplinaria requiere combinar conocimientos médicos, estadísticos, informáticos, éticos y organizacionales para validar cada resultado.

El procesamiento del lenguaje natural también ofrece aplicaciones relevantes en salud mental. Swaminathan et al. (2023) desarrollaron un sistema orientado a detectar rápidamente mensajes vinculados con situaciones de crisis. Este tipo de herramienta integra informática, psicología, lingüística y atención sanitaria para reconocer señales que podrían requerir intervención. Sin embargo, su uso exige controlar falsos positivos, privacidad y posibles errores de interpretación. La eficacia del sistema debe evaluarse junto con protocolos humanos que determinen cómo actuar ante una alerta y cómo proteger a las personas involucradas.

En educación, la inteligencia artificial puede apoyar revisión documental, análisis de información, producción académica y organización de procesos investigativos. Cedeño et al. (2024) señalan que su incorporación en el ámbito universitario amplía posibilidades de búsqueda, análisis y escritura científica. No obstante, una aplicación educativa rigurosa debe considerar factores pedagógicos, institucionales y sociales. La tecnología no puede evaluarse únicamente por su eficiencia operativa, porque también influye en la autonomía

del estudiante, la integridad académica y las capacidades del profesorado para supervisar su uso.

Gallent et al. (2023) advierten que la inteligencia artificial generativa plantea desafíos éticos y de integridad académica en la educación superior. Esta observación muestra que las aplicaciones interdisciplinarias deben integrar conocimientos pedagógicos, tecnológicos y normativos desde el inicio. Un sistema capaz de generar textos o retroalimentación puede apoyar el aprendizaje, pero también facilitar dependencia, atribuciones incorrectas o prácticas poco transparentes. Su implementación requiere criterios claros sobre autoría, verificación, responsabilidad y uso formativo de los contenidos producidos mediante herramientas generativas.

Tabla 6
Aplicaciones interdisciplinarias de inteligencia artificial en campos científicos seleccionados

Campo	Aplicación principal	Disciplinas que convergen	Aporte científico	Riesgo crítico
Salud	Diagnóstico, predicción y síntesis clínica	Medicina, informática, estadística y ética	Mejora del análisis y apoyo a decisiones	Sesgos, opacidad y errores clínicos
Educación	Análisis académico y generación de contenidos	Pedagogía, lingüística, informática y psicología	Personalización y apoyo investigativo	Dependencia e integridad académica
Ingeniería	Simulación, automatización y mantenimiento predictivo	Ingeniería, ciencia de datos y gestión	Optimización de procesos complejos	Dependencia técnica y fallas sistémicas
Ciencias ambientales	Modelado, monitoreo y predicción	Ecología, geografía, informática y estadística	Análisis integrado de sistemas naturales	Datos incompletos y simplificación ecológica
Investigación transversal	Integración de datos y descubrimiento de patrones	Metodología, computación y campos aplicados	Generación de conocimiento convergente	Incompatibilidad conceptual y metodológica

Nota. La tabla sintetiza aplicaciones interdisciplinarias y muestra que su valor depende de integrar competencias científicas, técnicas y éticas, además de controlar los riesgos específicos de cada campo.

En ingeniería, los gemelos digitales constituyen una de las aplicaciones más representativas. Estas arquitecturas permiten reproducir virtualmente objetos, procesos o sistemas para analizar su funcionamiento y anticipar posibles cambios. Qian et al. (2022)

explican que una réplica cibernética puede apoyar monitoreo, simulación y toma de decisiones. Su desarrollo exige combinar ingeniería, modelado, sensores, informática y análisis de datos. La calidad de la aplicación depende de la correspondencia entre el modelo digital y el sistema físico, así como de la actualización continua de la información.

Las aplicaciones industriales de gemelos digitales muestran que la inteligencia artificial puede intervenir en mantenimiento, diseño, producción y evaluación de escenarios. Singh et al. (2022) identifican usos en diferentes sectores donde la simulación favorece optimización y predicción. Sin embargo, trasladar estas capacidades a la investigación científica requiere controlar supuestos, parámetros y márgenes de error. Una simulación puede orientar decisiones experimentales, pero no reemplaza la validación empírica. La colaboración entre ingenieros, analistas y especialistas del dominio resulta necesaria para interpretar correctamente los resultados obtenidos.

En ciencias ambientales, la inteligencia artificial puede integrar imágenes, sensores, series temporales y registros territoriales para reconocer patrones y anticipar variaciones. Los algoritmos inteligentes amplían la capacidad de analizar sistemas complejos donde interactúan múltiples factores naturales y humanos (Hua et al., 2023). No obstante, la modelización ambiental debe considerar incertidumbre, cambios de escala y calidad desigual de los datos. Una predicción precisa en condiciones controladas puede perder validez cuando se aplica a territorios con dinámicas ecológicas, sociales o climáticas diferentes.

La investigación ambiental también puede beneficiarse de modelos de simulación y representaciones digitales de procesos físicos. Estas herramientas permiten explorar escenarios antes de intervenir sobre ecosistemas o infraestructuras sensibles. Sin embargo, su interpretación exige conocimientos de ecología, ingeniería, estadística y ciencias sociales, porque las decisiones ambientales afectan simultáneamente recursos, comunidades y actividades productivas. La integración interdisciplinaria evita que el análisis se limite a variables técnicas e incorpora consecuencias territoriales, institucionales y distributivas dentro de la valoración científica (Morin, 1999).

Las aplicaciones interdisciplinarias de inteligencia artificial alcanzan mayor valor cuando fortalecen la comprensión del problema y no solo la eficiencia del procesamiento. En salud, educación, ingeniería y ciencias ambientales, los sistemas inteligentes pueden apoyar análisis, predicción, simulación y organización de evidencia. Sin embargo, su

incorporación requiere trazabilidad, interpretabilidad y validación contextual. Contreras et al. (2024) señalan que el uso científico de estas tecnologías debe acompañarse de responsabilidad y supervisión. La innovación sostenible surge de combinar especialización disciplinar, cooperación metodológica y control humano.

2.6 Barreras epistemológicas, metodológicas e institucionales de la interdisciplinariedad

Las barreras epistemológicas de la interdisciplinariedad aparecen cuando las disciplinas sostienen concepciones distintas sobre qué constituye conocimiento válido, cómo debe producirse y qué criterios permiten justificarlo. Cabrera y Cepeda (2022) señalan que la epistemología orienta la construcción científica mediante principios de validez, fundamento y alcance. En proyectos asistidos por inteligencia artificial, estas diferencias pueden intensificarse si los resultados algorítmicos son aceptados como evidencia equivalente en todos los campos, sin considerar las particularidades conceptuales y explicativas de cada tradición disciplinar.

Las ciencias naturales suelen privilegiar medición, experimentación y explicación causal, mientras las ciencias sociales y humanas incorporan interpretación, historicidad y análisis contextual. Bautista (2022) destaca que la investigación cualitativa se fundamenta en la comprensión de significados, experiencias y realidades situadas. La dificultad surge cuando un equipo interdisciplinario intenta traducir todos los hallazgos a una única lógica analítica. La inteligencia artificial puede integrar información heterogénea, pero no debe homogeneizar los supuestos epistemológicos que sostienen cada tipo de evidencia.

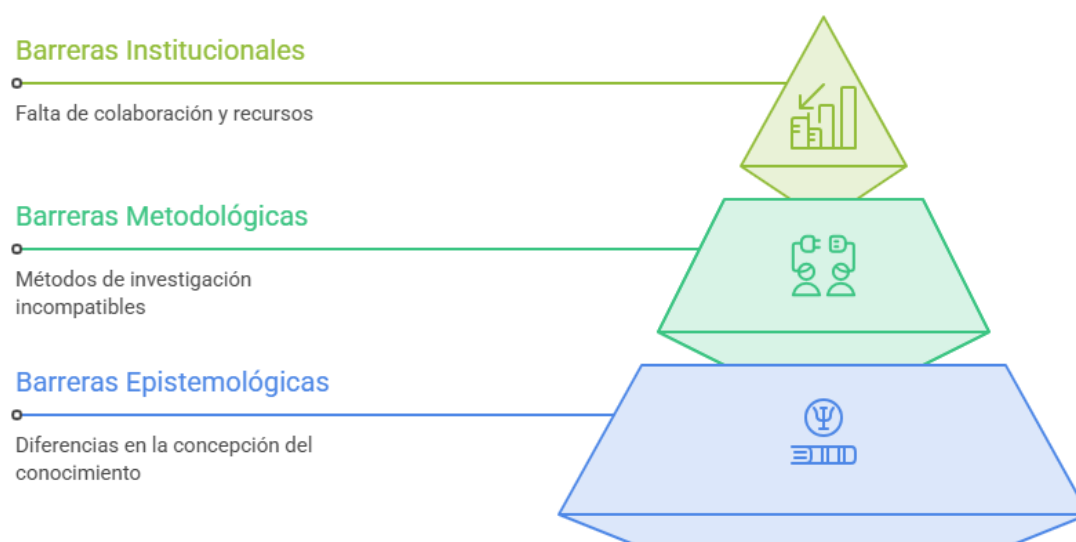
Otra barrera consiste en confundir integración con acumulación. Reunir especialistas, bases de datos y técnicas distintas no garantiza una investigación interdisciplinaria. Duedra et al. (2020) cuestionan la compartimentalización del conocimiento y proponen una ecología de saberes capaz de vincular ciencia básica, aplicada, extensión y divulgación. Sin acuerdos sobre el problema, los conceptos y la interpretación, los aportes permanecen aislados. Los sistemas inteligentes pueden facilitar conexiones documentales, aunque la articulación científica depende de una construcción teórica compartida.

Las barreras metodológicas se manifiestan cuando los métodos empleados responden a preguntas, escalas o unidades de análisis incompatibles. Falcón y Serpa (2021) explican que los métodos teóricos y empíricos cumplen funciones diferentes dentro del proceso

investigativo. En un estudio interdisciplinario, combinar análisis estadístico, interpretación cualitativa, simulación y minería textual exige definir cómo se relacionarán los resultados. Sin esta planificación, la inteligencia artificial puede producir asociaciones técnicamente correctas, pero desconectadas de la lógica metodológica general del proyecto.

La integración de datos también enfrenta problemas de calidad, comparabilidad y procedencia. Los registros pueden utilizar escalas distintas, contener valores faltantes o responder a criterios de clasificación propios de cada disciplina. Cuando estas diferencias no se documentan, los algoritmos pueden identificar patrones originados por inconsistencias técnicas y no por relaciones científicas reales. La reproducibilidad exige registrar transformaciones, decisiones y parámetros utilizados durante el procesamiento, especialmente en investigaciones basadas en aprendizaje automático (Kapoor y Narayanan, 2023; Semmelrock et al., 2025).

Figura 4
Jerarquía de barreras interdisciplinarias



Nota. La figura organiza las barreras interdisciplinarias en niveles epistemológicos, metodológicos e institucionales. Elaboración propia.

El lenguaje especializado constituye una barrera adicional. Un mismo término puede expresar significados diferentes entre disciplinas, mientras conceptos equivalentes pueden aparecer bajo denominaciones distintas. Garzón et al. (2025) destacan que el procesamiento del lenguaje natural mejora la interacción entre seres humanos y sistemas inteligentes. Sin embargo, la traducción automatizada no resuelve por sí sola los desacuerdos conceptuales. Los equipos deben construir glosarios compartidos, definir

categorías operativas y revisar las equivalencias propuestas para evitar interpretaciones ambiguas o relaciones semánticas artificiales.

Las barreras institucionales se relacionan con estructuras universitarias organizadas por departamentos, carreras y líneas de investigación separadas. Estas divisiones influyen en la asignación de recursos, reconocimiento académico y evaluación de resultados. Schuff y Hubert (2024) muestran que las agendas de investigación y financiamiento condicionan las prioridades científicas. Cuando los incentivos favorecen productos disciplinares tradicionales, los equipos interdisciplinarios enfrentan dificultades para sostener proyectos de largo plazo, compartir infraestructura y distribuir de manera equitativa los méritos obtenidos.

El financiamiento puede orientar la interdisciplinariedad hacia áreas consideradas estratégicas, aunque también limitar la autonomía de los equipos. Fochler et al. (2025) analizan cómo los recursos económicos influyen en la gobernanza de la investigación universitaria. La dependencia de convocatorias, empresas o prioridades gubernamentales puede definir problemas, métodos y resultados esperados. En proyectos con inteligencia artificial, esta situación se intensifica por los costos de infraestructura, datos y personal especializado, lo que puede concentrar capacidades en pocas instituciones y reducir la diversidad científica.

La desigualdad tecnológica constituye otra barrera institucional importante. No todas las universidades cuentan con acceso equivalente a plataformas, capacidad computacional, bases de datos o especialistas en inteligencia artificial. Ross et al. (2022) advierten que la ciencia abierta puede reproducir ventajas acumulativas y amenazas a la equidad. La interdisciplinariedad puede quedar restringida a instituciones con mayores recursos, mientras otros equipos participan como proveedores de datos o contextos locales sin intervenir plenamente en las decisiones metodológicas, interpretativas y autorales.

Las diferencias en responsabilidad y autoridad también generan conflictos. Los expertos técnicos pueden controlar los modelos y procedimientos computacionales, mientras los especialistas disciplinares poseen el conocimiento necesario para interpretar los resultados. Contreras et al. (2024) señalan que el uso científico de inteligencia artificial exige transparencia, supervisión y responsabilidad. La ausencia de protocolos claros puede dificultar la identificación de quién seleccionó los datos, modificó los parámetros

o validó las conclusiones. Una gobernanza colaborativa debe distribuir funciones sin fragmentar la responsabilidad científica final.

Superar estas barreras requiere acuerdos epistemológicos, metodológicos e institucionales contruidos desde el inicio del proyecto. Morin (1999) sostiene que el pensamiento complejo debe relacionar conocimientos diversos sin eliminar sus diferencias. La inteligencia artificial puede apoyar integración, comunicación y análisis, pero no reemplaza la deliberación entre disciplinas. La interdisciplinariedad se fortalece cuando existen objetivos comunes, documentación compartida, reconocimiento equitativo y espacios institucionales estables. Sin estas condiciones, la tecnología puede ampliar capacidades operativas mientras mantiene intactas las fragmentaciones tradicionales de la ciencia.



Capítulo

3.

Tecnologías emergentes para la investigación científica

Franklin Santiago Sosa Cifuentes

3.1 Aprendizaje automático y aprendizaje profundo en el análisis científico

El aprendizaje automático permite que los sistemas identifiquen regularidades en los datos y ajusten su desempeño mediante experiencia computacional, sin depender exclusivamente de reglas programadas de forma explícita. En la investigación científica, esta capacidad favorece la clasificación, la predicción, la detección de anomalías y el análisis de relaciones complejas. Hua et al. (2023) señalan que los algoritmos inteligentes amplían las posibilidades del cómputo científico. Su utilidad, sin embargo, depende de la calidad de los datos, de la selección metodológica y de la interpretación realizada por el investigador.

El aprendizaje profundo constituye una rama del aprendizaje automático basada en redes neuronales con múltiples capas capaces de construir representaciones jerárquicas. Esta arquitectura resulta especialmente útil para analizar imágenes, señales, lenguaje y conjuntos de datos de elevada dimensionalidad. A diferencia de modelos más simples, puede aprender características relevantes directamente desde la información disponible. No obstante, esa capacidad incrementa las exigencias de procesamiento y dificulta comprender cómo se generan determinadas predicciones, lo que plantea desafíos de interpretabilidad y control científico (Joyce et al., 2023; Hu, 2021).

La elección entre aprendizaje automático y aprendizaje profundo debe responder al problema científico, al volumen de información y al tipo de evidencia disponible. Los modelos tradicionales suelen funcionar adecuadamente con datos estructurados, muestras moderadas y variables previamente definidas, mientras las redes profundas demandan conjuntos extensos y recursos computacionales elevados. Krenn et al. (2022) sostienen que el rendimiento algorítmico adquiere valor científico cuando se relaciona con explicaciones comprensibles. Por ello, la complejidad técnica no debe considerarse automáticamente sinónimo de superioridad metodológica.

En ciencias de la salud, estos enfoques permiten analizar historiales clínicos, imágenes, textos médicos y señales fisiológicas. Swaminathan et al. (2023) aplicaron procesamiento del lenguaje natural para detectar mensajes asociados con crisis de salud mental, mostrando el potencial de los modelos computacionales para apoyar intervenciones rápidas. Sin embargo, una clasificación incorrecta puede producir consecuencias sensibles. La evaluación debe considerar precisión, tasas de error, representatividad de

los datos y capacidad de integrar los resultados dentro de decisiones profesionales supervisadas.

El aprendizaje profundo puede descubrir patrones difíciles de identificar mediante procedimientos estadísticos convencionales, pero esa ventaja no elimina el riesgo de reproducir sesgos presentes en los datos. Omiye et al. (2023) evidencian que los grandes modelos de lenguaje pueden propagar enfoques médicos basados en categorías raciales. Este hallazgo demuestra que un alto rendimiento promedio puede ocultar efectos desiguales entre grupos. La evaluación científica debe incorporar análisis desagregados, revisión de categorías y controles que permitan detectar resultados potencialmente discriminatorios o metodológicamente injustificados.

Tabla 7
Comparación entre aprendizaje automático y aprendizaje profundo en el análisis científico

Criterio	Aprendizaje automático	Aprendizaje profundo
Base metodológica	Algoritmos que aprenden relaciones a partir de variables definidas	Redes neuronales con múltiples capas de representación
Tipo de datos	Principalmente datos estructurados y tabulares	Imágenes, audio, texto, señales y datos de alta dimensionalidad
Preparación de variables	Requiere selección y diseño previo de características	Puede aprender características directamente desde los datos
Volumen de datos	Puede funcionar con conjuntos moderados	Generalmente necesita grandes cantidades de información
Recursos computacionales	Exigencia baja o media según el modelo	Exigencia elevada de procesamiento y almacenamiento
Interpretabilidad	Suele ofrecer mayor facilidad de explicación	Presenta mayor opacidad en decisiones internas
Aplicaciones científicas	Clasificación, regresión, predicción y detección de anomalías	Visión artificial, lenguaje natural y reconocimiento de patrones complejos
Riesgo principal	Sobreajuste, sesgos y selección inadecuada de variables	Opacidad, dependencia de datos masivos y altos costos técnicos

Nota. La comparación muestra que ambos enfoques son complementarios. Su elección debe fundamentarse en el problema científico, la naturaleza de los datos, la necesidad de explicación y los recursos disponibles.

La preparación de datos constituye una fase determinante para ambos enfoques. Los registros deben revisarse, limpiarse, transformar valores y documentar criterios de inclusión o exclusión. Cuando estas operaciones se realizan sin control, el modelo puede

aprender errores, duplicaciones o patrones irrelevantes. Kapoor y Narayanan (2023) advierten que la fuga de datos puede generar resultados artificialmente favorables y afectar la reproducibilidad. Por ello, los conjuntos de entrenamiento, validación y prueba deben mantenerse separados mediante procedimientos definidos antes de evaluar el rendimiento final.

La interpretabilidad adquiere especial importancia cuando los resultados algorítmicos influyen en conclusiones científicas o decisiones aplicadas. Zárate (2022) explica que la explicabilidad permite comprender los fundamentos de una decisión producida por inteligencia artificial. En aprendizaje automático, algunos modelos facilitan identificar la contribución de las variables, mientras las redes profundas presentan estructuras internas más difíciles de examinar. La selección debe considerar si el objetivo principal es maximizar precisión, comprender mecanismos o equilibrar ambas necesidades dentro del diseño investigativo.

La reproducibilidad exige documentar arquitectura, parámetros, versiones de software, procedimientos de entrenamiento y condiciones computacionales. Semmelrock et al. (2025) identifican barreras técnicas y organizacionales que dificultan repetir estudios basados en aprendizaje automático. Estas limitaciones se intensifican en aprendizaje profundo por la cantidad de configuraciones posibles y la dependencia de infraestructura especializada. Compartir código y datos puede favorecer la verificación, aunque también deben registrarse decisiones intermedias, semillas aleatorias y criterios empleados para seleccionar el modelo presentado como resultado final.

Ciobanu et al. (2024) señalan que la investigación médica basada en aprendizaje automático necesita combinar reproducibilidad e interpretabilidad para alcanzar resultados confiables. Este principio puede extenderse a cualquier campo científico. Un modelo no debe validarse únicamente mediante una métrica general, sino también mediante pruebas externas, análisis de estabilidad y revisión de su comportamiento frente a datos diferentes. La generalización constituye una condición esencial, porque un sistema puede funcionar adecuadamente en una muestra específica y fracasar al trasladarse a otros contextos.

El aprendizaje automático y el aprendizaje profundo amplían la capacidad científica para analizar datos complejos, pero no reemplazan la formulación teórica ni el juicio metodológico. Su incorporación debe justificarse por la naturaleza del problema y no por

la novedad de la tecnología. Contreras et al. (2024) destacan que el uso de inteligencia artificial en investigación exige transparencia, supervisión y responsabilidad. La innovación se produce cuando los modelos permiten descubrir conocimiento verificable, comprensible y reproducible, sin ocultar limitaciones ni trasladar decisiones críticas exclusivamente al sistema computacional.

3.2 Inteligencia artificial generativa aplicada a la producción académica

La inteligencia artificial generativa ha ampliado las posibilidades de la producción académica al intervenir en tareas de ideación, organización, síntesis y reformulación textual. A diferencia de otros sistemas orientados principalmente a clasificar o predecir, estos modelos producen contenidos nuevos a partir de patrones aprendidos en grandes volúmenes de datos. Su incorporación modifica los flujos tradicionales de escritura científica, aunque no reemplaza la responsabilidad intelectual del investigador. Díaz (2024) destaca que estas herramientas pueden apoyar diversas fases de la actividad investigativa cuando se utilizan con criterios metodológicos definidos.

En las etapas iniciales, los sistemas generativos pueden ayudar a ordenar ideas, proponer esquemas y explorar distintas formas de delimitar un tema. Esta capacidad resulta útil para transformar notas dispersas en una estructura preliminar y comparar posibles secuencias argumentativas. Sin embargo, las sugerencias obtenidas no deben aceptarse como decisiones científicas definitivas. La pertinencia del enfoque, la profundidad conceptual y la coherencia con el problema dependen del juicio del investigador, quien debe verificar que la organización propuesta responda a los objetivos del trabajo (Cedeño et al., 2024).

La generación de borradores constituye una de sus aplicaciones más visibles. Los modelos pueden producir párrafos, transiciones, resúmenes y versiones alternativas de una explicación, lo que reduce parte del esfuerzo mecánico de redacción. Ros y Pérez (2023) señalan que herramientas como ChatGPT pueden apoyar la escritura de artículos científicos, pero no deben reconocerse como autoras. Esta distinción es fundamental porque el sistema no asume responsabilidad sobre la veracidad, originalidad ni consecuencias del contenido, funciones que corresponden exclusivamente a quienes firman la publicación.

La síntesis automatizada puede facilitar la revisión de documentos extensos y la identificación preliminar de ideas centrales. No obstante, resumir evidencia científica

exige preservar matices, métodos, limitaciones y niveles de certeza. Tang et al. (2023) evaluaron modelos de lenguaje aplicados a la síntesis de evidencia médica y mostraron que su desempeño debe examinarse cuidadosamente. Una respuesta fluida puede omitir resultados contradictorios, exagerar conclusiones o mezclar afirmaciones procedentes de fuentes distintas, por lo que siempre requiere contraste directo con los textos originales.

Los modelos generativos también pueden apoyar la adaptación lingüística de contenidos académicos, mediante corrección gramatical, simplificación de expresiones o reformulación según diferentes públicos. Esta función favorece la claridad y la divulgación, especialmente cuando el investigador necesita comunicar hallazgos fuera de su disciplina. Sin embargo, una modificación estilística puede alterar conceptos técnicos, reducir precisión o introducir interpretaciones no previstas. La supervisión humana debe asegurar que cada versión conserve el significado original, la terminología especializada y la correspondencia con la evidencia científica presentada (Castillo, 2023).

Gao et al. (2023) demostraron que los resúmenes científicos generados por ChatGPT pueden parecer suficientemente plausibles para dificultar su identificación por revisores humanos. Este resultado evidencia que la calidad superficial del lenguaje no garantiza autenticidad ni validez académica. Un texto coherente puede contener datos inexistentes, métodos improbables o conclusiones sin respaldo. Por ello, la evaluación de producciones generativas debe superar criterios estilísticos y centrarse en la exactitud de las afirmaciones, la trazabilidad de las fuentes y la consistencia entre resultados, argumentos y referencias citadas.

Uno de los riesgos más graves corresponde a las alucinaciones, entendidas como contenidos formulados con seguridad aparente, pero carentes de sustento verificable. Estas pueden manifestarse mediante referencias inexistentes, autores incorrectos, datos inventados o interpretaciones que exceden el alcance de una publicación. Gilbert et al. (2024) proponen modelos aumentados con fuentes externas para reducir respuestas médicas no fundamentadas. En la producción académica, esta limitación obliga a comprobar manualmente cada cita, concepto, resultado y dato antes de incorporarlo a un manuscrito científico.

La integridad académica también se ve comprometida cuando el uso de inteligencia artificial generativa permanece oculto o sustituye tareas intelectuales que debían realizar los autores. Gallent et al. (2023) analizan los desafíos éticos que estas herramientas

introducen en la educación superior, particularmente en relación con autoría y autenticidad. La asistencia puede ser legítima cuando mejora la expresión o facilita la organización, pero resulta problemática si encubre plagio, atribuye conocimientos no adquiridos o impide determinar cuál fue la contribución real de cada participante.

La transparencia exige informar el uso de herramientas generativas cuando su intervención haya influido significativamente en la elaboración, revisión o transformación del manuscrito. Penabad et al. (2024) proponen orientaciones para reportar el empleo de inteligencia artificial en revistas científico-académicas. Declarar la herramienta, su función y el alcance de su participación permite evaluar el proceso con mayor claridad. Este reporte no traslada responsabilidad al sistema, sino que facilita reconocer qué tareas fueron asistidas y cuáles decisiones conceptuales, metodológicas y argumentativas permanecieron bajo control humano.

Los sesgos presentes en los datos de entrenamiento pueden trasladarse a los textos académicos generados. Omiye et al. (2023) evidencian que los grandes modelos de lenguaje reproducen planteamientos médicos asociados con categorías raciales. En otros campos, pueden privilegiar perspectivas dominantes, invisibilizar contribuciones regionales o reforzar estereotipos culturales y disciplinares. La revisión crítica debe analizar no solo la corrección formal, sino también las omisiones, supuestos y jerarquías incorporadas en el contenido, especialmente cuando se abordan poblaciones, territorios o conocimientos históricamente subrepresentados.

La protección de datos representa otra condición relevante cuando se introducen manuscritos inéditos, entrevistas, registros clínicos o información institucional en plataformas generativas. Sánchez (2024) advierte que estas tecnologías plantean retos para la protección de datos personales debido al tratamiento y posible reutilización de la información. Antes de emplearlas, el investigador debe revisar condiciones de privacidad, restricciones institucionales y sensibilidad de los contenidos. La eficiencia de una herramienta no justifica exponer datos confidenciales ni materiales sujetos a consentimiento, propiedad intelectual o compromisos editoriales.

La inteligencia artificial generativa puede fortalecer la producción académica cuando se utiliza como apoyo controlado y no como sustituto del razonamiento científico. Su contribución resulta valiosa para organizar ideas, revisar estilo, sintetizar información y preparar versiones preliminares, siempre que exista verificación documental y

supervisión permanente. Contreras et al. (2024) sostienen que el uso responsable de inteligencia artificial en investigación exige transparencia, ética y control humano. La calidad editorial depende finalmente de la exactitud, originalidad, coherencia y responsabilidad asumidas por los autores.

3.3 Procesamiento del lenguaje natural y minería de textos científicos

El procesamiento del lenguaje natural permite que los sistemas computacionales analicen, interpreten y generen información expresada mediante lenguaje humano. En la investigación científica, esta tecnología facilita el examen automatizado de artículos, resúmenes, informes, entrevistas y otros documentos académicos. Su aplicación reduce el tiempo destinado a tareas de clasificación y búsqueda, pero no sustituye la lectura crítica. Garzón et al. (2025) destacan que la interacción entre personas y sistemas inteligentes depende de capacidades lingüísticas capaces de reconocer significados, contextos y relaciones presentes en los textos.

La minería de textos científicos comprende procedimientos destinados a extraer información relevante, identificar regularidades y organizar contenidos dentro de colecciones documentales extensas. A diferencia de una búsqueda convencional basada en palabras clave, puede reconocer entidades, temas, relaciones semánticas y variaciones terminológicas. Estas funciones favorecen la exploración de campos con alta producción académica y vocabularios especializados. La integración de redes neuronales y grafos de conocimiento amplía las posibilidades de recuperación documental al conectar conceptos que no siempre aparecen expresados mediante los mismos términos (Polo y Casique, 2025).

Una de sus aplicaciones principales consiste en clasificar publicaciones según temas, enfoques metodológicos o áreas disciplinares. Los algoritmos pueden representar cada documento mediante características lingüísticas y agrupar textos con semejanzas semánticas. Esta capacidad facilita la construcción de corpus, la depuración de resultados y el reconocimiento de líneas de investigación. Sin embargo, una clasificación automatizada puede ignorar matices conceptuales o asignar categorías inadecuadas cuando los textos presentan terminología ambigua. La revisión especializada continúa siendo necesaria para comprobar la pertinencia científica de los grupos obtenidos (Mata et al., 2024).

El reconocimiento de entidades permite localizar autores, instituciones, métodos, variables, poblaciones y conceptos dentro de documentos científicos. Garzón et al. (2025) señalan que los sistemas de procesamiento del lenguaje natural han mejorado su capacidad para interpretar información y sostener interacciones más eficientes. En minería textual, esta función ayuda a estructurar datos procedentes de textos no organizados y facilita comparaciones entre estudios. No obstante, las entidades identificadas deben normalizarse, porque abreviaturas, variantes lingüísticas o nombres semejantes pueden generar duplicaciones y relaciones documentales incorrectas.

Figura 5

¿Qué enfoque tecnológico priorizar para el análisis de documentos académicos?



Nota. La figura compara el procesamiento del lenguaje natural y la minería de textos para el análisis de documentos académicos. Elaboración propia.

El modelado de temas constituye otra técnica empleada para descubrir núcleos conceptuales dentro de grandes colecciones. A partir de patrones de coaparición, los algoritmos estiman conjuntos de palabras que representan asuntos recurrentes en la literatura. Este procedimiento puede mostrar tendencias, transformaciones y áreas emergentes, aunque los temas generados no poseen significado científico autónomo. Su interpretación depende del conocimiento disciplinar, de la composición del corpus y de las decisiones analíticas adoptadas. La visualización mediante grafos puede facilitar la comprensión de estas relaciones temáticas (Ávila, 2024).

La minería de textos también apoya la identificación de brechas de investigación mediante el análisis de preguntas, limitaciones y recomendaciones presentes en publicaciones anteriores. Arias y Artigas (2022) explican que una brecha científica debe

conducir a problemas relevantes, delimitados y metodológicamente abordables. Los sistemas pueden señalar temas poco relacionados, poblaciones escasamente estudiadas o contradicciones entre resultados. Sin embargo, la baja frecuencia documental no demuestra por sí sola la existencia de un vacío significativo. La pertinencia debe evaluarse mediante criterios teóricos, sociales y metodológicos.

Los modelos de lenguaje permiten sintetizar evidencia distribuida entre numerosos documentos y producir representaciones breves de sus contenidos. Tang et al. (2023) evaluaron esta capacidad en la síntesis de evidencia médica, donde la precisión resulta especialmente importante. Aunque estas herramientas pueden reducir el tiempo de revisión, también pueden omitir limitaciones, fusionar resultados incompatibles o exagerar conclusiones. La síntesis automatizada debe considerarse una aproximación preliminar que requiere contraste con los artículos originales, verificación de citas y revisión del nivel real de evidencia disponible.

En salud mental, el procesamiento del lenguaje natural puede reconocer señales lingüísticas asociadas con situaciones de riesgo. Swaminathan et al. (2023) desarrollaron un sistema para detectar rápidamente mensajes relacionados con crisis y facilitar intervenciones oportunas. Esta aplicación demuestra que la minería textual trasciende la organización documental y puede apoyar decisiones sensibles. Sin embargo, los errores de clasificación pueden generar alarmas injustificadas o impedir una respuesta necesaria. Su implementación exige privacidad, evaluación contextual, supervisión profesional y protocolos claros para actuar ante los resultados emitidos.

La calidad del corpus determina en gran medida la confiabilidad de los patrones obtenidos. Una colección sesgada por idioma, región, disciplina o disponibilidad de acceso puede ofrecer una representación incompleta del conocimiento científico. Los algoritmos aprenderán las regularidades presentes en esos documentos, incluidas sus omisiones y desigualdades. Omiye et al. (2023) muestran que los grandes modelos de lenguaje pueden reproducir enfoques médicos basados en categorías raciales. Por ello, todo análisis textual debe documentar criterios de selección, cobertura y representatividad de las fuentes.

La reproducibilidad requiere registrar las etapas de recopilación, limpieza, tokenización, representación y análisis de los textos. También deben documentarse versiones de modelos, parámetros, diccionarios, reglas lingüísticas y criterios utilizados para excluir

documentos. Semmelrock et al. (2025) identifican la falta de documentación como una barrera recurrente en investigaciones basadas en aprendizaje automático. Sin estos registros, los resultados pueden depender de configuraciones particulares difíciles de reconstruir. La transparencia metodológica permite evaluar si los patrones lingüísticos encontrados son estables y científicamente interpretables.

El procesamiento del lenguaje natural y la minería de textos amplían la capacidad para explorar literatura, organizar información y descubrir relaciones dentro de corpus científicos extensos. Su valor no reside únicamente en automatizar la lectura, sino en facilitar preguntas y análisis que serían difíciles de realizar manualmente. Krenn et al. (2022) sostienen que la comprensión científica exige vincular resultados con explicaciones y estructuras conceptuales. En consecuencia, los hallazgos textuales deben interpretarse con conocimiento disciplinar, validación documental y supervisión humana permanente.

3.4 Visión artificial para el análisis de imágenes, muestras y fenómenos experimentales

La visión artificial permite que los sistemas inteligentes extraigan información de imágenes, videos y registros visuales obtenidos durante procesos científicos. Su aplicación incluye clasificación, segmentación, detección de objetos, reconocimiento de patrones y seguimiento de cambios temporales. En investigación, estas funciones facilitan el análisis de muestras biológicas, materiales, estructuras y fenómenos experimentales. El aprendizaje profundo ha ampliado esta capacidad mediante redes neuronales capaces de representar características complejas, aunque su rendimiento depende de la calidad de las imágenes, del etiquetado y de la correspondencia entre los datos y el problema estudiado (Hua et al., 2023).

A diferencia de la inspección visual exclusivamente humana, la visión artificial puede procesar grandes cantidades de imágenes bajo criterios uniformes y repetibles. Esto resulta útil cuando el análisis manual es lento, subjetivo o sensible a variaciones entre observadores. Sin embargo, la automatización no elimina la necesidad de experiencia disciplinar. Krenn et al. (2022) señalan que la comprensión científica requiere relacionar los resultados computacionales con explicaciones y estructuras conceptuales. Una clasificación visual solo adquiere valor cuando puede interpretarse dentro del fenómeno investigado.

En el análisis de muestras, los sistemas pueden identificar formas, texturas, colores y distribuciones espaciales difíciles de cuantificar mediante observación directa. Estas capacidades apoyan tareas de conteo, medición, diferenciación y detección de alteraciones. La precisión depende de que las condiciones de captura, iluminación, resolución y enfoque se mantengan controladas. Una imagen técnicamente deficiente puede inducir errores, aunque el modelo sea avanzado. Por ello, la calidad metodológica debe abarcar tanto el algoritmo como la preparación de la muestra y el procedimiento de adquisición visual (Ciobanu et al., 2024).

La segmentación de imágenes permite separar regiones relevantes y delimitar estructuras dentro de una muestra o fenómeno experimental. Esta operación facilita medir superficies, identificar bordes y comparar cambios entre condiciones. En medicina, por ejemplo, puede apoyar el análisis de imágenes clínicas, aunque las decisiones derivadas requieren revisión profesional. Bienefeld et al. (2023) destacan que los sistemas explicables deben responder a las necesidades de quienes utilizan sus resultados. La utilidad científica aumenta cuando el modelo permite comprender qué características visuales influyeron en su clasificación.

La detección automática de anomalías constituye otra aplicación significativa. Los algoritmos pueden señalar irregularidades, defectos o variaciones inesperadas dentro de secuencias visuales. Esta capacidad ayuda a identificar eventos que podrían pasar inadvertidos en revisiones manuales extensas. Sin embargo, una anomalía visual no equivale necesariamente a un fenómeno científicamente relevante. Puede deberse a ruido, cambios de iluminación, errores de calibración o diferencias en el procesamiento. La interpretación requiere contrastar la señal con controles experimentales y criterios disciplinares previamente establecidos (Mata et al., 2024).

Tabla 8
Aplicaciones de la visión artificial en el análisis científico

Función	Aplicación científica	Aporte principal	Condición metodológica	Riesgo crítico
Clasificación	Diferenciación de muestras, materiales o imágenes clínicas	Organización rápida y uniforme	Etiquetas válidas y clases bien definidas	Confusión entre categorías semejantes
Segmentación	Delimitación de regiones y estructuras	Medición espacial precisa	Calidad de imagen y	Pérdida de detalles relevantes

Función	Aplicación científica	Aporte principal	Condición metodológica	Riesgo crítico
			criterios de borde	
Detección de objetos	Identificación y conteo de elementos	Automatización de observaciones repetitivas	Resolución y escala controladas	Omisiones o falsos positivos
Seguimiento temporal	Observación de cambios y movimientos	Análisis dinámico de fenómenos	Captura continua y sincronizada	Interpretación incorrecta de trayectorias
Detección de anomalías	Reconocimiento de defectos o eventos inusuales	Identificación temprana de desviaciones	Datos de referencia representativos	Confundir ruido con hallazgos
Reconstrucción visual	Modelado de estructuras y escenas	Representación detallada del fenómeno	Calibración y múltiples perspectivas	Distorsiones del modelo reconstruido

Nota. La tabla sintetiza funciones de la visión artificial y evidencia que su aporte depende de protocolos de captura, validación disciplinar y control de errores.

El seguimiento temporal permite analizar movimientos, transformaciones y cambios progresivos dentro de experimentos registrados mediante video o secuencias de imágenes. Los sistemas pueden estimar trayectorias, velocidad, frecuencia y variación morfológica, generando datos difíciles de obtener manualmente. Esta capacidad resulta relevante en biología, ingeniería y ciencias naturales. No obstante, cualquier interrupción en la captura, cambio de escala o pérdida de fotogramas puede alterar las mediciones. La trazabilidad del procedimiento debe incluir dispositivos, parámetros de registro, software y criterios de procesamiento empleados.

La integración con laboratorios automatizados amplía aún más el alcance de la visión artificial. Fushimi et al. (2025) describen sistemas autónomos capaces de coordinar instrumentos y apoyar experimentos biotecnológicos. En estos entornos, las imágenes pueden utilizarse para monitorear procesos, evaluar resultados y activar ajustes automáticos. Sin embargo, una decisión basada en interpretación visual debe contar con límites operativos y mecanismos de intervención humana. Si el sistema clasifica incorrectamente una condición, el error puede propagarse hacia fases posteriores del experimento.

La validación de modelos visuales exige separar adecuadamente los datos de entrenamiento, ajuste y prueba. Kapoor y Narayanan (2023) advierten que la fuga de datos puede producir estimaciones artificialmente elevadas y comprometer la reproducibilidad.

En visión artificial, este problema puede surgir cuando imágenes semejantes de una misma muestra aparecen en distintos conjuntos. La evaluación debe utilizar datos independientes, condiciones variadas y métricas complementarias. También conviene analizar errores por categoría, contexto o tipo de imagen, en lugar de depender únicamente de un indicador promedio.

Los sesgos pueden originarse en una representación desigual de muestras, equipos, poblaciones o condiciones experimentales. Un modelo entrenado con imágenes obtenidas bajo protocolos homogéneos puede perder precisión cuando se aplica en otros entornos. Omiye et al. (2023) muestran que los sistemas inteligentes pueden reproducir sesgos presentes en sus datos y categorías. En visión artificial, esta limitación obliga a diversificar el conjunto de entrenamiento, documentar su procedencia y verificar el rendimiento en grupos o condiciones distintas antes de generalizar los resultados.

La visión artificial fortalece el análisis científico cuando convierte información visual en evidencia cuantificable, reproducible e interpretable. Su valor no reside únicamente en automatizar la observación, sino en ampliar la capacidad para medir fenómenos complejos y comparar resultados bajo criterios consistentes. Contreras et al. (2024) señalan que el uso de inteligencia artificial en investigación exige transparencia, supervisión y responsabilidad. La aplicación rigurosa requiere controlar la adquisición de imágenes, validar los modelos y mantener la interpretación final bajo responsabilidad del equipo científico.

3.5 Internet de las cosas, sensores inteligentes y recopilación automatizada de datos

El internet de las cosas integra dispositivos físicos, sensores, redes de comunicación y plataformas de análisis capaces de recopilar e intercambiar datos de manera continua. En la investigación científica, esta infraestructura permite observar fenómenos en tiempo real y ampliar la cobertura espacial y temporal de los registros. Su valor no reside únicamente en automatizar mediciones, sino en generar flujos de información útiles para analizar procesos dinámicos. Singh et al. (2022) muestran que la conexión entre sistemas físicos y digitales fortalece el monitoreo y la optimización en distintos sectores.

Los sensores inteligentes constituyen el componente básico de estos entornos porque transforman condiciones físicas, químicas o biológicas en datos procesables. Pueden medir temperatura, humedad, presión, movimiento, composición, vibración o variaciones

ambientales con diferentes niveles de precisión. Cuando incorporan capacidades de procesamiento, filtran señales, detectan cambios y transmiten información relevante sin depender de una revisión manual constante. Sin embargo, la validez científica de sus registros exige calibración, mantenimiento, control de errores y documentación de las condiciones bajo las cuales fueron obtenidos.

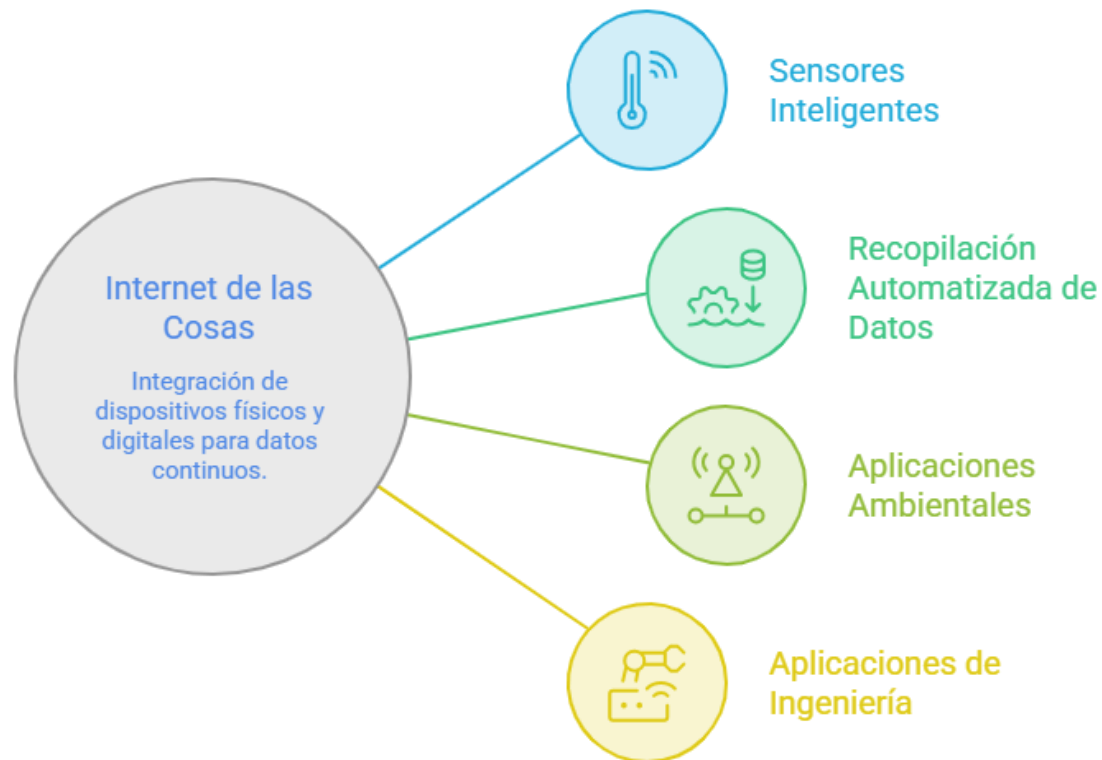
La recopilación automatizada permite sustituir observaciones esporádicas por registros continuos, lo que favorece el análisis de tendencias, eventos inusuales y relaciones temporales. Esta capacidad resulta especialmente importante cuando los fenómenos cambian con rapidez o presentan variaciones que podrían pasar inadvertidas mediante mediciones periódicas. Hua et al. (2023) destacan que los algoritmos inteligentes amplían el análisis de sistemas complejos. No obstante, un mayor volumen de datos también incrementa la necesidad de depuración, almacenamiento y selección de información científicamente pertinente.

En investigaciones ambientales, los sensores conectados pueden registrar variables climáticas, calidad del aire, humedad del suelo, niveles de agua o cambios en ecosistemas. La integración de estos datos facilita el estudio de procesos que combinan dimensiones territoriales y temporales. Sin embargo, las condiciones del entorno pueden afectar el funcionamiento de los dispositivos mediante corrosión, interferencias, pérdida de energía o alteraciones en la conectividad. La interpretación debe considerar estos factores para evitar que fallas técnicas se confundan con cambios reales del fenómeno observado (Morin, 1999).

En ingeniería, el internet de las cosas permite monitorear equipos, infraestructuras y procesos productivos mediante registros continuos. Los sensores pueden detectar vibraciones, temperatura, consumo energético o variaciones operativas que anticipen fallas. Singh et al. (2022) explican que estas tecnologías se articulan con gemelos digitales para representar el comportamiento de sistemas físicos. Esta integración favorece simulación, mantenimiento predictivo y evaluación de escenarios. Sin embargo, la calidad de la representación depende de la confiabilidad de los datos y de su actualización constante.

Figura 6

Revelando el impacto del internet de las cosas



Nota. La figura muestra el impacto del internet de las cosas mediante sensores inteligentes, recopilación automatizada de datos y aplicaciones ambientales y de ingeniería. Elaboración propia.

Los gemelos digitales utilizan información procedente de sensores para mantener una representación dinámica del objeto o proceso estudiado. Qian et al. (2022) señalan que estas réplicas cibernéticas permiten monitorear condiciones, simular respuestas y apoyar decisiones. En investigación, la conexión entre dispositivos físicos y modelos digitales facilita comparar el comportamiento real con escenarios previstos. No obstante, si los sensores transmiten datos incompletos o sesgados, la réplica puede ofrecer una representación engañosa. La validación debe abarcar tanto el modelo como la infraestructura de captura.

En laboratorios automatizados, los sensores inteligentes permiten controlar variables experimentales, registrar cambios y activar respuestas programadas. Fushimi et al. (2025) describen un sistema autónomo orientado a apoyar investigaciones biotecnológicas mediante la coordinación de instrumentos y procesos. Estos entornos pueden aumentar la precisión y continuidad de los experimentos, pero también requieren protocolos de seguridad y supervisión. Una lectura incorrecta puede modificar condiciones experimentales y propagar errores dentro de la secuencia, por lo que deben existir mecanismos de alerta e intervención humana.

La calidad del dato constituye una condición central en la recopilación automatizada. Los registros pueden verse afectados por ruido, pérdida de señal, deriva del sensor, duplicaciones o desajustes temporales. Antes del análisis, deben evaluarse consistencia, completitud, precisión y correspondencia con el fenómeno medido. Semmelrock et al. (2025) identifican la documentación insuficiente como una barrera para la reproducibilidad en investigaciones basadas en aprendizaje automático. En sistemas conectados, registrar versiones, calibraciones, ubicaciones y frecuencias de medición resulta igualmente necesario.

La interoperabilidad representa otro desafío metodológico. Los dispositivos pueden utilizar protocolos, formatos y unidades diferentes, dificultando la integración de datos procedentes de múltiples fuentes. Los sistemas científicos requieren estándares que permitan combinar registros sin perder contexto ni significado. La normalización debe conservar información sobre procedencia, escala, precisión y condiciones de captura. Cuando estas características se eliminan durante la integración, los algoritmos pueden detectar patrones originados por incompatibilidades técnicas y no por relaciones reales entre las variables analizadas (Ávila, 2024).

La recopilación continua también plantea riesgos para la privacidad cuando los sensores registran información sobre personas, comportamientos o ubicaciones. Andrade et al. (2024) analizan la relación entre inteligencia artificial y protección de datos personales desde una perspectiva legal. En salud, educación o ciencias sociales, los sistemas conectados deben limitar la información recolectada, proteger su transmisión y definir quién puede acceder a ella. La utilidad científica no justifica una vigilancia desproporcionada ni el uso secundario de datos sin conocimiento o autorización de los participantes.

La ciberseguridad condiciona la confiabilidad de toda infraestructura conectada. Un acceso no autorizado puede alterar registros, interrumpir mediciones o manipular decisiones automatizadas. Estas amenazas afectan tanto la protección de la información como la validez de los resultados científicos. La seguridad debe integrarse desde el diseño mediante autenticación, cifrado, control de accesos y respaldo de datos. Sánchez (2024) advierte que las tecnologías inteligentes amplían los desafíos asociados con el tratamiento y la protección de información personal, especialmente cuando intervienen plataformas externas.

El internet de las cosas y los sensores inteligentes amplían la capacidad para observar fenómenos con continuidad, detalle y alcance. Su incorporación favorece investigaciones más dinámicas, conectadas y sensibles a cambios temporales. Sin embargo, la automatización no garantiza por sí sola datos confiables. Contreras et al. (2024) destacan que el uso científico de inteligencia artificial exige transparencia, supervisión y responsabilidad. La calidad final depende de calibrar los dispositivos, documentar los flujos de información, proteger los datos y mantener control humano sobre las decisiones derivadas.

3.6 Computación en la nube, ciencia de datos y procesamiento de alto rendimiento

La computación en la nube proporciona infraestructura remota para almacenar, procesar y compartir datos científicos sin depender exclusivamente de equipos instalados en una institución. Su flexibilidad permite ajustar recursos según el tamaño del proyecto, ejecutar análisis simultáneos y facilitar el acceso distribuido a herramientas especializadas. Esta capacidad resulta especialmente útil cuando la investigación maneja grandes volúmenes de información o requiere colaboración entre equipos ubicados en distintos lugares. Sin embargo, su adopción exige valorar seguridad, costos, conectividad, interoperabilidad y dependencia de proveedores tecnológicos.

La ciencia de datos articula estadística, programación, gestión de información y conocimiento disciplinar para extraer significado de conjuntos complejos. No se limita a aplicar algoritmos, pues comprende recopilación, limpieza, integración, exploración, modelado, interpretación y comunicación de resultados. Hua et al. (2023) señalan que los algoritmos inteligentes amplían las posibilidades del cómputo científico y permiten abordar problemas de elevada complejidad. La calidad del análisis depende de que cada transformación sea metodológicamente justificable y conserve la relación entre los datos procesados y el fenómeno investigado.

El procesamiento de alto rendimiento utiliza múltiples unidades computacionales para ejecutar cálculos intensivos en tiempos menores que los sistemas convencionales. Esta capacidad favorece simulaciones, análisis genómicos, entrenamiento de modelos profundos y procesamiento masivo de imágenes o textos. Hartung (2025) relaciona la automatización y los modelos con capacidad de acción con nuevas formas de descubrimiento científico. No obstante, incrementar la potencia computacional no corrige

errores en el diseño, datos deficientes o supuestos inadecuados. La infraestructura debe seleccionarse según necesidades científicas concretas y no por sofisticación tecnológica.

La combinación de nube, ciencia de datos y alto rendimiento permite construir flujos de trabajo escalables. Los datos pueden almacenarse en repositorios compartidos, procesarse mediante recursos distribuidos y analizarse con herramientas reproducibles. En bioinformática, esta articulación resulta relevante porque los estudios genómicos producen archivos extensos y requieren múltiples etapas computacionales. Loza et al. (2026) destacan la importancia de la reproducibilidad en flujos aplicados a genómica y oncología clínica. Documentar cada operación permite reconstruir resultados y detectar diferencias entre ejecuciones.

Tabla 9
Comparación de infraestructuras computacionales aplicadas a la investigación científica

Criterio	Computación local	Computación en la nube	Procesamiento de alto rendimiento
Infraestructura	Equipos propios de capacidad limitada	Recursos remotos ajustables bajo demanda	Clústeres y unidades de cálculo coordinadas
Escalabilidad	Depende de ampliaciones físicas	Flexible según almacenamiento y procesamiento requerido	Elevada para cálculos intensivos y paralelos
Acceso	Vinculado al equipo o red institucional	Distribuido mediante conexión autorizada	Generalmente restringido a centros especializados
Aplicaciones	Análisis moderados y tareas cotidianas	Colaboración, almacenamiento y flujos escalables	Simulación, genómica y entrenamiento de modelos complejos
Costos	Inversión inicial y mantenimiento local	Pago según uso y servicios contratados	Infraestructura, energía y personal especializado
Riesgo principal	Obsolescencia y capacidad insuficiente	Dependencia del proveedor y protección de datos	Complejidad técnica y acceso desigual
Reproducibilidad	Depende del registro manual del entorno	Facilita entornos compartidos y automatizables	Exige documentar configuraciones y condiciones de ejecución

Nota. La tabla compara tres modalidades de infraestructura y evidencia que su elección debe considerar escala del análisis, sensibilidad de los datos, capacidad institucional y exigencias de reproducibilidad.

La nube también fortalece el trabajo colaborativo al permitir que investigadores accedan a datos, modelos y resultados mediante entornos compartidos. Ávila (2024) muestra que la organización de datos abiertos mediante ontologías y grafos de conocimiento favorece su representación y reutilización. Estas posibilidades requieren metadatos claros, control de versiones y criterios comunes de acceso. Cuando los archivos se modifican sin trazabilidad o se almacenan sin descripción suficiente, la disponibilidad técnica no garantiza comprensión, interoperabilidad ni aprovechamiento científico por parte de otros equipos.

La reproducibilidad constituye una exigencia central en entornos computacionales complejos. Semmelrock et al. (2025) identifican barreras asociadas con falta de documentación, acceso limitado a recursos y diferencias entre configuraciones. Un mismo código puede generar resultados distintos si cambian las bibliotecas, versiones, parámetros o capacidades del equipo. Por ello, los proyectos deben conservar información sobre sistemas operativos, dependencias, semillas aleatorias y secuencias de procesamiento. Los entornos virtualizados pueden mejorar la repetición, aunque no sustituyen una descripción metodológica completa y comprensible.

La protección de datos adquiere especial importancia cuando la información científica se almacena o procesa en servicios externos. Andrade et al. (2024) destacan que el uso de inteligencia artificial debe sustentarse en bases legales que protejan los datos personales. Los equipos deben determinar dónde se encuentran los servidores, quién administra los accesos y qué usos secundarios podrían realizarse. En investigaciones con información clínica, educativa o social, la escalabilidad tecnológica debe acompañarse de cifrado, anonimización, restricciones de descarga y procedimientos institucionales de autorización.

El acceso desigual a infraestructura avanzada puede profundizar diferencias entre instituciones y regiones. Los centros con mayores recursos disponen de capacidad para entrenar modelos, almacenar grandes colecciones y ejecutar simulaciones complejas, mientras otros dependen de plataformas externas o colaboraciones asimétricas. Ross et al. (2022) advierten que las dinámicas de ventaja acumulativa también afectan la equidad en la ciencia abierta. Democratizar el procesamiento científico requiere formación, infraestructura compartida y modelos de colaboración que permitan participar en las decisiones analíticas, no únicamente aportar datos.

La computación en la nube, la ciencia de datos y el procesamiento de alto rendimiento amplían la escala y velocidad de la investigación, pero su valor depende de la coherencia entre infraestructura, método y problema científico. Estas tecnologías facilitan análisis complejos, colaboración y automatización, aunque también introducen riesgos de opacidad, dependencia y desigualdad. Semmelrock et al. (2025) señalan que la reproducibilidad requiere acciones técnicas y organizacionales coordinadas. Una implementación rigurosa debe garantizar documentación, seguridad, acceso responsable y supervisión humana durante todo el flujo computacional.

3.7 Gemelos digitales, robótica científica y laboratorios automatizados

Los gemelos digitales, la robótica científica y los laboratorios automatizados representan una convergencia tecnológica orientada a transformar la experimentación. Estas herramientas conectan sistemas físicos, modelos computacionales, sensores y mecanismos de control para observar, simular y modificar procesos con menor intervención manual. Su aporte no consiste únicamente en acelerar tareas repetitivas, sino en ampliar la capacidad de explorar escenarios, mantener registros continuos y ajustar parámetros experimentales. Las aplicaciones de estas tecnologías abarcan distintos sectores científicos e industriales, aunque requieren validación, interoperabilidad y supervisión especializada (Singh et al., 2022).

Un gemelo digital es una representación computacional vinculada con un objeto, proceso o sistema físico mediante datos que permiten actualizar su estado. Qian et al. (2022) explican que esta réplica cibernética integra arquitectura, comunicación y modelado para apoyar monitoreo, simulación y toma de decisiones. En investigación científica, puede utilizarse para comparar condiciones, anticipar comportamientos y evaluar modificaciones antes de aplicarlas sobre el sistema real. Su confiabilidad depende de la correspondencia entre la representación digital, los datos disponibles y las condiciones efectivamente observadas.

La simulación mediante gemelos digitales permite estudiar fenómenos cuya experimentación directa puede ser costosa, lenta, peligrosa o difícil de repetir. Los investigadores pueden modificar variables, comparar escenarios y estimar posibles respuestas sin alterar inmediatamente el objeto físico. Esta capacidad favorece diseños experimentales más informados y reduce algunos riesgos operativos. Sin embargo, las predicciones permanecen condicionadas por supuestos, simplificaciones y calidad de los

datos utilizados. Una representación digital no sustituye la validación empírica, pues puede omitir relaciones relevantes presentes en el sistema real (Qian et al., 2022).

La robótica científica incorpora dispositivos programables capaces de manipular muestras, trasladar materiales, ejecutar mediciones y repetir protocolos con elevada precisión. Su integración en laboratorios disminuye la variabilidad asociada con determinadas operaciones manuales y permite mantener actividades durante periodos prolongados. Hartung (2025) vincula la automatización, la inteligencia artificial y los modelos con capacidad de acción con nuevas posibilidades para el descubrimiento científico. No obstante, la autonomía robótica debe acompañarse de límites operativos, mecanismos de seguridad y procedimientos que permitan interrumpir acciones ante resultados inesperados.

Los laboratorios automatizados combinan robótica, instrumentos analíticos, software y sistemas inteligentes dentro de flujos experimentales coordinados. Fushimi et al. (2025) desarrollaron un sistema autónomo destinado a apoyar la investigación biotecnológica mediante la integración de equipos y procesos. Estas arquitecturas pueden ejecutar protocolos, registrar resultados y modificar condiciones según criterios previamente establecidos. La automatización favorece continuidad y estandarización, pero también puede propagar errores si una lectura incorrecta activa decisiones posteriores. Por ello, cada etapa necesita controles, registros y validaciones independientes.

Los laboratorios autónomos avanzan más allá de la automatización de tareas aisladas porque conectan planificación, ejecución, análisis y ajuste dentro de ciclos iterativos. Tobias y Wahab (2025) examinan estas infraestructuras como laboratorios capaces de orientar experimentalmente sus siguientes acciones. El sistema puede seleccionar condiciones prometedoras a partir de resultados previos y dirigir nuevas pruebas. Esta capacidad acelera la exploración de espacios experimentales complejos, aunque plantea interrogantes sobre control científico, responsabilidades institucionales y criterios utilizados para priorizar determinadas rutas de investigación frente a otras.

La integración entre gemelos digitales y robótica permite que una representación computacional influya directamente sobre operaciones físicas. Los datos obtenidos por sensores actualizan el modelo, mientras sus predicciones pueden orientar movimientos, ajustes o nuevas mediciones. Singh et al. (2022) muestran que los gemelos digitales se aplican en diferentes industrias para monitorear y optimizar sistemas. En el ámbito

científico, este circuito favorece experimentos adaptativos. Sin embargo, cualquier discrepancia entre el modelo y el sistema físico puede desencadenar decisiones inadecuadas y afectar la validez de los resultados.

La reproducibilidad constituye una ventaja potencial de los laboratorios automatizados, debido a que cada operación puede registrarse mediante protocolos digitales. Tiempos, temperaturas, movimientos, parámetros y resultados quedan asociados con secuencias susceptibles de revisión. No obstante, la existencia de registros automáticos no garantiza que el experimento pueda repetirse en otro entorno. También deben documentarse versiones de software, configuraciones, instrumentos, calibraciones y criterios de decisión. Las investigaciones basadas en sistemas inteligentes enfrentan barreras técnicas y organizacionales que dificultan su reproducción cuando estos elementos no se comunican adecuadamente (Semmelrock et al., 2025).

La autonomía tecnológica también introduce problemas de explicabilidad. Un laboratorio puede seleccionar un nuevo experimento mediante un modelo cuyo razonamiento interno resulte difícil de interpretar. Zárate (2022) señala que la explicabilidad permite comprender los fundamentos de las decisiones algorítmicas y evaluar sus consecuencias. En investigación, esta exigencia es relevante porque los científicos deben justificar por qué se modificaron determinadas variables o se descartaron ciertas alternativas. La eficiencia operativa pierde valor si el proceso experimental no puede reconstruirse, comprenderse y someterse a evaluación crítica.

Los gemelos digitales, la robótica científica y los laboratorios automatizados pueden ampliar la escala, precisión y adaptabilidad de la experimentación. Su contribución más profunda consiste en integrar observación, simulación, ejecución y análisis dentro de sistemas capaces de aprender de resultados sucesivos. Sin embargo, la responsabilidad científica permanece en los equipos humanos que diseñan los protocolos, validan los modelos e interpretan los hallazgos. La automatización responsable exige transparencia, supervisión y control institucional para evitar que la autonomía técnica debilite la trazabilidad y la confiabilidad del conocimiento producido (Contreras et al., 2024).



Capítulo

4.

Gestión del conocimiento mediante inteligencia artificial

Delly Maribel San Martin Torres

4.1 Fundamentos de la gestión inteligente del conocimiento científico

La gestión inteligente del conocimiento científico comprende el conjunto de procesos destinados a identificar, organizar, relacionar, preservar y reutilizar información mediante sistemas capaces de aprender de los datos. Su finalidad no consiste únicamente en almacenar documentos, sino en convertir contenidos dispersos en estructuras útiles para la investigación. Ávila (2024) muestra que los grafos de conocimiento permiten representar relaciones entre datos abiertos, conceptos y metadatos. Esta capacidad favorece una visión más integrada del conocimiento y facilita su recuperación dentro de entornos científicos complejos.

El conocimiento científico posee características que dificultan su gestión: se produce en múltiples formatos, circula en distintas plataformas y utiliza lenguajes especializados. Artículos, bases de datos, protocolos, códigos, imágenes y registros experimentales requieren criterios de organización diferenciados. Tapia y Recalde (2024) destacan la importancia de mejorar la gestión documental mediante procedimientos sistemáticos. La inteligencia artificial puede apoyar clasificación, búsqueda y actualización, pero su eficacia depende de una arquitectura informacional clara y de estándares que permitan interpretar correctamente cada recurso disponible.

La gestión inteligente combina herramientas de automatización con decisiones humanas sobre relevancia, calidad y contexto. Los algoritmos pueden clasificar documentos, sugerir vínculos y detectar duplicaciones, aunque no determinan por sí mismos qué información posee mayor valor científico. Cedeño et al. (2024) señalan que la inteligencia artificial fortalece distintas fases de la investigación universitaria, incluida la organización de información. Sin embargo, la selección final debe considerar la pertinencia de las fuentes, la validez metodológica y la relación de cada contenido con los objetivos del proyecto.

Los metadatos constituyen la base para describir y localizar recursos científicos. Incluyen información sobre autoría, fecha, tipo de documento, palabras clave, procedencia, versión y condiciones de acceso. Cuando se estructuran adecuadamente, permiten que los sistemas identifiquen relaciones entre publicaciones, datos y proyectos. Ávila (2024) destaca el valor de ontologías y modelos vinculados para representar datos de investigación. Sin metadatos consistentes, la automatización pierde precisión y puede generar asociaciones incompletas, duplicadas o difíciles de interpretar.

Los grafos de conocimiento amplían esta lógica al representar entidades y vínculos dentro de una red semántica. Un artículo puede conectarse con autores, métodos, variables, instituciones y conjuntos de datos, creando rutas de exploración más ricas que una clasificación lineal. Polo y Casique (2025) proponen integrar grafos y redes neuronales para fortalecer la recuperación documental. Esta combinación permite identificar relaciones implícitas y mejorar la navegación entre fuentes, aunque exige controlar la calidad de los nodos y la validez de los enlaces generados.

La minería de textos cumple una función central porque transforma documentos no estructurados en información analizable. Los sistemas pueden extraer conceptos, reconocer entidades, agrupar temas y detectar tendencias en grandes corpus científicos. Garzón et al. (2025) destacan el potencial del procesamiento del lenguaje natural para mejorar la interacción entre personas y sistemas inteligentes. Aplicado a la gestión del conocimiento, facilita búsquedas más precisas y conexiones semánticas, pero requiere revisión experta para evitar interpretaciones superficiales o errores derivados de ambigüedades terminológicas.

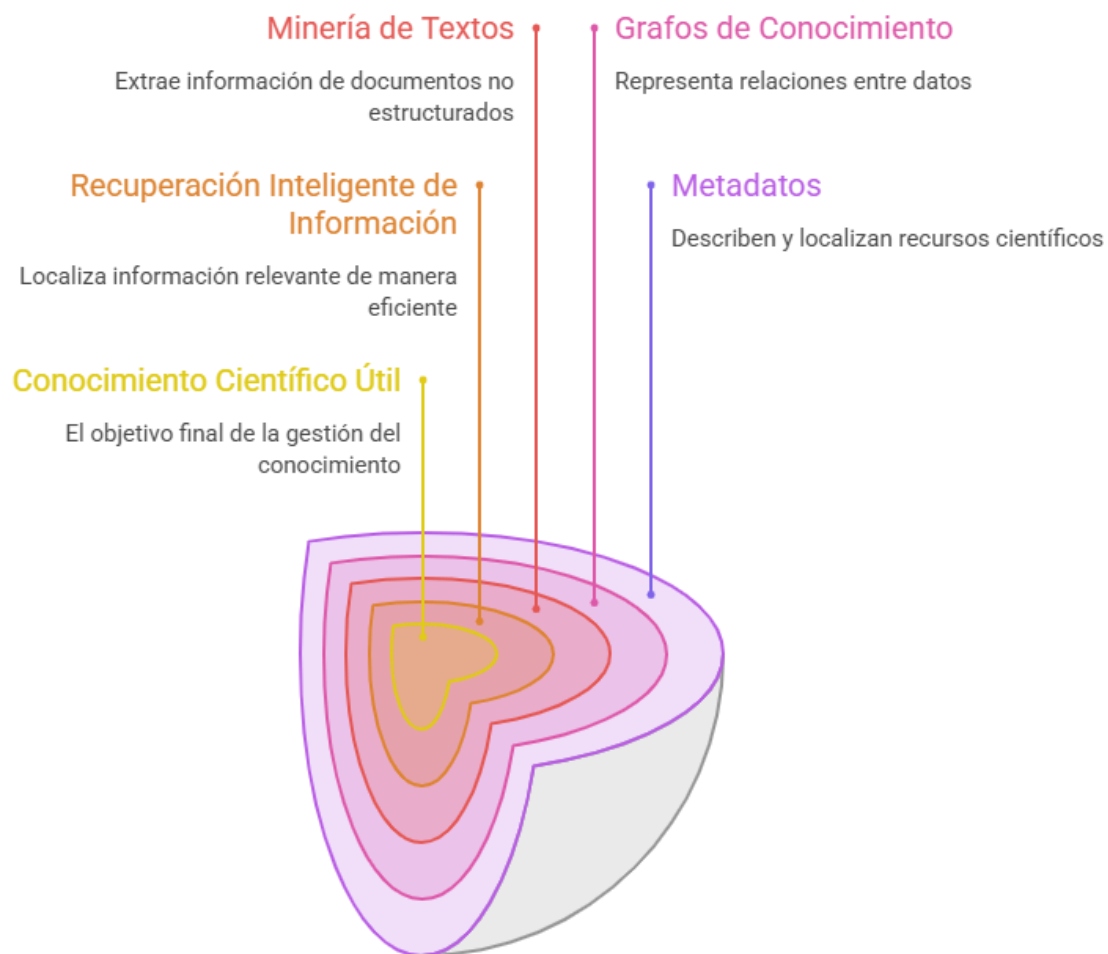
La recuperación inteligente de información supera la búsqueda basada exclusivamente en coincidencias literales. Los sistemas pueden considerar contexto, relaciones conceptuales y similitud semántica para localizar documentos relevantes (Giordano, 2022). Polo y Casique (2025) muestran que la combinación de grafos de conocimiento y redes neuronales favorece esta tarea. Sin embargo, una búsqueda más amplia no siempre produce mejores resultados. Es necesario establecer criterios de inclusión, pertinencia y calidad para evitar que la abundancia de información dificulte la identificación de evidencia realmente útil para la investigación.

La preservación del conocimiento también forma parte de esta gestión. Los recursos científicos deben conservarse con información suficiente para permitir su consulta, reutilización y verificación futura. Esto incluye versiones de documentos, códigos, datos, parámetros y decisiones metodológicas. Semmelrock et al. (2025) señalan que la falta de documentación constituye una barrera importante para la reproducibilidad en investigaciones basadas en aprendizaje automático. Una gestión inteligente debe registrar no solo el resultado final, sino también el proceso que permitió producirlo.

La reutilización del conocimiento depende de la interoperabilidad entre repositorios, plataformas y formatos. Los sistemas deben poder intercambiar información sin perder

significado, contexto o trazabilidad. Ávila (2024) muestra que las ontologías y los estándares de metadatos facilitan esta conexión. Sin embargo, la interoperabilidad técnica no garantiza comprensión científica. También se requiere compatibilidad conceptual y claridad sobre las condiciones de producción de los datos. De lo contrario, un recurso puede ser accesible, pero resultar inapropiado para responder nuevas preguntas de investigación.

Figura 7
Gestión inteligente del conocimiento científico



Nota. La figura sintetiza los componentes que facilitan la gestión inteligente y el aprovechamiento del conocimiento científico. Elaboración propia.

La gestión inteligente del conocimiento científico alcanza valor cuando mejora la capacidad de localizar, comprender y reutilizar información confiable. No se limita a incorporar tecnologías, sino que exige criterios de calidad, trazabilidad y responsabilidad. Tapia y Recalde (2024) resaltan que la gestión documental debe organizarse de forma sistemática para sostener procesos institucionales. En el ámbito científico, esta exigencia

implica combinar automatización, estándares, supervisión humana y políticas claras de acceso, preservación y uso responsable del conocimiento producido.

4.2 Búsqueda y recuperación automatizada de literatura académica

La búsqueda automatizada de literatura académica utiliza algoritmos capaces de localizar, clasificar y priorizar documentos científicos dentro de grandes colecciones digitales. Su principal ventaja consiste en ampliar la cobertura y reducir el tiempo invertido en tareas repetitivas de exploración. Sin embargo, la velocidad no garantiza pertinencia ni calidad. Polo y Casique (2025) destacan que la integración de grafos de conocimiento y redes neuronales puede fortalecer la recuperación documental, siempre que los resultados se interpreten de acuerdo con el problema científico y los criterios metodológicos establecidos.

Los sistemas tradicionales de búsqueda se apoyan principalmente en coincidencias de palabras clave, operadores booleanos y filtros bibliográficos. Las herramientas inteligentes incorporan análisis semántico, similitud contextual y reconocimiento de relaciones entre conceptos. Esto permite recuperar documentos que no utilizan exactamente los mismos términos empleados por el investigador. Garzón et al. (2025) señalan que el procesamiento del lenguaje natural mejora la comunicación entre personas y sistemas inteligentes. En este ámbito, facilita consultas más flexibles, aunque puede introducir resultados ambiguos cuando los conceptos poseen significados disciplinares diferentes.

La formulación de la estrategia de búsqueda continúa siendo una responsabilidad central del investigador. La inteligencia artificial puede sugerir términos, sinónimos, categorías y relaciones conceptuales, pero no determina por sí sola cuáles son científicamente adecuados. Una consulta mal delimitada puede producir resultados abundantes y poco útiles. La identificación de brechas depende de preguntas relevantes, criterios de inclusión explícitos y revisión crítica de los antecedentes, no únicamente de la cantidad de documentos recuperados (Arias y Artigas, 2022).

Los grafos de conocimiento permiten representar vínculos entre autores, publicaciones, instituciones, métodos, variables y conjuntos de datos. Ávila (2024) muestra que estas estructuras favorecen la visualización de relaciones dentro de datos abiertos de investigación. En la recuperación académica, pueden revelar conexiones temáticas, redes de colaboración y trayectorias conceptuales que no son evidentes en una lista

convencional de resultados. Su utilidad depende de la calidad de los metadatos y de la consistencia con que se describen las entidades incorporadas al sistema.

La minería de textos complementa la búsqueda al analizar contenidos completos, resúmenes, palabras clave y referencias bibliográficas. Esta técnica permite agrupar publicaciones, reconocer temas recurrentes y detectar tendencias en corpus extensos. Mata et al. (2024) identifican múltiples usos de la inteligencia artificial en el desarrollo de investigaciones científicas, incluida la organización de información. No obstante, el agrupamiento automatizado puede simplificar debates complejos o reunir documentos metodológicamente incompatibles. Por ello, los resultados deben considerarse una ayuda preliminar y no una clasificación definitiva.

Tabla 10
Comparación entre búsqueda académica convencional y recuperación automatizada

Criterio	Búsqueda académica convencional	Recuperación automatizada con inteligencia artificial
Base de consulta	Palabras clave, operadores y filtros	Significado semántico, contexto y relaciones conceptuales
Cobertura	Dependiente de términos definidos manualmente	Amplía términos y conexiones entre documentos
Organización de resultados	Listados ordenados por relevancia o fecha	Agrupamientos temáticos, redes y recomendaciones
Velocidad de revisión	Requiere lectura y depuración manual intensiva	Automatiza clasificación y priorización preliminar
Interpretación	Depende directamente del investigador	Combina análisis algorítmico y revisión especializada
Riesgo principal	Omisión de documentos por diferencias terminológicas	Inclusión de resultados irrelevantes o relaciones artificiales
Exigencia metodológica	Estrategia reproducible y filtros claros	Documentación de consultas, modelos y criterios de selección

Nota. La tabla evidencia que la recuperación automatizada amplía la exploración documental, pero no sustituye la delimitación metodológica ni la evaluación crítica de la literatura localizada.

Los sistemas de recomendación académica pueden sugerir documentos a partir del historial de consulta, las referencias utilizadas o la similitud temática. Esta función facilita el descubrimiento de publicaciones relacionadas, especialmente en campos con alta producción científica. Sin embargo, también puede reforzar circuitos cerrados de información y priorizar trabajos ya visibles. Ross et al. (2022) advierten que las dinámicas de ventaja acumulativa pueden afectar la equidad en la ciencia abierta. Una recuperación rigurosa debe evitar depender exclusivamente de recomendaciones personalizadas.

La cobertura lingüística y geográfica condiciona la calidad de los resultados automatizados. Los sistemas entrenados o alimentados principalmente con literatura en inglés pueden subrepresentar publicaciones latinoamericanas, regionales o procedentes de revistas con menor visibilidad internacional. Esta limitación afecta la diversidad de perspectivas y puede distorsionar la identificación de tendencias. La selección del corpus debe documentar idiomas, bases consultadas, periodos y criterios de acceso, de modo que las omisiones puedan reconocerse y no se interpreten como ausencia real de conocimiento científico (Zemelman, 2021).

La evaluación de relevancia requiere revisar título, resumen, metodología, resultados y correspondencia con la pregunta investigativa. Los sistemas pueden asignar puntuaciones o jerarquías, pero estas dependen de criterios técnicos que no siempre reflejan calidad científica. Penabad et al. (2024) destacan la necesidad de transparencia en el uso y reporte de inteligencia artificial dentro de revistas académico científicas. Aplicado a la búsqueda, este principio implica declarar herramientas, consultas, filtros y decisiones de exclusión cuando influyan de manera significativa en la construcción del corpus.

La reproducibilidad de la búsqueda automatizada exige conservar las expresiones utilizadas, fechas de consulta, bases, versiones de herramientas y criterios de selección. Semmelrock et al. (2025) señalan que la documentación insuficiente dificulta reproducir investigaciones basadas en aprendizaje automático. En revisiones de literatura, esta exigencia es igualmente relevante porque los sistemas pueden actualizar resultados, modificar recomendaciones o cambiar sus modelos internos. Registrar el procedimiento permite evaluar si el corpus obtenido responde a una estrategia científica transparente y no a una selección algorítmica imposible de reconstruir.

4.3 Clasificación, síntesis y análisis semántico de información científica

La clasificación automatizada de información científica permite organizar artículos, datos, informes y otros recursos según temas, métodos, disciplinas o niveles de relevancia. Esta función resulta especialmente útil ante el crecimiento sostenido de la producción académica y la dispersión de contenidos entre múltiples repositorios. Los sistemas inteligentes pueden reconocer semejanzas lingüísticas y conceptuales que no siempre son evidentes mediante búsquedas convencionales. Sin embargo, toda clasificación depende de categorías previamente definidas o aprendidas, por lo que debe

revisarse para evitar agrupamientos imprecisos o científicamente inconsistentes (Mata et al., 2024).

El análisis semántico amplía esta capacidad al examinar el significado contextual de palabras, expresiones y relaciones presentes en los textos. A diferencia de una coincidencia literal, permite identificar conceptos equivalentes, sinónimos y asociaciones temáticas entre documentos. Garzón et al. (2025) señalan que el procesamiento del lenguaje natural fortalece la interacción entre personas y sistemas inteligentes mediante una interpretación más compleja del lenguaje. En investigación, esta función facilita la exploración de corpus amplios, aunque requiere control disciplinar para evitar equivalencias conceptuales incorrectas.

Los grafos de conocimiento ofrecen una estructura adecuada para representar entidades científicas y sus vínculos. Autores, métodos, variables, instituciones y resultados pueden integrarse como nodos conectados mediante relaciones explícitas. Ávila (2024) muestra que estas arquitecturas permiten visualizar y organizar datos abiertos de investigación de manera estructurada. Su utilidad no se limita a la representación gráfica, pues también favorece consultas complejas y recorridos semánticos. No obstante, la confiabilidad del grafo depende de la calidad de los metadatos y de la validez de las relaciones establecidas.

Polo y Casique (2025) proponen combinar grafos de conocimiento y redes neuronales para mejorar la recuperación documental. Esta integración puede fortalecer la clasificación científica al relacionar patrones lingüísticos con estructuras conceptuales más amplias. Los sistemas pueden identificar documentos próximos por contenido, aunque utilicen vocabularios distintos. Aun así, la proximidad semántica no garantiza equivalencia metodológica ni teórica. El investigador debe comprobar si las publicaciones agrupadas responden realmente a preguntas semejantes, utilizan conceptos compatibles y producen evidencia comparable.

La síntesis automatizada busca condensar información extensa sin perder sus elementos centrales. Puede aplicarse a artículos individuales, conjuntos de publicaciones o evidencias procedentes de varias fuentes. Tang et al. (2023) evaluaron modelos de lenguaje orientados a resumir evidencia médica y mostraron que su utilidad debe acompañarse de revisión especializada. Una síntesis puede presentar fluidez y coherencia, pero omitir limitaciones, confundir resultados o exagerar conclusiones. Por ello, la salida

generada debe contrastarse con los documentos originales antes de incorporarse a una investigación.

La síntesis científica exige distinguir entre descripción, interpretación y valoración de evidencia. Los modelos generativos pueden mezclar estos niveles y presentar como conclusión una relación que solo fue mencionada en los textos analizados. Gao et al. (2023) demostraron que los resúmenes científicos generados artificialmente pueden parecer plausibles incluso para revisores humanos. Este riesgo obliga a verificar cada afirmación, especialmente cuando el sistema no muestra con claridad de qué fuente obtuvo una idea. La trazabilidad documental debe acompañar todo proceso de condensación automatizada.

La clasificación temática también puede apoyar la identificación de tendencias y vacíos de investigación. Los algoritmos agrupan publicaciones según patrones recurrentes y permiten observar áreas con alta concentración o escasa presencia documental. Sin embargo, una baja frecuencia no demuestra por sí sola que exista una brecha relevante. Arias y Artigas (2022) sostienen que los vacíos deben vincularse con problemas científicos delimitados y justificables. El análisis automatizado funciona como una herramienta exploratoria, pero la pertinencia final depende del examen teórico, metodológico y contextual realizado por el investigador.

La calidad de los resultados está condicionada por la composición del corpus. Si predominan ciertos idiomas, regiones, disciplinas o enfoques, los sistemas reproducirán esa distribución en sus clasificaciones y síntesis. Ross et al. (2022) advierten que la ciencia abierta puede mantener dinámicas de ventaja acumulativa y desigualdad. En el análisis semántico, esto puede traducirse en mayor visibilidad para perspectivas ya dominantes y menor presencia de conocimientos regionales. Documentar la procedencia y cobertura de las fuentes resulta indispensable para interpretar adecuadamente los patrones obtenidos.

La ambigüedad terminológica constituye otro desafío importante. Un mismo concepto puede variar entre disciplinas, mientras expresiones semejantes pueden designar fenómenos distintos. Garzón et al. (2025) destacan que el procesamiento del lenguaje natural mejora la comprensión contextual, pero no elimina todas las dificultades semánticas. Los equipos deben definir vocabularios, revisar categorías y validar los agrupamientos propuestos. Sin esta intervención, la automatización puede unir textos por

semejanza lingüística y ocultar diferencias epistemológicas relevantes para la interpretación científica.

La reproducibilidad requiere registrar criterios de clasificación, versiones de modelos, parámetros, reglas de preprocesamiento y decisiones de depuración. Semmelrock et al. (2025) identifican la falta de documentación como una barrera recurrente en investigaciones basadas en aprendizaje automático. En análisis textual, pequeñas variaciones en tokenización, diccionarios o configuraciones pueden modificar los resultados. La transparencia metodológica permite examinar cómo se construyeron las categorías y si la síntesis obtenida puede reproducirse con el mismo corpus y las mismas condiciones de procesamiento.

La explicabilidad también influye en la confianza depositada en los resultados. Zárate (2022) sostiene que comprender los fundamentos de una decisión algorítmica facilita su evaluación. En clasificación científica, esto implica conocer qué términos, relaciones o características llevaron al sistema a ubicar un documento dentro de determinada categoría. Cuando el modelo funciona como una caja negra, resulta difícil distinguir entre una asociación válida y un patrón espurio. Los investigadores deben priorizar métodos que permitan justificar, revisar y corregir decisiones automatizadas.

La clasificación, síntesis y análisis semántico fortalecen la gestión del conocimiento cuando transforman grandes colecciones documentales en estructuras comprensibles y útiles para la investigación. Su aporte principal consiste en facilitar conexiones, reducir carga operativa y apoyar la interpretación de tendencias. No obstante, estas funciones requieren validación humana, trazabilidad y control de sesgos. Krenn et al. (2022) señalan que la comprensión científica depende de vincular resultados con explicaciones y marcos conceptuales. La automatización aporta organización; el sentido final corresponde al investigador.

4.4 Sistemas de recomendación para la identificación de fuentes y tendencias de investigación

Los sistemas de recomendación académica emplean algoritmos para identificar publicaciones, autores, temas y vínculos potencialmente relevantes para una investigación. Su funcionamiento se apoya en similitudes entre documentos, patrones de consulta, redes de citación y relaciones semánticas. Estas herramientas pueden ampliar el acceso a literatura especializada y reducir el tiempo dedicado a búsquedas manuales. Polo

y Casique (2025) muestran que la combinación de grafos de conocimiento y redes neuronales fortalece la recuperación documental al conectar contenidos que no siempre comparten los mismos términos.

La recomendación basada en contenido analiza características de los documentos, como palabras clave, resúmenes, métodos y conceptos principales. A partir de estas propiedades, sugiere publicaciones semejantes a las ya consultadas por el investigador. Esta estrategia facilita la profundización temática y la identificación de antecedentes cercanos. Sin embargo, también puede limitar la exploración a perspectivas conocidas y reforzar una lectura excesivamente homogénea. La calidad de las sugerencias depende de la representación semántica utilizada y de la diversidad real del corpus disponible (Garzón et al., 2025).

Los enfoques colaborativos utilizan patrones de comportamiento de usuarios con intereses semejantes para recomendar fuentes. Si varios investigadores consultan o citan documentos relacionados, el sistema puede sugerirlos a otros perfiles cercanos. Esta lógica favorece el descubrimiento de literatura relevante, aunque también introduce riesgos de popularidad acumulativa. Ross et al. (2022) advierten que la ciencia abierta puede reproducir ventajas previas y amenazas a la equidad. Un sistema basado en frecuencia puede fortalecer la visibilidad de trabajos ya reconocidos y relegar investigaciones menos citadas.

Los grafos de conocimiento permiten superar parcialmente estas limitaciones al representar relaciones entre autores, instituciones, conceptos, métodos y conjuntos de datos. Ávila (2024) explica que estas estructuras facilitan la visualización de conexiones dentro de datos abiertos de investigación. En recomendación académica, pueden mostrar rutas alternativas entre campos y revelar vínculos interdisciplinarios. No obstante, la existencia de una conexión en el grafo no garantiza relevancia científica. Cada recomendación debe revisarse según el problema, el enfoque metodológico y la calidad de la fuente sugerida.

La identificación de tendencias se apoya en el análisis de frecuencia, crecimiento temporal, coaparición temática y cambios en redes de publicación. Los sistemas pueden detectar asuntos emergentes y áreas con mayor actividad académica. Mata et al. (2024) reconocen que la inteligencia artificial puede apoyar la organización y análisis de información científica. Sin embargo, una tendencia documental no equivale

necesariamente a un avance sólido. Puede reflejar modas, disponibilidad de financiamiento o concentración editorial, por lo que debe interpretarse con cautela y contraste crítico.

Figura 8

Desafíos en los sistemas de recomendación académica



Nota. La figura presenta los principales desafíos de los sistemas de recomendación académica relacionados con tendencias, relevancia, sesgos y limitación de perspectivas. Elaboración propia.

La recomendación también puede contribuir a detectar vacíos de investigación cuando identifica combinaciones temáticas poco exploradas o relaciones débiles entre campos. Arias y Artigas (2022) sostienen que una brecha debe traducirse en problemas relevantes, delimitados y científicamente abordables. Por ello, la ausencia de conexiones en una red no demuestra por sí sola una oportunidad válida. El investigador debe examinar si el vacío responde a limitaciones del corpus, diferencias terminológicas, falta de acceso o una necesidad real de conocimiento.

Los sistemas de recomendación pueden favorecer la interdisciplinariedad al sugerir literatura procedente de áreas distintas. Esta función resulta valiosa cuando un problema exige integrar enfoques, datos y métodos heterogéneos. Krenn et al. (2022) señalan que

la comprensión científica requiere vincular resultados con explicaciones y estructuras conceptuales. La recomendación interdisciplinaria adquiere valor cuando ayuda a construir esas conexiones, pero puede generar asociaciones superficiales si solo se basa en proximidad lingüística o coincidencias estadísticas entre documentos.

La personalización mejora la pertinencia de las sugerencias, aunque puede crear burbujas informativas. Un sistema que aprende del historial de consultas tenderá a reforzar temas, autores y enfoques ya utilizados. Esto reduce la exposición a perspectivas críticas, regiones menos visibles o marcos teóricos alternativos. Zemelman (2021) destaca la importancia del pensamiento epistémico para cuestionar esquemas cerrados de interpretación. En consecuencia, los investigadores deben combinar recomendaciones personalizadas con búsquedas deliberadamente amplias y criterios de diversidad académica.

La transparencia del sistema influye en la confianza que puede depositarse en sus resultados. Zárate (2022) sostiene que la explicabilidad permite comprender los fundamentos de las decisiones algorítmicas. En recomendación académica, esto implica conocer por qué una fuente fue sugerida, qué variables influyeron y qué datos alimentaron el modelo. Cuando estas razones permanecen ocultas, el investigador no puede evaluar adecuadamente la pertinencia ni detectar sesgos. Las interfaces deberían mostrar criterios de similitud, redes de relación y limitaciones de cobertura.

La reproducibilidad exige registrar consultas, filtros, fechas, versiones de herramientas y criterios utilizados para aceptar o descartar recomendaciones. Semmelrock et al. (2025) advierten que la documentación insuficiente dificulta repetir investigaciones basadas en aprendizaje automático. Los sistemas académicos pueden modificar sus algoritmos o actualizar sus bases, generando resultados diferentes ante una misma consulta. Mantener un registro del proceso permite reconstruir cómo se formó el corpus y valorar qué influencia tuvo la recomendación automatizada en la revisión de literatura.

Los sistemas de recomendación fortalecen la identificación de fuentes y tendencias cuando se utilizan como instrumentos exploratorios sujetos a control metodológico. Su valor reside en ampliar conexiones, facilitar descubrimientos y organizar grandes volúmenes de información. Sin embargo, no sustituyen la evaluación crítica de la calidad, la pertinencia ni la diversidad de las publicaciones. Penabad et al. (2024) destacan la

necesidad de transparentar el uso de inteligencia artificial en entornos académico científicos. La selección final continúa siendo responsabilidad del investigador.

4.5 Repositorios inteligentes, datos abiertos y preservación del conocimiento

Los repositorios inteligentes constituyen infraestructuras digitales diseñadas para almacenar, describir, relacionar y recuperar productos científicos mediante funciones automatizadas. A diferencia de un archivo documental convencional, integran metadatos, sistemas de búsqueda semántica, clasificación automática y mecanismos de recomendación. Su propósito es facilitar el acceso organizado a artículos, bases de datos, códigos, informes y materiales complementarios. Ávila (2024) muestra que la representación mediante ontologías y grafos de conocimiento fortalece la conexión entre recursos y mejora su reutilización dentro de entornos académicos.

La inteligencia aplicada a los repositorios permite reconocer relaciones entre autores, temas, proyectos, instituciones y conjuntos de datos. Estas conexiones amplían la capacidad de localizar información más allá de palabras clave aisladas. Polo y Casique (2025) proponen integrar grafos de conocimiento y redes neuronales para mejorar la recuperación documental. En este tipo de sistemas, la calidad de las recomendaciones depende de la consistencia de los metadatos y de la precisión con que se representen las entidades. Una relación técnicamente posible no siempre implica pertinencia científica.

Los datos abiertos favorecen la transparencia, la reutilización y la verificación de los resultados científicos. Su disponibilidad permite que otros investigadores examinen procedimientos, desarrollen análisis secundarios o comparen hallazgos en contextos distintos. Sin embargo, abrir datos no significa publicarlos sin preparación. Deben acompañarse de descripciones, formatos comprensibles, criterios de calidad y condiciones claras de uso. Ross et al. (2022) advierten que la ciencia abierta puede ampliar oportunidades, aunque también reproducir desigualdades cuando el acceso técnico y las capacidades de reutilización se distribuyen de manera desigual.

La apertura debe equilibrarse con la protección de información sensible. Algunos conjuntos contienen datos personales, clínicos, educativos o institucionales que no pueden difundirse sin restricciones. Andrade et al. (2024) analizan la necesidad de sustentar legalmente el tratamiento de datos personales en sistemas de inteligencia artificial. Por ello, los repositorios deben incorporar niveles de acceso, anonimización, licencias y mecanismos de autorización. La transparencia científica no debe confundirse

con exposición indiscriminada, especialmente cuando la reutilización puede afectar derechos, confidencialidad o compromisos adquiridos con los participantes.

La preservación digital busca garantizar que los productos científicos permanezcan accesibles, interpretables y utilizables a lo largo del tiempo. Esto exige conservar archivos, versiones, formatos, metadatos, relaciones y condiciones de producción. Tapia y Recalde (2024) destacan la importancia de una gestión documental sistemática para mantener orden, continuidad y disponibilidad de la información. En investigación, la preservación debe incluir no solo publicaciones finales, sino también datos, códigos, instrumentos, protocolos y decisiones metodológicas necesarias para comprender cómo se obtuvieron los resultados.

Tabla 11
Componentes de un repositorio inteligente orientado a la preservación científica

Componente	Función principal	Aporte a la investigación	Riesgo si se gestiona inadecuadamente
Metadatos estructurados	Describir recursos y facilitar su localización	Mejora búsqueda, citación y reutilización	Pérdida de contexto y duplicación
Grafos de conocimiento	Relacionar autores, datos, métodos y proyectos	Permite navegación semántica y análisis de redes	Asociaciones imprecisas o incompletas
Control de versiones	Registrar modificaciones y estados previos	Favorece trazabilidad y reproducibilidad	Confusión sobre el archivo utilizado
Datos abiertos	Facilitar acceso y análisis secundario	Amplía verificación y colaboración	Uso indebido o interpretación descontextualizada
Preservación digital	Mantener accesibilidad a largo plazo	Protege la memoria científica institucional	Obsolescencia y pérdida de información
Gestión de permisos	Regular acceso según sensibilidad	Protege privacidad y confidencialidad	Exposición no autorizada o restricciones excesivas
Formatos interoperables	Permitir intercambio entre plataformas	Facilita integración y reutilización	Dependencia tecnológica y aislamiento documental

Nota. La tabla sintetiza los componentes que permiten combinar acceso, organización, seguridad y preservación en repositorios científicos asistidos por inteligencia artificial.

El control de versiones ocupa un lugar central porque los datos y documentos científicos cambian durante el desarrollo de una investigación. Corregir errores, actualizar registros o modificar códigos puede alterar los resultados obtenidos. Los repositorios deben

identificar qué versión se utilizó en cada análisis y conservar los estados anteriores cuando resulten relevantes. Semmelrock et al. (2025) señalan que la reproducibilidad en investigaciones basadas en aprendizaje automático depende de una documentación detallada. Sin control de versiones, resulta difícil reconstruir el flujo de trabajo y verificar conclusiones.

La interoperabilidad permite que los repositorios intercambien información sin perder significado, contexto o trazabilidad. Ávila (2024) destaca el valor de modelos como DDI-RDF y DataCite Ontology para representar datos de investigación. Estos recursos favorecen conexiones entre plataformas, proyectos y publicaciones, pero requieren estándares compartidos y descripciones coherentes. La interoperabilidad no se limita a transferir archivos; implica que otros sistemas puedan interpretar correctamente autores, variables, métodos, licencias y relaciones. Cuando estas condiciones faltan, los datos quedan técnicamente accesibles, pero científicamente difíciles de reutilizar.

La preservación también enfrenta riesgos de obsolescencia tecnológica. Formatos, programas y plataformas pueden dejar de utilizarse, impidiendo abrir archivos o ejecutar códigos desarrollados años atrás. La gestión inteligente debe prever migraciones, documentación de dependencias y conservación de entornos computacionales. Loza et al. (2026) destacan la necesidad de reproducibilidad en flujos bioinformáticos aplicados a genómica y oncología clínica. En este tipo de investigaciones, preservar únicamente los resultados finales resulta insuficiente si no se mantienen las condiciones necesarias para reconstruir el análisis.

Los repositorios inteligentes pueden fortalecer la memoria institucional al reunir productos generados por equipos, laboratorios y proyectos en una estructura compartida. Esta función evita dispersión, pérdida de información y repetición innecesaria de esfuerzos. Cedeño et al. (2024) señalan que la inteligencia artificial puede apoyar la organización de procesos investigativos en el ámbito universitario. Sin embargo, la sostenibilidad del repositorio depende de políticas, personal especializado, recursos técnicos y responsabilidades claras. Una plataforma sin mantenimiento puede convertirse rápidamente en un archivo incompleto y desactualizado.

La preservación del conocimiento científico requiere equilibrar apertura, seguridad, calidad e interoperabilidad. Los repositorios inteligentes aportan mecanismos para organizar relaciones, facilitar búsquedas y sostener la reutilización de recursos, pero su

eficacia depende de decisiones institucionales y metodológicas coherentes. Penabad et al. (2024) destacan la importancia de transparentar el uso de inteligencia artificial en entornos académico científicos. Esa transparencia debe extenderse a la gestión documental, indicando cómo se clasifican, recomiendan, modifican y conservan los materiales depositados.

4.6 Grafos de conocimiento, redes de citación y visualización de dominios científicos

Los grafos de conocimiento representan información científica mediante entidades conectadas por relaciones explícitas y semánticamente definidas. Autores, publicaciones, instituciones, conceptos, métodos, datos y proyectos pueden organizarse como nodos dentro de una estructura que permite explorar vínculos complejos. Ávila (2024) señala que la integración de ontologías y estándares de metadatos favorece la representación de datos abiertos de investigación. Esta arquitectura supera la organización lineal de los repositorios tradicionales y facilita consultas orientadas a comprender cómo se relacionan distintos componentes de un dominio científico.

La construcción de un grafo comienza con la identificación de entidades relevantes y la definición de los vínculos que pueden establecerse entre ellas. Un artículo puede relacionarse con sus autores, referencias, palabras clave, metodología, conjunto de datos e institución de procedencia. Estas conexiones requieren vocabularios controlados y reglas consistentes para evitar duplicaciones o ambigüedades. Cuando los metadatos son incompletos, el grafo puede representar relaciones fragmentadas y generar interpretaciones incorrectas. La calidad de la estructura depende tanto de la tecnología como del rigor documental aplicado (Ávila, 2024).

Las ontologías aportan una base conceptual para definir categorías, propiedades y relaciones dentro del grafo. Su función consiste en establecer qué significa cada entidad y bajo qué condiciones puede vincularse con otras. Esto permite que los sistemas interpreten la información más allá de coincidencias textuales. Polo y Casique (2025) destacan que la combinación de grafos de conocimiento y redes neuronales mejora la recuperación documental. No obstante, la utilidad semántica depende de que las categorías utilizadas reflejen adecuadamente las particularidades teóricas y metodológicas del campo estudiado.

Las redes de citación muestran cómo las publicaciones se conectan mediante referencias bibliográficas. Cada documento funciona como un nodo y cada cita establece un vínculo dirigido hacia otra obra. El análisis de estas redes permite reconocer trabajos influyentes, comunidades temáticas y trayectorias de desarrollo científico. Sin embargo, la cantidad de citas no equivale necesariamente a calidad, validez o vigencia. Las citaciones también responden a dinámicas disciplinares, visibilidad editorial y ventajas acumulativas que pueden favorecer a determinados autores o instituciones (Ross et al., 2022).

El análisis de co-citación examina la frecuencia con que dos documentos aparecen citados conjuntamente por publicaciones posteriores. Esta relación puede indicar proximidad conceptual, pertenencia a una tradición teórica o participación en un mismo campo de discusión. Su representación permite identificar bases intelectuales compartidas y cambios en la estructura de una disciplina. Aun así, la co-citación no demuestra acuerdo entre los trabajos vinculados, pues dos estudios pueden citarse juntos precisamente por sostener posiciones contradictorias. La interpretación requiere revisar el contexto específico de cada referencia.

El acoplamiento bibliográfico sigue una lógica diferente, porque relaciona documentos que comparten referencias dentro de sus respectivas bibliografías. Cuanto mayor es el número de fuentes comunes, mayor puede ser su proximidad temática o metodológica. Esta técnica facilita identificar frentes de investigación recientes, debido a que no depende de citas recibidas posteriormente. Sin embargo, su eficacia está condicionada por la cobertura de la base documental y por las prácticas de citación de cada disciplina. Las conexiones calculadas deben complementarse con lectura crítica y análisis conceptual.

La visualización de dominios científicos transforma redes complejas en mapas que muestran agrupamientos, proximidades y posiciones relativas. Ávila (2024) sostiene que los grafos favorecen la representación de relaciones dentro de datos abiertos de investigación. Los mapas resultantes pueden revelar comunidades temáticas, áreas centrales, zonas periféricas y conexiones interdisciplinarias. No obstante, la disposición espacial responde a algoritmos específicos y no reproduce literalmente la estructura del conocimiento. La cercanía visual debe interpretarse como una representación analítica, no como prueba automática de relación causal o equivalencia conceptual.

Los mapas de coocurrencia de términos permiten examinar conceptos que aparecen conjuntamente en títulos, resúmenes o palabras clave. Esta técnica ayuda a reconocer

vocabularios dominantes, temas emergentes y transformaciones terminológicas. El procesamiento del lenguaje natural puede mejorar la identificación de equivalencias y relaciones contextuales entre expresiones científicas (Garzón et al., 2025). Sin embargo, los resultados dependen de la limpieza del corpus, la normalización de sinónimos y la resolución de términos ambiguos. Sin estos controles, una misma idea puede fragmentarse en varios nodos.

La inteligencia artificial amplía el análisis de grafos al detectar comunidades, predecir vínculos y recomendar conexiones entre entidades. Polo y Casique (2025) proponen integrar redes neuronales con grafos de conocimiento para fortalecer la recuperación de información documental. Estas técnicas pueden descubrir relaciones implícitas entre publicaciones y áreas científicas, aunque también generan asociaciones difíciles de explicar. Una predicción algorítmica debe distinguirse de una relación documental comprobada. El investigador necesita verificar si el vínculo sugerido posee fundamento teórico, metodológico o empírico dentro del dominio analizado.

Las redes de citación también pueden evidenciar desigualdades en la producción y circulación del conocimiento. Autores, revistas e instituciones con mayor visibilidad tienden a recibir más citas, lo que refuerza su posición central dentro de los mapas científicos. Ross et al. (2022) advierten que las dinámicas de ventaja acumulativa pueden amenazar la equidad en la ciencia abierta. Por ello, la posición periférica de una publicación no debe interpretarse automáticamente como falta de relevancia. Puede reflejar barreras lingüísticas, regionales, económicas o editoriales que limitan su integración.

La reproducibilidad de estos análisis exige documentar las bases consultadas, periodos cubiertos, criterios de depuración, unidades analizadas y algoritmos de visualización. Pequeñas modificaciones en los datos o parámetros pueden cambiar la forma de las comunidades y la centralidad de los nodos. Semmelrock et al. (2025) identifican la documentación insuficiente como una barrera frecuente en investigaciones basadas en aprendizaje automático. En los estudios de dominios científicos, también deben conservarse versiones del corpus y reglas utilizadas para unificar autores, instituciones, términos y referencias duplicadas.

Los grafos de conocimiento, las redes de citación y la visualización de dominios científicos permiten comprender cómo se organiza, conecta y transforma la producción

académica. Su valor reside en convertir colecciones dispersas en estructuras explorables que apoyan la identificación de tendencias, comunidades y relaciones interdisciplinarias. Sin embargo, los mapas no sustituyen la interpretación teórica ni la evaluación de las fuentes. Krenn et al. (2022) sostienen que la comprensión científica requiere vincular resultados con explicaciones conceptuales. La visualización orienta el análisis; el investigador construye su significado.

4.7 Transferencia, reutilización y circulación del conocimiento asistida por inteligencia artificial

La transferencia del conocimiento científico comprende los procesos mediante los cuales los resultados de investigación se comunican, adaptan y aplican en nuevos contextos. La inteligencia artificial puede acelerar esta circulación al clasificar contenidos, identificar destinatarios y transformar información especializada en formatos más accesibles. Su utilidad no depende solo de difundir con rapidez, sino de conservar la precisión conceptual y metodológica del conocimiento original. Cedeño et al. (2024) señalan que estas tecnologías fortalecen distintas fases de la actividad investigativa, incluida la organización y comunicación académica.

La reutilización implica emplear datos, métodos, modelos o resultados previos para responder nuevas preguntas científicas. Esta práctica puede reducir duplicaciones, ampliar comparaciones y favorecer investigaciones acumulativas. Los sistemas inteligentes facilitan la localización de recursos relacionados y la identificación de posibles usos secundarios. Sin embargo, un material disponible no siempre resulta adecuado para otro estudio. Deben revisarse su procedencia, calidad, condiciones de producción y compatibilidad con el nuevo problema, especialmente cuando los contextos, poblaciones o escalas analíticas son diferentes (Ávila, 2024).

Los repositorios inteligentes desempeñan una función central al conectar publicaciones, conjuntos de datos, códigos y proyectos mediante metadatos y relaciones semánticas. Polo y Casique (2025) muestran que la integración de grafos de conocimiento y redes neuronales mejora la recuperación documental. Esta capacidad favorece la circulación porque permite encontrar recursos mediante conceptos relacionados y no únicamente por coincidencias literales. No obstante, la reutilización depende de que los materiales estén correctamente descritos, versionados y acompañados de información suficiente para comprender su alcance y sus limitaciones.

La inteligencia artificial también puede adaptar contenidos científicos a públicos diversos mediante resúmenes, traducciones, visualizaciones y reformulaciones. Esta función facilita el diálogo entre disciplinas, instituciones y sectores sociales. Garzón et al. (2025) destacan el potencial del procesamiento del lenguaje natural para mejorar la comunicación entre personas y sistemas inteligentes. Sin embargo, la simplificación puede eliminar matices, alterar conceptos o presentar resultados preliminares como conclusiones firmes. Toda adaptación debe preservar la fidelidad al contenido original y señalar los límites de la evidencia comunicada.

La transferencia hacia ámbitos aplicados requiere algo más que difundir publicaciones. El conocimiento debe relacionarse con necesidades institucionales, capacidades técnicas y condiciones sociales concretas. Duedra et al. (2020) proponen superar la compartimentalización entre ciencia básica, aplicada, extensión y divulgación. La inteligencia artificial puede apoyar este tránsito mediante análisis de pertinencia, identificación de actores y organización de evidencias. Aun así, la adopción final depende de procesos de interpretación, negociación y adaptación que no pueden resolverse exclusivamente mediante automatización.

Tabla 12
Modalidades de transferencia y reutilización del conocimiento asistidas por inteligencia artificial

Modalidad	Función de la inteligencia artificial	Requisito científico	Riesgo principal
Recuperación de recursos	Localiza datos, artículos, códigos y proyectos relacionados	Metadatos completos y criterios de pertinencia	Reutilización descontextualizada
Adaptación de contenidos	Resume, traduce y reformula información especializada	Verificación de precisión conceptual	Simplificación o distorsión
Reutilización de datos	Identifica conjuntos compatibles con nuevas preguntas	Revisión de calidad, licencia y procedencia	Uso indebido o sesgo
Transferencia interdisciplinaria	Relaciona conceptos y métodos entre campos	Compatibilidad epistemológica y metodológica	Equivalencias artificiales
Difusión social	Personaliza formatos para distintos públicos	Claridad sobre límites y nivel de evidencia	Exageración de resultados

Modalidad	Función de la inteligencia artificial	Requisito científico	Riesgo principal
Transferencia institucional	Organiza evidencia para apoyar decisiones	Evaluación contextual y supervisión humana	Dependencia de recomendaciones algorítmicas

Nota. La tabla sintetiza formas de circulación del conocimiento y muestra que la asistencia inteligente debe acompañarse de trazabilidad, validación y adaptación contextual.

Los sistemas de recomendación pueden ampliar la circulación al sugerir fuentes y resultados a investigadores con intereses relacionados. Sin embargo, la personalización puede reforzar circuitos cerrados de información y reducir la exposición a perspectivas alternativas. Ross et al. (2022) advierten que las dinámicas de ventaja acumulativa afectan la equidad en la ciencia abierta. Cuando los algoritmos priorizan popularidad, citación o visibilidad previa, pueden dejar fuera contribuciones regionales o publicaciones menos reconocidas. La transferencia responsable requiere criterios de diversidad y cobertura.

La ciencia abierta fortalece la reutilización al facilitar acceso a datos, publicaciones y materiales metodológicos. No obstante, la apertura técnica no garantiza que todos los actores puedan aprovecharlos en igualdad de condiciones. Se requieren competencias, infraestructura y tiempo para interpretar y transformar esos recursos. Ávila (2024) muestra que los estándares y ontologías favorecen la conexión entre datos de investigación, pero su eficacia depende de capacidades institucionales. La circulación del conocimiento debe considerar no solo disponibilidad, sino también comprensibilidad, interoperabilidad y posibilidades reales de uso.

La protección de datos condiciona la reutilización cuando la información contiene elementos personales, clínicos, educativos o institucionales. Andrade et al. (2024) destacan la necesidad de sustentar legalmente el tratamiento de datos en sistemas de inteligencia artificial. Antes de transferir o reutilizar un conjunto, deben revisarse consentimiento, anonimización, seguridad y restricciones de acceso. Un dato útil para la investigación original puede generar riesgos distintos al emplearse en otro contexto. La reutilización responsable exige evaluar nuevamente su pertinencia y las consecuencias potenciales.

La circulación asistida por inteligencia artificial también requiere transparencia sobre las transformaciones aplicadas. Los resúmenes, traducciones, clasificaciones y recomendaciones deben poder rastrearse hasta los materiales originales. Penabad et al.

(2024) resaltan la importancia de reportar el uso de inteligencia artificial en entornos científico-académicos. Este principio permite distinguir entre contenido original, procesamiento automatizado e interpretación humana. Sin trazabilidad, resulta difícil identificar errores, corregir distorsiones o determinar qué parte del conocimiento fue modificada durante su transferencia hacia nuevos usuarios o contextos.

La transferencia, reutilización y circulación del conocimiento alcanzan valor científico cuando amplían su alcance sin debilitar su calidad ni su contexto. La inteligencia artificial facilita conexiones, adaptaciones y búsquedas, pero no garantiza por sí sola una aplicación pertinente. Semmelrock et al. (2025) destacan que la reproducibilidad depende de documentación, acceso y condiciones organizacionales. De manera semejante, la reutilización responsable exige conservar metadatos, versiones, restricciones y decisiones metodológicas. La tecnología amplía la circulación; la gobernanza científica define sus límites y responsabilidades.



Capítulo

5.

Calidad, ética e integridad en la investigación con inteligencia artificial

Carlos Alejandro Segarra Ramos

5.1 Calidad, confiabilidad y validación de resultados generados con inteligencia artificial

La calidad de los resultados generados con inteligencia artificial depende de la correspondencia entre el problema científico, los datos utilizados, el modelo seleccionado y los criterios de evaluación. Un resultado preciso en apariencia puede ser metodológicamente débil si se obtuvo con información incompleta, variables mal definidas o procedimientos opacos. Ciobanu et al. (2024) destacan que la investigación basada en aprendizaje automático debe atender simultáneamente reproducibilidad e interpretabilidad. La confiabilidad no surge del rendimiento técnico aislado, sino de la coherencia integral del proceso investigativo.

La validación debe comenzar antes de ejecutar el modelo, mediante la revisión de la calidad de los datos y de sus condiciones de producción. Registros duplicados, valores faltantes, errores de etiquetado o muestras poco representativas pueden alterar significativamente los resultados. Los sistemas inteligentes aprenden de las regularidades disponibles, incluidas sus inconsistencias y sesgos. Por ello, documentar procedencia, criterios de inclusión, transformaciones y exclusiones constituye una exigencia central para valorar la solidez de las conclusiones obtenidas (Simmelrock et al., 2025; Hu, 2021).

Kapoor y Narayanan (2023) advierten que la fuga de datos representa una amenaza importante para la validez de investigaciones basadas en aprendizaje automático. Este problema aparece cuando información destinada a la evaluación final influye directa o indirectamente en el entrenamiento. Como consecuencia, el modelo puede mostrar un rendimiento artificialmente elevado que no se mantiene ante datos nuevos. Para evitarlo, los conjuntos de entrenamiento, validación y prueba deben separarse correctamente y administrarse mediante procedimientos definidos antes del análisis definitivo.

La validación científica no consiste únicamente en confirmar que un modelo alcanza resultados favorables. Desde las propuestas de Popper y Kuhn, las teorías y explicaciones científicas permanecen abiertas a contrastación, revisión y eventual sustitución cuando la evidencia muestra sus limitaciones (Humpiri et al., 2021). En consecuencia, los resultados producidos mediante inteligencia artificial deben someterse a pruebas críticas capaces de revelar errores, condiciones de fracaso y límites de generalización.

La confiabilidad también exige seleccionar métricas adecuadas para el tipo de problema científico. La precisión global puede ocultar errores importantes cuando existen clases

desbalanceadas, poblaciones heterogéneas o consecuencias distintas según el tipo de fallo. En tales casos, conviene analizar sensibilidad, especificidad, estabilidad y desempeño por subgrupos. Un indicador elevado no demuestra por sí solo que el sistema sea útil o justo. La evaluación debe relacionar cada métrica con el propósito del estudio y con las decisiones que podrían derivarse de sus resultados (Kuppler et al., 2022).

La validación interna permite examinar el comportamiento del modelo dentro del conjunto de datos disponible, pero no garantiza su funcionamiento en otros contextos. Ciobanu et al. (2024) señalan que los modelos aplicados en medicina necesitan pruebas rigurosas de reproducibilidad e interpretación. Esta exigencia resulta aplicable a cualquier disciplina. Un sistema puede aprender patrones específicos de una institución, población o equipo de medición y perder rendimiento cuando se enfrenta a condiciones distintas. La validación externa es esencial para evaluar generalización y transferibilidad.

La interpretabilidad contribuye a determinar si las relaciones identificadas por el sistema poseen sentido científico. Joyce et al. (2023) relacionan la inteligencia artificial explicable con transparencia, interpretabilidad y comprensión. Conocer qué variables influyen en una predicción permite detectar asociaciones espurias, dependencias indebidas o decisiones difíciles de justificar. No todos los modelos ofrecen el mismo nivel de explicación, pero el investigador debe describir sus fundamentos, limitaciones y criterios de uso. La opacidad reduce la posibilidad de revisar críticamente los resultados y establecer confianza.

La explicabilidad no debe confundirse con una presentación simplificada del resultado. Zárate (2022) sostiene que explicar una decisión algorítmica implica hacer comprensibles sus fundamentos y efectos. Una visualización atractiva o una lista de variables importantes puede ser insuficiente si no muestra cómo interactúan los factores o bajo qué condiciones cambia la predicción. La validación requiere comprobar que la explicación sea estable, relevante y coherente con el conocimiento disciplinar. De lo contrario, puede ofrecer una apariencia de transparencia sin resolver la incertidumbre metodológica.

Los modelos generativos plantean desafíos adicionales porque pueden producir respuestas fluidas, convincentes y formalmente correctas sin respaldo suficiente. Gao et al. (2023) demostraron que resúmenes científicos generados por ChatGPT pueden parecer plausibles incluso para revisores humanos. Esta capacidad obliga a distinguir calidad lingüística de validez científica. Cada afirmación, dato y referencia debe contrastarse con

fuentes verificables. La confiabilidad no puede evaluarse solo por coherencia textual, sino por la correspondencia entre el contenido generado, la evidencia disponible y el alcance real de las fuentes.

Gilbert et al. (2024) proponen modelos aumentados orientados a reducir alucinaciones y mejorar la curación de información médica. Este enfoque muestra que conectar la generación con fuentes identificables puede fortalecer la confiabilidad. Sin embargo, la recuperación documental no elimina todos los errores, porque el sistema puede seleccionar evidencia poco pertinente, interpretar incorrectamente un texto o combinar afirmaciones incompatibles. La validación humana continúa siendo necesaria para revisar la calidad de las fuentes, el contexto de las citas y la fidelidad de la síntesis producida.

Los sesgos afectan la calidad cuando el modelo ofrece rendimientos desiguales entre grupos o reproduce categorías problemáticas. Omiye et al. (2023) evidencian que los grandes modelos de lenguaje pueden propagar enfoques médicos asociados con categorías raciales. Un sistema puede parecer confiable en términos generales y, al mismo tiempo, producir resultados perjudiciales para poblaciones específicas. La validación debe incluir análisis desagregados, revisión de representatividad y examen de las consecuencias prácticas. La calidad científica también comprende equidad, pertinencia contextual y ausencia de discriminación injustificada.

La reproducibilidad requiere conservar datos, código, versiones, parámetros, semillas aleatorias y condiciones computacionales. Semmelrock et al. (2025) identifican barreras técnicas y organizacionales que dificultan repetir investigaciones basadas en aprendizaje automático. Una descripción general del modelo no basta si otros equipos no pueden reconstruir el flujo analítico. La documentación debe abarcar decisiones de limpieza, criterios de selección, ajustes realizados y procedimientos para elegir el resultado final. Sin esta trazabilidad, la validación independiente se vuelve limitada y la confianza depende excesivamente de los autores originales.

La calidad, confiabilidad y validación de resultados generados con inteligencia artificial exigen una evaluación continua y no una comprobación aislada al final del estudio. Contreras et al. (2024) destacan que el uso científico de estas tecnologías debe acompañarse de transparencia, supervisión y responsabilidad. La validez depende de datos adecuados, métricas pertinentes, pruebas externas, interpretabilidad y documentación completa. La inteligencia artificial puede ampliar la capacidad analítica,

pero corresponde al investigador demostrar que sus resultados son sólidos, reproducibles, comprensibles y científicamente justificables.

5.2 Alucinaciones, errores y límites de los sistemas generativos

Las alucinaciones de los sistemas generativos consisten en producir afirmaciones plausibles, pero incorrectas, incompletas o carentes de respaldo verificable. Su gravedad aumenta en la investigación científica porque pueden presentarse como datos, referencias, resultados o explicaciones aparentemente confiables. Gao et al. (2023) demostraron que los resúmenes generados por ChatGPT pueden parecer suficientemente creíbles para dificultar su identificación por parte de revisores humanos. La fluidez lingüística, por tanto, no constituye evidencia de exactitud ni reemplaza la comprobación documental de cada contenido producido.

Los errores generativos se originan en la forma probabilística con que estos modelos producen texto. El sistema selecciona secuencias lingüísticas coherentes según patrones aprendidos, pero no verifica necesariamente la verdad de cada afirmación. Por ello, puede combinar conceptos incompatibles, alterar cifras, atribuir ideas a autores incorrectos o inventar referencias. La respuesta generada debe considerarse una propuesta textual preliminar y no una fuente científica autónoma. Su uso riguroso exige contrastar afirmaciones, fechas, nombres y conclusiones con los documentos originales (Castillo, 2023).

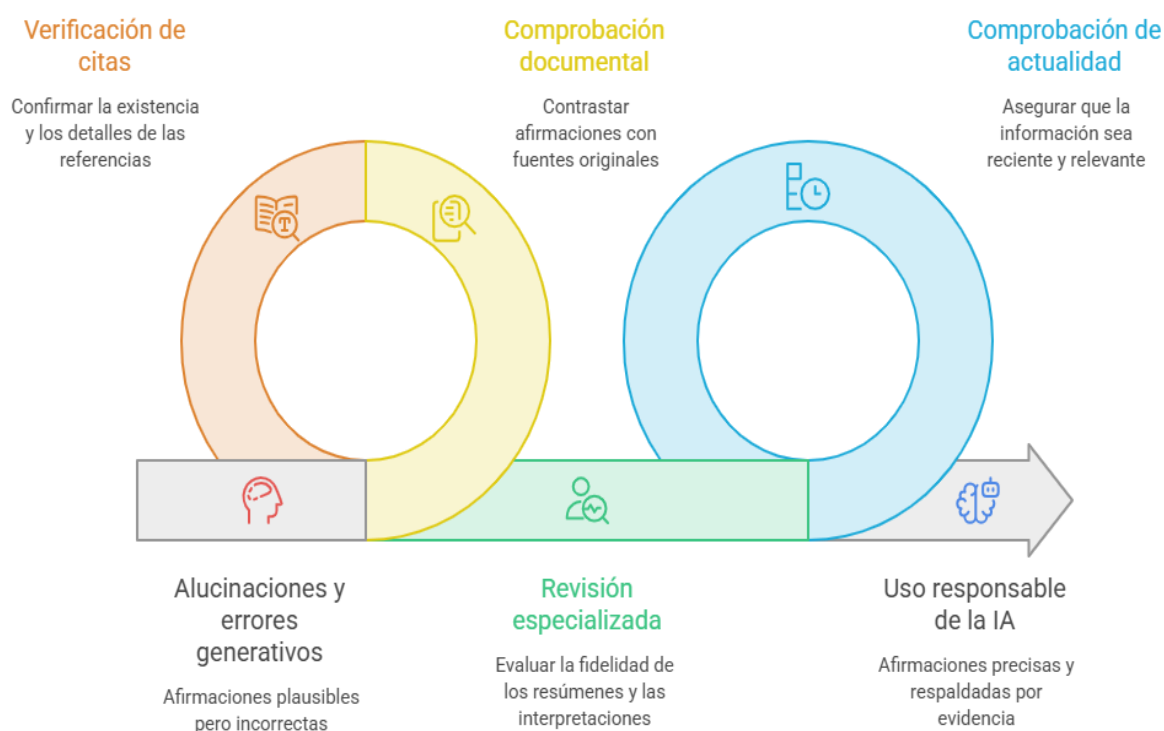
Una limitación frecuente consiste en fabricar citas bibliográficas con apariencia académica. Los modelos pueden generar títulos verosímiles, nombres de revistas, años, volúmenes o identificadores digitales que no corresponden a publicaciones reales. Este problema resulta especialmente peligroso cuando el investigador incorpora la referencia sin comprobarla. Penabad et al. (2024) destacan la necesidad de orientar y reportar el uso de inteligencia artificial en revistas científico-académicas. La verificación debe incluir la existencia de la fuente, su autoría, los datos editoriales y la correspondencia entre el texto citado y la idea atribuida.

Los sistemas también pueden distorsionar fuentes existentes mediante resúmenes inexactos o interpretaciones que exceden los hallazgos originales. Tang et al. (2023) evaluaron modelos de lenguaje en la síntesis de evidencia médica y mostraron que estas herramientas requieren revisión especializada. Una respuesta puede omitir limitaciones, confundir resultados preliminares con conclusiones firmes o eliminar diferencias entre

estudios. El riesgo no desaparece, aunque las referencias sean auténticas, porque la fidelidad de la síntesis depende de conservar métodos, contextos, niveles de certeza y alcances expresamente señalados por los autores.

La desactualización del conocimiento constituye otro límite relevante. Un modelo puede responder a partir de información incompleta, no incorporar publicaciones recientes o desconocer modificaciones institucionales y normativas. Incluso cuando utiliza materiales actuales, no siempre distingue entre versiones, correcciones o retractaciones. La producción académica requiere comprobar la vigencia de la evidencia antes de utilizarla. Gilbert et al. (2024) plantean el uso de modelos aumentados con recuperación de información para reducir respuestas médicas no fundamentadas, aunque esta estrategia tampoco elimina la necesidad de validar las fuentes recuperadas.

Figura 9
Mitigando los riesgos de los sistemas generativos



Nota. La figura presenta medidas para mitigar los riesgos de los sistemas generativos mediante la verificación de citas, fuentes, actualidad y revisión especializada. Elaboración propia.

Las alucinaciones pueden adquirir mayor apariencia de certeza cuando el usuario formula preguntas ambiguas o proporciona contexto insuficiente. El sistema tiende a completar vacíos mediante asociaciones plausibles, aunque no disponga de evidencia suficiente para responder con precisión. Cedeño et al. (2024) reconocen que la inteligencia artificial

puede apoyar la investigación universitaria, pero su aprovechamiento depende de capacidades críticas y metodológicas. Elaborar instrucciones claras, delimitar el propósito y aportar materiales verificables reduce errores, aunque no garantiza por sí solo que la respuesta sea correcta o científicamente pertinente.

Los sesgos también limitan la confiabilidad de los sistemas generativos. Omiye et al. (2023) evidencian que los grandes modelos de lenguaje pueden reproducir planteamientos médicos asociados con categorías raciales. Este problema muestra que una respuesta puede ser lingüísticamente coherente y, al mismo tiempo, incorporar supuestos discriminatorios o científicamente cuestionables. La revisión debe examinar no solo errores factuales, sino también categorías, omisiones y jerarquías presentes en el contenido. La calidad exige valorar cómo se representan poblaciones, regiones, disciplinas y perspectivas minoritarias.

La opacidad dificulta comprender por qué el sistema produjo una respuesta específica o seleccionó determinadas relaciones entre conceptos. Zárate (2022) sostiene que la explicabilidad permite examinar los fundamentos de una decisión algorítmica y valorar sus efectos. En los modelos generativos, esta exigencia resulta compleja porque no siempre es posible reconstruir el origen preciso de cada afirmación. La falta de trazabilidad limita la evaluación científica y obliga a conservar una separación clara entre texto generado, fuente consultada, interpretación humana y decisión final del investigador.

La supervisión humana constituye la principal barrera frente a la incorporación de errores en productos académicos. Castillo (2023) subraya que el uso de ChatGPT en escritura científica requiere control especializado. Revisar gramática o coherencia no es suficiente; también deben verificarse conceptos, evidencia, citas, datos y alcance de las conclusiones. El investigador conserva la responsabilidad por todo contenido presentado, aunque haya sido generado o reformulado por una herramienta. Delegar la redacción no implica transferir autoría, juicio metodológico ni obligación de responder por posibles inexactitudes.

Los límites de los sistemas generativos no impiden su uso, pero obligan a integrarlos dentro de procedimientos de verificación explícitos. Pueden apoyar organización, síntesis y reformulación, siempre que sus resultados sean tratados como insumos revisables. Contreras et al. (2024) señalan que el uso científico de inteligencia artificial exige

transparencia, responsabilidad y supervisión. La práctica rigurosa requiere contrastar con fuentes originales, declarar intervenciones significativas y descartar cualquier contenido que no pueda validarse. La utilidad tecnológica depende finalmente del control crítico ejercido por el investigador.

5.3 Sesgos algorítmicos y desigualdades en la producción científica

Los sesgos algorítmicos surgen cuando los sistemas de inteligencia artificial reproducen, amplifican o generan diferencias injustificadas a partir de los datos, categorías y decisiones incorporadas en su diseño. En la producción científica, estos sesgos pueden afectar la selección de literatura, la clasificación de evidencias, la interpretación de resultados y la visibilidad de determinados grupos. Omiye et al. (2023) demostraron que los grandes modelos de lenguaje pueden propagar enfoques médicos asociados con categorías raciales, incluso cuando sus respuestas parecen técnicamente coherentes y lingüísticamente convincentes.

Una fuente frecuente de sesgo se encuentra en los datos de entrenamiento. Si una base representa de manera desigual idiomas, regiones, poblaciones o disciplinas, el modelo aprenderá una visión parcial del conocimiento disponible. Esta limitación puede generar menor precisión en contextos poco representados y reforzar la centralidad de perspectivas dominantes. La calidad algorítmica no debe evaluarse únicamente mediante resultados generales, porque un rendimiento promedio satisfactorio puede ocultar errores sistemáticos sobre grupos específicos o áreas científicas periféricas (Kuppler et al., 2022).

Las desigualdades también se manifiestan en la selección automatizada de fuentes. Los sistemas de recomendación suelen priorizar documentos citados, consultados o publicados en espacios de alta visibilidad, lo que favorece dinámicas acumulativas. Ross et al. (2022) advierten que la ciencia abierta puede reproducir ventajas previas y amenazas a la equidad. En consecuencia, investigaciones valiosas procedentes de revistas regionales, idiomas distintos del inglés o instituciones con menor infraestructura pueden quedar relegadas, aunque resulten pertinentes para comprender problemas científicos situados.

El sesgo lingüístico afecta especialmente la recuperación, clasificación y síntesis de literatura. Los modelos entrenados con predominio de determinados idiomas suelen interpretar con mayor precisión sus estructuras y vocabularios, mientras reducen la visibilidad de otras tradiciones académicas. Zemelman (2021) destaca la importancia de

un pensamiento epistémico capaz de cuestionar marcos cerrados y categorías impuestas. Aplicado a la inteligencia artificial, este criterio exige revisar qué lenguas, conceptos y perspectivas quedan ausentes cuando los sistemas organizan automáticamente el conocimiento científico disponible.

Las categorías utilizadas para etiquetar datos también pueden introducir desigualdades. Una clasificación aparentemente neutral puede incorporar divisiones históricas, simplificaciones culturales o conceptos que no se ajustan a todas las poblaciones. Kontiainen et al. (2022) proponen una agenda interdisciplinaria para estudiar la equidad algorítmica desde experiencias relacionadas con el acceso a la justicia. Esta perspectiva muestra que los problemas de sesgo no se resuelven únicamente ajustando métricas, porque también requieren examinar el origen, significado y consecuencias institucionales de las categorías empleadas.

Tabla 13
Sesgos algorítmicos y sus efectos en la producción científica

Tipo de sesgo	Origen principal	Efecto sobre la investigación	Estrategia de control
Sesgo de representación	Datos desiguales por población, idioma o región	Resultados poco generalizables y mayor error en grupos subrepresentados	Diversificar fuentes y evaluar desempeño por subgrupos
Sesgo de selección	Criterios restrictivos o recomendaciones basadas en popularidad	Invisibilización de literatura periférica o emergente	Combinar búsqueda automatizada con revisión deliberadamente amplia
Sesgo de etiquetado	Categorías ambiguas, históricas o culturalmente inadecuadas	Clasificaciones distorsionadas y conclusiones discriminatorias	Validación interdisciplinaria de categorías
Sesgo de medición	Variables o indicadores que no representan adecuadamente el fenómeno	Interpretaciones parciales o asociaciones espurias	Revisar validez, contexto y condiciones de recolección
Sesgo de automatización	Confianza excesiva en resultados algorítmicos	Reducción del juicio crítico y aceptación de errores	Mantener supervisión humana y trazabilidad
Sesgo de visibilidad	Priorización de autores, revistas o instituciones dominantes	Reproducción de desigualdades en circulación y citación	Incorporar diversidad geográfica, lingüística e institucional

Nota. La tabla sintetiza fuentes de sesgo que pueden afectar la producción científica y destaca la necesidad de combinar controles técnicos, revisión contextual y participación interdisciplinaria.

La evaluación de equidad requiere analizar cómo se distribuyen errores, beneficios y riesgos entre distintos grupos. Kuppler et al. (2022) sostienen que una predicción considerada justa no conduce automáticamente a una decisión socialmente justa. En investigación, un modelo puede aplicar el mismo procedimiento a todos los casos y, aun así, producir efectos desiguales debido a diferencias históricas o contextuales. Las métricas deben complementarse con análisis sustantivos sobre quiénes resultan favorecidos, excluidos o expuestos a conclusiones científicas incorrectas.

La concentración tecnológica constituye otra fuente de desigualdad. Instituciones con mayor financiamiento pueden acceder a infraestructura, datos, especialistas y modelos avanzados, mientras universidades con menos recursos dependen de plataformas externas. Fochler et al. (2025) muestran que el financiamiento influye en la gobernanza y orientación de la investigación universitaria. Esta diferencia no solo afecta la capacidad de producir resultados, sino también quién define agendas, diseña modelos y establece estándares para evaluar conocimiento científico generado mediante inteligencia artificial.

La opacidad dificulta detectar y corregir sesgos porque impide comprender qué variables, datos o relaciones influyeron en una decisión. Zárate (2022) explica que la explicabilidad permite examinar los fundamentos y efectos de los sistemas algorítmicos. Cuando un modelo clasifica literatura, selecciona casos o produce inferencias sin mostrar sus criterios, resulta complejo identificar exclusiones o asociaciones problemáticas. La transparencia debe incluir procedencia de los datos, reglas de procesamiento, limitaciones conocidas y análisis comparativos del comportamiento del sistema en diferentes contextos.

Reducir los sesgos algorítmicos exige combinar estrategias técnicas, epistemológicas e institucionales. Contreras et al. (2024) señalan que el uso científico de inteligencia artificial requiere responsabilidad, transparencia y supervisión. La mitigación debe incluir revisión de representatividad, evaluación por grupos, documentación de categorías, participación de especialistas y apertura a perspectivas regionales. Ningún modelo es completamente neutral, pero sus límites pueden identificarse y gestionarse mediante procesos críticos que impidan convertir desigualdades históricas en resultados científicos aparentemente objetivos.

5.4 Transparencia, explicabilidad y reproducibilidad de los procesos investigativos

La transparencia en la investigación asistida por inteligencia artificial implica comunicar de manera suficiente cómo se seleccionaron los datos, qué herramientas se utilizaron y cuáles decisiones influyeron en los resultados. No se limita a mencionar el nombre del modelo, pues también exige describir versiones, parámetros, criterios de limpieza, transformaciones y condiciones de ejecución. Penabad et al. (2024) destacan la necesidad de reportar el uso de inteligencia artificial en publicaciones científico-académicas. Esta información permite valorar el alcance real de la intervención tecnológica y facilita una revisión metodológica más rigurosa.

La explicabilidad se relaciona con la posibilidad de comprender por qué un sistema produjo una determinada clasificación, predicción o recomendación. Zárate (2022) sostiene que explicar una decisión algorítmica exige hacer visibles sus fundamentos, efectos y límites. En investigación, este criterio permite detectar asociaciones espurias, dependencia excesiva de ciertas variables o comportamientos difíciles de justificar. La utilidad de una explicación depende de su pertinencia para el campo científico, porque una descripción técnica detallada puede resultar insuficiente si no conecta con los conceptos y preguntas del estudio.

La transparencia y la explicabilidad cumplen funciones distintas, aunque complementarias. Un proceso puede estar bien documentado y continuar siendo difícil de interpretar, especialmente cuando utiliza modelos complejos. Del mismo modo, una explicación parcial puede resultar comprensible sin revelar información suficiente sobre los datos o las condiciones del análisis. Joyce et al. (2023) relacionan la inteligencia artificial explicable con transparencia, interpretabilidad y comprensibilidad. La investigación rigurosa debe atender simultáneamente estos componentes para evitar que la claridad comunicativa oculte limitaciones metodológicas relevantes.

La reproducibilidad permite que otros investigadores reconstruyan un procedimiento y comprueben si obtienen resultados semejantes bajo condiciones equivalentes. Semmelrock et al. (2025) identifican barreras técnicas, documentales y organizacionales que dificultan reproducir investigaciones basadas en aprendizaje automático. Compartir únicamente el artículo final suele ser insuficiente, porque los resultados pueden depender de código, datos, configuraciones, versiones y decisiones intermedias. La documentación

debe abarcar todo el flujo de trabajo, desde la selección inicial de información hasta la evaluación y presentación del modelo definitivo.

Kapoor y Narayanan (2023) advierten que la fuga de datos puede generar resultados artificialmente favorables y contribuir a la crisis de reproducibilidad en estudios basados en aprendizaje automático. Este problema demuestra que la transparencia debe incluir la separación entre conjuntos de entrenamiento, validación y prueba. También deben registrarse los procedimientos de preprocesamiento y cualquier ajuste realizado después de observar resultados preliminares. Sin esta información, otros equipos no pueden determinar si el rendimiento reportado refleja capacidad real de generalización o una ventaja metodológica indebida.

La reproducibilidad computacional exige conservar versiones de software, bibliotecas, sistemas operativos, semillas aleatorias y condiciones de hardware. Pequeñas diferencias en estos elementos pueden modificar el comportamiento de un modelo, especialmente en aprendizaje profundo. Loza et al. (2026) destacan la importancia de documentar los flujos bioinformáticos aplicados a genómica y oncología clínica. Este principio se extiende a otros campos donde los análisis dependen de múltiples herramientas encadenadas. Preservar el entorno de ejecución facilita reconstruir resultados y localizar el origen de posibles discrepancias.

La interpretabilidad también debe adaptarse a las necesidades de quienes utilizarán los resultados. Bienefeld et al. (2023) señalan que la inteligencia artificial explicable debe conectar las necesidades de los profesionales clínicos con los objetivos de los desarrolladores. En investigación interdisciplinaria, esta exigencia implica producir explicaciones comprensibles para especialistas con diferentes formaciones. Un indicador técnico puede ser suficiente para un analista de datos, pero no para un profesional que debe interpretar el resultado dentro de una decisión científica, clínica, educativa o institucional.

La transparencia sobre el origen y la calidad de los datos resulta indispensable para evaluar sesgos y límites de generalización. Un modelo entrenado con información procedente de una sola región, población o institución puede perder validez en otros contextos. Ciobanu et al. (2024) destacan la necesidad de combinar reproducibilidad e interpretabilidad en investigaciones médicas basadas en aprendizaje automático. Documentar cobertura, criterios de inclusión, variables y distribución de los datos permite

reconocer qué poblaciones están representadas y cuáles conclusiones no deberían extenderse más allá del ámbito analizado.

Los modelos generativos plantean dificultades adicionales porque no siempre permiten identificar la fuente exacta de cada afirmación. Gao et al. (2023) demostraron que los textos científicos generados pueden parecer plausibles, aunque contengan problemas de autenticidad o exactitud. La transparencia exige distinguir entre información recuperada de documentos, contenido generado probabilísticamente y decisiones realizadas por el investigador. Cuando una herramienta interviene, en síntesis, redacción o interpretación, su uso debe declararse y cada afirmación debe verificarse directamente con las fuentes originales disponibles.

La explicabilidad no debe presentarse como una garantía absoluta de corrección. Un sistema puede ofrecer una justificación comprensible y continuar basándose en relaciones sesgadas o metodológicamente débiles. Omiye et al. (2023) evidencian que los modelos de lenguaje pueden reproducir planteamientos médicos asociados con categorías raciales. Una explicación convincente podría normalizar esas relaciones si no se examinan sus fundamentos históricos y científicos. Por ello, la evaluación debe integrar análisis técnico, conocimiento disciplinar y revisión ética de las consecuencias producidas por el modelo.

La apertura de datos y código puede fortalecer la reproducibilidad, aunque no siempre es posible compartir todos los materiales. La privacidad, la confidencialidad y la propiedad intelectual pueden exigir restricciones. Andrade et al. (2024) destacan que el tratamiento de datos personales mediante inteligencia artificial debe sostenerse en bases legales y medidas de protección. En estos casos, la transparencia puede garantizarse mediante descripciones detalladas, protocolos de acceso controlado, datos anonimizados o materiales simulados. El objetivo es permitir evaluación sin comprometer derechos ni obligaciones institucionales.

La transparencia, explicabilidad y reproducibilidad deben incorporarse desde el diseño de la investigación y no como requisitos añadidos al finalizar el estudio. Contreras et al. (2024) señalan que la responsabilidad científica en el uso de inteligencia artificial requiere supervisión y claridad sobre los procedimientos. Planificar registros, controles, documentación y formas de reporte facilita la auditoría y fortalece la confianza en los hallazgos. Un proceso investigativo es más sólido cuando sus decisiones pueden comprenderse, reconstruirse y someterse a evaluación crítica independiente.

5.5 Integridad académica, autoría y responsabilidad científica

La integridad académica comprende principios y prácticas orientados a garantizar honestidad, rigor, transparencia y responsabilidad en la producción del conocimiento. Cuando intervienen sistemas de inteligencia artificial, estas exigencias adquieren nuevas dimensiones porque parte del contenido puede ser generado, reformulado o analizado mediante herramientas automatizadas. Gallent et al. (2023) advierten que la inteligencia artificial generativa plantea retos para la autenticidad y la integridad en educación superior. Su uso legítimo depende de que no sustituya de forma encubierta las contribuciones intelectuales exigidas al investigador.

La autoría científica no se limita a redactar un manuscrito, sino que implica realizar aportes intelectuales sustantivos y asumir responsabilidad pública por el contenido publicado. Ros y Pérez (2023) sostienen que ChatGPT puede apoyar la escritura científica, pero no debe considerarse autor, porque carece de capacidad para responder por los argumentos, datos o consecuencias de una publicación. Los sistemas generativos tampoco pueden aprobar la versión final, resolver conflictos éticos ni garantizar la exactitud de las afirmaciones producidas durante el proceso académico.

La asistencia tecnológica puede intervenir en tareas como corrección lingüística, elaboración de esquemas, síntesis preliminar y reformulación de párrafos. Estas funciones no eliminan la autoría humana cuando los investigadores mantienen control sobre las decisiones conceptuales, metodológicas y argumentativas. Sin embargo, ocultar una intervención significativa puede generar una representación engañosa del proceso de producción. Penabad et al. (2024) destacan la necesidad de reportar el uso de inteligencia artificial en revistas científico-académicas, especialmente cuando la herramienta influye de manera relevante en la elaboración o transformación del manuscrito.

La responsabilidad científica permanece en las personas que diseñan el estudio, seleccionan los datos, interpretan los resultados y firman la publicación. Castillo (2023) subraya que la supervisión humana es indispensable cuando ChatGPT se utiliza como apoyo en escritura científica. Esta obligación incluye verificar afirmaciones, referencias, cifras y conclusiones antes de difundirlas. El hecho de que un error provenga de un sistema automatizado no reduce la responsabilidad de los autores, pues ellos deciden incorporar, modificar o rechazar el contenido generado dentro de la obra académica.

Las alucinaciones generativas representan una amenaza directa para la integridad cuando se incorporan contenidos sin comprobación. Gao et al. (2023) demostraron que los resúmenes científicos producidos por ChatGPT pueden parecer convincentes incluso para revisores humanos. Esta capacidad aumenta el riesgo de utilizar referencias inexistentes, métodos inventados o conclusiones sin sustento. La verificación debe realizarse directamente en las fuentes originales, no mediante nuevas consultas al mismo sistema. Una redacción formalmente sólida carece de valor académico si no existe evidencia verificable que respalde sus afirmaciones.

Tabla 14
Responsabilidades en el uso académico de inteligencia artificial

Actor	Responsabilidad principal	Práctica de integridad requerida	Conducta inadecuada
Autor principal	Garantizar exactitud y coherencia integral	Verificar contenidos, datos, citas y conclusiones	Incorporar respuestas generadas sin revisión
Coautores	Responder por sus contribuciones y la versión final	Revisar el manuscrito y declarar aportes reales	Aceptar autoría sin participación sustantiva
Equipo investigador	Documentar decisiones metodológicas y tecnológicas	Registrar herramientas, funciones y alcance de uso	Ocultar intervenciones automatizadas relevantes
Institución	Establecer normas y formar a investigadores	Crear protocolos de uso, supervisión y protección	Delegar toda responsabilidad al usuario individual
Revista científica	Definir criterios editoriales transparentes	Exigir declaraciones y verificar cumplimiento	Aplicar reglas ambiguas o inconsistentes
Revisor	Evaluar rigor, evidencia y transparencia	Examinar coherencia y posibles contenidos no verificables	Juzgar únicamente por estilo o fluidez textual
Sistema de IA	Apoyar tareas automatizadas	No aplica responsabilidad moral ni autoría	Ser presentado como responsable o autor

Nota. La tabla distingue responsabilidades humanas e institucionales y evidencia que los sistemas de inteligencia artificial pueden asistir el trabajo científico, pero no asumir autoría ni responder por los contenidos publicados.

La atribución adecuada también exige diferenciar entre asistencia técnica y contribución intelectual. Una herramienta puede corregir gramática o sugerir una estructura, mientras el investigador aporta la idea, el análisis y la interpretación. Cuando el sistema produce fragmentos extensos o participa en síntesis de evidencia, su intervención debe

documentarse con mayor detalle. Penabad et al. (2024) proponen informar la herramienta utilizada y la función desempeñada. Esta transparencia permite evaluar el proceso sin atribuir al sistema una responsabilidad que corresponde a los autores y a la institución.

El uso de inteligencia artificial puede favorecer prácticas deshonestas cuando se emplea para encubrir plagio, fabricar resultados o simular competencias que el autor no posee. Gallent et al. (2023) señalan que estas tecnologías desafían los mecanismos tradicionales de evaluación de autenticidad. Sin embargo, la respuesta no debe limitarse a detectores automáticos, cuya precisión puede ser insuficiente. Las instituciones necesitan fortalecer formación ética, evaluación crítica y diseño de actividades que permitan observar procesos, decisiones y evidencias, en lugar de juzgar únicamente el producto textual final.

La responsabilidad científica incluye proteger datos personales, documentos inéditos y materiales confidenciales introducidos en plataformas generativas. Sánchez (2024) analiza los desafíos que la inteligencia artificial generativa plantea para la protección de datos personales. Un investigador no debe cargar entrevistas, registros clínicos, manuscritos en revisión o información institucional sin conocer las condiciones de almacenamiento y reutilización. La integridad abarca tanto la honestidad autoral como el cumplimiento de compromisos de confidencialidad, consentimiento y seguridad asumidos durante la recopilación y tratamiento de la información.

Los sesgos generados o reproducidos por una herramienta también forman parte de la responsabilidad autoral. Omiye et al. (2023) muestran que los modelos de lenguaje pueden propagar enfoques médicos asociados con categorías raciales. Incorporar estas respuestas sin revisión puede reforzar conceptos discriminatorios bajo una apariencia de objetividad científica. Los autores deben examinar categorías, omisiones, lenguaje y consecuencias potenciales, especialmente cuando trabajan con poblaciones históricamente excluidas. La supervisión ética exige revisar no solo la exactitud factual, sino también la justicia y pertinencia de las representaciones utilizadas.

La integridad académica en investigaciones asistidas por inteligencia artificial requiere normas claras, documentación continua y responsabilidad humana indelegable. Contreras et al. (2024) señalan que el uso científico de estas tecnologías debe sostenerse en transparencia, ética y supervisión. Las herramientas pueden apoyar escritura, análisis y organización, pero no asumir compromiso intelectual ni responder ante errores. La confianza científica depende de autores capaces de declarar sus procedimientos, verificar

sus resultados, reconocer los límites tecnológicos y mantener control crítico sobre cada contenido presentado como conocimiento válido.

5.6 Privacidad, protección de datos y seguridad de la información científica

La privacidad en la investigación científica asistida por inteligencia artificial se refiere al control que las personas mantienen sobre la recopilación, uso, almacenamiento y circulación de sus datos. Este principio adquiere especial relevancia cuando se procesan historias clínicas, entrevistas, registros educativos, información biométrica o documentos institucionales. Andrade et al. (2024) señalan que la aplicación de inteligencia artificial debe sostenerse en fundamentos legales que protejan los datos personales. La eficiencia analítica no justifica una recolección excesiva ni usos distintos de aquellos inicialmente autorizados.

La protección de datos comienza con la minimización de la información recopilada. Cada variable debe responder a una necesidad científica concreta y guardar relación con los objetivos del estudio. Sánchez (2024) advierte que la inteligencia artificial generativa plantea retos debido al volumen de datos que puede procesar y a la dificultad para controlar su reutilización. Recopilar información innecesaria amplía los riesgos de exposición, reidentificación y tratamiento secundario. Por ello, el diseño metodológico debe limitar los datos a lo estrictamente pertinente.

El consentimiento informado debe explicar de manera comprensible cómo intervendrán los sistemas inteligentes en el tratamiento de la información. No basta con indicar que los datos serán utilizados con fines académicos si también serán cargados en plataformas externas, procesados por algoritmos o reutilizados en fases posteriores. Andrade et al. (2024) destacan la necesidad de vincular el tratamiento automatizado con bases legales claras. Los participantes deben conocer el propósito, los riesgos, las condiciones de almacenamiento y las posibilidades de acceso o eliminación de su información.

La anonimización busca reducir la posibilidad de relacionar los datos con una persona identificable. Sin embargo, eliminar nombres o números de identificación puede resultar insuficiente cuando persisten combinaciones de variables capaces de revelar identidades. Edad, ubicación, ocupación y antecedentes clínicos pueden permitir reidentificación al cruzarse con otras fuentes. Sánchez (2024) subraya que los entornos generativos amplían los desafíos asociados con el control de información personal. La evaluación del riesgo debe considerar tanto el conjunto original como sus posibles vinculaciones externas.

La seguridad de la información científica comprende medidas técnicas, organizacionales y humanas destinadas a prevenir accesos no autorizados, pérdida, alteración o divulgación. El cifrado, la autenticación y el control de permisos son importantes, pero no suficientes si no existen protocolos institucionales. Los equipos deben definir quién puede consultar, modificar, transferir o descargar los datos. La responsabilidad no puede recaer únicamente en especialistas informáticos, porque cada investigador interviene en decisiones que afectan la confidencialidad y la integridad de los registros (Contreras et al., 2024).

Las plataformas de inteligencia artificial generativa introducen riesgos particulares cuando se utilizan para resumir entrevistas, analizar documentos o corregir manuscritos inéditos. El investigador puede desconocer dónde se almacenan los contenidos, durante cuánto tiempo permanecen disponibles o si se emplean para mejorar otros modelos. Sánchez (2024) analiza estas dificultades en relación con la protección de datos personales. Antes de introducir información sensible, deben revisarse las políticas del servicio, las restricciones institucionales y la posibilidad de trabajar con versiones anonimizadas o entornos controlados.

Figura 10

Protección de datos científicos en la investigación de IA



Nota. La figura presenta los principales riesgos para la protección de datos científicos en investigaciones que utilizan inteligencia artificial. Elaboración propia.

La confidencialidad también protege información científica que no corresponde directamente a datos personales. Protocolos experimentales, resultados preliminares, bases inéditas, secretos técnicos y manuscritos en evaluación pueden verse comprometidos si se cargan en servicios externos. Esta exposición puede afectar propiedad intelectual, prioridad científica y acuerdos de colaboración. Penabad et al. (2024) destacan la importancia de reportar y regular el uso de inteligencia artificial en entornos académico científicos. Las instituciones deben diferenciar entre materiales públicos, restringidos, confidenciales y sujetos a revisión editorial.

La integridad de los datos exige garantizar que la información no haya sido alterada de manera accidental o maliciosa. Registros incompletos, archivos modificados y cambios no documentados pueden afectar la validez de los resultados. La trazabilidad debe conservar versiones, responsables, fechas y transformaciones realizadas durante el procesamiento. Semmelrock et al. (2025) señalan que la falta de documentación dificulta

la reproducibilidad de investigaciones basadas en aprendizaje automático. Proteger la integridad significa mantener tanto la seguridad técnica como la posibilidad de reconstruir el flujo analítico.

La computación en la nube y los repositorios compartidos facilitan colaboración, almacenamiento y procesamiento, pero distribuyen la información entre infraestructuras que pueden estar fuera del control directo del equipo. La gestión debe considerar ubicación de servidores, condiciones contractuales, copias de respaldo y mecanismos de recuperación. Andrade et al. (2024) relacionan la protección de datos con la necesidad de establecer bases legales y responsabilidades claras. La conveniencia operativa no debe prevalecer sobre la evaluación de riesgos, especialmente cuando se procesan datos sensibles o identificables.

Los incidentes de seguridad requieren planes de respuesta previamente definidos. El equipo debe saber cómo actuar ante accesos indebidos, pérdida de dispositivos, filtraciones o alteración de archivos. Esto incluye contener el incidente, evaluar el alcance, documentar lo ocurrido y aplicar medidas correctivas. Contreras et al. (2024) sostienen que la investigación con inteligencia artificial exige responsabilidad y supervisión. Una respuesta improvisada puede agravar los daños y dificultar la rendición de cuentas, mientras un protocolo claro protege a participantes, instituciones y proyectos científicos.

La privacidad, la protección de datos y la seguridad deben integrarse desde la planificación del estudio y mantenerse durante todo el ciclo de vida de la información. No son requisitos secundarios, sino condiciones para la legitimidad científica y la confianza social. Sánchez (2024) advierte que los sistemas generativos amplían los riesgos asociados con tratamiento, circulación y reutilización de datos personales. Una investigación responsable limita la recopilación, controla accesos, documenta transformaciones y evita trasladar información sensible a herramientas cuyos mecanismos no puedan evaluarse.



Capítulo

6.

Gobernanza y futuro de la investigación científica con inteligencia artificial

Cynthia Shakira Enríquez Fierro

6.1 Gobernanza institucional de la inteligencia artificial en la investigación científica

La gobernanza institucional de la inteligencia artificial en la investigación científica comprende normas, responsabilidades, procedimientos y mecanismos de supervisión destinados a orientar su uso dentro de universidades, centros y laboratorios. No se limita a autorizar herramientas, sino que debe establecer criterios sobre datos, modelos, transparencia, autoría, seguridad y rendición de cuentas. Contreras et al. (2024) señalan que la aplicación científica de estas tecnologías exige responsabilidad y control humano. Una gobernanza sólida traduce esos principios en decisiones operativas verificables y coherentes.

El primer componente corresponde a la definición de políticas institucionales claras. Estas deben precisar qué usos están permitidos, cuáles requieren autorización y en qué situaciones la inteligencia artificial no resulta adecuada. También deben diferenciar asistencia lingüística, análisis automatizado, generación de contenidos y toma de decisiones científicas. Penabad et al. (2024) destacan la necesidad de establecer orientaciones para el uso y reporte de inteligencia artificial en revistas académico científicas. Este criterio puede extenderse a proyectos, tesis, artículos, laboratorios y procesos de evaluación internos.

La distribución de responsabilidades evita que los errores se atribuyan de forma imprecisa a la tecnología. Investigadores, directores de proyecto, unidades de ética, especialistas en datos y autoridades institucionales deben conocer sus funciones. El investigador conserva la responsabilidad sobre los resultados, mientras la institución debe proporcionar normas, formación y mecanismos de control. Castillo (2023) sostiene que la supervisión humana es indispensable cuando herramientas generativas apoyan la escritura científica. Esta obligación requiere procedimientos que permitan revisar decisiones, corregir fallas y documentar intervenciones relevantes.

La gobernanza de datos ocupa un lugar central porque la inteligencia artificial depende de información cuya calidad, procedencia y sensibilidad pueden variar. Las instituciones deben regular recolección, acceso, almacenamiento, anonimización, reutilización y eliminación. Andrade et al. (2024) analizan la necesidad de sustentar legalmente el tratamiento de datos personales en sistemas inteligentes. Una política institucional debe establecer niveles de acceso, responsables de custodia y condiciones para utilizar

plataformas externas. Sin estas medidas, la innovación puede comprometer privacidad, confidencialidad y validez científica.

La evaluación de modelos debe incorporarse a la gobernanza antes de su implementación. No basta con comprobar que una herramienta funciona; también es necesario examinar precisión, sesgos, interpretabilidad, generalización y efectos sobre distintos grupos. Ciobanu et al. (2024) destacan que la reproducibilidad y la interpretabilidad son condiciones esenciales en investigaciones basadas en aprendizaje automático. Las instituciones pueden crear comités técnicos o protocolos de revisión para valorar si el modelo resulta apropiado para el problema, los datos y el nivel de riesgo involucrado.

Tabla 15
Componentes de la gobernanza institucional de inteligencia artificial en investigación

Componente	Finalidad institucional	Responsable principal	Riesgo ante su ausencia
Política de uso	Definir aplicaciones permitidas, restringidas y prohibidas	Autoridades y unidades de investigación	Uso improvisado o inconsistente
Gobernanza de datos	Regular acceso, protección, reutilización y conservación	Custodios de datos y comités de ética	Filtraciones, usos indebidos y pérdida de confianza
Evaluación de modelos	Revisar precisión, sesgos, explicabilidad y validez	Especialistas técnicos y disciplinares	Resultados poco confiables o discriminatorios
Transparencia	Documentar herramientas, versiones y funciones realizadas	Equipos investigadores	Procesos imposibles de auditar
Responsabilidad	Asignar funciones y mecanismos de rendición de cuentas	Dirección institucional y responsables de proyecto	Dilución de obligaciones
Formación	Desarrollar competencias técnicas, metodológicas y éticas	Unidades académicas y de capacitación	Dependencia acrítica de herramientas
Supervisión continua	Revisar cambios, incidentes y efectos durante el proyecto	Comités institucionales	Persistencia de errores y riesgos no detectados

Nota. La tabla sintetiza los elementos mínimos para una gobernanza institucional capaz de coordinar innovación, control metodológico, protección de datos y responsabilidad científica.

La transparencia institucional exige registrar qué herramientas se emplearon, con qué finalidad y bajo qué condiciones. Semmelrock et al. (2025) identifican la documentación insuficiente como una barrera para la reproducibilidad en investigaciones basadas en aprendizaje automático. Las políticas deben exigir información sobre versiones, parámetros, datos, modificaciones y criterios de selección. Este registro permite reconstruir procesos, evaluar resultados y responder ante auditorías o controversias. La transparencia no debe depender únicamente de la voluntad individual del investigador, sino formar parte de los procedimientos institucionales ordinarios.

La formación constituye otro pilar de la gobernanza. Los investigadores necesitan comprender capacidades, límites, riesgos y obligaciones asociados con los sistemas utilizados. Cedeño et al. (2024) señalan que la inteligencia artificial puede fortalecer la investigación universitaria, aunque su aprovechamiento requiere competencias adecuadas. Las instituciones deben ofrecer capacitación diferenciada para estudiantes, docentes, técnicos y autoridades, incluyendo verificación de resultados, protección de datos, integridad académica y evaluación de sesgos. Sin formación, las normas pueden existir formalmente y aplicarse de manera deficiente.

La supervisión institucional debe mantenerse durante todo el ciclo de investigación. Los modelos, plataformas y condiciones de uso pueden cambiar después de la aprobación inicial, alterando riesgos y resultados. Fochler et al. (2025) muestran que la gobernanza universitaria está influida por recursos, incentivos y estructuras de decisión. Por ello, las instituciones necesitan mecanismos para revisar incidentes, actualizar políticas y corregir prácticas. La gobernanza eficaz no consiste en una autorización única, sino en un proceso continuo de evaluación, aprendizaje y ajuste.

La gobernanza institucional alcanza legitimidad cuando equilibra innovación, autonomía científica y protección de derechos. Un control excesivamente rígido puede frenar usos valiosos, mientras una regulación débil favorece opacidad, desigualdad y dependencia tecnológica. Meskó y Topol (2023) subrayan la necesidad de supervisión regulatoria para los grandes modelos de lenguaje en salud. En otros campos, la misma lógica exige estructuras proporcionadas al nivel de riesgo. Las instituciones deben promover innovación responsable mediante normas claras, revisión interdisciplinaria y rendición de cuentas.

6.2 Regulación y marcos normativos para el uso responsable de sistemas inteligentes

La regulación de los sistemas inteligentes aplicados a la investigación científica debe establecer límites, responsabilidades y condiciones de uso proporcionales a los riesgos involucrados. No todas las aplicaciones requieren el mismo nivel de control, pues una herramienta de apoyo lingüístico no produce las mismas consecuencias que un modelo utilizado para clasificar pacientes o recomendar decisiones. Meskó y Topol (2023) sostienen que los grandes modelos de lenguaje en salud requieren supervisión regulatoria específica. Este principio puede extenderse a otros campos mediante esquemas diferenciados según finalidad, impacto y sensibilidad.

Un marco normativo responsable debe comenzar por definir con claridad el objeto regulado. La inteligencia artificial generativa, el aprendizaje automático, la automatización experimental y los sistemas de apoyo a decisiones presentan capacidades y riesgos distintos. Tratar estas tecnologías como equivalentes puede producir normas demasiado amplias o insuficientes. La regulación debe considerar la función concreta del sistema, el tipo de datos procesados y el grado de autonomía concedido. Una clasificación precisa facilita establecer obligaciones coherentes de documentación, validación, supervisión y rendición de cuentas (Contreras et al., 2024).

La protección de datos constituye uno de los ejes centrales de cualquier marco regulatorio. Andrade et al. (2024) señalan que el tratamiento de información personal mediante inteligencia artificial debe sostenerse en bases legales claras. En investigación, esto implica regular recolección, almacenamiento, acceso, reutilización y eliminación de datos. Las instituciones deben garantizar que la información utilizada corresponda con los objetivos autorizados y que no se transfiera a plataformas externas sin evaluación previa. La innovación tecnológica no puede justificar prácticas incompatibles con privacidad, confidencialidad o consentimiento.

Los sistemas generativos plantean desafíos específicos porque pueden procesar documentos inéditos, entrevistas, bases institucionales y otros contenidos sensibles. Sánchez (2024) advierte que estas herramientas amplían los riesgos relacionados con circulación, reutilización y control de datos personales. Por ello, los marcos normativos deben establecer restricciones para cargar información confidencial en plataformas cuyos mecanismos de almacenamiento no sean transparentes. También deben exigir medidas de

anonimización, seguridad y evaluación contractual cuando los servicios tecnológicos pertenezcan a proveedores externos con condiciones de uso propias.

La transparencia regulatoria exige que los investigadores declaren qué herramientas utilizaron, con qué finalidad y en qué etapas del estudio intervinieron. Penabad et al. (2024) proponen orientaciones para el uso y reporte de inteligencia artificial en revistas científico-académicas. Estas prácticas permiten distinguir entre asistencia técnica, generación de contenido, análisis automatizado e interpretación humana. Un marco normativo institucional debe transformar esta declaración en un requisito verificable, evitando que el reporte dependa únicamente de decisiones voluntarias o de criterios editoriales variables entre publicaciones.

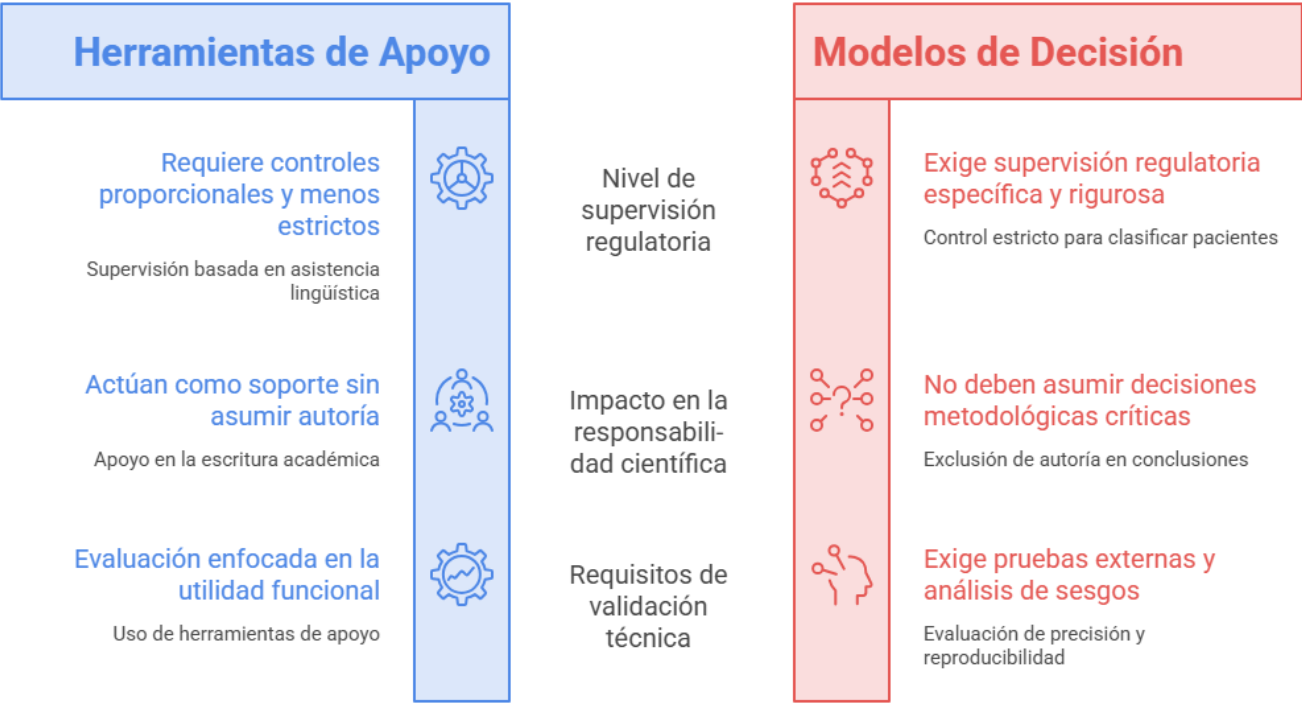
La responsabilidad científica debe permanecer claramente atribuida a personas e instituciones. Los sistemas inteligentes no pueden responder por errores, daños, sesgos o incumplimientos éticos. Ros y Pérez (2023) sostienen que herramientas como ChatGPT pueden apoyar la escritura, pero no asumir autoría. De forma equivalente, tampoco deben considerarse responsables de decisiones metodológicas o conclusiones científicas. La regulación debe identificar quién diseña el estudio, valida los resultados, autoriza el uso de datos y aprueba la publicación, evitando que la intervención tecnológica diluya obligaciones humanas.

La validación previa constituye otra exigencia normativa esencial. Antes de aplicar un modelo en contextos científicos sensibles, deben evaluarse precisión, generalización, interpretabilidad y comportamiento frente a datos diferentes. Ciobanu et al. (2024) destacan que la reproducibilidad y la interpretación son condiciones indispensables en investigaciones médicas basadas en aprendizaje automático. Los marcos regulatorios pueden exigir pruebas externas, documentación de errores y análisis por grupos. Esta evaluación permite determinar si el sistema es adecuado para la finalidad propuesta y bajo qué restricciones puede emplearse.

La equidad algorítmica debe formar parte del control normativo, especialmente cuando los resultados afectan poblaciones o decisiones institucionales. Kontiainen et al. (2022) proponen abordar la justicia algorítmica desde una perspectiva interdisciplinaria, vinculando el funcionamiento técnico con sus consecuencias sociales. La regulación debe exigir revisión de categorías, representatividad de los datos y distribución de errores. No basta con aplicar una misma regla a todos los casos, pues un procedimiento formalmente

uniforme puede producir efectos desiguales sobre grupos históricamente subrepresentados o expuestos a discriminación (Kuppler et al., 2022).

Figura 11
¿Qué enfoque regulatorio es más adecuado para el uso de sistemas inteligentes en la investigación científica?



Nota. La figura compara el nivel de regulación requerido para las herramientas de apoyo y los modelos de decisión utilizados en la investigación científica. Elaboración propia.

La explicabilidad debe exigirse de manera proporcional al impacto del sistema. Zárate (2022) señala que comprender los fundamentos de una decisión algorítmica permite evaluar sus efectos y límites. En investigaciones de bajo riesgo, una descripción general del modelo puede resultar suficiente; en aplicaciones clínicas, jurídicas o institucionales, se requieren explicaciones más detalladas y comprensibles. La regulación debe evitar exigencias meramente formales y garantizar que la información proporcionada permita revisar la lógica del sistema, detectar errores y cuestionar decisiones automatizadas.

Los marcos normativos también deben establecer procedimientos para incidentes, actualizaciones y cambios tecnológicos. Un modelo validado puede modificar su comportamiento después de una nueva versión, cambios en los datos o ajustes realizados por el proveedor. Semmelrock et al. (2025) advierten que la reproducibilidad depende de documentar configuraciones, versiones y condiciones de ejecución. La supervisión regulatoria no debe limitarse a una aprobación inicial, sino incorporar seguimiento

continuo, registro de modificaciones y mecanismos para suspender el uso cuando aparezcan riesgos no previstos.

La regulación responsable debe equilibrar protección, innovación y autonomía científica. Normas excesivamente rígidas pueden impedir aplicaciones valiosas, mientras reglas ambiguas favorecen prácticas opacas y desiguales. Meskó y Topol (2023) plantean que la supervisión de modelos generativos debe responder a sus riesgos reales y a los contextos donde se aplican. Un marco normativo sólido combina clasificación de usos, protección de datos, transparencia, validación, equidad y responsabilidad humana. Su finalidad no es frenar la inteligencia artificial, sino garantizar que contribuya a una ciencia confiable y socialmente legítima.

6.3 Propiedad intelectual y uso legítimo de datos y contenidos científicos

La propiedad intelectual en la investigación asistida por inteligencia artificial plantea interrogantes sobre autoría, titularidad, reutilización y transformación de contenidos científicos. Los sistemas generativos pueden producir textos, imágenes, códigos o síntesis a partir de patrones aprendidos en grandes volúmenes de información, pero no asumen responsabilidad jurídica ni académica por esos resultados. Ros y Pérez (2023) sostienen que herramientas como ChatGPT pueden apoyar la escritura científica, aunque no deben reconocerse como autoras. La atribución y la responsabilidad permanecen vinculadas con las personas que dirigen, revisan y validan el trabajo.

El uso legítimo de contenidos científicos exige distinguir entre consulta, análisis, reproducción, adaptación y apropiación. Una herramienta puede ayudar a resumir o reformular un artículo, pero el investigador debe preservar la atribución correspondiente y evitar presentar como propia una contribución intelectual ajena. La inteligencia artificial no elimina las obligaciones relacionadas con citación, reconocimiento de fuentes y respeto a los derechos de los autores. La integridad académica depende de identificar con claridad qué ideas proceden de publicaciones previas y cuáles corresponden al análisis original desarrollado en el estudio (Gallent et al., 2023).

Penabad et al. (2024) destacan la necesidad de declarar el uso de inteligencia artificial en revistas científico-académicas cuando su intervención sea significativa. Este principio también protege la transparencia sobre la elaboración de contenidos y permite valorar la contribución humana real. Informar la herramienta empleada, su función y el alcance de su participación reduce ambigüedades sobre autoría y proceso creativo. No obstante,

declarar el uso tecnológico no sustituye la obligación de verificar originalidad, exactitud, pertinencia y correspondencia con las fuentes consultadas.

Los datos científicos pueden estar sujetos a condiciones de acceso, licencias, acuerdos institucionales o restricciones derivadas de su naturaleza sensible. La disponibilidad técnica de un conjunto no implica autorización irrestricta para copiarlo, combinarlo o introducirlo en plataformas externas. Andrade et al. (2024) señalan que el tratamiento de información mediante inteligencia artificial debe sostenerse en fundamentos legales claros. Antes de reutilizar datos, el investigador debe revisar su procedencia, finalidad original, permisos, restricciones y posibles obligaciones frente a participantes, instituciones o responsables del repositorio.

La reutilización legítima requiere conservar información sobre procedencia, versión y condiciones de uso. Un conjunto de datos puede haber sido recopilado para responder una pregunta específica y resultar inadecuado para otros propósitos. Además, las transformaciones realizadas mediante sistemas inteligentes pueden modificar categorías, formatos o relaciones presentes en el material original. La trazabilidad permite identificar qué parte fue obtenida de una fuente, cuál fue procesada automáticamente y qué decisiones correspondieron al equipo investigador. Esta documentación fortalece transparencia, reproducibilidad y responsabilidad científica (Semmelrock et al., 2025).

Los contenidos introducidos en plataformas generativas pueden incluir manuscritos inéditos, protocolos, bases de datos, códigos o resultados preliminares. Sánchez (2024) advierte que la inteligencia artificial generativa amplía los retos relacionados con tratamiento, circulación y control de datos personales. Una preocupación semejante se aplica a materiales científicos confidenciales o protegidos. Antes de utilizar servicios externos, deben examinarse sus condiciones de almacenamiento, reutilización y acceso. La comodidad de una herramienta no justifica comprometer prioridad científica, confidencialidad editorial o derechos de terceros.

La generación automática de textos plantea el riesgo de incorporar fragmentos excesivamente próximos a materiales existentes sin que el usuario pueda reconocer su origen. Aunque el sistema no reproduzca de forma literal una fuente identificable, puede combinar estructuras, conceptos y expresiones aprendidas durante su entrenamiento. Por ello, la revisión de originalidad debe abarcar tanto coincidencias textuales como apropiaciones conceptuales no atribuidas. Castillo (2023) subraya que la supervisión

humana es indispensable en la escritura científica asistida por inteligencia artificial, especialmente para controlar exactitud y responsabilidad.

La propiedad intelectual también se relaciona con los productos creados por equipos interdisciplinarios que combinan aportes humanos y automatizados. La distribución de créditos debe basarse en contribuciones reales a la conceptualización, metodología, análisis e interpretación, no en el uso de una herramienta determinada. Los sistemas inteligentes pueden facilitar tareas, pero no participar en decisiones éticas ni responder ante controversias. La transparencia sobre funciones y aportes permite diferenciar asistencia técnica, contribución intelectual y responsabilidad autoral dentro de la producción académica (Penabad et al., 2024).

La ciencia abierta favorece el acceso y la reutilización de publicaciones, datos y materiales, pero no elimina la necesidad de respetar licencias, atribuciones y condiciones de uso. Ross et al. (2022) advierten que la apertura puede coexistir con desigualdades y dinámicas de ventaja acumulativa. Un recurso abierto puede ser reutilizable bajo determinadas condiciones, aunque siga requiriendo reconocimiento de autoría y preservación de su contexto. La apertura responsable debe facilitar circulación sin convertir el conocimiento compartido en un material desprovisto de procedencia, límites o responsabilidades.

El uso legítimo de datos y contenidos científicos exige combinar propiedad intelectual, integridad académica, protección de información y transparencia tecnológica. Contreras et al. (2024) señalan que la investigación con inteligencia artificial debe sostenerse en responsabilidad, supervisión y criterios éticos. Las instituciones necesitan políticas que regulen atribución, reutilización, licencias, carga de materiales y declaración de herramientas. La innovación no debe construirse mediante apropiación opaca o uso indiscriminado de recursos, sino sobre prácticas que reconozcan autoría, procedencia, consentimiento y obligaciones científicas.

6.4 Ciencia abierta, innovación colaborativa y democratización del conocimiento

La ciencia abierta promueve el acceso, intercambio y reutilización de publicaciones, datos, métodos y recursos científicos. Su propósito es ampliar la transparencia y favorecer que el conocimiento pueda ser examinado, reproducido y aprovechado por comunidades diversas. La inteligencia artificial fortalece este enfoque al facilitar clasificación, búsqueda semántica, conexión entre repositorios y análisis de grandes colecciones. Sin

embargo, la apertura no debe reducirse a disponibilidad técnica, pues también exige calidad documental, interoperabilidad, condiciones de uso claras y capacidades efectivas para interpretar los materiales compartidos (Ross et al., 2022).

Los datos abiertos constituyen una base importante para la innovación colaborativa porque permiten realizar análisis secundarios, comparar resultados y formular nuevas preguntas. Ávila (2024) muestra que los grafos de conocimiento y las ontologías favorecen la representación estructurada de datos de investigación. Estas herramientas facilitan vincular publicaciones, autores, proyectos y conjuntos de información. No obstante, la reutilización requiere metadatos completos, control de versiones y descripción de las condiciones de producción, ya que un archivo accesible puede resultar científicamente inadecuado si carece de contexto metodológico.

La innovación colaborativa surge cuando distintos equipos comparten recursos, capacidades y conocimientos para resolver problemas que exceden las posibilidades de una sola institución. La inteligencia artificial puede apoyar esta cooperación mediante plataformas de análisis, recuperación documental y organización distribuida de información. Duedra et al. (2020) proponen superar la compartimentalización entre ciencia básica, aplicada, extensión y divulgación. Desde esta perspectiva, la colaboración adquiere valor cuando conecta producción científica, necesidades sociales y aplicaciones concretas, sin diluir las responsabilidades ni las contribuciones de los participantes.

Los repositorios inteligentes amplían la circulación del conocimiento al relacionar documentos mediante similitud semántica, redes de citación y estructuras conceptuales. Polo y Casique (2025) señalan que la integración de grafos de conocimiento y redes neuronales mejora la recuperación de información documental. Esta capacidad puede facilitar el descubrimiento de trabajos procedentes de áreas distintas y fortalecer proyectos interdisciplinarios. Sin embargo, los algoritmos de recomendación deben revisarse para evitar que prioricen exclusivamente publicaciones populares, instituciones dominantes o contenidos con mayor presencia digital.

La democratización del conocimiento supone reducir barreras económicas, geográficas, lingüísticas y tecnológicas que limitan la participación científica. El acceso abierto puede ampliar oportunidades, pero no elimina por sí solo las desigualdades existentes. Ross et al. (2022) advierten que la ciencia abierta también puede reproducir ventajas acumulativas cuando ciertos actores concentran infraestructura, visibilidad y

competencias para reutilizar información. Por ello, democratizar exige combinar disponibilidad de recursos con formación, conectividad, herramientas accesibles y reconocimiento de conocimientos producidos fuera de los centros científicos dominantes.

Tabla 16

Dimensiones de la ciencia abierta y la democratización del conocimiento asistidas por inteligencia artificial

Dimensión	Aporte de la inteligencia artificial	Condición para la democratización	Riesgo principal
Acceso a literatura	Mejora búsqueda, clasificación y recomendación	Cobertura lingüística y regional diversa	Invisibilización de fuentes periféricas
Datos abiertos	Facilita integración, análisis y reutilización	Metadatos, licencias y formatos interoperables	Uso descontextualizado o indebido
Colaboración científica	Conecta equipos, proyectos y recursos	Participación equitativa en decisiones y créditos	Relaciones asimétricas entre instituciones
Divulgación	Adapta contenidos a distintos públicos	Verificación de precisión y claridad	Simplificación excesiva de resultados
Reproducibilidad	Organiza código, datos y flujos de trabajo	Documentación completa y acceso sostenible	Imposibilidad de reconstruir procesos
Formación	Apoya aprendizaje y acceso a herramientas	Desarrollo de competencias críticas y técnicas	Dependencia tecnológica
Visibilidad científica	Identifica relaciones y tendencias	Inclusión de revistas, idiomas y regiones diversas	Concentración algorítmica de reconocimiento

Nota. La tabla sintetiza las oportunidades y condiciones necesarias para que la inteligencia artificial contribuya a una ciencia abierta, colaborativa y socialmente accesible.

La desigualdad en el financiamiento condiciona quiénes pueden participar plenamente en ecosistemas de innovación abierta. Ugarte y Parra (2021) relacionan los recursos disponibles con la producción científica, evidenciando que la capacidad investigativa depende también de condiciones institucionales. El acceso a publicaciones no compensa la falta de infraestructura, personal especializado o recursos computacionales. Las políticas de ciencia abierta deben acompañarse de financiamiento, repositorios sostenibles y mecanismos de cooperación que permitan a instituciones con menores capacidades intervenir en el diseño, análisis e interpretación de los proyectos.

La inteligencia artificial generativa puede favorecer la divulgación al reformular contenidos técnicos, producir resúmenes y adaptar explicaciones para públicos no especializados. Cedeño et al. (2024) reconocen que estas herramientas amplían posibilidades de organización y comunicación en la investigación universitaria. Sin embargo, la democratización no consiste en simplificar indiscriminadamente. Cada adaptación debe conservar la precisión, distinguir resultados de interpretaciones y señalar incertidumbres. La revisión humana resulta indispensable para evitar que una comunicación accesible transforme hallazgos parciales en afirmaciones exageradas o científicamente imprecisas.

La apertura también debe respetar privacidad, propiedad intelectual y condiciones éticas de reutilización. No todos los datos pueden difundirse sin restricciones, especialmente cuando contienen información personal, clínica o institucional. Andrade et al. (2024) destacan que el tratamiento de datos mediante inteligencia artificial requiere bases legales y medidas de protección. La ciencia abierta responsable utiliza licencias, anonimización, accesos controlados y protocolos de uso. Democratizar el conocimiento no significa eliminar toda limitación, sino establecer condiciones transparentes que equilibren acceso científico, protección de derechos y responsabilidad institucional.

La ciencia abierta y la innovación colaborativa alcanzan legitimidad cuando amplían la participación sin reproducir desigualdades tecnológicas o epistémicas. Morin (1999) plantea que el conocimiento debe reconocer relaciones, diversidad y complejidad. Aplicado a estos procesos, ello exige integrar infraestructuras, capacidades humanas, marcos éticos y políticas de inclusión. La inteligencia artificial puede facilitar circulación, conexión y reutilización, pero la democratización depende de decisiones institucionales que garanticen acceso sostenible, diversidad de perspectivas, reconocimiento equitativo y control crítico sobre los sistemas empleados.

6.5 Financiamiento, infraestructura y desigualdades en el acceso a la inteligencia artificial

El financiamiento condiciona la capacidad de las instituciones para incorporar inteligencia artificial en sus procesos de investigación. La adquisición de infraestructura, licencias, almacenamiento, conectividad y personal especializado requiere inversiones sostenidas que no siempre están disponibles. Ugarte y Parra (2021) relacionan el financiamiento con la producción científica, mostrando que los recursos influyen en la

continuidad y alcance de las actividades académicas. En este marco, las diferencias presupuestarias pueden traducirse en capacidades desiguales para acceder, adaptar y evaluar sistemas inteligentes.

La infraestructura tecnológica comprende equipos de procesamiento, redes, servidores, repositorios, sistemas de respaldo y plataformas de análisis. Su disponibilidad determina el tipo de proyectos que una institución puede ejecutar y la complejidad de los modelos que puede utilizar. Hartung (2025) señala que la automatización y los modelos con capacidad de acción amplían las posibilidades del descubrimiento científico. Sin embargo, estas oportunidades dependen de entornos técnicos capaces de sostener experimentación continua, procesamiento intensivo y almacenamiento seguro de grandes volúmenes de información.

La concentración de infraestructura en pocas universidades, empresas o centros de investigación genera relaciones asimétricas. Las instituciones con mayores recursos pueden entrenar modelos, administrar bases extensas y definir estándares técnicos, mientras otras dependen de servicios externos. Fochler et al. (2025) muestran que el financiamiento influye en la gobernanza de la investigación universitaria. Esta dependencia puede limitar la autonomía científica, porque los proveedores tecnológicos establecen condiciones de acceso, costos, formatos y restricciones que afectan la orientación de los proyectos y el control sobre los datos.

El acceso desigual también se manifiesta entre regiones y países. Instituciones ubicadas fuera de los principales centros científicos pueden enfrentar conectividad limitada, costos elevados y menor disponibilidad de especialistas. Ross et al. (2022) advierten que las dinámicas de ventaja acumulativa pueden reproducirse incluso dentro de la ciencia abierta. En consecuencia, la inteligencia artificial puede ampliar la productividad de actores ya consolidados y profundizar brechas existentes si no se crean mecanismos de cooperación, infraestructura compartida y fortalecimiento de capacidades locales.

La dependencia de plataformas comerciales introduce desafíos económicos y metodológicos. Los costos pueden cambiar, los servicios pueden desaparecer y las condiciones de uso pueden modificarse sin control de los equipos científicos. Además, los modelos cerrados dificultan conocer datos de entrenamiento, criterios internos y limitaciones técnicas. Semmelrock et al. (2025) señalan que la reproducibilidad depende del acceso a herramientas, configuraciones y documentación. Cuando una investigación

se apoya en servicios opacos o inaccesibles, otros equipos pueden encontrar dificultades para repetir o auditar los resultados.

Figura 12
Desafíos en el acceso a la IA



Nota. La figura presenta los principales desafíos que limitan el acceso equitativo a la inteligencia artificial en la investigación científica. Elaboración propia.

La infraestructura no se limita a equipos físicos. También incluye conocimientos, protocolos, soporte técnico y capacidades institucionales para seleccionar y gestionar herramientas. Cedeño et al. (2024) destacan que la inteligencia artificial puede fortalecer la investigación universitaria, aunque su aprovechamiento exige preparación adecuada. Una institución puede adquirir tecnología avanzada y obtener resultados limitados si sus investigadores no comprenden cómo formular consultas, evaluar modelos, proteger datos o interpretar resultados. La inversión debe articular recursos técnicos con formación metodológica y ética.

Las desigualdades también aparecen dentro de los equipos de investigación. Quienes dominan programación, análisis de datos o infraestructura computacional pueden concentrar decisiones relevantes sobre modelos y procedimientos. Esta situación puede reducir la participación de especialistas disciplinares en etapas centrales del análisis. Contreras et al. (2024) sostienen que el uso científico de inteligencia artificial requiere responsabilidad, transparencia y supervisión. La colaboración debe distribuir funciones de manera clara y asegurar que la capacidad técnica no sustituya la deliberación científica ni concentre autoridad de forma injustificada.

El financiamiento competitivo puede orientar proyectos hacia áreas consideradas rentables o estratégicas, dejando en segundo plano problemas locales, sociales o humanísticos. Schuff y Hubert (2024) analizan cómo las agendas de investigación y financiamiento influyen en las prioridades científicas. En inteligencia artificial, esta presión puede favorecer aplicaciones con alta visibilidad tecnológica, aunque su pertinencia social sea limitada. La asignación de recursos debería considerar impacto científico, necesidades territoriales, sostenibilidad y contribución al desarrollo institucional, no únicamente novedad o rendimiento económico esperado.

La ciencia abierta puede reducir algunas barreras mediante repositorios, código compartido, datos accesibles y herramientas reutilizables. Sin embargo, la disponibilidad de recursos no garantiza su aprovechamiento efectivo. Ross et al. (2022) advierten que las capacidades para participar en entornos abiertos se distribuyen de manera desigual. Ejecutar modelos, adaptar código y gestionar grandes bases requiere infraestructura y formación. Las políticas de apertura deben acompañarse de soporte técnico, documentación comprensible y mecanismos de colaboración que permitan una participación científica más equitativa.

Las estrategias de infraestructura compartida pueden fortalecer a instituciones con recursos limitados. Redes universitarias, laboratorios regionales y plataformas cooperativas permiten distribuir costos, ampliar acceso y evitar duplicaciones. Estas iniciativas requieren acuerdos sobre gobernanza, seguridad, mantenimiento y propiedad de los resultados. Fochler et al. (2025) muestran que la organización del financiamiento afecta la autonomía y las decisiones científicas. Un esquema colaborativo debe asegurar participación real de todas las instituciones y evitar que los actores con mayor capacidad controlen datos, agendas o beneficios.

Reducir las desigualdades de acceso a la inteligencia artificial exige políticas sostenidas de financiamiento, formación e infraestructura. La solución no consiste únicamente en distribuir herramientas, sino en crear condiciones para utilizarlas de manera autónoma, crítica y responsable. Ugarte y Parra (2021) evidencian que la producción científica depende de recursos que permitan continuidad y desarrollo institucional. Una gobernanza equitativa debe promover capacidades locales, cooperación entre centros y acceso transparente a tecnologías, evitando que la innovación amplíe brechas preexistentes en el sistema científico.

6.6 Nuevos paradigmas y perspectivas futuras de la investigación científica

La investigación científica se orienta hacia paradigmas más automatizados, conectados y adaptativos, capaces de integrar datos, modelos y decisiones dentro de ciclos continuos de descubrimiento. Hartung (2025) plantea que la convergencia entre inteligencia artificial, modelos con capacidad de acción y automatización de laboratorios está configurando nuevas formas de producción científica. Este cambio no elimina los fundamentos tradicionales de la investigación, pero modifica la velocidad, escala y organización de los procesos. El desafío consiste en incorporar estas capacidades sin debilitar la reflexión teórica ni la validación empírica (Lopardo, 2024).

Los laboratorios autónomos representan una de las expresiones más visibles de este nuevo paradigma. Pueden planificar experimentos, coordinar instrumentos, analizar resultados y seleccionar condiciones para pruebas posteriores. Fushimi et al. (2025) muestran que estas arquitecturas ya se aplican en investigación biotecnológica mediante sistemas capaces de integrar equipos y flujos operativos. Su desarrollo abre posibilidades para explorar espacios experimentales complejos, aunque exige criterios de supervisión, trazabilidad y control institucional que permitan comprender cada decisión y evitar la propagación automatizada de errores.

La investigación futura también se caracterizará por una interacción más estrecha entre investigadores y sistemas inteligentes. En lugar de limitarse a ejecutar instrucciones, algunos modelos podrán proponer alternativas, evaluar resultados y reorganizar tareas según objetivos científicos. Tobias y Wahab (2025) analizan las implicaciones tecnológicas y políticas de los laboratorios autónomos, destacando que su expansión plantea preguntas sobre control, responsabilidad y orientación de las agendas. La autonomía técnica deberá mantenerse subordinada a decisiones humanas capaces de justificar prioridades, límites y consecuencias científicas.

Los gemelos digitales ampliarán la capacidad para estudiar sistemas físicos, biológicos e industriales mediante representaciones dinámicas conectadas con datos reales. Qian et al. (2022) explican que estas réplicas cibernéticas permiten monitorear, simular y anticipar comportamientos bajo distintas condiciones. En el futuro, su integración con sensores, aprendizaje automático y robótica favorecerá experimentos híbridos entre entornos digitales y físicos. Sin embargo, la utilidad de estas representaciones dependerá de su

actualización, precisión y correspondencia con los fenómenos, evitando confundir una simulación sofisticada con la realidad estudiada.

Otro paradigma emergente se relaciona con la producción de conocimiento mediante modelos generativos especializados y conectados con fuentes verificables. Gilbert et al. (2024) proponen sistemas aumentados capaces de actuar como curadores de información médica con menor riesgo de alucinación. Esta tendencia puede extenderse hacia herramientas científicas orientadas a sintetizar evidencia, formular hipótesis y organizar literatura. Aun así, la conexión con fuentes no garantiza exactitud absoluta. Los investigadores deberán evaluar selección documental, fidelidad interpretativa y pertinencia de cada respuesta antes de incorporarla al proceso científico.

La comprensión científica continuará siendo una exigencia central frente al aumento de la capacidad predictiva. Krenn et al. (2022) sostienen que la inteligencia artificial puede contribuir al entendimiento cuando sus resultados se relacionan con explicaciones y estructuras conceptuales. Un modelo capaz de predecir con precisión no necesariamente revela los mecanismos del fenómeno analizado. Los paradigmas futuros deberán equilibrar rendimiento, interpretabilidad y construcción teórica, evitando que la ciencia se limite a producir correlaciones útiles, pero insuficientes para explicar causalidad, contexto y alcance.

La ciencia del futuro será más interdisciplinaria, aunque esta convergencia exigirá superar diferencias epistemológicas y metodológicas. Morin (1999) plantea que el conocimiento complejo debe articular relaciones, incertidumbres y niveles diversos de realidad. La inteligencia artificial puede conectar bases, lenguajes y métodos procedentes de distintas áreas, pero no resolver automáticamente sus incompatibilidades. Los equipos deberán construir acuerdos sobre conceptos, evidencia y responsabilidad. La integración será valiosa cuando genere explicaciones más amplias sin reducir todas las disciplinas a una única lógica computacional.

La reproducibilidad adquirirá mayor importancia a medida que los procesos científicos dependan de modelos, plataformas y flujos automatizados. Semmelrock et al. (2025) identifican barreras técnicas y organizacionales que afectan la repetición de investigaciones basadas en aprendizaje automático. En escenarios futuros, será necesario preservar no solo datos y código, sino también versiones de modelos, parámetros, entornos, decisiones automáticas y registros de interacción humana. La documentación

deberá evolucionar para representar procesos dinámicos donde varias herramientas intervienen sucesivamente en la producción de un resultado científico.

Las perspectivas futuras también estarán condicionadas por la distribución desigual de infraestructura, financiamiento y competencias. Ross et al. (2022) advierten que la ciencia abierta puede reproducir ventajas acumulativas, incluso cuando amplía el acceso a publicaciones y datos. La inteligencia artificial podría concentrar la producción científica en instituciones con mayor capacidad computacional y control tecnológico. Evitar esta tendencia requerirá infraestructuras compartidas, formación especializada y políticas de cooperación que permitan a universidades con menores recursos participar en el diseño, interpretación y gobernanza de sistemas avanzados.

Los nuevos paradigmas científicos no deberían evaluarse únicamente por su novedad tecnológica, sino por la calidad y legitimidad del conocimiento que permiten generar. Contreras et al. (2024) señalan que la investigación con inteligencia artificial exige transparencia, supervisión y responsabilidad. El futuro dependerá de combinar automatización, explicación, apertura y protección de derechos dentro de marcos institucionales sólidos. La inteligencia artificial ampliará la capacidad de descubrir, simular y conectar información, pero la orientación ética y epistemológica continuará siendo una responsabilidad humana indelegable.

Bibliografía

- Andrade, D., Toapanta, M., Baño, M., y Gómez, E. (2024). Un enfoque de la inteligencia artificial para la protección de datos personales sustentando en la base legal. *LATAM Revista Latinoamericana de Ciencias Sociales y Humanidades*, 5(4), 3808-3820. <https://doi.org/10.56712/latam.v5i4.2530>
- Arias, F., y Artigas, W. (2022). Cómo plantear problemas científicos relevantes identificando brechas de investigación. *Mujer Andina*, 1(1), 76-82. <https://doi.org/10.36881/ma.v1i1.644>
- Ávila, E. (2024). Visualización de datos abiertos de investigación mediante grafos de conocimiento: metodología y aplicación de DDI-RDF y DataCite Ontology. *Scire: Representación Y organización Del Conocimiento*, 30(2), 59-71. <https://doi.org/10.54886/scire.v30i2.4963>
- Barradas, J. (2023). Inteligencia artificial como elemento transformador de la investigación científica. *Entrelíneas*, 2(1), 113-122. <https://doi.org/10.56368/Entrelineas213>
- Bautista, N. (2022). *Proceso de la investigación cualitativa: epistemología, metodología y aplicaciones*. Editorial El Manual Moderno. <https://books.google.es/books?hl=es&lr=&id=yr2CEAAAQBAJ&oi=fnd&pg=PT4&dq=epistemologia+de+La+investigaci%C3%B3n+b%C3%A1sica&ots=1zM-pYLUJs&sig=VH5QaLnoknbVQ1dYtOpvdDPHxy8#v=onepage&q=epistemologia%20de%20La%20investigaci%C3%B3n%20b%C3%A1sica&f=false>
- Bienefeld, N., Boss, J., Lüthy, R., Brodbeck, D., Azzati, J., Blaser, M., . . . Keller, E. (2023). Solving the explainable AI conundrum by bridging clinicians' needs and developers' goals. *npj Digital Medicine*, 6. <https://doi.org/10.1038/s41746-023-00837-4>
- Cabrera, S., y Cepeda, J. (2022). La epistemología, guía para el conocimiento científico. *Portal De La Ciencia*, 3(2), 123–133. <https://doi.org/10.51247/pdlc.v3i2.317>
- Callirgos, L. A., Gamarra, P. C., y Cisneros, J. D. (2022). *La investigación científica. Una aventura epistémica, creativa e intelectual*. Religacion Press.

- Castelo, E. M. (2025). Problemas de la investigación tecnológica y su aplicación en la generación de innovación. *Journal of Economic and Social Science Research*, 5(1), 146-160. <https://doi.org/10.55813/gaea/jessr/v5/n1/166>
- Castillo, W. (2023). The importance of human supervision in the use of ChatGPT as a support tool in scientific writing. *Metaverse Basic and Applied Research*, 2. <https://doi.org/10.56294/mr202329>
- Cedeño, J., Maitta, I., Vélez, M., y Palomeque, J. (2024). Investigación universitaria con inteligencia artificial. *Revista Venezolana De Gerencia*, 29(106), 817-830. <https://doi.org/10.52080/rvgluz.29.106.23>
- Ciobanu, O., Aicher, A., Kernbach, J., Regli, L., Serra, C., y Staartjes, V. (2024). A critical moment in machine learning in medicine: on reproducible and interpretable learning. *Acta Neurochirurgica*, 166. <https://doi.org/10.1007/s00701-024-05892-8>
- Contreras, L., Puma, I., Morales, J., Gil, C., y Chalco, N. (2024). Ética en el uso de la inteligencia artificial en la investigación científica: desafíos y consideraciones. *Aula virtual*, 5(12), 1-22. <https://doi.org/10.5281/zenodo.14653199>
- Díaz, L. (2024). El uso de la inteligencia artificial en la investigación científica. *Revista Historia De La Educación Latinoamericana*, 26(43), 253-272. <https://doi.org/10.19053/uplc.01227238.18014>
- Duedra, M. J., Tobes, P. G., Montes, M. M., y Galliari, F. (2020). Ciencia básica, ciencia aplicada, extensión y divulgación: de la compartimentalización a la ecología de saberes o transdisciplinariedad. *Trayectorias Universitarias*, 6(11), 1-13. <https://doi.org/10.24215/24690090e038>
- Falcón, A. L., y Serpa, G. R. (2021). Acerca de los métodos teóricos y empíricos de investigación: significación para la investigación educativa. *Revista Conrado*, 17(S3), 22-31. <https://conrado.ucf.edu.cu/index.php/Conrado/article/view/2133>
- Fochler, M., Janger, J., Stampfer, M., Strassnig, M., Benner, M., Charos, A., . . . Unterlass, F. (2025). Just following the money? How research funding shapes the governance of university research. *Science and Public Policy*, 1-16. <https://doi.org/10.1093/scipol/scaf056>

- Fushimi, K., Nakai, Y., Nishi, A., Suzuki, R., Ikegami, M., Nimura, R., . . . Hasunuma, T. (2025). Development of the autonomous lab system to support biotechnology research. *Scientific Reports*, 15, 1-13. <https://doi.org/10.1038/s41598-025-89069-y>
- Gallent, C., Zapata, A., y Ortego, J. (2023). El impacto de la inteligencia artificial generativa en educación superior: una mirada desde la ética y la integridad académica. *RELIEVE. Revista Electrónica de Investigación y Evaluación Educativa*, 29(2), 1-21. <https://doi.org/10.30827/relieve.v29i2.29134>
- Gao, C., Howard, F., Markov, N., Dyer, E., Ramesh, S., Luo, Y., y Pearson, A. (2023). Comparing scientific abstracts generated by ChatGPT to real abstracts with detectors and blinded human reviewers. *NPJ digital medicine*, 6. <https://doi.org/10.1038/s41746-023-00819-6>
- García, J., Carhuas, L., Gonzales, E., y Barrios, C. (2023). Importancia de la Gnoseología y la Epistemología en el proceso de investigación. *Delectus*, 6(2), 77-85. <https://doi.org/10.36996/delectus.v6i2.213>
- Garzón, M., Campo, G. D., y Loor, B. (2025). Análisis sistemático sobre la eficiencia comunicativa entre humanos y sistemas de inteligencia artificial mediante procesamiento del lenguaje natural. *Universitas-XX. Revista de ciencias sociales y humanas*(42). <https://doi.org/10.17163/uni.n42.2025.07>
- Gilbert, S., Kather, J., y Hogan, A. (2024). Augmented non-hallucinating large language models as medical information curators. *npj Digital Medicine*, 7(100). <https://doi.org/10.1038/s41746-024-01081-0>
- Giordano, P. (2022). El constructivismo operativo de Luhmann y su reflexión sobre el conocimiento científico. *Estudios sociológicos*, 40(120), 725-754. <https://doi.org/10.24201/es.2022v40n120.2186>
- Hartung, T. (2025). AI, agentic models and lab automation for scientific discovery. The beginning of scAInce. *Frontiers in Artificial Intelligence*, 8. <https://doi.org/10.3389/frai.2025.1649155>
- Hu, X. (2021). Hempel on scientific understanding. *Studies in History and Philosophy of Science*, 88, 164-171. <https://doi.org/10.1016/j.shpsa.2021.05.009>

- Hua, C., Cao, X., Liao, B., y Li, S. (2023). Advances on intelligent algorithms for scientific computing: An overview. *Frontiers in Neurorobotics*, 17. <https://doi.org/10.3389/fnbot.2023.1190977>
- Humpiri, J., Humpri, F. d., y Mamani , E. (2021). Teorías científicas. Las propuestas de Popper y Kuhn sobre investigaciones científicas. *Horizontes Revista de Investigación en Ciencias de la Educación*, 5(17), 277-296. <https://doi.org/10.33996/revistahorizontes.v5i17.171>
- Joyce, D., Kormilitzin, A., Smith, K., y Cipriani, A. (2023). Explainable artificial intelligence for mental health through transparency and interpretability for understandability. *npj Digital Medicine*, 6. <https://doi.org/10.1038/s41746-023-00751-9>
- Kapoor, S., y Narayanan, A. (2023). Leakage and the reproducibility crisis in machine-learning-based science. *Patterns*, 4(9). <https://doi.org/10.1016/j.patter.2023.100804>
- Kontinen, L., Koulou, R., y Sankari, S. (2022). Research agenda for algorithmic fairness studies: Access to justice lessons for interdisciplinary research. *Frontiers in Artificial Intelligence*, 5. <https://doi.org/10.3389/frai.2022.882134>
- Krenn, M., Pollice, R., Guo, S., Aldeghi, M., Cervera, A., Friederich, P., . . . Aspuru, A. (2022). On scientific understanding with artificial intelligence. *Nature Reviews Physics*, 4, 761-769. <https://doi.org/10.1038/s42254-022-00518-3>
- Kuppler, M., Kern, C., Bach, R., y Kreuter, F. (2022). From fair predictions to just decisions? Conceptualizing algorithmic fairness and distributive justice in the context of data-driven decision-making. *Frontiers in Sociology*, 7. <https://doi.org/10.3389/fsoc.2022.883999>
- Lopardo, H. Á. (2024). Ciencia básica, ciencia aplicada y técnica. *Acta Bioquímica Clínica Latinoamericana*, 58(1). <https://www.redalyc.org/journal/535/53576131002/53576131002.pdf>
- Loza, J., Riera, V., Miranda, M., y Lopez, D. (2026). Reproducibilidad en los flujos de trabajo bioinformáticos aplicados a genómica y oncología clínica. *La Ciencia al Servicio de la Salud y Nutrición*, 16(2), 155-162. <https://doi.org/10.47187/cssn.Vol16.Iss2.453>

- Mata, K., Sancán, V., Káiser, I., y Kaiser, R. (2024). Una revisión sistemática del uso de la Inteligencia artificial en el desarrollo de investigaciones científicas. *Reincisol*, 3(6), 1642-1660. [https://doi.org/10.59282/reincisol.V3\(6\)1642-1660](https://doi.org/10.59282/reincisol.V3(6)1642-1660)
- Meskó, B., y Topol, E. (2023). The imperative for regulatory oversight of large language models (or generative AI) in healthcare. *npj Digital Medicine*, 6, 1-6. <https://doi.org/10.1038/s41746-023-00873-0>
- Morin, E. (1999). *Los siete saberes necesarios para la educación del futuro*. UNESCO. <http://users.df.uba.ar/solari/Docencia/Complejos/morin.pdf>
- Omiye, J., Lester, J., Spichak, S., Rotemberg, V., y Daneshjou, R. (2023). Large language models propagate race-based medicine. *npj Digital Medicine*, 6. <https://doi.org/10.1038/s41746-023-00939-z>
- Penabad, L., Morera, M., y Penabad, M. (2024). Guía para uso y reporte de inteligencia artificial en revistas científico-académicas. *Revista Electrónica Educare*, 28(S), 1-41. <https://doi.org/10.15359/ree.28-s.19830>
- Polo, L., y Casique, R. (2025). Propuesta metodológica para la recuperación de información documental: integración de grafos de conocimiento y redes neuronales. *Investigación Bibliotecológica: Archivonomía, Bibliotecología e Información*, 39(105), 141-163. <https://doi.org/10.22201/iibi.24488321xe.2025.105.59051>
- Qian, C., Liu, X., Ripley, C., Qian, M., Liang, F., y Yu, W. (2022). Digital twin—Cyber replica of physical things: Architecture, applications and future research directions. *Future Internet*, 14(2). <https://doi.org/10.3390/fi14020064>
- Ros, P., y Pérez, Á. (2023). ChatGPT: una novedosa herramienta de escritura para artículos científicos, pero no un autor (por el momento). *Revista de Neurología*, 76(8). <https://doi.org/10.33588/rn.7608.2023066>
- Ross, T., Reichmann, S., Cole, N. L., Fessl, A., Klebel, T., y Pontika, N. (2022). Dynamics of cumulative advantage and threats to equity in open science: A scoping review. *Royal Society Open Science*, 9(1), 211032. <https://doi.org/10.1098/rsos.211032>

- Sánchez, M. (2024). Inteligencia artificial generativa y los retos en la protección de datos personales. *Estudios en derecho a la información*, 9(18), 179-205. <https://doi.org/10.22201/ijj.25940082e.2024.18.18852>
- Schuff, P., y Hubert, M. (2024). Agendas de investigación y financiación de la ciencia y la tecnología. *Revista Iberoamericana De Ciencia, Tecnología Y Sociedad - CTS*, 20(59), 155–177. <https://doi.org/10.52712/issn.1850-0013-536>
- Semmelrock, H., Ross, T., Kopeinik, S., Theiler, D., Haberl, A., Thalmann, S., y Kowald, D. (2025). Reproducibility in machine-learning-based research: Overview, barriers and drivers. *AI Magazine*, 46(2), 1-19. <https://doi.org/10.1002/aaai.70002>
- Singh, M., Srivastava, R., Fuenmayor, E., Kuts, V., Qiao, Y., Murray, N., y Devine, D. (2022). Applications of digital twin across industries: A review. *Applied Sciences*, 12(11). <https://doi.org/10.3390/app12115727>
- Swaminathan, A., López, I., Garcia, R., Heist, T., McClintock, T., Caoili, K., . . . Nock, M. (2023). Natural language processing system for rapid detection and intervention of mental health crisis chat messages. *npj Digital Medicine*, 6. <https://doi.org/10.1038/s41746-023-00951-3>
- Tang, L., Sun, Z., Idnay, B., Nestor, J., Soroush, A., Elias, P., . . . Peng, Y. (2023). Evaluating large language models on medical evidence summarization. *npj Digital Medicine*, 6. <https://doi.org/10.1038/s41746-023-00896-7>
- Tapia, R., y Recalde, A. (2024). Mejora de gestión documental: revisión sistemática. Ingenium et Potentia. *Revista Electrónica Multidisciplinaria de Ciencias Básicas, Ingeniería y Arquitectura*, 6(11), 16-26. <https://doi.org/10.35381/i.p.v6i11.4076>
- Tobias, A., y Wahab, A. (2025). Autonomous “self-driving” laboratories: A review of technology and policy implications. *Royal Society Open Science*, 12(7). <https://doi.org/10.1098/rsos.250646>
- Ugarte, E., y Parra, G. (2021). La importancia del financiamiento sobre la producción científica en México. *Investigación bibliotecológica*, 35(87), 187-202. <https://doi.org/10.22201/iibi.24488321xe.2021.87.58330>

Zárate, L. O. (2022). Explicabilidad (de la inteligencia artificial). *EUNOMÍA. Revista en Cultura de la Legalidad*(22), 328-344.
<https://doi.org/10.20318/eunomia.2022.6819>

Zemelman, H. (2021). Pensar Teórico y Pensar Epistémico: los retos de las Ciencias Sociales latinoamericanas. *Espacio Abierto*, 30(3), 234-244.
<https://www.redalyc.org/journal/122/12268654011/12268654011.pdf>

INNOVACIÓN e INTELIGENCIA ARTIFICIAL EN LA INVESTIGACIÓN CIENTÍFICA

Innovación e Inteligencia Artificial en la Investigación Científica analiza cómo los sistemas inteligentes están transformando la producción de conocimiento, desde la búsqueda documental y el diseño metodológico hasta el análisis de datos, la experimentación y la comunicación académica.

Con una perspectiva crítica, la obra examina sus principales aplicaciones, oportunidades y riesgos, destacando la importancia de la calidad de los datos, la supervisión humana, la ética, la transparencia y la responsabilidad científica. Un libro dirigido a investigadores, docentes, estudiantes y profesionales interesados en comprender el papel de la inteligencia artificial en el desarrollo de una ciencia más rigurosa, colaborativa e innovadora.



CIDPROS

Centro de innovación y desarrollo profesional

ISBN: 978-9907-9562-7-6

