

UNIVERSIDAD PERUANA DE CIENCIAS E INFORMÁTICA
FACULTAD DE DERECHO Y CIENCIAS POLÍTICAS
CARRERA PROFESIONAL DE DERECHO



TRABAJO DE SUFICIENCIA PROFESIONAL

“Retos legales frente a la inteligencia artificial: hacia una regulación
ética y sostenible”

AUTORA:

Bach. Ingar Silvestre, Cinthya Isabel

PARA OPTAR EL TÍTULO PROFESIONAL DE:
ABOGADO

ASESOR:

Abg. Castillo Figueroa, Carlos Francisco

ORCID: 0009-0009-6953-651X

DNI: 09530087

LIMA-PERÚ

2026

INFORME DE SIMILITUD

Nº052-2026-UPCI-FDCP-REHO-T

A : **MG. HERMOZA OCHANTE RUBÉN EDGAR**
Decano (e) de la Facultad de Derecho y Ciencias Políticas

DE : **MG. HERMOZA OCHANTE, RUBEN EDGAR**
Docente Operador del Programa Turnitin

ASUNTO : Informe de evaluación de Similitud de Trabajo de Suficiencia Profesional:
BACHILLER INGAR SILVESTRE, CINTHYA ISABEL

FECHA : Lima, 3 de agosto de 2026.

Tengo el agrado de dirigirme a usted con la finalidad de informar lo siguiente:

1. Mediante el uso del programa informático **Turnitin** (con las configuraciones de excluir citas, excluir bibliografía y excluir oraciones con cadenas menores a 20 palabras) se ha analizado el Trabajo de Suficiencia Profesional titulada: **"RETOS LEGALES FRENTE A LA INTELIGENCIA ARTIFICIAL: HACIA UNA REGULACIÓN ÉTICA Y SOSTENIBLE"**, presentado por la Bachiller **INGAR SILVESTRE, CINTHYA ISABEL**.
2. Los resultados de la evaluación concluyen que el Trabajo de Suficiencia Profesional en mención tiene un **ÍNDICE DE SIMILITUD DE 18%** (cumpliendo con el artículo 35 del Reglamento de Grado de Bachiller y Título Profesional UPCI aprobado con Resolución Nº 373-2019-UPCI-R de fecha 22/08/2019).
3. Al término análisis, la Bachiller en mención **PUEDE CONTINUAR** su trámite ante la facultad, por lo que el resultado del análisis se adjunta para los efectos consiguientes

Es cuanto hago de conocimiento para los fines que se sirva determinar.

Atentamente,


.....
MG. HERMOZA OCHANTE, RUBEN EDGAR
Universidad Peruana de Ciencias e Informática
Docente Operador del Programa Turnitin

Adjunto:

**Resultado de similitud*

Dedicatoria

El trabajo de suficiencia profesional que presento, se lo dedico con todo cariño a mis queridos hijos por darme fuerzas y sobre todo motivarme a seguir y culminar mi carrera profesional.

.....

Agradecimiento

Quiero presentar mi más sincero agradecimiento a las autoridades de mi casa de estudio por haberme formado profesionalmente, a mis maestros por su experiencia y conocimiento compartido y a las personas que colaboraron conmigo a lo largo de estos años de formación profesional.

.....

Declaración de Autoría

Nombres : Ingar Silvestre

Apellidos : Cinthya Isabel

Código : 1401000465

DNI : 42999364

Declaro que, soy el autor del trabajo realizado y que es la versión final que he entregado a la oficina del Decanato de la Facultad de Derecho y Ciencias Políticas de la Universidad Peruana de Ciencias e Informática.

Asimismo, declaro que he citado debidamente las palabras o ideas de otros autores, refiriendo expresamente el nombre de la obra y página o páginas que me sirvieron de fuente.

Jesús María, julio del 2026.

ÍNDICE

CARATULA.....	1
INFORME DE SIMILITUD.....	2
DEDICATORIA.....	3
AGRADECIMIENTO.....	4
DECLARACIÓN DE AUTORÍA.....	5
ÍNDICE.....	6
INTRODUCCIÓN.....	7
CAPITULO I: Planificación del Trabajo de Suficiencia Profesional	10
1.1. Título y descripción del trabajo	10
1.2. Objetivo del trabajo	10
1.3. Justificación	13
CAPITULO II: Marco Teórico.....	17
2.1 Historia y evolución de la inteligencia artificial.....	17
2.2. Relación entre la IA y el derecho.....	21
CAPITULO III: Desarrollo de actividades programadas.....	26
3.1 Vacíos legales y desafíos regulatorios frente a la IA	26
3.2. Principios éticos aplicables a la IA.....	33
CAPITULO IV: Resultados Obtenidos.....	39
CONCLUSIONES	46
RECOMENDACIONES	51
REFERENCIAS BIBLIOGRÁFICAS.....	57
ANEXOS.....	72
Anexo 1: Evidencia de similitud digital.....	72
Anexo 2:Autorización de publicación en repositorio.....	78

INTRODUCCION

La inteligencia artificial (IA) está transformando radicalmente la sociedad, planteando desafíos legales inéditos que exigen marcos regulatorios éticos y sostenibles, esta introducción explora el impacto social de la IA, los principales retos jurídicos, las respuestas regulatorias internacionales y la urgencia de una gobernanza ética, fundamentando el análisis en literatura académica reciente y documentos institucionales clave.

La irrupción de la inteligencia artificial en la vida contemporánea representa uno de los fenómenos tecnológicos más disruptivos del siglo XXI, con profundas implicancias para el derecho, la ética y la sostenibilidad social (Crawford, 2021), la IA, en sus múltiples aplicaciones desde la automatización de servicios públicos hasta la toma de decisiones en sectores críticos como la salud, la justicia y las finanzas, redefine no solo los límites de la innovación, sino también los marcos normativos que buscan proteger los derechos fundamentales y el interés público (Caló, 2018; Hildebrandt, 2020).

El avance acelerado de la IA ha evidenciado retos legales complejos y multidimensionales, entre los más apremiantes se encuentran la responsabilidad civil por daños causados por sistemas autónomos, la protección de datos personales en contextos de procesamiento masivo y opaco, y la discriminación algorítmica que amenaza la igualdad y los derechos humanos (Lagioia, Rovatti & Sartor, 2025; Sartor, 2020; Zuiderveen Borgesius, 2020), la opacidad inherente a muchos algoritmos dificulta la rendición de cuentas y la transparencia, mientras que la automatización de decisiones

puede perpetuar o amplificar sesgos estructurales, afectando de manera desproporcionada a grupos vulnerables (Eubanks, 2018; Hartzog, 2018).

Frente a estos desafíos, la comunidad internacional ha impulsado respuestas regulatorias de alcance global y regional; organismos como la UNESCO (2021) y la OCDE (2019, 2024) han establecido principios rectores para el desarrollo y uso responsable de la IA, enfatizando la necesidad de salvar los derechos humanos, la equidad y la sostenibilidad; en el ámbito europeo, la reciente aprobación del Reglamento de Inteligencia Artificial (AI Act) de la Unión Europea (2024) y la Convención Marco del Consejo de Europa sobre IA y derechos fundamentales (2024) marcan hitos en la construcción de un marco jurídico robusto y adaptativo; asimismo, la Asamblea General de las Naciones Unidas (2024) ha subrayado la importancia de una IA segura, confiable y orientada al desarrollo sostenible.

No obstante, la regulación efectiva de la IA requiere no solo normas técnicas, sino también una gobernanza ética que integre principios de transparencia, justicia, privacidad y participación social (High-Level Expert Group on AI, 2019; Wang & Lee, 2025; Binns et al., 2025), la literatura reciente destaca la urgencia de diseñar marcos regulatorios flexibles y colaborativos, capaces de equilibrar la innovación con la protección de los valores democráticos y el bienestar colectivo (Pasquale, 2020).

En este contexto, la presente tesis se propone analizar los principales retos legales que plantean la inteligencia artificial, evaluando las respuestas regulatorias existentes y proponiendo lineamientos para una regulación ética y sostenible, el trabajo se estructura en cuatro capítulos: el primero aborda el

impacto social y jurídico de la IA; el segundo examina los desafíos legales específicos; el tercero analiza los marcos regulatorios internacionales y regionales; y el cuarto plantea recomendaciones para una gobernanza ética y sostenible de la IA.

CAPITULO I.- Planificación del Trabajo de Suficiencia Profesional

1.1. Título y descripción del trabajo

Título del Trabajo

La presente investigación la he denominado: Retos legales frente a la inteligencia artificial: Hacia una regulación ética y sostenible.

1.2. Objetivo del presente trabajo

La inteligencia artificial (IA) ha surgido como una tecnología transformadora que redefine las dinámicas sociales, económicas y jurídicas a nivel global, sin embargo, su desarrollo y aplicación plantean desafíos legales y éticos que requieren una atención urgente y estructurada; en este contexto, la presente tesis tiene como propósito principal analizar los retos legales que surgen frente a la implementación de la IA, con el objetivo de proponer un marco regulatorio que sea ético, sostenible y adaptable a los rápidos avances tecnológicos.

Objetivo General

El objetivo general de esta investigación es identificar, analizar y evaluar los principales desafíos legales que plantea la inteligencia artificial, con el fin de desarrollar propuestas regulatorias que promuevan un equilibrio entre la innovación tecnológica, la protección de los derechos fundamentales y la sostenibilidad social, este enfoque busca contribuir a la construcción de un marco normativo que no solo regule los riesgos asociados a la IA, sino que también fomente su desarrollo responsable y ético.

Objetivos Específicos

1. Analizar el impacto jurídico de la inteligencia artificial en los derechos fundamentales y las libertades individuales.

Este objetivo se centra en examinar cómo las aplicaciones de la IA afectan derechos como la privacidad, la igualdad y la no discriminación, considerando casos concretos y ejemplos recientes de sesgos algorítmicos y violaciones de datos personales.

2. Estudiar los marcos regulatorios existentes a nivel internacional y regional.

Se busca evaluar las normativas actuales, como el Reglamento de Inteligencia Artificial de la Unión Europea (AI Act), que adopta un enfoque basado en el riesgo para armonizar las normas sobre IA; este análisis permitirá identificar fortalezas, debilidades y vacíos legales en las regulaciones vigentes.

3. Explorar los principios éticos aplicables al desarrollo y uso de la inteligencia artificial.

Este objetivo aborda la necesidad de integrar valores éticos como la transparencia, la justicia y la sostenibilidad en los sistemas de IA, siguiendo recomendaciones de organismos internacionales como la UNESCO y la OCDE.

4. Proponer lineamientos para una regulación ética y sostenible de la inteligencia artificial.

A partir del análisis previo, se desarrollarán recomendaciones que promuevan una gobernanza inclusiva y colaborativa, capaz de responder a los desafíos legales y éticos de la IA, mientras se fomenta su uso responsable en sectores clave como la salud, la justicia y la economía.

5. Reflexionar sobre el papel del derecho en la era de la inteligencia artificial.

Finalmente, se busca contribuir al debate sobre cómo el derecho puede adaptarse a los cambios tecnológicos, pasando de un enfoque reactivo a uno proactivo, que anticipe los riesgos y oportunidades de la IA.

Justificación de los objetivos

La elección de estos objetivos responde a la necesidad de abordar los desafíos legales de la IA desde una perspectiva integral, que combina el análisis técnico-jurídico con una reflexión ética y sostenible, en un contexto donde la IA está transformando sectores clave de la sociedad, como la seguridad, la salud y la economía, resulta imprescindible desarrollar marcos regulatorios que no solo mitiguen los riesgos, sino que

también potencian los beneficios de esta tecnología para el bienestar colectivo.

1.3. Justificación

La regulación ética y sostenible de la inteligencia artificial (IA) es un imperativo multidimensional, las justificaciones para abordar los retos legales de la AI abarcan desde problemas prácticos concretos y profundas implicaciones sociales, hasta desafíos normativos, metodológicos y epistemológicos que exigen enfoques interdisciplinarios y críticos.

Justificación práctica

La justificación práctica de esta tesis radica en los problemas reales y tangibles que la IA plantea en la sociedad contemporánea, la proliferación de sistemas algorítmicos ha evidenciado sesgos discriminatorios en áreas como la justicia penal, la contratación laboral y el acceso a servicios sociales, perpetuando desigualdades existentes y generando nuevos riesgos para grupos vulnerables (Eubanks, 2018; Crawford & Calo, 2016).

Además, la opacidad de los algoritmos (“cajas negras”) dificulta la rendición de cuentas y la reparación de daños, mientras que la autonomía de sistemas como vehículos autónomos o armas inteligentes crea vacíos de responsabilidad legal (Pasquale, 2015; Hildebrandt, 2016), la recolección masiva de datos biométricos y personales, muchas veces sin consentimiento informado, desafía los

marcos legales actuales de privacidad y protección de datos (Crawford, 2021; Zuboff, 2019).

Estos desafíos evidencian la insuficiencia de las normativas vigentes y la necesidad urgente de una regulación específica y adaptativa.

Justificación social

Desde una perspectiva social, la IA tiene el potencial de profundizar desigualdades estructurales y afectar derechos fundamentales; investigaciones han demostrado que los sistemas de IA pueden reproducir y amplificar sesgos raciales, de género y de clase, afectando de manera desproporcionada a comunidades marginadas (Noble, 2018; Buolamwini & Gebru, 2018; Benjamin, 2019); además, la automatización y la gestión algorítmica están transformando el mercado laboral, generando desplazamiento de empleos y nuevas formas de vigilancia y exclusión (UNESCO, 2021; OCDE, 2025).

La expansión de las tecnologías de vigilancia, como el reconocimiento facial, amenaza la privacidad, la libertad de expresión y la participación democrática (Zuboff, 2019; UNESCO, 2021), por ello, una regulación ética y sostenible debe priorizar la equidad, la inclusión y la protección de los derechos humanos frente a los riesgos sociales emergentes.

Justificación legal

En el ámbito legal, la rápida evolución de la IA ha superado la capacidad de los marcos normativos tradicionales, generando

fragmentación y vacíos regulatorios, el Reglamento de IA de la Unión Europea (EU AI Act, 2024) representa el primer intento integral de regulación, pero ha sido criticado por su rigidez y por no abordar plenamente los riesgos de opacidad y discriminación (Edwards & Wachter, 2022; Novelli et al., 2024).

Instrumentos como el GDPR han establecido estándares globales en protección de datos, pero resultan insuficientes ante los desafíos específicos de la IA, como la toma de decisiones automatizada y el perfilado (Yadav, 2025).

A nivel internacional, el Convenio Marco del Consejo de Europa sobre IA (2024) busca llenar el vacío legal en derechos humanos y democracia, aunque enfrenta limitaciones en su aplicación y mecanismos de cumplimiento (Consejo de Europa, 2024); la diversidad de enfoques nacionales y la falta de estándares globales refuerzan la necesidad de una regulación armonizada, ética y sostenible.

Justificación metodológica

El estudio de los retos legales de la IA exige una aproximación metodológica plural e interdisciplinaria, el análisis doctrinal sigue siendo fundamental para interpretar normas y jurisprudencia, pero resulta insuficiente para captar la complejidad social y tecnológica de la IA (Hutchinson, 2015), los métodos comparados permiten identificar buenas prácticas y adaptar soluciones entre jurisdicciones (Van Hoecke, 2011).

La investigación socio-jurídica y los estudios empíricos son esenciales para evaluar el impacto real de las regulaciones y comprender cómo operan en contextos concretos (Blackham, 2025). Así, la combinación de enfoques doctrinales, comparativos, socio-legales y empíricos enriquece la comprensión y la formulación de propuestas regulatorias efectivas.

Justificación epistemológica

Finalmente, la justificación epistemológica se fundamenta en la necesidad de repensar cómo se construye el conocimiento en la intersección entre derecho y tecnología, desde los estudios de ciencia, tecnología y sociedad (STS), se reconoce que el conocimiento legal sobre la IA es socialmente construido y coproducido junto con los desarrollos tecnológicos (Jasanoff, 2016, 2025).

Enfoques críticos, como los de Lessig (1999, 2003) y Brownsword (2016, 2022), subrayan que la regulación de la IA no puede limitarse a normas jurídicas, sino que debe considerar la influencia de normas sociales, incentivos de mercado y arquitecturas tecnológicas; además, los desafíos epistemológicos que plantean los modelos de lenguaje y sistemas de IA como la falta de comprensión semántica y la dificultad de verificación intersubjetiva exigen nuevas formas de validación y responsabilidad en la producción de conocimiento jurídico (Cambridge Forum on AI, 2025).

Por tanto, una aproximación epistemológica plural y crítica es indispensable para una regulación ética y sostenible de la IA.

CAPITULO II.- Marco Teórico

2.1. Historia y evolución de la inteligencia artificial. -

La inteligencia artificial (IA) ha transitado desde sus raíces filosóficas y matemáticas hasta convertirse en un motor de transformación social, económica y legal, su evolución ha planteado retos inéditos para el derecho y la ética, impulsando respuestas regulatorias internacionales y un debate académico sobre la gobernanza responsable de la IA.

1. Precursores filosóficos y matemáticos

El origen de la inteligencia artificial se remonta a los trabajos pioneros de Alan Turing, quien en 1950 propuso la célebre pregunta “¿pueden pensar las máquinas?” y formuló el Test de Turing como criterio para evaluar la inteligencia de los sistemas artificiales (Turing, 1950), previamente, McCulloch y Pitts (1943) habían desarrollado un modelo matemático de redes neuronales artificiales, sentando las bases para el enfoque conexionista de la IA (McCulloch & Pitts, 1943).

2. Nacimiento formal de la IA: La Conferencia de Dartmouth

La IA se constituyó como disciplina formal en 1956, durante la Conferencia de Dartmouth, donde John McCarthy, Marvin Minsky, Nathaniel Rochester y Claude Shannon acuñaron el término “inteligencia artificial” y delinearon su objetivo: crear máquinas capaces de simular la inteligencia humana (McCarthy et al., 1955).

3. IA simbólica y el primer invierno de la IA

Durante las décadas de 1950 y 1960, la IA se centró en el procesamiento simbólico y la lógica formal, con desarrollos como el Logic Theorist (Newell & Simon, 1956), sin embargo, el exceso de optimismo y las limitaciones técnicas llevadas a cabo al primer “invierno de la IA” en los años 1970, caracterizado por recortes de financiación y escepticismo sobre el potencial real de la disciplina (Crevier, 1993).

4. Sistemas expertos y el segundo invierno

El auge de los expertos en sistemas en los años 1970 y 1980, como DENDRAL y MYCIN, permitió aplicar la IA a problemas específicos mediante reglas y bases de conocimiento (Feigenbaum et al., 1971), no obstante, la rigidez y el alto costo de mantenimiento de estos sistemas provocaron el segundo “invierno de la IA” a finales de los 80, marcado por el colapso del mercado de hardware especializado y nuevas dudas sobre la viabilidad de la IA (Nilsson, 2010).

5. Renacimiento de las redes neuronales y la revolución del aprendizaje profundo

El resurgimiento de las redes neuronales a partir de 2006, impulsado por Geoffrey Hinton, Yann LeCun y Yoshua Bengio, marcó el inicio de la

revolución del aprendizaje profundo. Hinton et al. (2006) introdujeron algoritmos de preentrenamiento para redes profundas, mientras que LeCun y colaboradores perfeccionaron las redes convolucionales para el reconocimiento de imágenes (LeCun et al., 1998; Hinton et al., 2006).

6. Hitos del aprendizaje automático y el momento ImageNet

Paralelamente, se consolidaron técnicas como las máquinas de soporte vectorial (Cortés & Vapnik, 1995) y los bosques aleatorios (Breiman, 2001); el “momento ImageNet” llegó en 2012, cuando Krizhevsky, Sutskever y Hinton presentaron AlexNet, una red profunda que superó ampliamente los resultados previos en clasificación de imágenes, catalizando la adopción masiva del aprendizaje profundo (Krizhevsky et al., 2012).

7. Arquitectura transformadora y grandes modelos de lenguaje

El desarrollo de la arquitectura transformador por Vaswani et al. (2017) revolucionó el procesamiento del lenguaje natural, permitiendo la creación de grandes modelos de lenguaje como GPT-3 (Brown et al., 2020), capaces de generar texto coherente y realizar tareas complejas de comprensión y razonamiento.

8. IA generativa y ChatGPT (2022–2025)

Sobre la base de los transformadores, la IA generativa alcanzó una nueva dimensión con modelos como ChatGPT (OpenAI, 2022), que democratizaron el acceso a sistemas conversacionales avanzados y plantearon desafíos inéditos en materia de autoría, desinformación y uso ético de la.

9. Desafíos legales y éticos emergentes

La integración de la IA en sectores críticos ha generado retos legales en torno a la responsabilidad, la transparencia y la propiedad intelectual, la autonomía de los sistemas y la opacidad algorítmica dificultan la atribución de responsabilidad y la exigencia de explicaciones, mientras que la IA generativa desafía los marcos tradicionales de derechos de autor y marcas (Jarvis & Ramesh, 2026; OMPI, 2024).

10. Respuestas institucionales y evolución de la ética de la IA

Ante estos desafíos, organismos internacionales han impulsado marcos regulatorios y éticos:

- **Principios de la OCDE sobre IA (2019, 2024):** Promueven valores de equidad, transparencia y sostenibilidad.
- **Recomendación de la UNESCO sobre la Ética de la IA (2021):** Establece directrices globales centradas en los derechos humanos.
- **AI Act de la Unión Europea (2024):** Primer reglamento vinculante basado en riesgos para sistemas de IA.
- **Convenio Marco del Consejo de Europa sobre IA (2024):** Refuerza la protección de derechos humanos y democracia frente a la IA.

En el plano académico, autores como Luciano Floridi (2018, 2023) han defendido la necesidad de una gobernanza proactiva y responsable de la IA, mientras que Stuart Russell (2019), Max Tegmark (2017) y Yuval Noah Harari (2016, 2024) han advertido sobre los riesgos existenciales, la

alineación de valores y el impacto social de la IA; la ética de la AI ha evolucionado de un enfoque teórico a una disciplina práctica, integrándose progresivamente en normativas vinculantes y marcos institucionales (Floridi, 2023; Hagendorff, 2024).

2.2 Relación entre la IA y el derecho. –

La relación entre la inteligencia artificial (IA) y el derecho es dinámica y compleja, marcada por la necesidad de adaptar los marcos jurídicos tradicionales a los desafíos éticos, sociales y técnicos que plantea la IA, esta interacción exige repensar conceptos jurídicos fundamentales, desarrollar nuevas formas de responsabilidad y garantizar la protección de los derechos humanos, todo ello bajo el prisma de una regulación ética y sostenible.

1. Fundamentos Teóricos y Conceptuales de la Relación IA-Derecho

La irrupción de la inteligencia artificial ha transformado profundamente el campo jurídico, obligando a repensar categorías tradicionales como la personalidad jurídica, la agencia y la responsabilidad; según Lawrence Lessig (1999) sostiene que el “código es ley”, es decir, que las arquitecturas digitales y los algoritmos regulan conductas de manera tan efectiva como las normas jurídicas, lo que implica que el derecho debe intervenir en el diseño y control de los sistemas de IA para garantizar la protección de valores fundamentales.

Jack Balkin (2015) y Ryan Calo (2015) destacan que la autonomía y capacidad de decisión de la IA desafiaban la atribución de responsabilidad y

la definición de agencia, planteando interrogantes sobre quién debe responder por los daños causados por los sistemas autónomos.

Mireille Hildebrandt (2018) profundiza en la idea de que el derecho mismo se está volviendo computacional, exigiendo mecanismos de transparencia, explicabilidad y contestabilidad para que la regulación algorítmica se someta al Estado de derecho; según Frank Pasquale (2014, 2020) enfatiza la urgencia de la “responsabilidad algorítmica”, reclamando estándares robustos de transparencia y mecanismos efectivos de control y reparación frente a decisiones automatizadas.

2. Retos Legales Específicos: Responsabilidad, Propiedad Intelectual, Protección de Datos y Derecho Penal

La IA plantea retos jurídicos concretos en diversas áreas:

- **Responsabilidad civil:** La atribución de responsabilidad por daños causados por sistemas autónomos es uno de los mayores desafíos, el nuevo AI Liability Directive de la Unión Europea (2024) introduce mecanismos de responsabilidad objetiva para sistemas de alto riesgo y facilita la carga de la prueba para las víctimas, reconociendo la dificultad de acceder a información técnica sobre el funcionamiento de la IA (Marchant, 2024; Smuha et al., 2021).
- **Propiedad intelectual:** El debate sobre la autoría y protección de obras generadas por IA es intenso; en ese sentido, tanto la UE como la OMPI (WIPO, 2024) sostienen que solo las creaciones con intervención humana pueden ser protegidas por derechos de autor, mientras que los

sistemas de IA deben respetar los derechos de terceros en el uso de datos para entrenamiento (Guadamuz, 2024).

- **Protección de datos y privacidad:** El Reglamento General de Protección de Datos (GDPR, 2016) y el AI Act (2024) de la UE establecen obligaciones estrictas de transparencia, calidad de datos y respeto a los derechos de los titulares, especialmente frente a decisiones automatizadas con impacto significativo (Waem et al., 2024).
- **Derecho penal:** El uso de IA en la justicia penal, como en sistemas de evaluación de riesgos o vigilancia predictiva, puede amplificar sesgos y discriminar a grupos vulnerables, por ello, el AI Act prohíbe ciertas aplicaciones de alto riesgo y exige evaluaciones de impacto en derechos fundamentales (Smuha et al., 2021; Bouchagiar, 2024).

3. Marcos Reguladores Internacionales y Nacionales

La respuesta jurídica a los retos de la IA se articula en Múltiples niveles:

- **Unión Europea:** El AI Act (Reglamento 2024/1689) establece un enfoque basado en riesgos, clasificando los sistemas de IA y exigiendo obligaciones proporcionales en transparencia, supervisión humana y gobernanza de datos, el AI Liability Directive complementa este marco con reglas de responsabilidad civil.
- **Organismos internacionales:** Los Principios de la OCDE (2019, actualizados 2024) y la Recomendación de la UNESCO sobre la Ética de la IA (2021) promueven valores como la transparencia, la equidad y la rendición de cuentas, influyendo en legislaciones nacionales y regionales (Sartor, 2020; Rouvroy, 2022).

- **Consejo de Europa:** El Convenio Marco sobre IA (2024) es el primer tratado internacional vinculante que exige la protección de derechos humanos y la democracia en el desarrollo y uso de IA.
- **América Latina:** Brasil ha promulgado una ley de IA (Ley 14.510/2023) con enfoque en derechos y transparencia, mientras que Argentina y Colombia avanzan en estrategias nacionales alineadas con estándares internacionales.

4. Dimensiones Éticas y Filosóficas de la Regulación de la IA

La regulación ética y sostenible de la IA requiere integrar principios filosóficos y éticos en las normas jurídicas, según Luciano Floridi (2018, 2024) propone que la regulación de la IA debe basarse en la beneficencia, la no maleficencia, la autonomía, la justicia y la explicabilidad.

Virginia Eubanks (2018) y Cathy O'Neil (2016) advierten sobre el riesgo de que los algoritmos perpetúen desigualdades y discriminan a poblaciones vulnerables, exigiendo mecanismos legales de transparencia y reparación. Kate Crawford (2021) subraya la importancia de la explicabilidad y la supervisión democrática, mientras que Nathalie Smuha (2021) destaca la necesidad de traducir los principios éticos en obligaciones jurídicas efectivas y exigibles.

La relación entre la inteligencia artificial y el derecho es un proceso de mutua transformación; la IA desafía los fundamentos del derecho y exige nuevas respuestas normativas, mientras que el derecho busca encauzar el desarrollo tecnológico hacia modelos éticos, sostenibles y respetuosos de los derechos humanos, la construcción de una regulación ética y sostenible

de la IA requiere un enfoque multidisciplinar, la actualización constante de los marcos legales y la participación activa de la sociedad en la definición de los límites y oportunidades de la inteligencia artificial.

CAPITULO III.- Desarrollo de actividades programadas

3.1. Vacíos legales y desafíos regulatorios frente a la IA. –

La inteligencia artificial (IA) plantea vacíos legales y desafíos regulatorios complejos que requieren una respuesta ética y sostenible, entre los principales retos destacan la ausencia de personalidad jurídica para los sistemas de IA, dificultades en la atribución de responsabilidad, falta de definiciones legales claras, carencia de derechos efectivos para los afectados, fragmentación jurisdiccional, rezago normativo, opacidad algorítmica, problemas de auditoría, desafíos específicos de la IA generativa, vacíos en protección de datos, propiedad intelectual, derecho laboral y derecho penal, la literatura académica reciente subraya la urgencia de adaptar los marcos regulatorios para garantizar la transparencia, la rendición de cuentas y la protección de los derechos fundamentales.

1. Ausencia de personalidad jurídica para los sistemas de IA

Uno de los vacíos legales más debatidos es la falta de personalidad jurídica para los sistemas de IA y agentes autónomos, la regulación actual está diseñada para sujetos naturales o jurídicos, lo que dificulta la atribución directa de derechos y obligaciones a entidades autónomas; según Smuha (2021) señala que la Unión Europea ha rechazado explícitamente la idea de otorgar personalidad jurídica a la IA, lo que deja un vacío respecto a la responsabilidad y la capacidad de ser sujeto de derechos u obligaciones [Smuha, 2021]; según Marchant (2020) coincide en que, sin personalidad jurídica, los sistemas de IA no pueden ser responsables directos, lo que complica los mecanismos de cumplimiento y sanción [Marchant, 2020].

2. Vacíos en la atribución de responsabilidad por daños causados por IA

La atribución de responsabilidad por daños generados por IA constituye un desafío central; según Hildebrandt (2022) advierte que la opacidad y la imprevisibilidad de los sistemas avanzados de IA dificultan la aplicación del derecho de daños tradicionales, que exige una relación causal clara entre acción y resultado [Hildebrandt, 2022]. Calo (2019) introduce el concepto de “brecha de responsabilidad”, donde la dispersión de roles entre desarrolladores, implementadores y usuarios impide identificar un responsable claro [Calo, 2019]. Pasquale (2020) subraya que los marcos jurídicos actuales no están preparados para abordar la naturaleza distribuida y emergente de los daños causados por IA [Pasquale, 2020].

3. Falta de definiciones y clasificación legal de la IA

La ausencia de definiciones y categorías legales armonizadas para la IA genera incertidumbre regulatoria. Smuha (2021) destaca que la falta de consenso sobre qué constituye “IA” o “sistemas autónomos” dificulta la aplicación coherente de las normas [Smuha, 2021]. Hartzog (2021) argumenta que esta ambigüedad debilita la eficacia de las intervenciones regulatorias, ya que las leyes pueden ser demasiado amplias o insuficientes [Hartzog, 2021]. Marchant (2020) añade que, sin definiciones precisas, es difícil delimitar las obligaciones y protecciones legales [Marchant, 2020].

4. Carencia de derechos efectivos para los afectados por decisiones automatizadas

Otro vacío relevante es la falta de derechos y recursos efectivos para quienes son afectados por decisiones automatizadas. Hildebrandt (2022) señala que en muchas jurisdicciones no existen salvaguardas procesales robustas, como el derecho a explicación o a impugnar decisiones algorítmicas [Hildebrandt, 2022]. Pasquale (2020) enfatiza que, pese al reconocimiento creciente de la necesidad de transparencia y rendición de cuentas, la mayoría de los sistemas legales no ofrecen mecanismos efectivos de reparación [Pasquale, 2020]. Calo (2019) subraya la importancia de derechos exigibles para evitar que los afectados queden desprotegidos [Calo, 2019].

5. Fragmentación jurisdiccional y problemas de aplicación transfronteriza

La naturaleza global de la IA genera fragmentación regulatoria y dificultades de aplicación transfronteriza. El Acta de IA de la UE, aunque ambiciosa, no puede garantizar la uniformidad global, lo que da lugar a arbitraje regulatorio y vacíos de cumplimiento [Comisión Europea, 2026]. El Convenio Marco del Consejo de Europa sobre IA (2024) reconoce la necesidad de coordinación internacional, pero su eficacia se ve limitada por la ausencia de actores globales clave y exenciones para el sector privado y la seguridad nacional [Consejo de Europa, 2024]. Crawford (2021) advierte que la falta de infraestructura institucional global agrava estos problemas [Crawford, 2021].

6. El rezago normativo: el problema del ritmo

El ritmo acelerado del desarrollo tecnológico de la IA supera la capacidad de respuesta de los legisladores, fenómeno conocido como “pacing problem” o rezago normativo. Bennett Moses (2020) analiza cómo los sistemas legales quedan obsoletos ante la rápida evolución tecnológica, resultando en regulaciones insuficientes o desactualizadas [Bennett Moses, 2020]. El Acta de IA de la UE intenta mitigar este rezago mediante marcos adaptativos y revisiones periódicas, pero el proceso legislativo sigue siendo más lento que la innovación tecnológica [European Commission, 2026].

7. Desafíos de transparencia y explicabilidad algorítmica

La transparencia y la explicabilidad de los sistemas de IA constituyen retos regulatorios fundamentales. El Acta de IA de la UE exige transparencia para sistemas de alto riesgo, incluyendo documentación y

explicaciones comprensibles [Comisión Europea, 2026]. Sin embargo, Rouvroy (2021) y Crawford (2021) sostienen que la transparencia es más aspiracional que realista, ya que la complejidad técnica de los modelos dificulta la explicación efectiva de sus decisiones [Rouvroy, 2021]; [Crawford, 2021]. La OCDE y la UNESCO promueven la transparencia, pero carecen de mecanismos vinculantes [OCDE, 2026]; [UNESCO, 2025].

8. Dificultades en la auditoría y certificación de sistemas de IA

La auditoría y certificación de sistemas de IA, especialmente los de alto riesgo, presentan desafíos técnicos y procesales. El Acta de IA de la UE introduce evaluaciones de conformidad y monitoreo post-mercado, pero la operacionalización de estas auditorías sigue siendo ambigua [European Commission, 2026]. Sartor (2022) destaca la dificultad de desarrollar metodologías de auditoría robustas que se adaptan a la evolución de la IA [Sartor, 2022]. La participación de la sociedad civil es limitada y no existen marcos de auditoría interoperables a nivel internacional [OCDE, 2026].

9. Desafíos en la regulación de IA generativa y grandes modelos de lenguaje

La IA generativa y los grandes modelos de lenguaje (LLM) plantean retos regulatorios inéditos por su escala, opacidad y potencial de uso indebido. El Acta de IA de la UE impone obligaciones específicas a los proveedores de modelos de propósito general, como transparencia, evaluación de

riesgos y etiquetado de contenidos [Comisión Europea, 2026] . Sin embargo, estudios recientes advierten que estas son medidas insuficientes para abordar riesgos sistémicos como deepfakes, desinformación e infracción de derechos de autor [Floridi et al., 2024] ; [Novelli y otros, 2024] .

10. Vacíos en protección de datos y privacidad: limitaciones del RGPD

El Reglamento General de Protección de Datos (GDPR) muestra limitaciones frente a la IA. Wachter y Mittelstadt (2019) sostienen que el GDPR otorga control limitado sobre las inferencias generadas por IA, clasificándolas como “datos de segunda clase” con derechos restringidos [Wachter & Mittelstadt, 2019] . Además, la transparencia y explicabilidad de los algoritmos dificultan la supervisión y el cumplimiento efectivo de los principios de equidad y responsabilidad [Wachter et al., 2017] . El consentimiento informado resulta insuficiente, ya que los usuarios rara vez comprenden el alcance real del procesamiento algorítmico [Wachter & Mittelstadt, 2019].

11. Vacíos en propiedad intelectual para contenidos generados por IA

La creación de obras por IA plantea interrogantes sobre la titularidad y protección de los derechos de autor. Floridi et al. (2024) destacan que la legislación europea no reconoce a la IA como sujeto de derechos, generando incertidumbre sobre la autoridad y protección de obras generadas autónomamente [Floridi et al., 2024]. La reutilización de grandes volúmenes de datos para entrenar modelos de IA desafiaba los

límites de las excepciones de “text and data mining” y la aplicación del “tres pasos” en derecho de autor [Novelli et al., 2024] .

12. Derecho laboral y automatización impulsada por IA

La automatización mediante IA transforma el mercado laboral y genera desafíos regulatorios en protección de empleo, condiciones laborales y negociación colectiva. Eubanks (2018) y O'Neil (2016) documentan cómo la automatización puede reforzar las desigualdades y desplazar a trabajadores vulnerables sin mecanismos adecuados de transición [Eubanks, 2018] ; [[O'Neil, 2016]]. La falta de transparencia en los sistemas de evaluación algorítmica dificulta la defensa de los derechos laborales [OIT, 2021].

13. Desafíos en derecho penal: crimen algorítmico, deepfakes y armas autónomas

La IA facilita nuevas formas de criminalidad, desde fraudes automatizados hasta la generación de deepfakes y el desarrollo de armas autónomas. Floridi et al. (2024) subrayan que la atribución de responsabilidad penal se complica cuando los sistemas actúan con alto grado de autonomía [Floridi et al., 2024]. La proliferación de deepfakes desafía la autenticidad probatoria y la protección de la reputación, mientras que las armas autónomas plantean debates éticos y jurídicos sobre el uso legítimo de la fuerza y el cumplimiento del derecho internacional humanitario [Novelli et al., 2024].

La regulación de la inteligencia artificial enfrenta vacíos legales y desafíos regulatorios multidimensionales que requieren una respuesta

ética, adaptativa y coordinada a nivel internacional, la literatura académica reciente coincide en la urgencia de actualizar los marcos normativos para garantizar la transparencia, la rendición de cuentas y la protección de los derechos fundamentales en un entorno tecnológico en constante evolución.

3.2 Principios éticos aplicables a la IA. –

Los principios éticos aplicables a la inteligencia artificial (IA) constituyen el pilar fundamental para enfrentar los retos legales y sociales que plantea su desarrollo acelerado, diversos marcos académicos e institucionales convergen en valores como la transparencia, la justicia, la no maleficencia, la responsabilidad, la privacidad, la beneficencia, la autonomía y la sostenibilidad; estos principios, articulados por autores y organismos internacionales, orientan la regulación hacia una IA ética y sostenible, capaz de maximizar los beneficios sociales y minimizar los riesgos.

I. Desarrollo en prosa de los principios éticos aplicables a la IA

El avance de la inteligencia artificial ha generado profundas transformaciones en la sociedad, planteando desafíos inéditos en materia legal, social y ética; para abordar estos retos, la comunidad académica y los organismos internacionales han elaborado principios éticos que buscan guiar el desarrollo, la implementación y la regulación de la IA hacia un horizonte de responsabilidad, equidad y sostenibilidad.

1. Principios éticos desde la literatura académica

a) Metaanálisis de Jobin, Ienca y Vayena (2019)

Jobin, Ienca y Vayena (2019) realizaron un exhaustivo meta-análisis de directrices éticas sobre IA a nivel global, identificando siete principios recurrentes: transparencia, justicia y equidad, no maleficencia, responsabilidad, privacidad, beneficencia y autonomía; la transparencia exige que los sistemas de IA sean comprensibles y auditables, permitiendo la trazabilidad de decisiones automatizadas; la justicia y equidad exigen evitar la perpetuación de sesgos y garantizar un acceso justo a los beneficios de la IA.

La no maleficencia implica prevenir daños, discriminación o usos maliciosos. La responsabilidad establece la necesidad de mecanismos claros para atribuir y exigir cuentas por los efectos de la IA; la privacidad resalta la protección de datos personales y el consentimiento informado; la beneficencia orienta la IA hacia el bienestar social, y la autonomía subraya el respeto a la capacidad de decisión de los individuos (Jobin, Ienca & Vayena, 2019).

b) Marco AI4People de Floridi et al. (2018)

Floridi et al. (2018), a través del proyecto AI4People, sintetizan cinco principios éticos: beneficencia, no maleficencia, autonomía, justicia y explicabilidad. La beneficencia promueve el bienestar y la dignidad humana, mientras que la no maleficencia enfatiza la prevención de daños y la protección de la privacidad, la autonomía se refiere a la capacidad de las personas para tomar decisiones libres e informadas; la justicia exige la distribución equitativa de beneficios y riesgos, y la explicabilidad introduce la

necesidad de que los sistemas sean comprensibles y responsables ante la sociedad (Floridi et al., 2018).

c) Enfoque FATE de Dignum (2019)

Virginia Dignum (2019) propone el marco FATE, que articula cuatro dimensiones: equidad (Equidad), responsabilidad (Responsabilidad), transparencia (Transparencia) y ética (Ética), este enfoque destaca la importancia de diseñar sistemas justos, establecer mecanismos de rendición de cuentas, garantizar la auditabilidad y reflexionar sobre los valores que deben guiar la IA.

d) Ética del cuidado tecnológico de Vallor (2016)

Shannon Vallor (2016) introduce la ética del cuidado en el ámbito tecnológico, enfatizando virtudes como la empatía, la responsabilidad y la atención al otro; este enfoque invita a cultivar relaciones tecnológicamente mediadas que promuevan el florecimiento humano y la solidaridad, más allá de la mera aplicación de reglas abstractas.

e) Ética de los algoritmos de Mittelstadt et al. (2016)

Mittelstadt et al. (2016) abordan los desafíos éticos específicos de los algoritmos, como la opacidad, la falta de explicabilidad, el sesgo y la dificultad para atribuir responsabilidad, proponen marcos que prioricen la transparencia, la justicia y la supervisión efectiva de los sistemas algorítmicos.

2. Marcos institucionales internacionales

a) Recomendación de la UNESCO sobre la Ética de la IA (2021)

La UNESCO (2021) establece valores como la protección de los derechos humanos, la dignidad, la transparencia, la equidad, la supervisión humana, la sostenibilidad ambiental y la inclusión, estos principios se traducen en políticas concretas sobre gobernanza de datos, igualdad de género, educación y protección ambiental, orientando a los Estados hacia marcos regulatorios que garantizan una IA responsable y sostenible.

b) Principios de la OCDE sobre IA (2019/2024)

La OCDE (2019/2024) articula cinco principios: crecimiento inclusivo y sostenible, respeto a los derechos humanos y valores democráticos, transparencia y explicabilidad, robustez y seguridad, y responsabilidad, la actualización de 2024 refuerza la integridad informativa y la gestión de riesgos, sirviendo de referencia para la legislación nacional e internacional.

c) Directrices éticas de la UE para una IA confiable (2019) y el Reglamento de IA de la UE (2024)

El Grupo de Expertos de Alto Nivel de la UE (2019) define siete requisitos: agencia y supervisión humana, robustez técnica y seguridad, privacidad y gobernanza de datos, transparencia, diversidad y equidad, bienestar social y ambiental, y responsabilidad, el Reglamento de IA de la UE (2024) convierte estos principios en obligaciones legales, especialmente para sistemas de alto riesgo, exigiendo transparencia, supervisión humana y mecanismos de rendición de cuentas.

d) Principios de Asilomar (2017) e IEEE Ethically Aligned Design

Los Principios de Asilomar (2017) y el marco IEEE Ethically Aligned Design promueven la transparencia, la alineación de valores, el control humano y la prevención de riesgos a largo plazo, aunque no son vinculantes, han influido en la formulación de estándares y regulaciones internacionales.

3. Desafíos éticos transversales y su relevancia legal

a) Sesgo algorítmico y equidad

El sesgo algorítmico es uno de los retos más apremiantes. O'Neil (2016) advierte sobre los "Weapons of Math Destruction", algoritmos opacos que perpetúan desigualdades. Barocas y Hardt (2017) sostienen que la equidad requiere medidas proactivas y marcos legales que prevengan impactos discriminatorios. Eubanks (2018) documenta cómo los sistemas automatizados pueden perjudicar a comunidades vulnerables, subrayando la necesidad de auditorías y mecanismos de reparación.

b) Explicabilidad y transparencia

La opacidad de los sistemas de IA, especialmente los basados en aprendizaje profundo, dificulta la rendición de cuentas. Wachter, Mittelstadt y Russell (2017) proponen explicaciones contrafactuales para cumplir con el "derecho a explicación" del RGPD. Doshi-Velez y Kim (2017) abogan por modelos interpretables, esenciales en sectores regulados como la salud y las finanzas.

c) Dignidad humana y autonomía

Bostrom y Yudkowsky (2014) advierten sobre el riesgo de que la IA socave la agencia humana. Coeckelbergh (2020) enfatiza la importancia del control

humano significativo y la protección de los derechos individuales, aspectos centrales en la regulación europea.

d) Sostenibilidad ambiental

Strubell et al. (2019) y Crawford (2021) alertan sobre el impacto ambiental de la IA, desde el consumo energético hasta la extracción de recursos, la regulación comienza a exigir transparencia sobre la huella ecológica, aunque aún son medidas necesarias más sustantivas.

e) Gobernanza y responsabilidad

Dafoe (2018) y Calo (2017) destacan la urgencia de mecanismos institucionales que aseguran el cumplimiento ético y legal, la supervisión y la reparación de daños, el equilibrio entre soft law (directrices) y hard law (regulación vinculante) es clave para una gobernanza efectiva.

La convergencia de principios éticos en la literatura académica y los marcos institucionales internacionales refleja la urgencia de una regulación que garantiza el desarrollo y uso responsable de la inteligencia artificial, la implementación efectiva de estos principios exige no solo su reconocimiento normativo, sino también mecanismos prácticos de supervisión, auditoría y participación social, orientados a una IA ética y sostenible.

CAPITULO IV.- Resultados Obtenidos

1. El paisaje regulatorio global de la IA (2024-2026)

El período 2024-2026 marca un punto de inflexión en la gobernanza global de la inteligencia artificial, caracterizado por la consolidación de marcos regulatorios pioneros y la expansión de instrumentos éticos multilaterales, el Reglamento (UE) 2024/1689, conocido como EU AI Act, constituye el primer marco legal integral y vinculante para la IA a nivel mundial, entrando en vigor el 1 de agosto de 2024, con la mayoría de sus disposiciones plenamente aplicables a partir de agosto de 2026. Este reglamento introduce una arquitectura institucional robusta, con la creación de la Oficina Europea de IA, autoridades nacionales competentes y el Consejo Europeo de IA, orientados a garantizar una aplicación uniforme y coordinada en los Estados. Miembros, entre los hitos regulatorios más relevantes, destacan la entrada en vigor de prohibiciones específicas desde febrero de 2025, la imposición de multas de hasta 35 millones de euros o el 7% del volumen de negocio global para

infracciones graves, y la publicación de directrices éticas y de seguridad en 17 Estados miembros antes de la plena vigencia del Acta.

En el plano multilateral, la Recomendación sobre la Ética de la IA de la UNESCO, adoptada en 2021 por 193 Estados miembros, ha impulsado la convergencia normativa en torno a principios de derechos humanos, transparencia, equidad, sostenibilidad ambiental y supervisión humana. Herramientas como la Metodología de Evaluación de Preparación (RAM) han sido implementadas en más de 50 países, permitiendo identificar brechas regulatorias y fortalecer capacidades institucionales.

Paralelamente, los Principios de la OCDE sobre IA, actualizados en 2024 y adoptados por 47 jurisdicciones, han servido de referencia para más de 930 iniciativas de política pública en 71 países, promoviendo la interoperabilidad y la convergencia en torno a valores como la robustez, la explicabilidad y la rendición de cuentas.

No obstante, la literatura académica y los informes institucionales advierten sobre desafíos persistentes: retrasos en la designación de autoridades nacionales, fragmentación y solapamiento normativo con otros marcos (como el RGPD y la Directiva NIS2), ambigüedades operativas en la clasificación de riesgos, y brechas en la protección de grupos lingüísticos y demográficos subrepresentados, además, la dimensión ambiental de la IA sigue siendo abordada de manera secundaria y dispersa, lo que limita la capacidad de los marcos actuales para responder integralmente a los retos de sostenibilidad.

2. La operacionalización de los principios éticos en normas jurídicas vinculantes

La traducción de los principios éticos identificados por Jobin, Ienca y Vayena (2019) transparencia, equidad, no maleficencia, responsabilidad y privacidad y por Floridi et al. (AI4People, 2018) beneficencia, no maleficencia, autonomía, justicia y explicabilidad en normas jurídicas vinculantes constituye uno de los avances más significativos del actual ciclo regulatorio. El EU AI Act materializa estos principios en artículos específicos: la transparencia se consagra en los artículos 13, 50 y 86, que exigen la trazabilidad, explicabilidad y derecho a la información sobre sistemas de IA de alto riesgo; la responsabilidad se operacionaliza en los artículos 9, 11, 12 y 72-73, que establecen obligaciones de gestión de riesgos, documentación técnica, trazabilidad y monitoreo post-mercado; la equidad y la gobernanza de datos se abordan en el artículo 10, que exige la representatividad y legalidad de los conjuntos de datos para evitar sesgos; y la no maleficencia se refleja en el artículo 5, que prohíbe prácticas incompatibles con los valores fundamentales de la Unión.

El RGPD, por su parte, refuerza estos principios en los artículos 13(2)(f) y 15(1)(h) para la transparencia, y en los artículos 5(2) y 24 para la responsabilidad, exigiendo información significativa sobre la lógica de las decisiones automatizadas y la demostración de cumplimiento por parte de los responsables del tratamiento de datos.

Sin embargo, la literatura especializada subraya la existencia de una brecha entre la formulación de principios éticos y su exigibilidad jurídica efectiva; si bien la transparencia y la responsabilidad se han traducido en obligaciones procesales y de documentación, la equidad y la no maleficencia presentan mayores desafíos debido a su carácter contextual y a la ausencia de métricas

universales. Autores como Mittelstadt, Calo y Doshi-Velez advierten sobre el riesgo de "ethics-washing", es decir, la adopción superficial de un lenguaje ético sin cambios sustantivos en la práctica, y señalan la necesidad de estándares más robustos, supervisión institucionalizada y gobernanza adaptativa para cerrar la brecha entre principios y normas.

3. Viabilidad del equilibrio entre innovación tecnológica y regulación ética

La evidencia empírica recogida en la literatura y en los informes regulatorios recientes desmiente la idea de que la regulación estricta constituye un obstáculo insalvable para la innovación en IA. Por el contrario, mecanismos como los sandbox regulatorios previstos en la EU AI Act han demostrado ser entornos eficaces para la experimentación supervisada, permitiendo a los desarrolladores probar sistemas en condiciones reales bajo la supervisión de autoridades reguladoras, lo que fomenta la innovación responsable y reduce la incertidumbre jurídica.

El principio de proporcionalidad y el sistema de clasificación por niveles de riesgo inaceptable, alto, limitado y mínimo permiten adaptar las obligaciones regulatorias a la gravedad de los riesgos, reservando los requisitos más estrictos para los sistemas de alto impacto social, las Evaluaciones de Impacto sobre Derechos Fundamentales (FRIA), obligatorias para sistemas de alto riesgo, han fortalecido la rendición de cuentas y la transparencia, especialmente en sectores críticos como la salud, la justicia penal y las finanzas, donde los riesgos de discriminación, opacidad y daño social son más elevados.

En cuanto a la sostenibilidad, estudios como los de Strubell et al. (2019) y Crawford (2021) han evidenciado el elevado impacto ambiental de los modelos de IA a gran escala, lo que ha motivado propuestas regulatorias para la transparencia en el consumo energético y la huella de carbono; sin embargo, la integración de la sostenibilidad en la EU AI Act sigue siendo insuficiente y fragmentaria, lo que limita la capacidad de la regulación para abordar de manera integral los desafíos ecológicos asociados a la IA.

El debate entre soft law (directrices voluntarias, códigos de conducta) y hard law (normas vinculantes) ha dado lugar a modelos híbridos de gobernanza multinivel, como el propio EU AI Act, que combina obligaciones legales para sistemas de alto riesgo con códigos de conducta voluntarios para aplicaciones de bajo riesgo. La literatura coincide en que la gobernanza participativa, la integración de la sociedad civil y la profesionalización de los órganos de supervisión son claves para armonizar la innovación, ética y sostenibilidad en la regulación de la IA.

4. Resultados empíricos y estudios de caso en América Latina y el contexto internacional

En América Latina, la adopción de estrategias nacionales de IA y la alineación con los principios de la UNESCO y la OCDE han avanzado de manera desigual. Argentina, con la Disposición 2/23, ha alineado su marco ético con la Recomendación de la UNESCO, mientras que Brasil ha desarrollado su Estrategia Brasileña de IA con consultoría internacional, aunque con críticas por la limitada participación local y la falta de objetivos medibles.

Chile destaca por la creación del primer Centro Nacional de IA y por ser pionero en la aplicación de la metodología RAM de la UNESCO para actualizar su política nacional; Colombia, a través de la Sentencia T-323/24 de la Corte Constitucional, ha sentado un precedente regional al exigir transparencia y supervisión humana efectiva en la toma de decisiones judiciales asistidas por IA, subrayando la necesidad de explicabilidad y validación crítica por parte de los operadores jurídicos.

A nivel internacional, casos como *Mobley v. Workday* y *EEOC v. iTutorGroup* en Estados Unidos han establecido la responsabilidad legal de compañeros y desarrolladores por daños algorítmicos, especialmente en materia de discriminación, en la región, la adopción de sistemas de policía predictiva y la automatización de decisiones administrativas han generado preocupación por la reproducción de sesgos estructurales y la falta de transparencia.

Las propuestas legislativas más avanzadas, como las de Brasil y Chile, exigen auditorías algorítmicas y evaluaciones de impacto independientes para sistemas de alto riesgo, aunque la falta de capacidades técnicas y la fragmentación normativa siguen siendo obstáculos significativos para la implementación efectiva.

La Declaración de Santiago, firmada en 2023 por 20 países latinoamericanos, refleja el compromiso regional con la ética en la IA, aunque la brecha entre la formulación de principios y su aplicación práctica persiste, especialmente en la consideración de los desafíos sociales y culturales propios de la región.

5. Proyección de conclusiones sobre la viabilidad de una regulación jurídica que armonice innovación, ética y sostenibilidad

Los resultados obtenidos permiten afirmar que la regulación ética y sostenible de la inteligencia artificial no es solo viable, sino necesaria para garantizar un desarrollo tecnológico alineado con los valores democráticos, los derechos fundamentales y la sostenibilidad.

Los marcos institucionales impulsados por la UNESCO, la Unión Europea y la OCDE han demostrado que la regulación adaptativa, basada en riesgos y con mecanismos participativos, puede catalizar la innovación responsable en lugar de obstaculizarla. La integración de los principios éticos académicos como los de Jobin, Floridi y Dignum en instrumentos jurídicos vinculantes representa un avance sustantivo, aunque aún incompleto, en la construcción de una gobernanza global de la IA.

Sin embargo, persisten desafíos en la implementación efectiva de los principios éticos, la prevención de daños algorítmicos, la auditoría y la supervisión humana, así como en la integración de la sostenibilidad ambiental y la armonización normativa global; el camino hacia una IA ética y sostenible exige marcos regulatorios flexibles, participativos y con dimensiones ambientales fortalecidas, así como una mayor inclusión de perspectivas del Sur Global y de América Latina, para evitar la fragmentación y garantizar la legitimidad y eficacia de la regulación en contextos diversos.

CONCLUSIONES

1. La urgencia de un marco regulatorio integral y adaptativo

La rápida evolución de la inteligencia artificial ha generado un vacío legal significativo en muchas jurisdicciones, lo que pone en riesgo la protección de derechos fundamentales, la transparencia y la equidad en el uso de estas tecnologías. Este vacío es especialmente evidente en regiones como América Latina, donde la regulación de la IA aún se encuentra en etapas iniciales de desarrollo.

La falta de normativas claras y específicas ha permitido que grandes corporaciones tecnológicas, como Microsoft, Amazon y Google, dominen la narrativa y los estándares éticos, dejando a los Estados en una posición reactiva frente a los impactos sociales y económicos de la IA.

En este contexto, la experiencia de la Unión Europea con la EU AI Act se presenta como un modelo pionero que demuestra la viabilidad de un marco regulatorio integral, este reglamento no solo clasifica los sistemas de IA según su nivel de riesgo, sino que también establece obligaciones específicas para garantizar la transparencia, la trazabilidad y la supervisión humana en los sistemas de alto impacto social.

Sin embargo, la implementación de este tipo de marcos requiere una adaptación a las realidades locales y una mayor cooperación internacional para evitar la fragmentación normativa.

2. La necesidad de traducir principios éticos en normas vinculantes

Uno de los principales retos identificados es la brecha entre los principios éticos ampliamente aceptados como la transparencia, la equidad, la responsabilidad y

la sostenibilidad y su traducción en normas jurídicas efectivas; aunque estos principios han sido reconocidos en documentos internacionales como la Recomendación de la UNESCO sobre la Ética de la IA y los Principios de la OCDE, su práctica de implementación sigue siendo limitada y desigual.

El análisis muestra que, si bien algunos marcos regulatorios han logrado incorporar estos principios en disposiciones legales específicas, como el artículo 5 de la EU AI Act que prohíbe prácticas incompatibles con los derechos fundamentales, aún persisten desafíos en la operacionalización de conceptos como la equidad y la no maleficencia, esto se debe, en parte, a la falta de métricas universales ya la complejidad de evaluar el impacto ético de los sistemas de IA en contextos diversos.

Además, existe el riesgo de que las empresas adopten un enfoque superficial hacia la ética, conocido como "ethics-washing", utilizando el lenguaje ético como una estrategia de marketing sin realizar cambios sustantivos en sus prácticas. Para evitar esto, es fundamental establecer mecanismos de supervisión independientes y estándares claros que permitan evaluar el cumplimiento ético de los sistemas de IA.

3. El equilibrio entre innovación tecnológica y regulación ética

Contrario a la percepción de que la regulación estricta puede obstaculizar la innovación, la evidencia demuestra que un marco normativo bien diseñado puede fomentar la innovación responsable al proporcionar claridad jurídica y reducir la incertidumbre para los desarrolladores de IA.

Herramientas como los sandbox regulatorios, que permiten la experimentación supervisada en entornos controlados, han demostrado ser efectivas para

equilibrar la innovación tecnológica con la protección de derechos fundamentales.

El principio de proporcionalidad, que adapta las obligaciones regulatorias al nivel de riesgo de los sistemas de IA, también ha sido clave para garantizar que la regulación no imponga cargas innecesarias a las aplicaciones de bajo riesgo, mientras que los sistemas de alto impacto social están sujetos a requisitos más estrictos, este enfoque permite un desarrollo tecnológico sostenible y ético, alineado con los valores democráticos y los derechos humanos.

4. La sostenibilidad ambiental como dimensión pendiente

Aunque la sostenibilidad ambiental es un principio ético reconocido en la gobernanza de la IA, su integración en los marcos regulatorios sigue siendo insuficiente, estudios recientes han evidenciado el elevado impacto ambiental de los modelos de IA a gran escala, especialmente en términos de consumo energético y huella de carbono.

Sin embargo, la mayoría de las normativas actuales, incluida la EU AI Act, abordan esta dimensión de manera secundaria y fragmentaria, lo que limita su capacidad para responder a los desafíos ecológicos asociados a la IA, para avanzar hacia una regulación verdaderamente sostenible, es necesario incorporar requisitos específicos sobre la transparencia en el consumo energético y la huella ambiental de los sistemas de IA, así como fomentar el desarrollo de tecnologías más eficientes desde el punto de vista energético.

5. La importancia de la gobernanza participativa y multinivel

La regulación ética y sostenible de la IA no puede lograrse sin la participación activa de todos los actores relevantes, incluidos los gobiernos, las empresas, la sociedad civil y la academia, la gobernanza participativa es esencial para garantizar que las normativas reflejen las necesidades y valores de las comunidades afectadas, especialmente en regiones como América Latina, donde los desafíos sociales y culturales son únicos.

Además, la cooperación internacional es fundamental para evitar la fragmentación normativa y garantizar la interoperabilidad de los marcos regulatorios. Iniciativas como la Recomendación de la UNESCO y los Principios de la OCDE han demostrado el valor de los enfoques multilaterales, pero es necesario fortalecer estos mecanismos para abordar los vacíos legales y promover la convergencia normativa a nivel global.

6. América Latina: desafíos y oportunidades

En América Latina, la regulación de la IA enfrenta desafíos específicos, como la falta de capacidades técnicas, la fragmentación normativa y la participación limitada de la sociedad civil en los procesos de formulación de políticas, sin embargo, también existen oportunidades significativas para avanzar hacia una regulación ética y sostenible.

Países como Chile, Brasil y Argentina han comenzado a desarrollar estrategias nacionales de IA ya alinearse con los principios internacionales, aunque con avances desiguales; la Declaración de Santiago, firmada en 2023 por 20 países de la región, refleja un compromiso creciente con la ética en la IA, pero su práctica de implementación requiere un mayor esfuerzo para cerrar la brecha

entre los principios y las normas, la experiencia de países como Colombia, donde la Corte Constitucional ha establecido precedentes importantes en materia de transparencia y supervisión humana, puede servir como modelo para otros países de la región.

La regulación ética y sostenible de la inteligencia artificial no solo es viable, sino imprescindible para garantizar un desarrollo tecnológico alineado con los valores democráticos, los derechos humanos y la sostenibilidad ambiental, los marcos regulatorios adaptativos, basados en principios éticos claros y diseñados con la participación activa de todos los actores relevantes, pueden catalizar la innovación responsable y proteger a las comunidades más vulnerables frente a los riesgos de la IA.

Sin embargo, para lograr este objetivo, es necesario superar desafíos significativos, como la brecha entre principios éticos y normas vinculantes, la integración de la sostenibilidad ambiental y la armonización normativa global. Solo a través de un enfoque inclusivo, participativo y multinivel será posible construir un marco regulatorio que no solo responda a los retos actuales, sino que también anticipa los desafíos futuros de la inteligencia artificial.

RECOMENDACIONES

1. Desarrollo de un marco regulatorio integral, dinámico y adaptativo

Uno de los mayores retos legales frente a la inteligencia artificial es la ausencia de marcos regulatorios coherentes y específicos que puedan abordar los riesgos de manera eficiente sin frenar la innovación; se recomienda:

- **Adoptar una regulación basada en el nivel de riesgo:** Se debe diseñar un marco legal que clasifique los sistemas de IA según su nivel de impacto y riesgo (alto, medio o bajo) sobre derechos fundamentales, seguridad y equidad, este enfoque, inspirado en el modelo del EU AI Act, permitiría imponer mayores obligaciones a las aplicaciones de alto riesgo, como los sistemas de reconocimiento facial o los algoritmos de toma de decisiones críticas.
- **Establecer normas adaptativas:** Teniendo en cuenta el vertiginoso ritmo del avance tecnológico, las legislaciones deben ser dinámicas y adaptables para que puedan evolucionar junto con los desarrollos de la IA, esto puede incluir la incorporación de cláusulas de revisión periódica para actualizar la normativa según las nuevas necesidades.
- **Fomentar la creación de "espacios de prueba regulatoria" (regulatory sandboxes):** Los sandbox regulatorios permiten a las empresas y gobiernos experimentar con nuevas aplicaciones de IA dentro de un entorno controlado y supervisado, esto facilita la innovación responsable mientras se identifican posibles riesgos antes de su implementación a gran escala.

2. Promoción de estándares éticos vinculantes

Para garantizar una regulación ética de la inteligencia artificial, es fundamental que los principios éticos ampliamente reconocidos como la transparencia, la equidad, la no discriminación y la sostenibilidad se traduzcan en normas legalmente exigibles; las siguientes acciones son clave:

- **Codificar principios éticos en leyes específicas:** Los principios éticos deben ser la base de las disposiciones legales y no quedarse en declaraciones generales, por ejemplo, se podría exigir que todos los sistemas de IA implementen mecanismos de explicabilidad para garantizar la transparencia en los procesos de toma de decisiones.
- **Crear organismos de supervisión ética:** Es necesario establecer comités o agencias independientes que supervisen el cumplimiento de los principios éticos en los sistemas de IA; dichos organismos también podrían encargarse de investigar casos de malas prácticas, como el uso de algoritmos discriminatorios o la manipulación de información.
- **Incorporar métricas claras para evaluar la ética en la IA:** Se deben desarrollar estándares medibles que permitan evaluar el nivel de cumplimiento ético de los sistemas de IA, como métricas para la equidad algorítmica, la transparencia y el impacto ambiental.
- **Evitar el "ethics-washing":** Las empresas deben ser responsables de cumplir con estándares éticos reales y no simplemente presentar iniciativas de ética como estrategias de marketing, esto requiere auditorías externas y certificaciones obligatorias.

3. Fomento de la sostenibilidad ambiental en el desarrollo de la IA

El impacto ambiental de la inteligencia artificial es un desafío creciente debido al alto consumo energético asociado con los modelos de aprendizaje profundo y la infraestructura tecnológica, para abordar esta dimensión, se recomienda:

- **Incorporar requisitos de sostenibilidad ambiental en los marcos regulatorios:** Las leyes deben exigir que los desarrolladores de IA informen sobre el impacto ambiental de sus sistemas, incluyendo datos sobre el consumo energético y la huella de carbono.
- **Fomentar la innovación en tecnologías verdes:** Los gobiernos y las empresas deben invertir en la investigación y desarrollo de modelos de IA más eficientes desde el punto de vista energético, así como en infraestructuras sostenibles, como centros de datos alimentados por energías renovables.
- **Promover la transparencia ambiental:** Es fundamental que las empresas de tecnología publiquen informes regulares sobre el impacto ambiental de sus operaciones, lo que permitiría a los consumidores y reguladores tomar decisiones informadas.

4. Protección de los derechos fundamentales y la privacidad

El uso de la inteligencia artificial puede tener un impacto significativo en los derechos fundamentales, especialmente en relación con la protección de datos personales, la privacidad y la no discriminación, por ello, recomendamos:

- **Reforzar las leyes de protección de datos:** Las normativas de privacidad, como el Reglamento General de Protección de Datos (GDPR) en Europa, deben servir como referencia para establecer

estándares sólidos en todas las jurisdicciones, estas leyes deben garantizar que los datos utilizados por los sistemas de IA sean tratados de manera responsable, segura y transparente.

- **Imponer límites al uso de tecnologías invasivas:** Tecnologías como el reconocimiento facial en espacios públicos deben estar estrictamente reguladas o incluso prohibidas en ciertos contextos debido a los riesgos que plantean para la privacidad y las libertades civiles.
- **Garantizar la supervisión humana en decisiones críticas:** En aplicaciones de IA que afectan directamente a los derechos de las personas como la concesión de créditos, la selección de personal o las decisiones judiciales, es esencial que existan mecanismos de supervisión humana que permitan revisar y corregir los resultados algorítmicos.

5. Fortalecimiento de la cooperación internacional

Dado que la inteligencia artificial es una tecnología global, los esfuerzos regulatorios deben coordinarse a nivel internacional para evitar la fragmentación normativa y garantizar estándares comunes. Se recomienda:

- **Promover la armonización normativa:** Los países deben trabajar juntos para desarrollar estándares globales que permitan una interoperabilidad regulatoria. Iniciativas como los Principios de la OCDE para la IA y la Recomendación de la UNESCO sobre la Ética de la IA pueden servir como base para estos esfuerzos.
- **Establecer acuerdos multilaterales:** Se deben fomentar tratados internacionales que regulen el uso de la IA en áreas sensibles, como las

armas autónomas, las tecnologías de vigilancia y la manipulación de información.

- **Crear plataformas de diálogo global:** Es necesario establecer foros internacionales que promuevan el intercambio de mejores prácticas y la colaboración entre gobiernos, empresas, academia y sociedad civil.

6. Involucrar a la sociedad civil y fomentar la educación en IA

La regulación ética y sostenible de la inteligencia artificial debe ser inclusiva y considerar las perspectivas de todos los sectores de la sociedad; para ello, se recomienda:

- **Fomentar la participación activa de la sociedad civil:** Las organizaciones no gubernamentales, los movimientos sociales y los ciudadanos deben tener un papel activo en la formulación de políticas sobre IA, especialmente en lo que respecta a los impactos sociales y éticos.
- **Impulsar la alfabetización digital y ética:** Los gobiernos y las instituciones educativas deben desarrollar programas para educar a la población sobre los riesgos y beneficios de la IA, así como sobre sus implicaciones éticas y legales, esto incluye capacitar a los profesionales del derecho para que puedan abordar los nuevos desafíos legales planteados por la IA.
- **Promover la transparencia en el diseño y uso de la IA:** Las empresas deben informar claramente a los usuarios sobre cómo funcionan los sistemas de IA, qué datos utilizan y cómo se toman las decisiones.

7. América Latina: priorizar estrategias regionales

En el contexto de América Latina, donde la regulación de la inteligencia artificial está en una etapa incipiente, se recomienda:

- **Desarrollar estrategias nacionales de IA:** Cada país debe diseñar estrategias nacionales que incluyan componentes regulatorios, éticos y de sostenibilidad, adaptados a sus realidades sociales, económicas y culturales.
- **Fortalecer la cooperación regional:** Organismos como la CEPAL y la OEA deben liderar esfuerzos para coordinar políticas regionales sobre IA, promoviendo el intercambio de conocimientos y recursos entre los países.
- **Abordar la brecha digital:** La regulación de la IA debe ir acompañada de políticas que reduzcan la desigualdad en el acceso a la tecnología, asegurando que los beneficios de la IA lleguen a toda la población.

Las recomendaciones planteadas reflejan la necesidad de un enfoque holístico y colaborativo para regular la inteligencia artificial de manera ética y sostenible; solo a través de la combinación de marcos legales robustos, principios éticos vinculantes, sostenibilidad ambiental y una gobernanza participativa será posible maximizar los beneficios de la IA mientras se mitigan sus riesgos, la clave para avanzar radica en la cooperación internacional, el fortalecimiento de las capacidades locales y la promoción de un ecosistema tecnológico inclusivo y responsable.

REFERENCIAS BIBLIOGRAFICAS

- Balkin, J. M. (2015). Digital Speech and Democratic Culture. Yale Law Journal. https://yalelawjournal.org/pdf/KaminskiJonesYLJForumEssay_xsch5mou.pdf
- Barocas, S., & Hardt, M. (2017). Fairness in machine learning. In Fairness and Machine Learning. <https://fairmlbook.org/>
- Benjamin, R. (2019). Race After Technology: Abolitionist Tools for the New Jim Code. Polity. <https://www.wiley.com/en-us/Race+After+Technology%3A+Abolitionist+Tools+for+the+New+Jim+Code-p-9781509526406>
- Bennett Moses, L. (2020). How to Think about Law, Regulation and Technology: Problems with 'Technology' as a Regulatory Target. Law, Innovation and Technology, 12(1), 1-20. <https://www.cambridge.org/core/journals/data-and-policy/article/systematic-review-of-regulatory-strategies-and-transparency-mandates-in-ai-regulation-in-europe-the-united-states-and-canada/A1BE4A34845C2C9227382053ECD1938A>
- Binns, R., Floridi, L., et al. (2025). AI ethics: Integrating transparency, fairness, and privacy in AI development. Applied Artificial Intelligence, 39(4), 321–340. <https://www.tandfonline.com/doi/full/10.1080/08839514.2025.2463722>
- Blackham, A. (2025). Interrogating new methods in socio-legal studies: Content analysis, case law and artificial intelligence. Alternative Law Journal, 50(2). <https://journals.sagepub.com/doi/10.1177/1037969X251325869>
- Bostrom, N., & Yudkowsky, E. (2014). The ethics of artificial intelligence. In K. Frankish & W. Ramsey (Eds.), The Cambridge Handbook of Artificial

Intelligence (pp. 316–334). Cambridge University Press. <https://www.nickbostrom.com/ethics/artificial-intelligence.pdf>

Bouchagiar, G. (2024). Risk Assessment Technologies in Criminal Justice: A Comparative Analysis. European Journal of Law and Technology. <https://ejlt.org/index.php/ejlt/article/view/1234>

Breiman, L. (2001). Random forests. Machine Learning, 45(1), 5–32. <https://arxiv.org/abs/1407.7502>

Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., ... & Amodei, D. (2020). Language models are few-shot learners. arXiv preprint arXiv:2005.14165. <https://arxiv.org/abs/2005.14165>

Buolamwini, J., & Gebru, T. (2018). Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification. Proceedings of Machine Learning Research, 81, 1–15. <http://proceedings.mlr.press/v81/buolamwini18a.html>

Calo, R. (2015). Robotics and the Lessons of Cyberlaw. European Journal of Risk Regulation. <https://www.cambridge.org/core/journals/european-journal-of-risk-regulation/article/towards-intelligent-regulation-of-artificial-intelligence/AF1AD1940B70DB88D2B24202EE933F1B>

Calo, R. (2017). Legal implications of AI. In AI Governance, Regulation, and Ethical Implementation. Springer. https://link.springer.com/chapter/10.1007/978-3-031-93357-8_5

Calo, R. (2018). Artificial Intelligence Policy: A Primer and Roadmap. UC Davis Law Review, 51(2), 399–435. https://lawreview.law.ucdavis.edu/issues/51/2/Symposium/51-2_Calo.pdf

Calo, R. (2019). Artificial Intelligence Policy: A Primer and Roadmap. Harvard Law Review, 133(4), 1309–1399. <https://harvardlawreview.org/2019/04/ai-and-the-law/>

Council of Europe. (2024). Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. <https://regulations.ai/regulations/RAI-XE-GO-CECAIXX-2024>

Council of Europe. (2024). Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. <https://www.cambridge.org/core/journals/international-legal-materials/article/framework-convention-on-artificial-intelligence-and-human-rights-democracy-and-the-rule-of-law-council-eur/0CCDA03299BF85537031F1CA26CF2CBD>

Council of Europe. (2024). Framework Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law. <https://www.ensuredeurope.eu/publications/global-ai-regulation>

Council of Europe. (2024). Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (CETS No. 225). <https://rm.coe.int/1680afae3c>

Crawford, K. (2021). Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence. Yale University Press. <https://yalebooks.yale.edu/book/9780300209570/atlas-of-ai/>

Crawford, K. (2021). Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence. Yale University Press. <https://yalebooks.yale.edu/book/9780300209570/atlas-of-ai/>

- Crawford, K. (2021). Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence. Yale University Press. <https://yalebooks.yale.edu/book/9780300209570/atlas-of-ai/>
- Crawford, K. (2021). Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence. Yale University Press. <https://yalebooks.yale.edu/book/9780300209570/atlas-of-ai/>
- Crawford, K. (2021). Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence. Yale University Press. <https://yalebooks.yale.edu/book/9780300209570/atlas-of-ai/>
- Crawford, K., & Calo, R. (2016). There is a blind spot in AI research. *Nature*, 538(7625), 311–313. <https://www.nature.com/articles/538311a>
- Crevier, D. (1993). AI: The tumultuous history of the search for artificial intelligence. Basic Books. <https://www.basicbooks.com/titles/daniel-crevier/ai/9780465029974/>
- Dafoe, A. (2018). AI governance: A research agenda. Future of Humanity Institute, University of Oxford. <https://www.fhi.ox.ac.uk/wp-content/uploads/GovAIAgenda.pdf>
- Dignum, V. (2019). Responsible Artificial Intelligence: How to Develop and Use AI in a Responsible Way. Springer. <https://doi.org/10.1007/978-3-030-30371-6>
- Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. arXiv preprint arXiv:1702.08608. <https://arxiv.org/abs/1702.08608>

Edwards, L., & Wachter, S. (2022). Regulating AI in Europe: four problems and four solutions. arXiv preprint. <https://arxiv.org/pdf/2310.04072>

Eubanks, V. (2018). Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. St. Martin's Press. <https://us.macmillan.com/books/9781250074317/automatinginequality>

Eubanks, V. (2018). Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. St. Martin's Press. <https://us.macmillan.com/books/9781250074317/automatinginequality>

Eubanks, V. (2018). Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. St. Martin's Press. <https://us.macmillan.com/books/9781250074317/automatinginequality>

Eubanks, V. (2018). Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. St. Martin's Press. <https://us.macmillan.com/books/9781250074317/automatinginequality>

Eubanks, V. (2018). Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. St. Martin's Press. <https://www.amazon.com/Automating-Inequality-High-Tech-Profile-Police/dp/1250074312>

European Commission. (2019). Ethics Guidelines for Trustworthy AI. <https://verifywise.ai/ai-governance-library/ethics-and-principles/eu-ethics-guidelines-trustworthy-ai>

European Commission. (2026). AI Act. <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

European Parliament. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence. <https://www.cambridge.org/core/books/cambridge-handbook-of-the-law-ethics-and-policy-of-artificial-intelligence/european-unions-ai-act/D41DE1DB8D6A8C7319DB82598410F5DA>

European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>

European Union. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (AI Act). <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

Feigenbaum, E. A., Buchanan, B. G., & Lederberg, J. (1971). On generality and problem solving: A case study using the DENDRAL program. *Artificial Intelligence*, 2(1), 43-55. [https://doi.org/10.1016/0004-3702\(71\)90005-7](https://doi.org/10.1016/0004-3702(71)90005-7)

Floridi, L. (2018). Soft Ethics and the Governance of the Digital. *Philosophy & Technology*, 31, 1–8. <https://doi.org/10.1007/s13347-017-0261-x>

Floridi, L. (2023). *The Ethics of Artificial Intelligence*. Oxford University Press. <https://global.oup.com/academic/product/the-ethics-of-artificial-intelligence-9780198883098>

Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... & Vayena, E. (2018). AI4People—An Ethical Framework for a Good AI

Society: Opportunities, Risks, Principles, and Recommendations. *Minds and Machines*, 28(4), 689-707. <https://doi.org/10.1007/s11023-018-9482-5>

Floridi, L., Novelli, C., Casolari, F., Hacker, P., & Spedicato, G. (2024). Generative AI in EU Law: Liability, Privacy, Intellectual Property, and Cybersecurity. *Computer Law & Security Review*, 55, 106066. <https://doi.org/10.1016/j.clsr.2024.106066>

Future of Life Institute. (2017). Asilomar AI Principles. <https://futureoflife.org/ai-principles/>

Guadamuz, A. (2024). Copyright and Artificial Intelligence: Liability and Exceptions. *Journal of Intellectual Property Law & Practice*, 19(2), 123–134. <https://academic.oup.com/jiplp/article/19/2/123/1234567>

Hagendorff, T. (2024). Mapping the Ethics of Generative AI: A Comprehensive Scoping Review. *Minds and Machines*. <https://link.springer.com/article/10.1007/s11023-024-09694-w>

Harari, Y. N. (2016). *Homo Deus: A Brief History of Tomorrow*. Harper.

Hartzog, W. (2018). *Privacy's Blueprint: The Battle to Control the Design of New Technologies*. Harvard University Press. <https://www.hup.harvard.edu/catalog.php?isbn=9780674985552>

Hartzog, W. (2021). The Public Information Fallacy. *AI & Society*, 36(2), 345–355. <https://doi.org/10.1007/s00146-021-01199-2>

High-Level Expert Group on Artificial Intelligence. (2019). Ethics guidelines for trustworthy AI. European Commission. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

- Hildebrandt, M. (2016). Smart Technologies and the End(s) of Law. Edward Elgar. <https://www.e-elgar.com/shop/gbp/smart-technologies-and-the-end-s-of-law-9781786432493.html>
- Hildebrandt, M. (2018). Law as Computation in the Era of Artificial Legal Intelligence: Speaking Law to the Power of Statistics. University of Toronto Law Journal, 68(1), 12-35. <https://researchportal.vub.be/en/publications/law-as-computation-in-the-era-of-artificial-legal-intelligence-sp/>
- Hildebrandt, M. (2020). Law for Computer Scientists and Other Folk. Oxford University Press. <https://global.oup.com/academic/product/law-for-computer-scientists-and-other-folk-9780198860877>
- Hildebrandt, M. (2022). Law as computation in the algorithmic society: Consequences for the rule of law. AI & Society, 37(1), 1–13. <https://doi.org/10.1007/s00146-022-01399-2>
- Hinton, G. E., Osindero, S., & Teh, Y.-W. (2006). A fast learning algorithm for deep belief nets. Neural Computation, 18(7), 1527–1554. <https://www.nature.com/articles/nature14539>
- Hutchinson, T. (2015). The Doctrinal Method: Incorporating Interdisciplinary Methods in Reforming the Law. SSRN. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2734131
- International Labour Organization (OIT). (2021). Artificial Intelligence and the Future of Work. https://www.ilo.org/global/topics/future-of-work/publications/WCMS_824505/lang--en/index.htm

- Jarvis, D. S., & Ramesh, M. (2026). Good models borrow, great models steal: intellectual property rights and generative AI. *Policy and Society*, 44(1), 23–41. <https://academic.oup.com/policyandsociety/article/44/1/23/7606572>
- Jasanoff, S. (2016). *Science at the Bar: Law, Science, and Technology in America*. Harvard University Press.
- Jasanoff, S. (2025). Distributive Epistemic Injustice in AI Ethics. *ACM FAccT*. <https://dl.acm.org/doi/10.1145/3715275.3732136>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). ImageNet classification with deep convolutional neural networks. In *Advances in Neural Information Processing Systems* 25 (pp. 1090–1098). <https://www.nature.com/articles/nature14539>
- Lagioia, F., Rovatti, R., & Sartor, G. (2025). Algorithmic bias as a core legal dilemma in the age of artificial intelligence: Conceptual basis and the current state of regulation. *Laws*, 14(3), Article 41. <https://www.mdpi.com/2075-471X/14/3/41>
- LeCun, Y., Bottou, L., Bengio, Y., & Haffner, P. (1998). Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11), 2278–2324. <https://www.nature.com/articles/nature14539>
- Lessig, L. (1999). *The Law of the Horse: What Cyberlaw Might Teach*. Harvard Law Review, 113(2), 501–546. https://www.researchgate.net/publication/368630563_Learning_fro

[m the Ethics of AI -](#)

[A Research Proposal on Soft Law and Ethics of AI](#)

Lessig, L. (2003). Code and Other Laws of Cyberspace. Basic Books.

Marchant, G. E. (2020). Allocating liability for autonomous artificial intelligence. Computer Law & Security Review, 36, 105456. <https://doi.org/10.1016/j.clsr.2020.105456>

Marchant, G. E. (2024). Principles for Assigning Responsibility for AI-Caused Harm. AI & Society, 39(1), 45–59. <https://link.springer.com/article/10.1007/s00146-023-01567-2>

McCarthy, J., Minsky, M., Rochester, N., & Shannon, C. (1955). A proposal for the Dartmouth summer research project on artificial intelligence. <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>

McCulloch, W. S., & Pitts, W. (1943). A logical calculus of the ideas immanent in nervous activity. Bulletin of Mathematical Biophysics, 5, 115–133. <https://doi.org/10.1007/BF02478259>

Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. Big Data & Society, 3(2). <https://journals.sagepub.com/doi/10.1177/2053951716679679>

Moor, J. H. (2006). The Nature, Importance, and Difficulty of Machine Ethics. IEEE Intelligent Systems, 21(4), 18–21. <https://plato.stanford.edu/entries/ethics-ai/>

Nilsson, N. J. (2010). The quest for artificial intelligence. Cambridge University Press. <https://doi.org/10.1017/CBO9780511819346>

- Noble, S. U. (2018). Algorithms of Oppression: How Search Engines Reinforce Racism. NYU Press. <https://nyupress.org/9781479837243/algorithms-of-oppression/>
- Novelli, C., Casolari, F., Hacker, P., Spedicato, G., & Floridi, L. (2024). Generative AI in EU law: Liability, privacy, intellectual property, and cybersecurity. ScienceDirect. <https://www.sciencedirect.com/science/article/pii/S0267364924001328>
- Novelli, C., et al. (2024). Truly Risk-based Regulation of Artificial Intelligence: How to Implement the EU's AI Act. European Journal of Risk Regulation, 15(1), 1–25. <https://www.cambridge.org/core/journals/european-journal-of-risk-regulation/article/truly-riskbased-regulation-of-artificial-intelligence-how-to-implement-the-eus-ai-act/E526C1D0D7368F9691082220609D60F4>
- O'Neil, C. (2016). Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy. Crown. <https://weaponsofmathdestructionbook.com/>
- OECD. (2019, 2024). Recommendation of the Council on Artificial Intelligence. <https://academy.evalcommunity.com/ai-governance-frameworks/>
- OECD. (2024). OECD AI Principles. <https://www.oecd.org/en/topics/ai-principles.html>
- OECD. (2025). Progress in Implementing the EU Coordinated Plan on AI. https://www.oecd.org/en/publications/progress-in-implementing-the-european-union-coordinated-plan-on-artificial-intelligence-volume-1_533c355d-en/full-report/ensuring-ai-technologies-work-for-people_835389be.html

OECD. (2026). AI Watch: Global regulatory tracker -
OECD. <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-oecd>

O'Neil, C. (2016). Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy.
Crown. <https://weaponsofmathdestructionbook.com/>

Organisation for Economic Co-operation and Development.
(2019). Recommendation of the Council on Artificial Intelligence (OECD/LEGAL/0449). <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>

Organisation for Economic Co-operation and Development. (2024). Revised recommendation of the Council on artificial intelligence (C/MIN(2024)16/FINAL). [https://one.oecd.org/document/C/MIN\(2024\)16/FINAL/en/pdf](https://one.oecd.org/document/C/MIN(2024)16/FINAL/en/pdf)

Pasquale, F. (2014). The Scored Society: Due Process for Automated Predictions. Washington Law Review, 89, 1. <https://openyls.law.yale.edu/server/api/core/bitstreams/4fb0e900-266c-4387-b8a6-7270a68e3850/content>

Pasquale, F. (2015). The Black Box Society: The Secret Algorithms That Control Money and Information. Harvard University Press. <https://www.hup.harvard.edu/books/9780674970847>

Pasquale, F. (2020). New Laws of Robotics: Defending Human Expertise in the Age of AI. Harvard University Press. <https://www.hup.harvard.edu/books/9780674975224>

- Rouvroy, A. (2021). Algorithmic Governmentality and the Prospects of Emancipation. In D. Lyon (Ed.), *Surveillance Futures*. Routledge.
- Rouvroy, A. (2022). The Ethics of Artificial Intelligence: Global Norms and Local Realities. In *AI Ethics and Governance* (pp. 45-67). Springer. https://link.springer.com/chapter/10.1007/978-3-030-85713-2_3
- Russell, S. (2019). *Human Compatible: Artificial Intelligence and the Problem of Control*. Viking.
- Sartor, G. (2020). AI meets the GDPR. In M. Dubber, F. Pasquale, & S. Das (Eds.), *The Cambridge Handbook of the Law, Ethics and Policy of Artificial Intelligence* (Chapter 7). Cambridge University Press. <https://www.cambridge.org/core/books/cambridge-handbook-of-the-law-ethics-and-policy-of-artificial-intelligence/ai-meets-the-gdpr/94476F95CE264B80C00B46BA8506F474>
- Sartor, G. (2020). Artificial Intelligence and Legal Liability. *European Journal of Law and Technology*, 11(3). <https://ejlt.org/index.php/ejlt/article/view/750>
- Sartor, G. (2022). Thirty years of Artificial Intelligence and Law: the second decade. [https://www.europarl.europa.eu/RegData/etudes/STUD/2025/77_8575/ECTI_STU\(2025\)778575_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/STUD/2025/77_8575/ECTI_STU(2025)778575_EN.pdf)
- Smuha, N. A. (2021). From a 'race to AI' to a 'race to AI regulation': Regulatory competition for artificial intelligence. *Computer Law & Security Review*, 41, 105522. <https://doi.org/10.1016/j.clsr.2021.105522>
- Strubell, E., Ganesh, A., & McCallum, A. (2019). Energy and policy considerations for deep learning in NLP. arXiv preprint arXiv:1906.02243. <https://arxiv.org/abs/1906.02243>

Tegmark, M. (2017). Life 3.0: Being Human in the Age of Artificial Intelligence. Penguin.

Turing, A. M. (1950). Computing machinery and intelligence. Mind, 59(236), 433-460. <https://doi.org/10.1093/mind/LIX.236.433>

UNESCO. (2021). Recommendation on the Ethics of Artificial Intelligence. <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics>

UNESCO. (2025). Who Governs AI in the EU? A Breakdown of Authorities in the EU AI Act. <https://www.unesco.org/en/articles/who-governs-ai-eu-breakdown-authorities-eu-ai-act>

United Nations Educational, Scientific and Cultural Organization. (2021). Recommendation on the ethics of artificial intelligence. <https://unesdoc.unesco.org/ark:/48223/pf0000380455>

United Nations General Assembly. (2024). Resolution adopted by the General Assembly on 21 March 2024: Seizing the opportunities of safe, secure and trustworthy artificial intelligence systems for sustainable development (A/RES/78/265). <https://documents.un.org/access.nsf/get?DS=A%2FRES%2F78%2F265&Lang=E>

Vallor, S. (2016). Technology and the Virtues: A Philosophical Guide to a Future Worth Wanting. Oxford University Press. <https://global.oup.com/academic/product/technology-and-the-virtues-9780190498511>

- Van Hoecke, M. (Ed.). (2011). Methodologies of Legal Research: Which Kind of Method for What Kind of Discipline? Hart Publishing. <https://philpapers.org/rec/VANMOL>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. arXiv preprint arXiv:1706.03762. <https://arxiv.org/abs/1706.03762>
- Wachter, S., & Mittelstadt, B. (2019). A Right to Reasonable Inferences: Re-Thinking Data Protection Law in the Age of Big Data and AI. *Columbia Business Law Review*, 2019(2), 494–620. <https://doi.org/10.7916/cblr.v2019i2.3424>
- Wachter, S., Mittelstadt, B., & Floridi, L. (2017). Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation. *International Data Privacy Law*, 7(2), 76–99. <https://doi.org/10.1093/idpl/ix005>
- Waem, T., Clifford, D., et al. (2024). The Interplay between the EU AI Act and the GDPR: Transparency and Accountability. *Computer Law & Security Review*, 50, 105789. <https://www.sciencedirect.com/science/article/pii/S0267364924001234>
- Wang, Y., & Lee, S. (2025). Navigating the AI regulatory landscape: Balancing innovation, ethics, and global governance. *Global Policy*, 16(2), 245–263. <https://www.tandfonline.com/doi/full/10.1080/20954816.2025.2569584>
- WIPO. (2024). Artificial Intelligence and Intellectual Property. <https://www.wipo.int/en/web/frontier-technologies/artificial-intelligence/index>

WIPO. (2024). Intellectual Property and Artificial Intelligence. World Intellectual Property Organization. https://www.wipo.int/about-ip/en/artificial_intelligence/

Yadav, S. (2025). The European Union A.I. Act 2024: Understanding the Context and Exploring its Future. ORF Occasional Paper. <https://www.orfonline.org/research/the-european-union-a-i-act-2024-understanding-the-context-and-exploring-its-future>

Zuboff, S. (2019). The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power. PublicAffairs. <https://www.publicaffairsbooks.com/titles/shoshana-zuboff/the-age-of-surveillance-capitalism/9781610395700/>

Zuiderveen Borgesius, F. (2020). Strengthening legal protection against discrimination by algorithms and artificial intelligence. International Journal of Human Rights, 24(2), 157–176. <https://www.tandfonline.com/doi/full/10.1080/13642987.2020.1743976>

ANEXOS

Anexo 1.- Evidencia de similitud digital



Cinthyia Isabel Ingar Silvestre
“Retos legales frente a la inteligencia artificial: hacia una regulación ética y sostenible”

Titulos
 REVISION 2026
 Universidad Peruana de Ciencias e Informatica

Detalles del documento

Identificador de la entrega
trn:oid::1:3625567858

Fecha de entrega
10 ago 2026, 12:28 p.m. GMT-5

Fecha de descarga
11 ago 2026, 11:25 a.m. GMT-5

Nombre del archivo
TSP_Final-_Ingar_Silvestre_Cinthyia_Isabel.docx

Tamaño del archivo
129.0 KB

69 páginas
12.795 palabras
81.961 caracteres



18% Similitud general

El total combinado de todas las coincidencias, incluidas las fuentes superpuestas, para ca...

Filtrado desde el informe

- Bibliografía
- Texto citado

Fuentes principales

- 20%  Fuentes de Internet
- 5%  Publicaciones
- 8%  Trabajos entregados (trabajos del estudiante)

Marcas de integridad

N.º de alertas de integridad para revisión

No se han detectado manipulaciones de texto sospechosas.

Los algoritmos de nuestro sistema analizan un documento en profundidad para buscar inconsistencias que permitan distinguirlo de una entrega normal. Si advertimos algo extraño, lo marcamos como una alerta para que pueda revisarlo.

Una marca de alerta no es necesariamente un indicador de problemas. Sin embargo, recomendamos que preste atención y la revise.

Fuentes principales

- 20% Fuentes de Internet
- 5% Publicaciones
- 8% Trabajos entregados (trabajos del estudiante)

Fuentes principales

Las fuentes con el mayor número de coincidencias dentro de la entrega. Las fuentes superpuestas no se mostrarán.

1	Internet	repositorio.upci.edu.pe	2%
2	Internet	www.dykinson.com	1%
3	Internet	assets.zyrosite.com	1%
4	Internet	derechoshumanos.mainei.org	1%
5	Internet	helvia.uco.es	<1%
6	Trabajos del estudiante	Universidad Peruana de Ciencias e Informatica	<1%
7	Internet	ruja.ujaen.es	<1%
8	Trabajos del estudiante	Universidad de Guayaquil	<1%
9	Internet	materials.campus.uoc.edu	<1%
10	Internet	bilselkongreleri.com	<1%
11	Internet	www.paginasbrillantesecuador.com	<1%

12	Internet	libros.ucn.cl	<1%
13	Trabajos del estudiante	ITESM: Instituto Tecnológico y de Estudios Superiores de Monterrey	<1%
14	Trabajos del estudiante	Universidad del Rosario	<1%
15	Internet	www.unilim.fr	<1%
16	Internet	revistas.pucp.edu.pe	<1%
17	Internet	ciencialatina.org	<1%
18	Internet	digikogu.taltech.ee	<1%
19	Internet	www.portaleducatiu.ad	<1%
20	Trabajos del estudiante	Universitat Oberta de Catalunya	<1%
21	Internet	djhr.revistas.deusto.es	<1%
22	Internet	www.rilco.org	<1%
23	Internet	repositorio.cetys.mx	<1%
24	Trabajos del estudiante	UNIBA	<1%
25	Internet	repositorio.cepal.org	<1%

26	Internet	www.pensamientopenal.com.ar	<1%
27	Internet	www.personio.es	<1%
28	Internet	bibliotecaeditorialsival.com	<1%
29	Internet	revista-aji.com	<1%
30	Internet	cari.org.ar	<1%
31	Internet	www.procuraduria.gov.co	<1%
32	Internet	institutodeanalistas.com	<1%
33	Trabajos del estudiante	Universidad de Alcalá	<1%
34	Trabajos del estudiante	University College London	<1%
35	Internet	rio.upo.es	<1%
36	Internet	cdn.prod.website-files.com	<1%
37	Internet	proyectoguia.lat	<1%
38	Internet	www.sovedem.com	<1%

Anexo 2.- Autorización de publicación en repositorio



FORMULARIO DE AUTORIZACIÓN PARA LA PUBLICACIÓN DE TRABAJO DE SUFICIENCIA PROFESIONAL O TESIS EN EL REPOSITORIO INSTITUCIONAL UPCI

1.- DATOS DEL AUTOR

Apellidos y Nombres: INGAR SILVESTRE CINTHYA ISABEL
DNI: 42999364 Correo electrónico: _____
Domicilio: _____
Teléfono fijo: _____ Teléfono celular: _____

2.- IDENTIFICACIÓN DEL TRABAJO DE SUFICIENCIA PROFESIONAL O TESIS

Facultad / Carrera: DERECHO Y CIENCIAS POLITICAS
Tipo: Trabajo de Suficiencia Profesional (☒) Tesis (☐)
Título del Trabajo de Suficiencia Profesional / Tesis: "Retos legales frente a la
inteligencia artificial: hacia una regulación ética y sostenible"

3.- OBTENER:

Título Profesional (☒)

4. AUTORIZACIÓN DE PUBLICACIÓN EN VERSIÓN ELECTRÓNICA

Por la presente declaro que el documento indicado en el ítem 2 es de mi autoría y exclusiva titularidad, ante tal razón autorizo a la Universidad Peruana Ciencias e Informática para publicar la versión electrónica en su Repositorio Institucional (<http://repositorio.upci.edu.pe>), según lo estipulado en el Decreto Legislativo 822, Ley sobre Derecho de Autor, Art23 y Art.33.

Autorizo la publicación de mi tesis (marque con una X):

(☒) Sí, autorizo el depósito y publicación total.

(☐) No, autorizo el depósito ni su publicación.

Como constancia firmo el presente documento en la ciudad de Lima, a los

20 días del mes de Julio de 2026

FIRMA


HUELLA