



**FUNDACIÓN UNIVERSITARIA SANITAS
CONSEJO DIRECTIVO**

ACUERDO No. 127 DE 2025

(Aprobado en la sesión del 20 de mayo de 2025, que consta en el Acta número 442 del Consejo Académico de la Fundación Universitaria Sanitas)

“POR EL CUAL SE DEFINEN LOS LINEAMIENTOS INSTITUCIONALES PARA EL ADECUADO USO DE LA INTELIGENCIA ARTIFICIAL (IA) EN LA FUNDACIÓN UNIVERSITARIA SANITAS – UNISANITAS”

EL CONSEJO ACADÉMICO DE LA FUNDACIÓN UNIVERSITARIA SANITAS, EN USO DE LAS ATRIBUCIONES QUE LE CONFIERE EL ESTATUTO GENERAL DE LA MISMA Y

CONSIDERANDO:

1. Que la Constitución Política en su Artículo 69°, consagró la Autonomía Universitaria, permitiendo que las Universidades puedan darse sus directivas y regirse por sus propios Estatutos de acuerdo con la Ley.
2. Que la Ley 30 de 1992 en su Artículo 1° señala que *“la Educación Superior es un proceso permanente que posibilita el desarrollo de las potencialidades del ser humano de una manera integral, se realiza con posterioridad a la educación media o secundaria y tiene por objeto el pleno desarrollo de los alumnos y su formación académica o profesional”*.
3. Que la Ley de Educación Superior (Ley 30 de 1992), en su Artículo 28° precisó que: *“La autonomía universitaria consagrada en la Constitución Política de Colombia y de conformidad con la presente Ley, reconoce a las universidades el derecho a darse y modificar sus estatutos, designar sus autoridades académicas y administrativas, crear, organizar y desarrollar sus programas académicos, definir y organizar sus labores formativas, académicas, docentes, científicas y culturales, otorgar los títulos correspondientes, seleccionar a sus profesores, admitir a sus alumnos y adoptar sus correspondientes regímenes, y establecer, arbitrar y aplicar sus recursos para el cumplimiento de su misión social y de su función institucional”*.

4. Que el Artículo 29° de la Ley 30, determina que *“la autonomía de las instituciones universitarias o escuelas tecnológicas y de las instituciones técnicas profesionales estará determinada por su campo de acción y de acuerdo con la presente Ley, en los siguientes aspectos: a) Darse y modificar sus estatutos; b) Designar sus autoridades académicas y administrativas; c) Crear, desarrollar sus programas académicos, lo mismo que expedir los correspondientes títulos; d) Definir y organizar sus labores formativas, académicas, docentes, científicas, culturales y de extensión; e) Seleccionar y vincular a sus docentes, lo mismo que a sus alumnos; f) Adoptar el régimen de alumnos y docentes; g) Arbitrar y aplicar sus recursos para el cumplimiento de su función institucional”*.
5. Que, en la misma Ley (Artículo 6°), se señala como una función de la Instituciones de Educación Superior *“(…) Profundizar en la formación integral de los colombianos dentro de las modalidades y calidades de la Educación Superior, capacitándolos para cumplir las funciones profesionales, investigativas y de servicio social que requiere el país (...)”*.
6. Que la Ley General de Educación (Ley 115 de 1994), en su Artículo 2°, declara que el Servicio Educativo comprende *“el conjunto de normas jurídicas, los programas curriculares, la educación por niveles y grados, la educación no formal, la educación informal, los establecimientos educativos, las instituciones sociales (estatales o privadas) con funciones educativas, culturales y recreativas, los recursos humanos, tecnológicos, metodológicos, materiales, administrativos y financieros, articulados en procesos y estructuras para alcanzar los objetivos de la educación”*.
7. Que la Fundación Universitaria Sanita (en adelante Unisanitas) en el marco de sus veinte años de fundación, con la participación de la comunidad universitaria, actualizó el Proyecto Educativo Institucional (Acuerdo 81 de 2023), de modo que responda a la evolución y a los desafíos trazados por la Institución en su Plan de Desarrollo Institucional relacionadas con sus diversos procesos misionales, estratégicos y de apoyo.
8. Que Unisanitas, en concordancia con su misión institucional y con los propósitos fundacionales que la orientan, ha adoptado como



lineamiento estratégico la implementación de propuestas innovadoras en los procesos de enseñanza y aprendizaje, en armonía con los estándares internacionales que rigen la educación superior.

9. Que, en virtud de lo anterior, Unisanitas ha optado por posicionarse como una institución de vanguardia, mediante la incorporación progresiva de tecnologías y metodologías pedagógicas contemporáneas, orientadas a la mejora continua de la calidad educativa.
10. Que, en desarrollo de dicha política institucional, Unisanitas ha definido el Aprendizaje Basado en Problemas – ABP como sistema didáctico, el cual estructura el proceso formativo a partir de la interacción entre estudiantes, saberes y profesores, de conformidad con los propósitos de formación establecidos en cada programa académico. Este enfoque permite al estudiante adelantar su proceso de aprendizaje en función del perfil profesional definido, y le facilita su desarrollo integral, en atención a las tendencias del contexto nacional e internacional, así como a las necesidades del sector productivo y de la sociedad en general.
11. Que de conformidad con lo dispuesto en el literal d) del Artículo 37º del Estatuto General de Unisanitas, es función del Consejo Académico *“decidir sobre los asuntos académicos que estatutariamente no sean de competencia de otros órganos o autoridades de la Institución.”*
12. Que el Consejo Académico, en su sesión del veinte (20) de mayo de dos mil veinticinco (2025), la cual consta en el Acta No. 442 de ese organismo, consideró necesario definir unos lineamientos institucionales que sirvan de marco normativo que guíe el uso ético, responsable y efectivo de herramientas de inteligencia artificial en Unisanitas, fomentando su integración en procesos educativos, administrativos y de investigación, con el fin de asegurar que su implementación sea propicia para el desarrollo académico, humano y social de los miembros de la comunidad universitaria, con respeto a los principios de equidad, transparencia y protección de datos.



En mérito de lo anterior,

ACUERDA:

Artículo 1º.- Aprobar los **LINEAMIENTOS INSTITUCIONALES PARA EL ADECUADO USO DE LA INTELIGENCIA ARTIFICIAL (IA) EN UNISANITAS**, que se incorporan y forman parte integrante del presente Acuerdo.

Artículo 2º.- Socializar los referidos lineamientos con toda la comunidad universitaria de Unisanitas, a efectos de su apropiación e implementación por parte de los miembros y/o diferentes grupos de interés de la misma.

Artículo 3º.- El presente acuerdo se complementa con el Estatuto General de Unisanitas, sus reglamentos estudiantiles y las políticas Institucionales de la misma, razón por la cual su implementación e interpretación deberá estar articulada y en plena coherencia con dicha normatividad, así como con las demás disposiciones legales que se ocupen de la materia.

Artículo 4º.- Corresponde al Consejo Académico de Unisanitas, modificar, complementar y/o adicionar los aspectos que considere pertinentes en relación con los lineamientos aprobados mediante el presente Acuerdo, así como interpretar o aclarar, en caso de duda, el alcance de cualquiera de las disposiciones contenidas en los mismos.

Artículo 5º.- El presente acuerdo rige a partir de su expedición y deroga las demás disposiciones que le sean contrarias.

Dado a los veintiocho (28) días del mes de mayo de dos mil veinticinco (2025).

COMUNIQUESE, PUBLÍQUESE Y CÚMPLASE,


MARIO ARTURO ISAZA RUGET
Presidente Consejo Académico


JOHNS STEVE NAVARRO LARA
Secretario General

LINEAMIENTOS INSTITUCIONALES PARA EL ADECUADO USO DE LA INTELIGENCIA ARTIFICIAL FUNDACIÓN UNIVERSITARIA SANITAS



Buenas prácticas para el uso responsable de la inteligencia artificial



En **UNISANITAS** promovemos el uso ético, responsable y consciente de la inteligencia artificial como herramienta de apoyo en la educación, la investigación y la gestión.

Esta guía reúne buenas prácticas para estudiantes, docentes e investigadores, con el fin de integrar la IA de forma segura, crítica y alineada con nuestros valores institucionales.



1. Usa la IA con propósitos claros

Explicita su uso en tus entregables

Informa cuándo, cómo y con qué fin se utiliza una herramienta de IA en cualquier actividad académica o institucional. La transparencia es la base de la confianza.



2. La IA apoya, pero tú decides

Nunca reemplaza el criterio humano

Toda decisión académica debe estar fundamentada en el juicio de las personas. La IA es un apoyo, no un sustituto del pensamiento crítico.



3. Protege la privacidad y los datos

Respetar la confidencialidad de la información

Asegura el uso ético y autorizado de los datos personales, académicos o de investigación. Usa plataformas seguras y confiables.



4. Evita sesgos y discriminación

Busca justicia y equidad tecnológica

Asegura que el uso de IA no genere exclusión ni refuerce desigualdades. Evalúa el impacto de los modelos y su neutralidad.



5. Verifica la información generada

Contrasta siempre con fuentes confiables

La IA puede generar errores o falsedades. Evalúa la validez y relevancia del contenido antes de usarlo.



6. Declara y cita el uso de IA

Reconoce su participación en tus trabajos

Menciona las herramientas utilizadas, su rol y el grado de aporte. Fomenta la integridad académica y evita el plagio.



7. Usa IA para aprender mejor, no solo para cumplir

Potencia tu comprensión y análisis

Empléala como un medio para profundizar el aprendizaje, no como atajo. El conocimiento significativo nace del proceso, no del resultado.



8. Comparte las buenas prácticas

Sé agente de cultura ética en IA

Invita a tus compañeros, docentes y colegas a reflexionar y actuar con responsabilidad en el uso de estas herramientas. Todos somos parte del cambio.

Te invitamos a leer el documento completo de Lineamientos Institucionales para el uso de la IA en UNISANITAS, donde encontrarás orientaciones éticas, normativas y pedagógicas para un uso responsable, seguro y transformador de la inteligencia artificial.

TABLA DE CONTENIDO

OBJETIVO GENERAL.....	3
ALCANCE Y PROCESO DE CONSTRUCCIÓN DEL DOCUMENTO.....	3
DEFINICIONES.....	3
INTRODUCCIÓN.....	6
MARCO TEÓRICO.....	7
MARCO NORMATIVO.....	8
PRINCIPIOS ÉTICOS.....	10
IA y ABP.....	12
IA EN INVESTIGACIÓN.....	14
Observatorio para la Implementación Tecnológica en Investigación.....	16
Marco Ético y Normativo de Mínimos Aceptables en Investigación.....	17
EVALUACIÓN DE RIESGOS Y ACCIONES INSTITUCIONALES.....	18
COMO ABORDAR UN CASO DE USO INDEBIDO DE IA.....	20
REGLAS PARA EL USO DE HERRAMIENTAS DE IA.....	20
EVALUACIÓN Y SEGUIMIENTO.....	21
AGRADECIMIENTOS.....	22
BIBLIOGRAFÍA.....	23

OBJETIVO GENERAL

Establecer un marco normativo que guíe el uso ético, responsable y efectivo de herramientas de inteligencia artificial en Unisanitas, fomentando su integración en procesos educativos, administrativos y de investigación, con el fin de asegurar que su implementación sea propicia para el desarrollo académico, humano y social de los miembros de la comunidad universitaria, con respeto a los principios de equidad, transparencia y protección de datos.

ALCANCE Y PROCESO DE CONSTRUCCIÓN DEL DOCUMENTO

Esta política aplica a todos los miembros de la comunidad académica de Unisanitas, incluyendo estudiantes, docentes, investigadores y personal administrativo que utilicen herramientas de inteligencia artificial en actividades relacionadas con la enseñanza, investigación, gestión administrativa y cualquier otro ámbito institucional. Abarca el uso, desarrollo, evaluación y regulación de tecnologías basadas en IA, garantizando que estas se empleen de manera ética, responsable y alineada con los principios institucionales.

Este documento ha sido elaborado con la participación de un equipo de expertos interdisciplinario de UNISANITAS, conformado por profesionales de diversas áreas del conocimiento. A través de un proceso de análisis y consulta, se han integrado perspectivas académicas, tecnológicas, éticas y administrativas para garantizar que los lineamientos aquí establecidos respondan a las necesidades y desafíos actuales en el uso de la inteligencia artificial en la educación, la investigación y la gestión institucional.

El proceso de construcción de este documento incluyó mesas de trabajo, revisión de normativas nacionales e internacionales, así como la adopción de principios de buenas prácticas en el uso responsable de la IA. Esta metodología colaborativa asegura que las directrices propuestas sean aplicables, pertinentes y alineadas con la misión y visión de Unisanitas, fomentando un desarrollo tecnológico ético y sostenible dentro de la comunidad académica.

DEFINICIONES

 **Inteligencia Artificial (IA):** Es un campo multidisciplinar, actualmente con fuertes aportes desde la computación y donde su objetivo es estudiar y desarrollar sistemas que piensan y actúan racionalmente a partir de asociar información y percibir contextos para crear nuevo conocimiento y tomar decisiones. Su desarrollo se basa principalmente en la implementación de algoritmos, modelos computacionales y lógicos matemáticos a partir de la inspiración de sistemas del mundo real que expresan inteligencia y/o grandes volúmenes de datos en su procesamiento para emular funciones cognitivas humanas, desempeñando procesos de “automatización de tareas que asociamos al pensamiento humano, actividades como toma de

decisiones, resolución de problemas y aprendizaje” (Bellman, 1978); generando así, valor en distintos sectores económicos y sociales. De esta manera, la IA impulsa la automatización de tareas, la optimización de procesos y el desarrollo de soluciones innovadoras en el marco de la transformación digital (Consejo Nacional de Política Económica y Social, 2019; Guío Español et al., 2021; UNESCO, 2021).

 **Inteligencia Humana (IH):** La inteligencia humana es uno de los constructos más complejos de la psicología y las disciplinas cognitivas. Hoy, a pesar de más de 100 años de investigación sistemática, esta conceptualización continúa en evolución, lo que refleja la profundidad del fenómeno en una diversidad de perspectivas teóricas que intentan comprenderlo (Sternberg, 2019). Hoy, el término “inteligencia” deriva etimológicamente del latín “intelligentia”, da cuenta de la capacidad de entender, comprender y establecer conexiones entre varias ideas (Simpson & Weiner, 1989). De manera histórica, es necesario relatar, que las primeras aproximaciones científicas al término “inteligencia” salen a la luz a finales del siglo XIX con los trabajos de Francis Galton y, posteriormente, Alfred Binet, ellos documentaron las pruebas para su medición (Nicolas et al., 2013); de manera subsecuente, durante el siglo XX, el concepto se transformó desde una concepción cuantitativa hasta modelos multidimensionales. La inteligencia se conceptualiza como una capacidad compleja y multifacética que, permite al individuo adaptarse a su entorno, razonar, resolver problemas, adquirir conocimientos y aprender de experiencias previas o imaginarias (Gottfredson, 1997). Esto, no constituye un constructo de facto, sino un conjunto de habilidades que son sinérgicas entre sí. Por su parte, Gardner propuso un modelo de inteligencias múltiples, documentando ocho formas de inteligencia: lógico matemática, lingüística, espacial, corporal, musical, kinestésica, interpersonal, intrapersonal y naturalista (Davis et al., 2011). Steinberg desarrolló la teoría triárquica que distingue la inteligencia práctica, creativa y analítica (Moustafa et al., 2018). Estos modelos contribuyeron de una manera contundente a una cosmogonía holística del fenómeno “inteligencia”.

 **Confidencialidad:** Principio que asegura que los datos personales, de investigación o información sensible sean protegidos frente a accesos no autorizados, manteniendo su privacidad y uso ético (Ley estatutaria 1581, 2012).

 **Generación de Contenido Sintético:** Producción de texto, imágenes, audio o video mediante herramientas de IA, que emulan el contenido creado por humanos (Franganillo, 2022, 2023).

 **Proporcionalidad:** En el contexto de la inteligencia artificial enfatiza que el uso de sistemas de IA debe limitarse a lo necesario para alcanzar un objetivo legítimo, evitando aplicaciones excesivas o innecesarias que puedan generar riesgos o daños. Este enfoque busca minimizar los riesgos y maximizar los beneficios, asegurando un equilibrio adecuado entre ambos factores (UNESCO, 2021).

 **Uso Responsable de la IA:** Mecanismos de control humano sobre los sistemas, evitando que las decisiones críticas se tomen de manera totalmente automatizada. Además, enfatiza la

necesidad de garantizar la seguridad de los datos y la integridad de las personas afectadas por la IA (Ministerio de Ciencia, Tecnología e Innovación, 2024).

-  **Ética en la inteligencia artificial:** La ética en la inteligencia artificial comprende un conjunto de principios y normativas orientadas a garantizar que su desarrollo y aplicación sean responsables, equitativos y transparentes. Incluye la protección de la privacidad, la equidad en la toma de decisiones y la supervisión humana en su implementación. Para ello, se requieren regulaciones que permitan la auditoría de los algoritmos y la mitigación de sesgos, asegurando que la IA se utilice en beneficio de la sociedad sin generar desigualdades ni riesgos injustificados, promoviendo la responsabilidad en su uso y desarrollo (Guío Español et al., 2021; Ministerio de Ciencia, Tecnología e Innovación, 2024; UNESCO, 2021).
-  **Inteligencia Computacional Bioinspirada:** Hace referencia al uso de técnicas inspiradas en procesos biológicos en el campo de la inteligencia artificial, específicamente en el área del aprendizaje evolutivo. Implica adoptar un enfoque basado en el aprendizaje y la cognición de un sistema o mecanismo biológico y utilizar a partir de su adaptación computacional buscar combinaciones óptimas de funciones en modelos empíricos (Triyason et al., 2015).
-  **Planeación y Optimización de soluciones:** La planificación, programación y optimización es una rama de la Inteligencia Artificial que estudia la búsqueda de la mejor solución o camino, a partir de elegir una secuencia de acciones para alcanzar un objetivo determinado desde un estado inicial determinado. Las acciones suelen basarse de forma esquemática en lenguajes e información del entorno (Ghallab et al., 1998).
-  **Inteligencia Basada en Conocimiento:** También denominados sistemas expertos, son una derivación de la inteligencia artificial que se ocupa de utilizar desarrollos computacionales para simular la inteligencia humana. Formalmente se refiere a herramientas computacionales que tienen como objetivo encapsular la experiencia de expertos humanos en un gestor de conocimiento accionable (Dawson et al., 2003). Estos sistemas están diseñados para brindar asesoramiento, tomar decisiones o resolver problemas en función del conocimiento almacenado en ellos.
-  **Aprendizaje Automático (Machine Learning):** Rama de la inteligencia artificial que desarrolla modelos matemáticos computacionales basados en la idea de que aprende siempre que puede utilizar una experiencia de manera que su rendimiento mejora sobre otras experiencias similares, formalmente Tom M. Mitchel define que “Un programa se dice que **aprende** a partir de una experiencia E respecto a una tarea específica T con un mediada de desempeño P à Si su desempeño en la tarea T , medido por P , mejora con la experiencia E ”, y este aprendizaje se da a partir del uso de datos y algoritmos que identifican patrones y hacen predicciones. Se usa en aplicaciones como reconocimiento de voz, diagnóstico médico y sistemas de recomendación (Mitchell, 1997).
-  **Redes Neuronales Artificiales:** Modelos computacionales bioinspirados en los procesos de aprendizaje y cognición del cerebro humano, compuestos por capas de nodos

interconectados que procesan información. Son la base de muchas técnicas de inteligencia artificial, como el aprendizaje profundo y la visión por computadora (Ministerio de Ciencia, Tecnología e Innovación, 2024).

 **Aprendizaje Profundo (Deep Learning):** Subárea del aprendizaje automático que utiliza redes neuronales con múltiples capas para analizar grandes volúmenes de datos y extraer patrones complejos. Es fundamental en aplicaciones como la visión artificial, el procesamiento de lenguaje natural y los asistentes virtuales (UNESCO, 2021).

 **Cuarta Revolución Industrial:** Transformación tecnológica caracterizada por la convergencia de la inteligencia artificial, el Internet de las cosas, la biotecnología, la robótica y otras innovaciones digitales, con un impacto profundo en la economía y la sociedad. Se basa en la automatización inteligente y la interconexión digital de sistemas (Consejo Nacional de Política Económica y Social, 2025).

INTRODUCCIÓN

En el marco de su compromiso con la excelencia académica, científica y humana, la Fundación Universitaria Sanitas (Unisanitas) reconoce el impacto creciente de la inteligencia artificial (IA) en los ámbitos educativos, profesionales y sociales. Como parte de su misión de formar profesionales íntegros y su visión de ser líder en la educación superior el área de las ciencias de la salud y aplicadas a esta, Unisanitas considera fundamental establecer lineamientos claros para el uso ético, responsable y crítico de las herramientas de IA en la Institución.

El rápido desarrollo de la IA en el mundo ha impulsado la necesidad de una regulación institucional que garantice su uso en beneficio de la comunidad académica, asegurando su integración de manera que potencie los procesos de aprendizaje, investigación y gestión administrativa, sin comprometer la equidad, la privacidad y la integridad académica.

Según el CONPES 3975 de 2019, la IA es un acelerador clave de la transformación digital en Colombia y debe alinearse con principios de seguridad, transparencia y equidad en el ámbito educativo (Consejo Nacional de Política Económica y Social, 2019).

En consonancia con recomendaciones internacionales como las establecidas por la UNESCO en su Recomendación sobre la Ética de la Inteligencia Artificial (UNESCO, 2021), estos lineamientos enfatizan la necesidad de minimizar los sesgos, promover el acceso equitativo a la tecnología y garantizar la protección de datos personales en su implementación. Así mismo, sigue las recomendaciones de universidades nacionales e internacionales, que han delineado marcos éticos para el uso de IA en la educación superior.

Unisanitas no solo busca adoptar la IA de manera estratégica, sino también liderar su implementación en el sector educativo en Colombia, garantizando su uso ético y alineado con los valores institucionales.

El presente documento establece los lineamientos para su aplicación en la docencia, la investigación y la gestión universitaria, asegurando que el aprovechamiento de la IA esté enmarcado en principios de justicia, equidad y desarrollo sostenible (UNESCO, 2021).

MARCO TEÓRICO

La Inteligencia artificial es una disciplina eminentemente científica en constante evolución en el campo de la cognición y la computación, su conceptualización y desarrollo se puede documentar desde la filosofía clásica hasta los avances en lógica del siglo XX. El concepto de máquinas pensantes tiene raíces en la historia de la filosofía de occidente donde varios pensadores exploran la idea que un artefacto pueda replicar el pensamiento del hombre, cuestionando así la naturaleza de la conciencia y la capacidad del ser humano de resolver problemas o lo que definiremos vagamente como inteligencia. Aristóteles (384-322 AC) desarrolló un sistema de lógica que influye directamente en el desarrollo de la Inteligencia artificial. Este pensador propuso el “silogismo” como estructura de razonamiento que permite inferencias a partir de premisas establecidas (Aristotle, IV a.C.). Este trabajo estableció las bases para el estudio de la lógica, posteriormente, durante el siglo XVII, René Descartes (1596-1650) estableció la idea del “dualismo cartesiano”, hablando del cuerpo y la mente como entidades separadas, influenció la visión que la inteligencia podría existir independientemente de la materia física, lo que posteriormente inspiró la computación basado en modelos abstractos (Descartes, 1641).

Después, en el siglo XVIII, el filósofo Gottfried Wilhelm Leibniz dedujo que el razonamiento humano se reduciría a cálculos lógicos; esto abrió las puertas para la lógica computacional. Este pensador tiene una obra cardinal “Dissertation on the Art of Combinations”, estableciendo que el pensamiento se puede representar mediante las combinaciones matemáticas (Leibniz, 1666).

El desarrollo de la Inteligencia artificial está ligado, necesariamente, a los avances en lógica y la teoría de la computación que proporcionaron un marco conceptual para la creación de sistemas que pudiesen emular el pensamiento de los seres humanos. A principios del siglo XX, George Boole postuló el álgebra que lleva su apellido; un sistema basado en valores binarios, verdadero y falso, estableciendo las bases conceptuales para el diseño de circuitos electrónicos y la computación digital (Boole, 1854). Para 1930, Kurt Gödel (1906-1978) introdujo “los teoremas de incompletitud” los cuales demostraron que, no todos los algoritmos matemáticos se pueden probar o se pueden derivar en su propio marco conceptual (Gödel, 1931), lo que tuvo un impacto cardinal en la teoría de la computación estableciendo un límite, complementariamente, donde las máquinas no podían resolver algunos problemas. De manera paralela, Alan Turing (1912-1954) conceptuó y, de hecho, en un esfuerzo heroico materializó la máquina de Turing, un modelo teórico que podía realizar cálculos con reglas predefinidas (A. M. Turing, 1937). Además, creó una base robusta para la “Church–Turing thesis” esta establece que, “cualquier problema computacional se puede resolver con un algoritmo ejecutado en una máquina de Turing”. Es importante de notar que, en su artículo, “Computing Machinery and Intelligence” (1950) introdujo el test de Turing en

el ánimo de evaluar la capacidad de un sistema para imitar el comportamiento humano, en especial, imitar la característica de resolver problemas (A. Turing, 1950).

El primer intento para construir máquinas con potencial inteligente; es decir, capaz de resolver problemas por sí mismas, iniciaron a mediados del siglo XX cuando existía ya, una capacidad computacional prudente para la implementación de algoritmos robustos. En este punto, es necesario anotar el trabajo de Norbert Wiener (1894–1964) quien cobra valor introduciendo el concepto de “cibernética” como un área de estudio que demostraba la comunicación y el control en los sistemas biológicos y mecánicos (Wiener, 1961). Este concepto es determinante para la Inteligencia artificial como concepto puesto que establece la “regulación de los sistemas de Inteligencia artificial por retroalimentación” lo que es fundamental para el aprendizaje automático. En 1951, Marvin Minsky y Dean Edmonds desarrollaron la primera red neuronal artificial llamada “Stochastic Neural Analog Reinforcement Calculator” (SNARC por sus siglas en inglés) en la Universidad de Harvard (Minsky, 1954), esto fue determinante en la objetivación del aprendizaje automático.

En 1951 Marvin Minsky, John McCarthy, Nathaniel Rochester y Claude Shannon encabezaron la conferencia de Dartmouth donde habló oficialmente del término de Inteligencia artificial, allí se estableció la teoría de que todo aprendizaje o cualquier otra característica de la inteligencia puede describirse con tanta precisión que una máquina puede recrearlo (McCarthy et al., 1955).

Los acontecimientos mencionados marcaron el comienzo de una disciplina que, con el tiempo, ha evolucionado de manera significativa, ampliando su influencia más allá del ámbito teórico y técnico. En la actualidad, la inteligencia artificial (IA) no solo representa un campo de estudio, sino también un motor clave en la transformación digital de diversos sectores, incluyendo la educación superior, la investigación y la gestión administrativa. Su desarrollo ha generado debates sobre su impacto en la equidad, la ética, la propiedad intelectual y la gobernanza, lo que ha llevado a la formulación de marcos regulatorios y normativos a nivel nacional e internacional (Consejo Nacional de Política Económica y Social, 2025).

MARCO NORMATIVO

Desde un enfoque global, la UNESCO, en su Recomendación sobre la Ética de la Inteligencia Artificial, establece que la IA debe desarrollarse y utilizarse bajo principios de transparencia, equidad, seguridad y minimización de riesgos. Además, resalta la importancia de garantizar un acceso equitativo a estas tecnologías, asegurando que los avances en IA no amplíen brechas preexistentes en educación y acceso al conocimiento (UNESCO, 2021).

A nivel nacional, Colombia ha avanzado en la formulación de políticas para la transformación digital y la adopción ética de la IA. Es así que el CONPES 3975 de 2019, reconoció la IA como un habilitador clave de la transformación digital, promoviendo el desarrollo de infraestructura tecnológica, la creación de marcos normativos y el fortalecimiento de capacidades en el uso de IA en diversos sectores; posteriormente, en 2025, se formuló la Política Nacional de Inteligencia

Artificial, la cual profundiza en la gobernanza, ética, seguridad y desarrollo sostenible de la IA en el país (Consejo Nacional de Política Económica y Social, 2019, 2025).

La Política Nacional de Inteligencia Artificial 2025 también destaca la necesidad de garantizar que la adopción de IA en Colombia tenga un enfoque inclusivo y con perspectiva de género, asegurando que el desarrollo y uso de estas tecnologías no refuercen desigualdades estructurales. Se enfatiza en la importancia de generar políticas de equidad en el acceso a la IA, promoviendo la participación de mujeres y poblaciones tradicionalmente excluidas de los desarrollos tecnológicos (Política Nacional de Inteligencia Artificial, 2024). Además, resalta la urgencia de una IA sostenible, con estrategias para mitigar su impacto ambiental a través del uso responsable de la energía y la reducción de la huella de carbono en los procesos de entrenamiento y aplicación de modelos de IA (Consejo Nacional de Política Económica y Social, 2025).

En el ámbito educativo, la UNESCO ha subrayado el potencial de la IA para mejorar la enseñanza y el aprendizaje, pero también advierte sobre los riesgos que conlleva no regularla adecuadamente. En la *“Guía para el uso de IA generativa en educación e investigación”* se enfatiza en la necesidad de desarrollar competencias digitales y pensamiento crítico en estudiantes y docentes, asegurando que la IA sea utilizada como un complemento para la enseñanza y no como un sustituto del rol humano en el proceso educativo (Miao, 2024). Se recomienda la creación de políticas de formación en ética de la IA para profesores, con el fin de preparar a las instituciones educativas para un uso informado y seguro de estas herramientas (Miao, 2021).

Otro aspecto clave en la adopción de IA en el ámbito universitario es la propiedad intelectual. Según el documento *Copyright and Artificial Intelligence*, la generación de contenido sintético mediante IA plantea interrogantes sobre la titularidad de los derechos de autor y la integridad académica. Es crucial establecer normativas claras que regulen el uso de IA en la producción de conocimiento, garantizando el reconocimiento adecuado de las fuentes y evitando el plagio académico (United States Copyright Office, 2025). En términos de gobernanza y regulación, la *Hoja de Ruta para la Adopción Ética y Sostenible de la Inteligencia Artificial en Colombia* establece que la IA debe implementarse bajo un enfoque de desarrollo sostenible y justicia social, asegurando que sus beneficios sean accesibles para toda la población y evitando impactos negativos en el empleo y la privacidad (Ministerio de Ciencia, Tecnología e Innovación, 2024). De igual forma, el Marco Ético de IA en Colombia enfatiza la importancia de establecer criterios de supervisión y control para garantizar que la IA sea utilizada de manera justa y segura (Guío Español et al., 2021).

En este contexto, los lineamientos de Unisanitas para el uso de IA buscan garantizar que su implementación en la educación, la investigación y la gestión institucional esté enmarcada en principios de transparencia, equidad, seguridad y supervisión humana, alineándose con las mejores prácticas nacionales e internacionales para la adopción ética y efectiva de la inteligencia artificial, sin pasar por alto los marcos teóricos y normativos establecidos en la actualización de su Proyecto Educativo Institucional (Acuerdo No. 081 del 1° de septiembre de 2023) y en su Política Institucional de Propiedad Intelectual (Acuerdo No. 065 del 12 de noviembre de 2020).

PRINCIPIOS ÉTICOS

En la institución, la implementación de los diferentes modelos de Inteligencia artificial se alineará a los principios éticos fundamentales que garanticen el respeto por la dignidad humana, los derechos de autor y la privacidad.

A continuación, establecemos lineamientos y principios para el uso de la Inteligencia artificial en esta institución (Figura 1), en el requerimiento de documentar y regular su aplicación; con un importante incentivo para su adopción y aplicación de manera responsable.

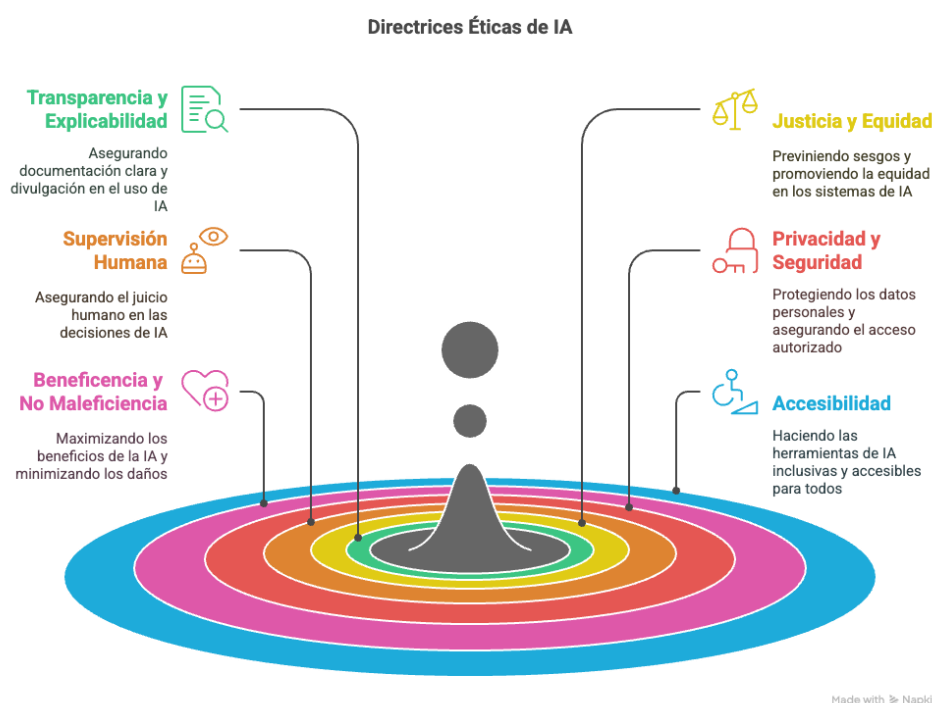


FIGURA 1. DIRECTRICES ÉTICAS DE LA IA EN UNISANITAS

Transparencia y Explicabilidad

Toda implementación, aplicación, integración y uso de inteligencia artificial en actividades académicas e investigativas, debe ser claramente documentada de manera visible. En este sentido, los usuarios están obligados a informar sobre el uso de herramientas de Inteligencia artificial en sus diferentes productos académicos, especificando su alcance, propósito y contribución. Esta directriz no tiene un carácter prohibitivo, por el contrario, tiene el ánimo de propender por la adopción responsable de las diferentes innovaciones tecnológicas que existentes y futuras.

Privacidad y Seguridad

El respeto por la privacidad y la seguridad de los datos de las personas debe prevalecer en el desarrollo y aplicación de la inteligencia artificial, garantizando su uso adecuado y autorizado de

los datos para los fines claramente establecidos en un proyecto o actividad específica. Adicionalmente, es esencial implementar medidas robustas para proteger la información sensible, evitando accesos no autorizados y garantizando su manejo seguro en todo momento.

Justicia y Equidad

El principio de justicia y equidad en la aplicación de la inteligencia artificial establece que su uso no debe causar daño, discriminación, exclusión ni inducir sesgos perjudiciales para las personas. Es fundamental que la implementación de la IA respete y proteja la dignidad humana, evitando cualquier uso malintencionado y promoviendo una distribución justa e igualitaria de sus beneficios. Esto implica desarrollar y utilizar sistemas de IA de manera que se garantice el acceso equitativo a sus ventajas, minimizando desigualdades y asegurando que todos los grupos sociales se beneficien de manera equilibrada.

Beneficencia y no Maleficencia

La IA debe usarse con el propósito fundamental de maximizar beneficios y minimizar cualquier riesgo o posible daño. Es necesario realizar evaluaciones de impacto y riesgos para asegurar que los efectos positivos prevalezcan sobre los negativos.

Supervisión Humana

Es necesario que las decisiones académicas no se releguen de una manera exclusiva o íntegra a los algoritmos de inteligencia artificial. Como se viene argumentando en este lineamiento, estas herramientas eficientizan los procesos humanos, razón por la cual todas las decisiones en lo académico deben estar fundamentadas en el juicio humano y contar, sin excepción, con la intervención activa de personas reales. La inteligencia artificial será un apoyo no un reemplazo del criterio o el juicio humano. En este sentido, es necesario que se adopte como norma ética el “**imperativo categórico**” Kantiano basado en la transparencia y la integridad. A saber, se debe de actuar de manera que las decisiones puedan ser universalizadas y entendidas por todos, respetando siempre la dignidad de los seres humanos, lo que implica que las decisiones que afectan a profesores, estudiantes y personal administrativo, deben de tomarse de manera justa, consciente y reflexiva. Teniendo esto en mente, se garantiza el ejercicio académico con rigor científico en su método (Kant, 1785).

Es cardinal que la implementación de estos aplicativos se utilice como una herramienta que complemente al ser humano y sus capacidades, sin que ello implique su reemplazo. Por lo tanto, aunque la inteligencia artificial puede generar análisis y contenidos de manera autónoma guiada por un prompt (órdenes o requerimientos dadas por un operador a un sistema, generalmente modelos de procesamiento de lenguaje) y su operador humano, es responsabilidad de los diferentes usuarios, llámese estudiantes, profesores, administrativos o contratistas, verificar la relevancia y exactitud de la información extractada de estos modelos tecnológicos; a saber, modelos de procesamiento de lenguaje o algoritmos de análisis. Asimismo, resulta indispensable implementar mecanismos claros para la revisión humana, con el fin de garantizar la integridad de los resultados y los datos a compartir, los cuales deben de ser objetivados en los diferentes textos, diapositivas, charlas, conferencias o presentaciones académicas donde se implementan estas herramientas.

Principio de Accesibilidad

Unisanitas fomentará la inclusión y acceso a las diferentes herramientas de Inteligencia artificial, en cuanto sus medios lo permitan; de esta manera, se implementarán las acciones pertinentes que garanticen que todos los miembros de la comunidad académica puedan tener las mismas oportunidades de capacitarse e implementar las diferentes tecnologías y aplicativos de inteligencia artificial en su quehacer académico y profesional, de una manera responsable y segura.

IAyABP

El Aprendizaje Basado en Problemas (ABP) o Problem-Based Learning (PBL) es una metodología de aprendizaje activo basado en el análisis, estudio y resolución de problemas como parte del abordaje de diferentes temas que pertenezcan a una malla curricular. En esta metodología los estudiantes trabajan en pequeños grupos, con el acompañamiento de un docente facilitador identifican lo que necesitan aprender y buscan activamente información para desarrollar soluciones. Este enfoque fomenta la autonomía, el pensamiento crítico y la aplicación práctica del conocimiento (Trullàs et al., 2022; Yew & Goh, 2016).

Es así como en Unisanitas, de conformidad con lo previsto en el Acuerdo No. 050 del 29 de agosto de 2017 (Sistema Institucional de Créditos Académicos), el ABP se desarrolla en cuatro (4) fases (Figura 2) que estructuran y orientan el proceso de enseñanza-aprendizaje:

Fase I: Planeación, consiste en la identificación de competencias, el análisis del problema y la formulación de objetivos de aprendizaje, estableciendo un plan de trabajo con actividades programadas y roles definidos.

Fase II: Abordaje y estudio del problema, los estudiantes, bajo la orientación del docente, realizan una búsqueda de la literatura, profundizan en los conceptos teóricos y metodológicos mediante seminarios, prácticas de laboratorio y simulaciones, con el fin de construir soluciones fundamentadas.

Fase III: Síntesis, se socializan los aprendizajes adquiridos, se presentan y analizan las soluciones propuestas, y se genera una discusión colectiva que permite consolidar el conocimiento, con la intervención del docente para clarificar aspectos clave.

Fase IV: Retroalimentación, permite evaluar el proceso de aprendizaje, identificar fortalezas y oportunidades de mejora, y establecer estrategias para optimizar futuras experiencias, aplicando autoevaluación, coevaluación y heteroevaluación.

Fases del Aprendizaje Basado en Problemas en Unisanitas



FIGURA 2. FASES DEL APRENDIZAJE BASADO EN PROBLEMAS EN UNISANITAS

Durante la revisión con el grupo de expertos, se determinó que la **Fase II: Abordaje y estudio del problema**, es el momento más adecuado para la implementación de la IA en el aula, en esta fase se visualiza el mayor beneficio al incorporar estas herramientas, debido a la gran cantidad de fuentes bibliográficas disponibles y su crecimiento acelerado, los estudiantes enfrentan una alta carga cognitiva y un considerable consumo de tiempo. Además, esta sobrecarga dificulta la selección de los recursos más pertinentes y actualizados, lo que refuerza la necesidad de herramientas que optimicen este proceso. Por la gran cantidad de información que deben seleccionar, depurar y sintetizar, se evidencia un nivel de aprendizaje superficial que resulta contraproducente frente a los propósitos de formación. El beneficio de incorporar herramientas de IA que aceleren la selección y la síntesis de la mejor información favorece el análisis crítico y el aprendizaje profundo.

Para aclarar el concepto, el aprendizaje profundo tiene que ver con el aprendizaje significativo, la reinterpretación, comprensión, conexión y aplicación de conocimientos en situaciones prácticas, todo en el marco de teoría constructivista del aprendizaje. Todo lo contrario, es el aprendizaje superficial que tiende a emplear estrategias de estudio superficial, como la memorización y la repetición, el conocimiento es estático en tanto el estudiante no reflexiona sobre la aplicación práctica, y solo actúa por motivaciones externas para el logro del aprendizaje (Ortega-Díaz & Hernández-Pérez, 2015).

La metodología de ABP se propone desarrollar habilidades como el pensamiento crítico, la capacidad de trabajar en equipo y las habilidades de comunicación. Estas habilidades principalmente se ejercitan durante la fase de búsqueda de la información, que debe ser realizada

a partir del trabajo en grupo, la selección crítica de las fuentes y la síntesis de la mejor información. Cuando la cantidad de información sobrepasa la capacidad de procesamiento de los estudiantes se incurre en un aprendizaje superficial, lo que significa que no logran vincular lo que están aprendiendo en una estructura conceptual mayor, tienden a aprender de manera pasiva y a menudo memorizan la información de determinado contenido (Santrock, 2006). El posible beneficio que tendría incorporar herramientas de inteligencia artificial se refleja en esta fase, acortando los tiempos de búsqueda e integración de las fuentes para permitir un análisis con mayor profundidad y una síntesis que promueva el aprendizaje significativo.

En el caso de los profesores, la utilización de herramientas de inteligencia artificial ha demostrado facilitar la creación de casos, problemas, escenarios y variantes de cada uno de estos, logrando el abordaje de una mayor cantidad de posibilidades de estudio. Adicionalmente, para la elaboración individualizada de evaluaciones y para el proceso de calificación. En estos casos, el beneficio se explica por un mayor nivel de realismo y detalle de los casos, así como un lenguaje más natural en los problemas generados por IA (Mool et al., 2024).

El proceso de incorporación del uso de las herramientas de inteligencia artificial debe hacerse de manera progresiva y paralela con capacitaciones dirigidas a los estudiantes de los primeros semestres. Estas capacitaciones se deben realizar de forma institucional y estandarizada para lograr el mejor aprovechamiento en la metodología de ABP. Es bien sabido que los estudiantes provienen en su mayoría de modelos de formación tradicionales, en donde el rol del docente es principal y las habilidades de autogestión del conocimiento que logran los estudiantes son mínimas. Incorporar el uso de estas herramientas en los ciclos de formación básica, cuando el acompañamiento docente es más amplio, favorecerá el entrenamiento y su utilización durante ciclo profesional cuando los estudiantes tienen menos acompañamiento y menos tiempo para aprendizaje autónomo.

La principal recomendación que se puede hacer a partir de los estudios de revisión sistemática sobre el uso de inteligencia artificial en educación médica y en educación en salud es su incorporación al currículo desde los primeros semestres de los programas para el entrenamiento y uso orientado de las herramientas de inteligencia artificial en un marco ético y responsable (Sapci & Sapci, 2020)

IA EN INVESTIGACIÓN

La investigación está experimentando una transformación cardinal con la integración de la inteligencia artificial en su ejercicio, lo que permite importantes adelantos en la recopilación, interpretación y análisis de datos; en su conjunto, constituyen un andamiaje fundamental en la construcción de conocimiento. En este sentido, de manera introductoria, vamos a reflexionar en el fundamento epistémico de la ciencia y su integración con este tipo de tecnologías (Hehl, 2016).

Karl Popper expuso que la realidad podría dividirse en tres mundos interconectados. En este nuevo contexto de “copilotaje” de la inteligencia artificial en el devenir de la mente humana, podríamos

teorizar acerca de una interacción integral de estos mundos. En el primero, el “mundo 1” de Popper, los sistemas de inteligencia artificial procesan datos empíricos provenientes de sensores, bases de datos y experimentos. En el “mundo 2”, los humanos, en este caso los investigadores, utilizan este tipo de tecnologías para el modelado de procesos cognitivos, la asistencia en la formulación de hipótesis y el análisis de volúmenes de información astronómicos que, incluso con las capacidades computacionales previas a este momento histórico, serían imposibles de procesar; y de manera final, en el “mundo 3” se contribuye a la expansión del conocimiento con el descubrimiento de patrones no objetivados anteriormente, generando modelos predictivos o nuevas conclusiones que pueden ser validadas mediante el método científico. Hay que mencionar que, para Popper, un principio fundamental del ejercicio de las ciencias, desde un punto de vista epistémico, es la falsabilidad como concepto angular. Esto se podría resumir de la siguiente manera, una teoría es científica en cuanto pueda ser refutada mediante evidencia (Popper, 1972).

Reflexionar sobre estas bases filosóficas de la ciencia nos permite contextualizar la aplicación de los modelos de inteligencia artificial a la investigación, estructurando el análisis del papel de estas herramientas en la generación de nuevo conocimiento, la interacción que sus resultados tienen en el pensamiento humano y el impacto que tienen en el método científico vigente, que, entre otras cosas, hace posible la existencia de esta tecnología. Los algoritmos computacionales pueden procesar grandes volúmenes de información en tiempos inalcanzables en un ejercicio análogo a lo clásicamente realizado en ciencia. Esto es un potenciador de eficiencia en la identificación de correlaciones y patrones, áreas fundamentales en la investigación científica.

Así, asistimos a un cambio en el paradigma para formulación de hipótesis y la extrapolación de tendencias, lo que permite refutar o formular teorías con un nivel de precisión sin precedentes. Teóricamente, podríamos plantear que, la inteligencia artificial no solo actúa como herramienta de apoyo, sino que podría desempeñar un papel protagónico en la generación de nuevo conocimiento. Cabe la salvedad, hasta la fecha de redacción de este documento, no hay evidencia fehaciente que este tipo de tecnologías tenga intencionalidad propia (keyBPS, 2022), los resultados de estos algoritmos dependen de los modelos subyacentes diseñados por humanos.

Una analogía bastante válida para referirnos a este tipo de tecnologías es una caja negra, en la cual, extraemos unos resultados sin saber muy bien de dónde vienen o cuál fue el mecanismo para llegar a ellos. Esto se debe a que los propietarios de este tipo de tecnologías, en su gran mayoría, reservan el derecho comercial de sus algoritmos y, de manera subsecuente, el modelo matemático tiene importantes brechas o “gaps” donde este procesamiento se realiza a una escala logarítmica, perdiéndose el detalle, el cómo, de la obtención de estas conclusiones.

Por lo tanto, desde el punto de vista ético, estamos ante un importante reto, si bien la inteligencia artificial podría generar nuevo conocimiento, éste se da a partir de los datos con los que se entrena. En este sentido, los datos sesgados pueden llevar a conclusiones falaces y, según lo anterior, el principio de reproducibilidad no se puede garantizar de manera fehaciente cuando se integran modelos provenientes de terceros en la interpretación de estos datos.

Con un nivel de precisión bastante alto, podemos predecir que la implementación de los diferentes modelos de inteligencia artificial en la academia transformará de manera rotunda la manera como se imparten y generan nuevos conocimientos. En este sentido, con la popularización de los modelos de aprendizaje y el procesamiento de datos a escalas masivas, se tejen oportunidades nunca vistas para optimizar el acceso y la generación de conocimientos. Esta clase de tecnologías mejora la eficiencia investigativa y personaliza la enseñanza; hasta ahora, de manera teórica, puesto que, debido a la prematuridad de su existencia, los estudios científicos al respecto aún son incipientes. Con esto en mente, no se puede menospreciar el hecho que estas innovaciones acarreen cambios epistémicos, éticos y metodológicos de los cuales es menester generar un sentido de responsabilidad y reflexión.

Observatorio para la Implementación Tecnológica en Investigación

El observatorio para la implementación tecnológica en investigación se constituirá como una entidad para la supervisión y promoción del uso responsable de la Inteligencia artificial en la Fundación Universitaria Sanitas, ejerciendo funciones cardinales que incluyen:

-  Asesoramiento, brindando orientación sobre la integración adecuada de la Inteligencia Artificial en los diferentes procesos.
-  Monitoreo, supervisando el uso e implementación de la Inteligencia Artificial con el fin de asegurar el cumplimiento de estos lineamientos y detectar posibles desviaciones.
-  Desarrollo, promoviendo iniciativas que exploren nuevas aplicaciones de Inteligencia Artificial y evaluando su alineación con la dinámica institucional.
-  Formación y divulgación, organizando cursos, talleres y seminarios que fomenten el conocimiento y la apropiación crítica de estas herramientas por parte de la comunidad universitaria.

Implementación de Inteligencia Artificial en Instituciones

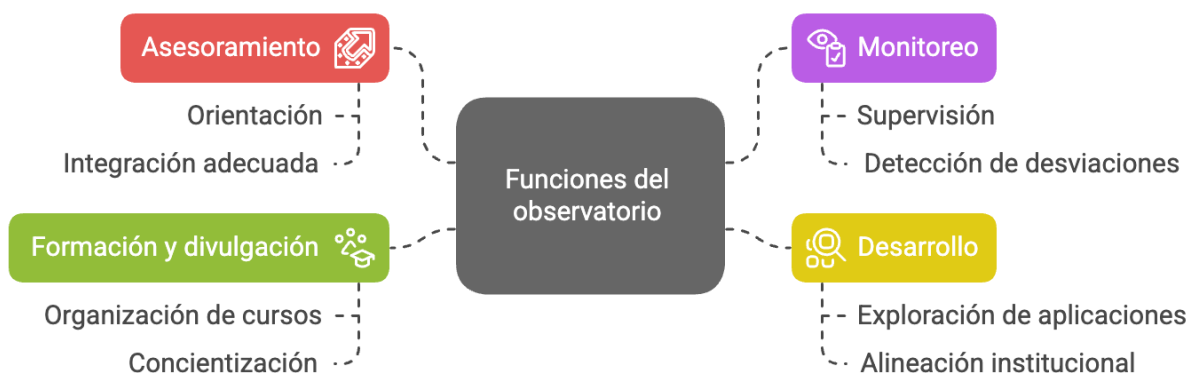


FIGURA 3. IMPLEMENTACIÓN DE INTELIGENCIA ARTIFICIAL EN INSTITUCIONES

Marco Ético y Normativo de Mínimos Aceptables en Investigación

Para garantizar el uso adecuado y responsable de la inteligencia artificial se dispone de un marco ético y normativo de mínimos aceptables.

-  Todo proyecto académico o investigativo que integre el uso de herramientas de inteligencia artificial en su construcción tendrá una sección que detalle su uso, especificando las herramientas empleadas, su propósito, justificación, el alcance de su contribución, ya sea significativa, íntegra o parcial, el manejo de los datos, mitigación de sesgos y supervisión o revisión humana. Complementario a esto, se exige que la adjudicación de los resultados obtenidos sea transparente, evitando la atribución a título propio los hallazgos que se hayan generado o coadyuvado por sistemas de Inteligencia artificial.
-  Todos los proyectos de investigación deben acogerse a los lineamientos establecidos por la Universidad y deben ser aprobados por el Comité de Ética en Investigación previo a su desarrollo.
-  Como parte del respeto a la propiedad intelectual, se establece que toda producción científica o académica sin distinción alguna deberá reconocer y citar el uso de estas herramientas, según los principios de responsabilidad académica e integridad científica. Se debe de garantizar, a criterio de cada docente o superior, el mecanismo para citar el uso de estas herramientas y la debida atribución de los aportes relacionados por la Inteligencia artificial.
-  En el marco del principio de Justicia, la dinámica académica debe garantizar que el uso de las diferentes herramientas de Inteligencia artificial en la investigación no favorezca a un grupo en detrimento de otro; por tanto, la organización en conjunto con El Observatorio para la Implementación Tecnológica en Investigación debe propender por la equidad en el acceso a

los diferentes aplicativos útiles según los medios de la Fundación Universitaria Sanitas para la comunidad.

-  En cuanto a la autonomía, se debe garantizar el respeto por la capacidad de cada investigador y participante en los proyectos para, según su criterio, decidir sobre el uso de las diferentes herramientas y, de esta manera, asegurar el aporte de información clara de riesgos, limitaciones y funcionamiento. Así las cosas, se fomentará una cultura de toma de decisiones informadas en investigación y que los involucrados ejerzan su autonomía sin sesgos o coacción derivados del uso de la Inteligencia artificial.
-  El principio de beneficencia se aplicará garantizando que la implementación de estos algoritmos tenga un impacto eminentemente positivo en el conocimiento y en la comunidad académica. La Inteligencia artificial debe de ser utilizada para mejorar la eficiencia y precisión de los procesos académicos, más no su detrimento o ejercicio malicioso, minimizando errores y optimizando los procesos, sin desplazar la supervisión humana ni comprometer la calidad e integridad académica.
-  Se implementarán medidas para minimizar cualquier daño potencial derivado del uso de Inteligencia artificial, especialmente en materia de sesgos algorítmicos. El respeto por la privacidad de datos y la prohibición de la manipulación individual de información generada por Inteligencia artificial se mantiene como un mínimo aceptable, anotando que, se subroga este lineamiento a la política de protección de datos del grupo Keralty.
-  La Fundación Universitaria Sanitas, se compromete a ofrecer programas formativos para los docentes, administrativos, estudiantes, investigadores sobre el uso ético y responsable de la Inteligencia artificial. Estas capacitaciones engranarán los aspectos técnicos, las consideraciones éticas y legales de la implementación de estos sistemas, área que está continuamente en discusión.
-  Se implementarán mecanismos de seguimiento para la evaluación del impacto de la inteligencia artificial en el ejercicio académico donde, el Observatorio para la Implementación Tecnológica en Investigación será el responsable de supervisar el cumplimiento de los lineamientos establecidos y asesorar la implementación de buenas prácticas.

EVALUACIÓN DE RIESGOS Y ACCIONES INSTITUCIONALES

RIESGOS IDENTIFICADOS	ACCIONES INSTITUCIONALES
Uso indebido de datos sensibles: El uso de IA puede implicar el procesamiento de información personal sin las medidas adecuadas de seguridad, exponiendo datos confidenciales de estudiantes, docentes e investigadores (Universidad Externado de Colombia, 2024).	Desarrollar políticas para la protección de datos sensibles: Implementar protocolos para garantizar que los datos personales sean resguardados y utilizados de manera responsable dentro de los marcos legales vigentes (Guío Español et al., 2021).
Generación de información errónea o sesgada: Los modelos de IA pueden producir resultados que refuercen prejuicios o difundan información incorrecta debido a	Implementar capacitaciones periódicas para el uso ético de la IA: Los estudiantes y docentes deben recibir formación continua sobre los riesgos y buenas prácticas en el uso

sesgos en los datos de entrenamiento (UNESCO, 2021; Universidad San Juan de la Cruz, 2024).	de IA en el ámbito educativo y administrativo (UNESCO, 2021; Universidad Externado de Colombia, 2024).
Dependencia excesiva de herramientas de IA: El uso inadecuado de IA en procesos educativos puede afectar el desarrollo del pensamiento crítico y la creatividad, generando una sobre dependencia en la tecnología (Política Para El Uso De Tecnologías Digitales E Inteligencia Artificial En La Universidad Del Pacífico, 2024; Universidad Externado de Colombia, 2024).	Realizar auditorías regulares sobre el uso de IA en la institución: Evaluar el impacto de las herramientas de IA en los procesos educativos, administrativos y de investigación para detectar posibles riesgos y corregir desviaciones (Política para el uso de Tecnologías Digitales e Inteligencia Artificial en la Universidad del Pacífico, 2024).
Impacto ambiental: El entrenamiento y uso de modelos de IA requieren un alto consumo energético, lo que contribuye a la huella de carbono y afecta la sostenibilidad (Consejo Nacional de Política Económica y Social, 2025).	Promover una IA sostenible: Establecer estrategias para reducir el impacto ambiental del uso de IA en la institución, priorizando tecnologías energéticamente eficientes y promoviendo el uso responsable de los recursos (Consejo Nacional de Política Económica y Social, 2025).
Riesgos en la propiedad intelectual: La generación de contenido sintético mediante IA plantea desafíos en la autoría y el reconocimiento de fuentes originales, afectando los derechos de propiedad intelectual y la integridad académica (United States Copyright Office, 2025).	Regular el uso de IA en la producción de contenido académico: Establecer normas que definan cómo deben citarse los contenidos generados por IA y garantizar que no se vulneren los derechos de propiedad intelectual (United States Copyright Office, 2025).
Desigualdad en el acceso a la IA: La falta de infraestructura y recursos puede limitar el acceso equitativo a las herramientas de IA, ampliando las brechas tecnológicas y educativas (Consejo Nacional de Política Económica y Social, 2025).	Garantizar acceso equitativo a la IA: Implementar iniciativas para reducir la brecha digital y asegurar que toda la comunidad académica tenga acceso a herramientas de IA, sin exclusión por factores socioeconómicos o geográficos (Consejo Nacional de Política Económica y Social, 2025; Universidad Externado de Colombia, 2024).
Falta de supervisión y evaluación adecuada del uso de IA.	Estos lineamientos podrán ser revisados periódicamente, por la supervisión y evaluación de IA en Unisanitas será realizada por el Consejo Académico y/o por los comités que se definan a nivel institucional para tal fin, para monitorear el uso de IA en la institución, garantizar su alineación con principios éticos

	y evaluar su impacto en la educación e investigación.
--	---

COMO ABORDAR UN CASO DE USO INDEBIDO DE IA

Se reconoce que no existe una herramienta de detección de uso indebido de la IA que sea 100% confiable. Las soluciones actuales pueden identificar indicios de generación artificial, pero no ofrecen una prueba concluyente. Por ello, es fundamental complementar la detección con un análisis crítico del trabajo del estudiante. Algunas estrategias incluyen:

-  **Comparación con trabajos previos:** Evaluar si el estilo, la profundidad del análisis y la coherencia del texto son consistentes con el desempeño habitual del estudiante.
-  **Patrones de lenguaje atípicos:** Identificar frases excesivamente formales o estructuradas de manera poco común para el estudiante.
-  **Errores inusuales:** La IA puede generar información errónea o inventada, por lo que se debe contrastar el contenido con fuentes verificables.
-  **Dificultad del estudiante para explicar su trabajo:** Si un estudiante no puede justificar o desarrollar las ideas que presenta en su escrito, puede ser una señal de uso indebido de IA.

Ante la sospecha de un uso indebido de herramientas de inteligencia artificial, se recomienda seguir los siguientes pasos:

1. **Notificación al estudiante:** Informar al estudiante sobre la situación y brindarle la oportunidad de presentar su perspectiva.
2. **Evaluación del caso:** Analizar la situación de manera individual y determinar la gravedad de la falta:
 - **Leve:** Se emitirá una advertencia y se permitirá la realización del trabajo nuevamente con supervisión.
 - **Moderada:** Se aplicará una reducción en la calificación del trabajo.
 - **Grave:** En casos de uso excesivo de IA generativa en trabajos o proyectos de investigación, el caso deberá ser evaluado en profundidad con la Unidad Académica o el área correspondiente para determinar las acciones pertinentes.
3. **Formación y orientación:** Brindar al estudiante información sobre el uso adecuado de la IA, enfatizando las normas éticas y académicas establecidas por la institución.

REGLAS PARA EL USO DE HERRAMIENTAS DE IA

1. **Uso por parte de los estudiantes:** Se permite el uso de herramientas de IA en la **Fase II del Aprendizaje Basado en Problemas (ABP): Abordaje y estudio del problema**, siempre que su aplicación sea responsable y alineada con los objetivos de aprendizaje; además de declarada.

2. **Regulación por los docentes:** Los profesores tienen la facultad de establecer normas sobre el uso de IA por parte de los estudiantes, determinando en qué contextos está permitido o restringido su uso.
3. **Uso en la enseñanza:** Los docentes pueden emplear herramientas de IA para planear, diseñar, desarrollar y evaluar procesos educativos, asegurando su correcta integración en el aula.
4. **Evaluación previa:** Antes de incorporar herramientas de IA en los cursos, los profesores deberán analizar el riesgo de desinformación, así como las implicaciones éticas y de derechos fundamentales. Se recomienda verificar si existen evaluaciones de impacto ético o normativo relacionadas con su implementación.
5. **Formación en buenas prácticas:** Cuando se requiera el uso de IA en actividades académicas, los docentes deberán orientar a los estudiantes en su uso adecuado, enfatizando prácticas responsables, la relación entre medios y fines, consideraciones éticas, integridad académica y los riesgos asociados.

EVALUACIÓN Y SEGUIMIENTO

Estos lineamientos podrán ser revisados periódicamente, por parte del Consejo Académico de Unisanitas y/o por los Comités que se definan a nivel institucional para tal fin, quienes evaluarán su impacto y propondrán actualizaciones, según las tendencias tecnológicas y las necesidades de la institución.

AGRADECIMIENTOS

El presente documento fue elaborado por un equipo interdisciplinario, quienes, desde su conocimiento y experiencia, aportaron para generar un marco normativo integral, alineado con los principios éticos, académicos y administrativos de la Fundación Universitaria Sanitas Unisanitas:

-  Joanna A. Acevedo T. Vicerrectora.
-  Johns Navarro L. Secretario General.
-  Jorge Martínez Bernal. Decano Facultad de Enfermería.
-  Rolando Salazar. Decano Facultad de Psicología, Ciencias Sociales y de la Educación.
-  Andrés Bula Calderón. Director Pregrado Medicina.
-  Johana Bolaños Macías. Directora Postgrados Facultad de Medicina.
-  María del Pilar García. Directora Bienestar Universitario.
-  Oscar Castro M. Coordinador Postgrados Facultad de Psicología, Ciencias Sociales y de la Educación.
-  Andrés Felipe González. Instructor Asistente.
-  Ingrid Milena Rodríguez Bedoya. Secretaria Ejecutiva del Comité de Ética en Investigación.
-  Jossie Murcia. Joven investigador del Instituto para la Transformación y el Desarrollo Sanitario.
-  María José Amaya. Unidad de Innovación y Desarrollo.
-  Valentina Ramón. Unidad de Innovación y Desarrollo.

BIBLIOGRAFÍA

-  Aristotle. (IV a.C.). The Organon. Cambridge, Harvard University Press. <http://archive.org/details/organoncooke01arisuoft>.
-  Bellman, R. (1978). An Introduction to Artificial Intelligence: Can Computers Think? Boyd & Fraser Publishing Company.
-  Boole, G. (1854). An Investigation of the Laws of Thought: On which are Founded the Mathematical Theories of Logic and Probabilities. Walton and Maberly.
-  Columbia University. (2024). Generative AI Policy. <https://provost.columbia.edu/content/office-senior-vice-provost/ai-policy>.
-  Consejo Nacional De Política Económica Y Social. (2019, noviembre 8). Documento CONPES 3975: Política nacional para la transformación digital e inteligencia artificial.
-  Consejo Nacional De Política Económica Y Social. (2025, febrero 14). Documento CONPES 4144: Política Nacional De Inteligencia Artificial.
-  Davis, K., Christodoulou, J., Seider, S., & Gardner, H. (2011). The theory of multiple intelligences. En The Cambridge handbook of intelligence (pp. 485–503). Cambridge University Press. <https://doi.org/10.1017/CBO9780511977244.025>
-  Dawson, J. F., Marvin, A. C., & Marshman, C. A. (2003). Electromagnetic Compatibility. En R. A. Meyers (Ed.), Encyclopedia of Physical Science and Technology (Third Edition) (pp. 261–275). Academic Press. <https://doi.org/10.1016/B0-12-227410-5/00210-6>
-  Descartes, R. (with Internet Archive). (1641). Meditaciones metafísicas. [Place of publication not identified]: [Createspace Independent Publishers]. http://archive.org/details/meditacionesmeta0000desc_l5m9
-  Equipo proyecto IA-Uniandes. (2024, octubre). Lineamientos para el uso de inteligencia artificial generativa (IAG) en la Universidad de los Andes.
-  Franganillo, J. (2022). Contenido generado por inteligencia artificial: Oportunidades y amenazas. Anuario ThinkEPI, 16. <https://doi.org/10.3145/thinkepi.2022.e16a24>
-  Franganillo, J. (2023). La inteligencia artificial generativa y su impacto en la creación de contenidos mediáticos. methaodos. Revista de ciencias sociales, 11(2), Article 2. <https://doi.org/10.17502/mrcs.v11i2.710>
-  Ghallab, M., Howe, A., Knoblock, C., McDermott, D., Ram, A., Veloso, M., Weld, D., & Wilkins, D. (1998, octubre). The Planning Domain Definition Language.
-  Gödel, K. (1931). Über formal unentscheidbare Sätze der Principia mathematica und verwandter Systeme. <https://www.semanticscholar.org/paper/Kurt-G%C3%B6del-%2C-%E2%80%98-%C3%9Cber-formal-unentscheidbare-S%C3%A4tze-I-Zach/2b10b241d9bae986bdd403f25c24bc33118bb63e>
-  Gottfredson, L. S. (1997). Mainstream science on intelligence: An editorial with 52 signatories, history and bibliography. Intelligence, 24(1), 13–23. [https://doi.org/10.1016/S0160-2896\(97\)90011-8](https://doi.org/10.1016/S0160-2896(97)90011-8)
-  Guío Español, A., Tamayo Uribe, E., & Gómez Ayerbe, P. (2021, mayo). Marco Ético Para La Inteligencia Artificial En Colombia.

-  Hehl, W. (2016). A General World Model with Poiesis: Poppers Three Worlds updated with Software (arXiv:1604.00360). arXiv. <https://doi.org/10.48550/arXiv.1604.00360>
-  Kant, I. (1785). Fundamentación para una metafísica de las costumbres.
-  keyBPS, F.-. (2022, noviembre 14). Por qué la inteligencia artificial no puede querer nada. key Business Process Solutions. <https://www.keybps.com/por-que-la-inteligencia-artificial-no-puedo-querer-nada>
-  Leibniz, G. W. (1666). Dissertation on the Art of Combinations. En G. W. Leibniz & L. E. Loemker (Eds.), *Philosophical Papers and Letters* (pp. 73–84). Springer Netherlands. https://doi.org/10.1007/978-94-010-1426-7_2
-  Ley estatutaria 1581 (2012). <https://www.funcionpublica.gov.co/eva/gestornormativo/norma.php?i=49981>
-  McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (1955). A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, August 31, 1955. *AI Magazine*, 27(4), Article 4. <https://doi.org/10.1609/aimag.v27i4.1904>
-  Miao, F. (2021). Inteligencia artificial y educación: Guía para las personas a cargo de formular políticas—UNESCO Digital Library. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000379376>
-  Miao, F. (2024). Guía para el uso de IA generativa en educación e investigación. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000389227>
-  Ministerio de Ciencia, Tecnología e Innovación. (2024, febrero). Hoja de ruta para el desarrollo y aplicación de la Inteligencia Artificial en Colombia.
-  Minsky, M. L. (1954). Theory of neural-analog reinforcement systems and its application to the brain-model problem [Thesis, Princeton University]. <https://ebooks.ub.rug.nl/337/>
-  Mitchell, T. M. (1997). *Machine Learning*. McGraw-Hill.
-  Mool, A., Schmid, J., Johnston, T., Thomas, W., Fenner, E., Lu, K., Gandhi, R., Western, A., Seabold, B., Smith, K., Patterson, Z., Feldt, H., Vollmer, D., Nallaveettil, R., Fanelli, A., Schmillen, L., Tischkau, S., Cianciolo, A. T., Benedict, P., & Selinfreund, R. (2024). Using Generative AI to Simulate Patient History-Taking in a Problem-Based Learning Tutorial: A Mixed-Methods Study (p. 2024.05.02.24306753). *medRxiv*. <https://doi.org/10.1101/2024.05.02.24306753>
-  Moustafa, N., Adi, E., Turnbull, B., & Hu, J. (2018). A New Threat Intelligence Scheme for Safeguarding Industry 4.0 Systems. *IEEE Access*, 6, 32910–32924. <https://doi.org/10.1109/ACCESS.2018.2844794>
-  Nicolas, S., Andrieu, B., Croizet, J.-C., Sanitioso, R. B., & Burman, J. T. (2013). Sick? Or slow? On the origins of intelligence as a psychological object. *Intelligence*, 41(5), 699–711. <https://doi.org/10.1016/j.intell.2013.08.006>
-  Ortega-Díaz, C., & Hernández-Pérez, A. (2015). Hacia El Aprendizaje Profundo En La Reflexión De La Práctica Docente. <https://www.redalyc.org/articulo.oa?id=46142596015>
-  Popper, K. R. (1972). *Objective Knowledge: An Evolutionary Approach*. Oxford University Press.
-  Santrock, J. W. (2006). *PSICOLOGIA DE LA EDUCACION*. McGraw-Hill Interamericana de España S.L.
-  Sapci, A. H., & Sapci, H. A. (2020). Artificial Intelligence Education and Tools for Medical and Health Informatics Students: Systematic Review. *JMIR Medical Education*, 6(1), e19285. <https://doi.org/10.2196/19285>

-  Simpson, J., & Weiner, E. (1989). *The Oxford English Dictionary: 20 Volume Set*. Oxford University Press.
-  Sternberg, R. J. (2019). *Human Intelligence: An Introduction*. Cambridge University Press.
-  Triyason, T., Valaisathien, S., Vanijja, V., Kanthamanon, P., & Chan, J. H. (2015). Chapter 5—VoIP Quality Prediction Model by Bio-Inspired Methods. En X.-S. Yang, S. F. Chien, & T. O. Ting (Eds.), *Bio-Inspired Computation in Telecommunications* (pp. 95–116). Morgan Kaufmann. <https://doi.org/10.1016/B978-0-12-801538-4.00005-7>
-  Trullàs, J. C., Blay, C., Sarri, E., & Pujol, R. (2022). Effectiveness of problem-based learning methodology in undergraduate medical education: A scoping review. *BMC Medical Education*, 22(1), 104. <https://doi.org/10.1186/s12909-022-03154-8>
-  Turing, A. (1950). *Computing Machinery and Intelligence* (1950). En B. J. Copeland (Ed.), *The Essential Turing* (p. 0). Oxford University Press. <https://doi.org/10.1093/oso/9780198250791.003.0017>
-  Turing, A. M. (1937). On Computable Numbers, with an Application to the Entscheidungsproblem. *Proceedings of the London Mathematical Society*, s2-42(1), 230–265. <https://doi.org/10.1112/plms/s2-42.1.230>
-  UNESCO. (2021, noviembre 23). Recomendación sobre la ética de la inteligencia artificial. https://unesdoc.unesco.org/ark:/48223/pf0000381137_spa
-  United States Copyright Office. (2025, enero). Copyright and Artificial Intelligence. Part 2: Copyrightability. <https://www.copyright.gov/ai/>
-  Universidad Externado de Colombia. (2024, septiembre). Lineamientos para el uso de IA.
-  Wiener, N. (1961). *Cybernetics or Control and Communication in the Animal and the Machine*. MIT Press. <https://mitpress.mit.edu/9780262537841/cybernetics-or-control-and-communication-in-the-animal-and-the-machine/>
-  Yew, E. H. J., & Goh, K. (2016). Problem-Based Learning: An Overview of its Process and Impact on Learning. *Health Professions Education*, 2(2), 75–79. <https://doi.org/10.1016/j.hpe.2016.01.004>



UNISANITAS
WWW.UNISANITAS.EDU.CO
2025