

Reacción de la tutela multinivel de protección de los derechos humanos ante la revolución de la inteligencia artificial: una aproximación desde España

ARTÍCULO EN PRENSA

Reaction of multi-level human rights protection to the artificial intelligence revolution: an approach from Spain

ARTICLE IN PRESS

Manuel Palomares-Herrera 

Universidad Internacional de la Rioja. España

manuel.palomares@unir.net

ORCID: <https://orcid.org/0000-0003-1520-5036>

<https://doi.org/10.18543/djhr.3376>

Fecha de recepción: 24.08.2025

Fecha de aceptación: 18.06.2026

Fecha de publicación en línea: julio 2026

Cómo citar / Citation: Palomares-Herrera, Manuel. 2026. «Reacción de la tutela multinivel de protección de los derechos humanos ante la revolución de la inteligencia artificial: una aproximación desde España». *Deusto Journal of Human Rights*. <https://doi.org/10.18543/djhr.3376>

Sumario: 1. Planteamiento. 2. Afectaciones constitucionales de la IA a los derechos humanos. 3. Alusiones comparadas en torno a la integración jurídica de la IA en clave de derechos humanos. 4. Reacciones genuinas de la tutela multinivel. 4.1. Desde las disposiciones normativas. 4.2. Desde la jurisprudencia y la doctrina. Conclusiones. Referencias.

Resumen: El artículo analiza la reacción del sistema multinivel de protección de los derechos humanos frente a la expansión de la inteligencia artificial, con especial atención al ámbito internacional, europeo y español. A partir de un enfoque jurídico-crítico, se examinan iniciativas normativas como la Resolución A/RES/78/265 de la ONU y el Reglamento (UE) 2024/1689, destacando un modelo regulatorio orientado a la contención del riesgo más que a la promoción tecnológica. El estudio identifica múltiples derechos potencialmente afectados —privacidad, igualdad, empleo o libertad— y advierte sobre los riesgos derivados de la opacidad algorítmica y los sesgos en los sistemas automatizados. Asimismo, se subraya el papel emergente de la

jurisprudencia y la doctrina como mecanismos de reacción ante la insuficiencia normativa. Desde una perspectiva comparada, se evidencia la diversidad de enfoques regulatorios a nivel global. Finalmente, el trabajo concluye que la clave reside en articular un marco jurídico adaptativo que garantice la protección efectiva de los derechos fundamentales en la era digital.

Palabras clave: Derechos humanos, inteligencia artificial, tutela multinivel, derecho constitucional, reacción legal, protección internacional, derechos fundamentales.

Abstract: This article analyzes the reaction of the multi-level human rights protection system to the expansion of artificial intelligence, with particular attention to the international, European, and Spanish contexts. Using a critical legal approach, it examines regulatory initiatives such as UN Resolution A/RES/78/265 and Regulation (EU) 2024/1689, highlighting a regulatory model geared toward risk containment rather than technological advancement. The study identifies multiple potentially affected rights —privacy, equality, employment, and freedom— and warns of the risks stemming from algorithmic opacity and biases in automated systems. It also underscores the emerging role of jurisprudence and legal scholarship as mechanisms for responding to regulatory inadequacies. From a comparative perspective, the diversity of regulatory approaches at the global level is evident. Finally, the article concludes that the key lies in establishing an adaptive legal framework that guarantees the effective protection of fundamental rights in the digital age.

Keywords: Human rights, artificial intelligence, multi-level protection, constitutional law, legal reaction, international protection, fundamental rights.

1. Planteamiento¹

Julio G. Verne, en boca del personaje Cyrus Smith en su obra *La Isla Misteriosa* (1874) afirmaba como ingeniero que “la civilización nunca retrocede, pues la ley de la necesidad siempre fuerza a ir hacia adelante” y es que, con la expansión de nuevas tecnologías, la nueva “ley de la necesidad” es la reacción internacional conjunta si se pretende avanzar con aquel impulso que indicaba aquel literato francés. Pues bien, en el año 2024 la Asamblea General de la Organización de Naciones Unidas (ONU), tras el descontrolado desarrollo tecnológico de herramientas de predictibilidad, de *big data* y con uso de algoritmos aprobó la Resolución A/RES/78/265 titulada “Aprovechar las oportunidades de sistemas seguros y fiables de inteligencia artificial para el desarrollo sostenible”², lo que supuso una histórica reacción donde muestra la preocupación de la sociedad internacional por estas afecciones en donde acota conceptos y reconoce preocupaciones de esta forma:

a efectos de esta resolución, se refieren a sistemas de inteligencia artificial en el ámbito no militar, cuyo ciclo de vida incluye las etapas de prediseño, diseño, desarrollo, evaluación, puesta a prueba, despliegue, utilización, venta, adquisición, explotación y retirada de servicio, y presentan las siguientes características: están centrados en las personas, son fiables, se pueden explicar, son éticos e inclusivos, respetan plenamente la promoción y la protección de los derechos humanos y el derecho internacional, mantienen la privacidad, están orientados al desarrollo sostenible y son responsables, tienen el potencial de acelerar y propiciar los avances hacia la consecución de los 17 Objetivos de Desarrollo Sostenible y el desarrollo sostenible en sus tres dimensiones —económica, social y ambiental— de manera equilibrada e integrada, fomentar la transformación digital, promover la paz, salvar las brechas digitales entre los países y dentro de ellos y promover y proteger el goce de los derechos humanos y las

¹ El presente capítulo ha sido elaborado en el marco del Grupo de Investigación “Derecho y Vulnerabilidad” (VULDER-UNIR) que comenzó en marzo de 2025 y del Proyecto de Innovación Docente Aplicada de la Universidad Internacional de la Rioja (UNIR). Referencia: PIDA 2526_050. “La aplicación de la Inteligencia Artificial a las asignaturas de la Facultad de Derecho de UNIR. Acrónimo: EDUCAIA- Educación en Derecho e Inteligencia Artificial”. 2025-26 y del Proyecto Precompetitivo de Investigación Referencia: PPI26_25. “El “*Compliance*” de la Inteligencia Artificial. La implementación de la Inteligencia artificial en Derecho a la luz de la Agenda 2030”. Acrónimo: CIAD: *Compliance*, Inteligencia Artificial, Derecho, Objetivos de Desarrollo Sostenible, Agenda 2030”, 2025-2027.

² Véase: https://digitallibrary.un.org/record/4043244/files/A_RES_78_265-ES.pdf.

libertades fundamentales para todos, manteniendo a la persona en el centro.

Pues bien, en el presente estudio examinamos esta resolución, las regionales y la interna para asociar la “revolución” (Floridi et al. 2018) de esta nueva tecnología disruptiva de transformación digital (automatismos humanoides, control facial, algoritmos de predicción de conductas en *big data*) con el importante papel que continúa desempeñando el sistema internacional de protección de los derechos humanos y centrado en el del desarrollo de la inteligencia artificial (IA) en nuestros días. Desde luego es una especialidad no cercada a disciplinas como la informática y la programación, sino que depende determinadamente su devenir de las ciencias jurídicas pues, como afirmaba el profesor Cumplido (2024, 73):

Nuestra tarea es determinar los límites que deben tener la inteligencia artificial y la robótica que, en términos generales, deben recoger el derecho constitucional, sin proscribir la investigación, sino ajustándola a los valores sobre la vida, la justicia, la paz, el desarrollo, etc. aceptados en las declaraciones universales sobre los derechos humanos.

Para ello, analizaremos cómo los ordenamientos jurídicos vigentes son afectados por el nuevo escenario, si pueden adaptarse y reformarse a las nuevas exigencias para continuar garantizando el principio de seguridad jurídica y la salvaguarda de derechos humanos —como el del honor, la intimidad o el libre desarrollo de la personalidad— o si necesitan previamente abordarse los desafíos éticos y jurídicos que plantean estas tecnologías emergentes antes de entrar a legislar.

En esta lid, y conscientes de que la regulación internacional va por detrás del progreso tecnológico y del desarrollo normativo interno, se evalúan las iniciativas actuales de organismos internacionales como la ONU o la Unión Europea (UE) y se proponen recomendaciones para fortalecer la protección de los derechos humanos en la era digital. Proponemos para ello algunas preguntas clave: ¿Hasta dónde alcanza la influencia de la A/RES/78/265? ¿pueden los marcos legales existentes adaptarse para abordar los desafíos éticos y jurídicos que plantean la IA y la robótica? ¿Qué nuevos instrumentos o mecanismos pueden ser necesarios para garantizar la protección de los derechos humanos en la era digital? ¿Cómo pueden los organismos internacionales colaborar eficazmente con gobiernos, empresas tecnológicas y la sociedad civil para desarrollar estándares y regulaciones apropiadas?

Para responder a estas preguntas, el estudio adopta un enfoque metodológico interdisciplinario, combinando el comentario jurídico de normativa vigente, documental, de revisión jurisprudencial y de revisión bibliográfica de los últimos referentes doctrinales en la materia para evaluar iniciativas actuales de las organizaciones internacionales y con el objetivo de proponer una reflexión jurídica y una *lege ferenda* para fortalecer la protección de los derechos humanos en este nuevo contexto tecnológico. De forma complementaria, incorporamos una sección de derecho comparado desde donde se puede tomar una perspectiva panorámica que nos sitúe en torno a cómo se está reaccionando por parte de distintos legisladores ante el mismo desafío.

Partimos de que la doctrina aún emergente se debate entre la regulación estricta que propone la implementación de salvaguardias adicionales, como evaluaciones de impacto de protección de datos y rendición de cuentas como Mantelero (2018)³. Por consiguiente, y a modo de estado de la cuestión, hemos de indicar que es un tema de estudio *mainstream* en donde se está derramando mucha tinta actualmente, donde podemos destacar en materia de robotización como parte de la IA el trabajo de Pineros (2024) sobre aplicaciones de IA a la Administración de Justicia, el de Méndez (2024, 82) que aborda la cacareada fiabilidad de los algoritmos en el nuevo paradigma respecto al encaje de la personalidad jurídica de los robots en una nueva “superficie legal” o el de sobre las aristas éticas de la IA desde el *soft law* del momento y sus posibles riesgos a los derechos humanos; no obstante, se facilita de la parte heurística de la investigación una práctica recopilación bibliográfica.

2. Afectaciones constitucionales de la IA a los derechos humanos

Si los derechos humanos se definen —art. 1 de la Declaración Universal de Derechos Humanos (DUDH) de 1948— como exigencias mínimas de dignidad, y entendemos que estos se convierten en fundamentales cuando se encuentran blindados en una constitución *erga homnes*, nos preguntamos si la aparición de una “constitución

³ Muchos expertos abogan por la implementación de medidas estrictas para mitigar el sesgo algorítmico. Esto incluye la auditoría regular de algoritmos, la transparencia en los datos y métodos utilizados, y la inclusión de diversos grupos en el proceso de desarrollo de IA (Binns 2018).

para robots” dirigidos por IA⁴ (Castellanos y Montero 2020) sería la antesala de una “pseudodignidad artificial” que prepare el camino a un estándar por explorar en su origen.

Ya la ciencia ficción coqueteó con ello al tratar las tres leyes de la robótica —o leyes de Asimov— aparecidas por primera vez en el relato “Círculo Vicioso” (*Runaround*) de 1942. Primera ley: Un robot no hará daño a un ser humano, ni por inacción permitirá que un ser humano sufra daño. Segunda ley: Un robot debe cumplir las órdenes dadas por los seres humanos, a excepción de aquellas que entren en conflicto con la primera ley. Tercera ley: Un robot debe proteger su propia existencia en la medida en que esta protección no entre en conflicto con la primera o con la segunda ley.⁵

Trasladado al caso nacional, también había cierta tendencia a acotar lo digital, y es que, cuando en la España de 1977 aún no habían proliferado el disquete y el Macintosh 128K en que se desconocía *scroll* infinito o la guerra de IA entre *OpenAI* y *DeepSeek* (Xu 2025), ya se predecía en nuestro país la futura influencia de la informática⁶ en los derechos humanos, en aquel momento con la intimidad y el honor, pero hoy puede afectar a varios factores que impactan en bienes jurídicos y empeños de los poderes públicos como el entretenimiento (art. 48 CE), la salud mental (art. 43 CE), el empleo (art. 40 CE), la educación (art. 27 CE), la igualdad (art. 14 CE), el libre desarrollo de la personalidad (art. 10 CE), la propiedad intelectual (art. 20 CE), etc.⁷

Como reflexionamos en este epígrafe, estamos en ciernes de una regulación robótica completa, y se está definiendo a pesar de que la AV

⁴ Firmada por el equipo de robótica de *Google DeepMind*. Véase: <https://deepmind.google/discover/blog/shaping-the-future-of-advanced-robotics/>

⁵ A esta pica en Flandes le siguieron los principios de robótica de EPSRC / AHRC, luego las Leyes de Satya Nadella o las leyes de la robótica de Tilden.

⁶ No deja de ser curioso que hace 46 años los padres de la Constitución Española (CE) de 1978 ya previesen de forma pionera (Palomares-Herrera 2024) el impacto *ad futurum* que iba a experimentar el sector tecnológico-digital cuando señala en su art. 18.4: “La ley limitará el uso de la informática para garantizar el honor y la intimidad personal y familiar de los ciudadanos y el pleno ejercicio de sus derechos” y desde luego, esta es la tónica de la ONU y de la UE.

⁷ Reconocidos los principales bienes jurídicos afectados en el art. 12 DUDH y el art. 17 del Pacto Internacional de Derechos Civiles y Políticos de 1966 en acciones como lo es la recopilación masiva de datos, el perfilado y toma de decisiones automatizada como, por ejemplo, para la selección de candidatos para empleos, la vigilancia omnipresente como vigilancias basadas en IA con reconocimiento facial o la inferencia de datos sensibles. *Verbi gratia*, la privacidad, dado que los sistemas de IA pueden analizar grandes volúmenes de datos de vigilancia, lo que permite a gobiernos y empresas rastrear y monitorear usuarios sin su consentimiento.

RES/78/265 de la ONU se refiere a la promoción de la IA, esencialmente; pero es evidente que lo hace desde el escepticismo y no desde el impulso, pues fuerza que esta herramienta evolucione en el canal de la "seguridad y la fiabilidad", por lo que se reconoce el temor de la existencia del uso inadecuado de la misma, y para ello la define y acota de forma transversal.⁸

Es más, ni siquiera la ONU al aprobar la Resolución A/RES/70/1 de 25 de septiembre de 2015 "Transformar nuestro mundo: la Agenda 2030 para el Desarrollo Sostenible"⁹ relaciona la igualdad de las naciones desde el acceso a la IA, pues el ODS 6, por ejemplo, establece como insumo imprescindible el agua, pero evaluemos las grandes diferencias educativas que existen entre un alumnado con acceso a internet e IA respecto a los que no, cuestiones que pueden perpetuar y amplificar sesgos existentes en los datos de entrenamiento de la IA, incluso derivando finalmente en decisiones discriminatorias en áreas de opinión, costumbres, exigencias democráticas o justicia. En este aspecto señalaba Pineros (2024, 75) que:

La configuración constitucional diseña una Justicia hecha por humanos para ser administrada por humanos. La irrupción de la inteligencia artificial en todos los sectores nos lleva a cuestionarnos cómo afectaría a la prestación de servicios públicos.

Esto expone el salto del laboratorio a lo cotidiano, de ahí que existan limitaciones de base creadas en los propios programadores para contemplar futuros conflictos.¹⁰ Lo mismo ocurre con la cuestión laboral en cuanto que la implementación de sistemas de IA en el lugar de trabajo puede llevar a la automatización de empleos que desaparecerían¹¹, lo que se enfrentaría a aquello del deber de trabajar

⁸ Méndez (2024, 84) añade que quizás se llegue tarde a legislar en, al menos, transparencia y justificación que es donde se tenían que haber plasmado todos los esfuerzos.

⁹ Véase: chrome-extension://efaidnbmninnibpcapjcgclcfndmkaj/https://unctad.org/system/files/official-document/ares70d1_es.pdf

¹⁰ Partamos de la base de que la delegación de decisiones críticas a sistemas de IA puede socavar la autonomía de los individuos si no tienen el control o la capacidad de impugnar dichas decisiones porque la IA puede ser utilizada para crear contenido falso o manipular la información (como *deepfakes*), afectando la capacidad de las personas para tomar decisiones informadas y erosionando la confianza pública.

¹¹ "La repercusión que tendrá la introducción de la robótica en los procesos de producción es, controvertida, La diversidad de resultados (del 9-54%) de empleos amenazados, refleja la complejidad de las opciones metodológicas y su diversa repercusión en los resultados de la investigación" (Goñi 2019, 62).

del art. 35 CE, afectando negativamente a los trabajadores y sus medios de vida en donde se perpetuará una suerte de (re)sentimiento “ludista” —del ludismo frente a la primera revolución industrial— contra la “máquina”, si no se implementan políticas adecuadas de transición laboral que eviten sistemas de IA opacos que censuren sistemas de escrutinio social de cómo se toman las decisiones, y que permita una fiscalización pública si se requiere.¹²

Autores como Corvalán (2017) o Valls (2021) conceptualizaban ya a la IA como a sistemas informáticos capaces de realizar tareas que normalmente requieren inteligencia humana, como el aprendizaje, la resolución de problemas y la toma de decisiones y tenía ese sentido como herramienta *ad hoc* que complementase tareas repetitivas pero que, desgraciadamente, ha llegado a otros escenarios que han movilizado a la doctrina por tener mayor implicación humana si lo analizamos desde la definición propuesta por el Grupo de Expertos de Alto Nivel sobre IA de la Comisión Europea (2021)¹³. Pues si la IA se alimenta del aprendizaje indiscriminado, podemos preguntarnos si aprende lo que debe y de quien debe —pues no todo lo aprendido ni toda conducta es deseable ser aprendida o difundida—.

Otro desafío detectado es que los sistemas de IA pueden perpetuar y amplificar discriminaciones existentes o crear nuevas formas de discriminación con sesgos en los datos de entrenamiento, la falta de diversidad en el desarrollo de IA con la subrepresentación de ciertos grupos en los equipos que desarrollan sistemas, la opacidad algorítmica o la discriminación por *proxy*, por lo que observamos se van desarrollando las normas de arriba abajo como veremos en el siguiente epígrafe. Y es que la comunidad internacional¹⁴ solo algunos países

¹² Académicos como Floridi et al. (2018) proponen la creación de organismos de supervisión (en no afectación negativa a los derechos humanos) independientes que auditen y evalúen regularmente los sistemas de IA, pero, según Larrondo y Grandi (2021, 21) esto en el sector privado quien debe de hacerlo es la propia empresa que la use.

¹³ “Sistemas de *software* (y posiblemente también de *hardware*) diseñados por humanos que, dado un objetivo complejo, actúan en la dimensión física o digital percibiendo su entorno mediante la adquisición de datos, interpretando los datos estructurados o no estructurados recopilados, razonando sobre el conocimiento o procesando la información derivada de estos datos y decidiendo las mejores acciones a tomar para lograr el objetivo dado”. Por tanto, abarca diversas técnicas y enfoques, incluyendo: aprendizaje automático (*machine learning*), algoritmos que mejoran automáticamente a través de la experiencia, aprendizaje profundo (*deep learning*) y una subclase del aprendizaje automático basada en redes neuronales artificiales entre otros. Véase: <https://digital-strategy.ec.europa.eu/es/policies/expert-group-ai>

¹⁴ Véase: <https://news.un.org/es/story/2024/03/1528511>

han comenzado a plantar estrategias nacionales como el Plan Nacional de IA de España o ENIA.¹⁵

No todo es negativo, el propósito fundamental del derecho como disciplina es la regulación, no la estigmatización, pues la IA tiene elementos positivos para la garantía de derechos humanos, como mejorar el diagnóstico y tratamiento de enfermedades, facilitando un acceso más equitativo a servicios de salud de calidad, detectando enfermedades como el cáncer en etapas tempranas (lo que puede salvar vidas) (Topol 2019), ayudando a estudiantes con discapacidades (De Asís y Ansuátegui 2020, 11) o aquellos en zonas rurales con acceso limitado a recursos educativos (Holmes et al. 2019), o mejorando los niveles de seguridad doméstica (O'Neil 2016).¹⁶

Por último, el siguiente nivel, cuando la IA se incorpora en mecanismos variados humanoides o no, pasamos a una robótica¹⁷ que salta a otras afectaciones que darían para varias monografías para abordarla en plenitud, pero ya encontramos en la década de 1950 precedentes como el de Alan Turing, la Conferencia de Dartmouth (Hanover) de 1956, los robots *Unimate* de 1961, pasando por el "invierno de la IA" de los 70 a los 80 hasta llegar a los actuales asistentes virtuales y *chatbots* basados en IA¹⁸, vehículos autónomos o los ADAS (sistemas de conducción asistida) que nos demuestran, en definitiva, que la regulación ha sido *a posteriori*.¹⁹

¹⁵ Véase: <https://planderecuperacion.gob.es/noticias/conoce-Estrategia-Nacional-Inteligencia-Artificial-ENIA-IA-prtr>. Aunque también podríamos destacar la normativa española sobre protección de datos tras puesta del reglamento europeo donde el GDPR de 2018 es pieza clave del marco regulatorio que afecta directamente el uso de IA en España pues establece estrictas reglas para la recopilación, procesamiento y almacenamiento de datos personales, elementos esenciales para el funcionamiento de muchos sistemas de IA (Reglamento (UE) 2016/679).

¹⁶ Pero a pesar de estos beneficios, no podemos obviar, *contrario sensu*, que las tecnologías de reconocimiento facial y vigilancia masiva pueden violar el derecho a la privacidad (Crawford y Calo 2016), que puede llevar a decisiones discriminatorias en áreas como el empleo, la justicia penal y los servicios financieros (Noble 2018), o puede generar noticias falsas y los *deepfakes* para engañar y manipular la opinión pública (Chesney y Citron 2019).

¹⁷ Según la Federación Internacional de Robótica, un robot es "un mecanismo programable accionado en dos o más ejes con un grado de autonomía, moviéndose dentro de su entorno, para realizar tareas previstas". Véase: <https://www.udit.es/grado-en-robotica-que-es-que-se-estudia-y-salidas-profesionales-en-2025/>

¹⁸ La oferta de la versión gratuita de los *chats* de AI como la estadounidense *OpenAI (ChatGPT)* y la china *DeepSeek* en 2020 y 2025 respectivamente, ha supuesto el acercamiento doméstico de la nueva tecnología, fechas que coinciden con la eclosión de las reacciones que estudiamos.

¹⁹ Es más, Grigore (2022, 168) sugiere que exista una "moratoria de su uso en espacios públicos, al menos hasta que las autoridades demuestren que no existen

3. Alusiones comparadas en torno a la integración jurídica de la IA en clave de derechos humanos

En un mundo desde años interconectado corresponde realizar una panorámica cenital sobre la forma en la que se aborda la integración jurídica de la IA en países de nuestro entorno inmediato en donde podemos partir de América Latina, que se encuentra en una fase de consolidación progresiva, caracterizada por la recepción —*mutatis mutandis*— del modelo garantista europeo y por el desafío estructural de superar la fragmentación normativa.

En este contexto, Chile se posiciona como referente regional al haber alineado su política pública con los estándares de la UE, destacando el Proyecto de Ley de 2024²⁰ y la institucionalización de una Comisión Asesora en materia ética ante la IA, lo que evidencia una aproximación *ex ante* orientada a la prevención de riesgos. Brasil, por su parte, avanza mediante el Proyecto de Ley 2338/2023²¹, cuyo eje vertebrador es la gobernanza de sistemas automatizados y la tutela de derechos *iura fundamentalia*, configurando un modelo que busca equilibrar innovación y responsabilidad jurídica en un modelo intermedio —*in statu nascendi*— en el que se sitúan Colombia, México o Perú.

Marcamos Colombia pues presenta una elevada densidad legislativa con iniciativas que oscilan entre el fomento de la innovación y la protección de datos personales —Decreto 1263 de 2022 y Ley 1581 de 2012 respectivamente— en donde destaca en materia de IA el “Proyecto de Ley 043 de 2025 por medio del cual se regula la inteligencia artificial en Colombia para garantizar su desarrollo ético y responsable y se dictan otras disposiciones”²². En el caso de México, a través de su Agenda Nacional de Inteligencia Artificial (ANIA)²³, apreciamos que se prioriza la ciberseguridad y la ética aplicada, mientras que Perú ha sido pionero en la promulgación de la Ley

problemas significativos con la precisión o los efectos discriminatorios, y que esos sistemas de IA cumplen con las normas aplicables a la protección de datos y a la privacidad”.

²⁰ Boletín 16821-19, con el proyecto ya en segundo trámite constitucional. Véase: <https://www.minciencia.gob.cl/areas/inteligencia-artificial/Inteligencia-Artificial/Proyecto-Ley-regula-sistemas-IA/>

²¹ N. de la Cámara de los Diputados: PL 2338/2023. Véase: <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233>

²² Véase: https://minciencias.gov.co/sites/default/files/upload/noticias/pl_ia_finalizado.pdf

²³ Sede virtual. Véase: <https://www.ania.org.mx/propuestadeagendanacionaldeia>

31.814²⁴, aunque permanece a la espera de desarrollos reglamentarios que permitan su implementación efectiva *in concreto*, especialmente en sectores estratégicos —por lo que la puesta en marcha puede atrasarse.

En el ámbito de Centroamérica y el Caribe, estimamos como abanderados los legislativos de países como la República Dominicana²⁵ o Costa Rica²⁶ pues han impulsado iniciativas recientes (2023–2026) orientadas a regular el desarrollo e implementación de sistemas inteligentes desde una perspectiva de fomento estatal *pro interventio publica* a la que quieren incorporar la IA que perciben como herramienta, no como algo que limitar, por lo que, en conjunto, el modelo latinoamericano se distingue por su impronta ético-humanista, sustentada en las recomendaciones de la UNESCO, lo que refleja una orientación *pro persona* en la configuración normativa frente a la norma comunitaria.

En el resto de parámetros no puede afirmarse lo mismo, dado que el informe ILIA-CEPAL²⁷ 2025 advierte —*prima facie*— sobre la persistencia de brechas estructurales como para dar eficacia a sus normas, la ausencia de marcos sancionadores suficientemente definidos (*vacatio sanctionis*) y por último, limitaciones en la capacidad institucional para la supervisión efectiva de la IA.

En contraste, el modelo de Estados Unidos se caracteriza por centrarse en la responsabilidad civil y la regulación sectorial, sin llegar aún a un marco penal específico (Lima et al. 2020), pero buscando un enfoque pragmático, que intenta evitar la descentralización (Tapia 2025) y orientado a la innovación donde se observa una dualidad normativa desde la Orden Ejecutiva de 11 de diciembre de 2025 del Presidente²⁸: mientras la administración federal intenta limitar

²⁴ Véase: <https://www.gob.pe/institucion/congreso-de-la-republica/normas-legales/4565760-31814>

²⁵ Informe de recomendaciones sobre el Proyecto de Ley Orgánica que regula los sistemas de Inteligencia Artificial (00495-2025-PLO-SE). Véase: https://procompetencia.gob.do/wpfd_file/informe-de-recomendaciones-sobre-el-proyecto-de-ley-organica-que-regula-los-sistemas-de-inteligencia-artificial-00495-2025-plo-se/

²⁶ En su Asamblea Legislativa, ahora disponibles los siguientes expedientes: Expediente 23.919 (Ley para la promoción responsable de la IA), Expediente 24.875 (Regulación de IA en procesos electorales), Expediente 24.484 (Ley para la implementación de Sistemas de IA en el sector público) y Expediente 23.771 (Regulación de la Inteligencia Artificial) en: https://www.asamblea.go.cr/sd/referencia_cedil/TA_14_2025/Vinculos/8_Exp23919.pdf

²⁷ Índice Latinoamericano de Inteligencia Artificial, véase: <https://www.cepal.org/es/publicaciones/82514-indice-latinoamericano-inteligencia-artificial-ilia-2025>

²⁸ La Orden ejecutiva de 11 de diciembre de 2025 comienza -en su Sección 1— refiriéndose a la precedente Orden Ejecutiva 14179 del 23 de enero de 2025 sobre la

regulaciones estatales consideradas gravosas²⁹ mediante órdenes ejecutivas, normas como la californiana Ley de la Asamblea n. 2885 de IA³⁰ o la Ley SB24-205 de Colorado³¹ que han desarrollado marcos regulatorios propios, cuya entrada en vigor en 2026 los posiciona como estándares *de facto*, todos muy restrictivos, lo que el Presidente Donald J. Trump viene procurando enfrentar para abrir la IA a la innovación en la norma.

Por la típica dicotomía tecnológica, hemos de acudir al Comité de Constitución y Derecho del Congreso Nacional del Pueblo de la República de China pues publicó pioneramente la Ley de Ciberseguridad de China (CSL 网络安全法), vigente desde 2017 y con enmiendas significativas a partir de enero de este 2026 donde se ha desarrollado un modelo regulatorio temprano, abanderado, previsor, prudente, específico y orientado al control estatal de *imperium technologicum*, centrado particularmente en la supervisión de contenidos (Quesada 2025).

Y es que, a diferencia del enfoque europeo, basado en la clasificación general de riesgos, el sistema asiático adopta regulaciones verticales para tecnologías específicas como las *deepfakes* y bulos de IA generativa imponiendo estándares de seguridad que obligan a las empresas a auditar sus datos de entrenamiento, con el objetivo de preservar el *ordo socialis*, trasladando una responsabilidad directa a los proveedores, quienes deben registrar formalmente sus algoritmos. Incluso los Estados más adelantados de la comunidad internacional —incluyendo la UE— fueron a la cola de China como se demostró con el Proceso de Hiroshima sobre la IA del G7 de 2023³² o

“Eliminación de Barreras al Liderazgo Estadounidense en Inteligencia Artificial” y, respecto de ella, dice: “revoqué el intento de mi predecesor de paralizar esta industria y ordené a mi Administración eliminar las barreras al liderazgo de Estados Unidos en IA. Mi Administración ya ha realizado un trabajo tremendo para avanzar en ese objetivo, incluyendo la actualización de los marcos regulatorios federales existentes para eliminar las barreras y fomentar la adopción de aplicaciones de IA en todos los sectores. Estos esfuerzos ya han brindado enormes beneficios al pueblo estadounidense y han dado lugar a billones de dólares de inversiones en todo el país” (Tapia 2025).

²⁹ Ante la inexistencia de una legislación federal integral —*lex generalis*—, estas normativas subnacionales obligan a las empresas a implementar mecanismos estrictos de gestión de riesgos, particularmente en lo relativo a la discriminación algorítmica.

³⁰ O la Ley SB 53 (TFAIA) de seguridad y transparencia, la SB 243 de protección de menores de la IA, o las AB 2602 y AB 1836 de protección de actores, ergo todas restrictivas de la IA. Véase: https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202320240AB2885

³¹ Véase: <https://leg.colorado.gov/bills/sb24-205>

³² Que adoptó once principios rectores internacionales orientados a promover una IA segura, fiable y alineada con valores democráticos *bonum commune*, y que implementó el Código Internacional de Conducta para sistemas avanzados de IA,

cuando se firmó —incluida España— la Declaración de Bletchley durante la Cumbre de Seguridad sobre la IA³³ en Reino Unido en noviembre de 2023.

En el modelo continental —fuera del paraguas comunitario— destaca el modelo germano, que se posiciona como referente con su Estrategia Nacional de Inteligencia Artificial (*KI-Strategie*, 2018, actualizada en 2020)³⁴, complementada por desarrollos normativos sectoriales como la adaptación de la Ley Federal de Protección de Datos (*Bundesdatenschutzgesetz, BDSG*)³⁵ que, aunque no es una única ley codificada, se caracteriza por su enfoque en la responsabilidad jurídica y la trazabilidad algorítmica, así como por la posible extensión de responsabilidad a operadores y desarrolladores (*Organisationsverschulden*) en donde resulta relevante que se hayan incluso planteado la posibilidad de atribuir responsabilidad penal indirecta por decisiones automatizadas a los empresarios (Januário 2023).

Por otro lado, nuestra vecina francesa destaca con una norma que, aunque no es monográfica de IA, sí se refería a algoritmos: la *Loi n. 2016-1321 pour une République numérique*.³⁶ Nos remitimos a esta norma que ya cumple una década pues repercutió en derechos fundamentales dado que albergaba la obligación de transparencia algorítmica en la Administración Republicana, ya que exige que las decisiones administrativas automatizadas sean explicables y accesibles para los ciudadanos, reconociendo el derecho de los ciudadanos a conocer las reglas que rigen los algoritmos utilizados por el Estado —anticipándose a principios que luego recogería el Reglamento Europeo que estudiaremos en el siguiente epígrafe.

Por todo ello, en definitiva, los modelos regulatorios estudiados pueden agruparse en tres bloques con enfoques iusfundamentales

concebido como un instrumento de *soft law* que establece directrices voluntarias para maximizar beneficios y mitigar riesgos, sujeto a revisión permanente. Véase: <https://digital-strategy.ec.europa.eu/es/library/hiroshima-process-international-guiding-principles-advanced-ai-system>

³³ Este documento constituye un compromiso multilateral para fomentar una IA segura, centrada en el ser humano y utilizada de manera responsable, garantizando *inter alia* el respeto a los derechos humanos, la seguridad global y la coherencia con los principios democráticos y de justicia. Véase: <https://www.gov.uk/government/publications/ai-safety-summit-2023-the-bletchley-declaration.es-419>.

³⁴ Véase: <https://www.ki-strategie-deutschland.de/>

³⁵ Véase: https://www.gesetze-im-internet.de/bdsg_2018/

³⁶ NOR: ECFI1524250L. Véase: <https://www.legifrance.gouv.fr/jorf/id/JORFTEXT000033202746>

diferenciados: En el ámbito europeo y su esfera de influencia (incluida América Latina), predomina un enfoque pro persona, donde la IA se somete a límites estrictos derivados de derechos fundamentales como la privacidad, la igualdad y la tutela judicial efectiva. Un segundo modelo estadounidense, donde la protección de derechos fundamentales se articula de forma más fragmentaria, centrada en la no discriminación y el debido proceso, especialmente en contextos de decisiones automatizadas y que, aunque menos sistemático, refuerza garantías como el derecho a impugnar decisiones algorítmicas (*audi alteram partem*), con un énfasis creciente en el *compliance* empresarial. Y un modelo chino que prioriza la estabilidad social y el control estatal sobre la protección individual, subordinando los derechos fundamentales a objetivos colectivos —que, si bien incorpora exigencias de seguridad y veracidad, la protección iusfundamental es más limitada y condicionada. En conjunto, podemos defender también la existencia de una tensión global entre innovación tecnológica y protección de derechos, donde Europa lidera un paradigma garantista que podría consolidarse como estándar internacional.

4. Reacciones genuinas de la tutela multinivel

4.1. Desde las disposiciones normativas

Hilando con el fin del anterior epígrafe, ofrecemos una identificación de reacciones del sistema de tutela multinivel de los derechos humanos, reacción que es necesaria pues, como diría Goñi (2019, 71):

Si no lo hacemos así pronto llegarán los desencantos y pasaremos de la utopía digital a la distopía. El espejismo “tecnoutópico”, sin unos valores y principios éticos que protejan a los individuos y a los grupos potencialmente vulnerables, nos puede llevar a una sociedad indeseable en sí misma.

Por ello, podemos ofrecer una concatenación de despliegues de reacciones normativas desde el sistema multinivel de protección de los derechos humanos.³⁷ En este sentido, hemos de comenzar por la ONU,

³⁷ Entendido como el conjunto de criterios instituciones, mecanismos y normas diseñados a nivel internacional, regional (comunitario) y nacional para promover y salvaguardar los derechos fundamentales V.G, en el contexto de la IA y la robótica, son particularmente relevantes los derechos relacionados con la privacidad —como

pues como hemos reseñado en su resolución sobre la IA, conminó a todos los Estados miembros y partes interesadas a que:

se abstengan de utilizar sistemas de inteligencia artificial que no puedan funcionar de conformidad con las normas internacionales de derechos humanos o que planteen riesgos indebidos para el disfrute de los derechos humanos.

Y en el mismo sentido han ido desplegándose las normas internas, más desde la contención que desde la promoción. Aparte de reacciones de organizaciones internacionales que se han pronunciado en torno a la IA como la ONU³⁸, el Consejo de Europa³⁹, la OCDE⁴⁰ o incluso la Unión Africana⁴¹, queremos destacar la UNESCO como antecedente no vinculante, pero sí relevante por abordar la IA con uno de los mayores consensos con la “Recomendación sobre la ética de la inteligencia artificial de 2021 (SHS/BIO/PI/2021/1)”⁴² por ser adoptado por 193 Estados miembros que se comprometen a que en la implementación de las IA nacionales públicas o privadas imperen principios fundamentales como la transparencia (Palomares-Herrera 2017 y 2022) y la equidad, recordando siempre la importancia de la supervisión humana de los sistemas de IA.⁴³

Resolución 68/167 de la Asamblea General de la ONU sobre el derecho a la privacidad en la era digital (2013)— la no discriminación, la libertad de expresión, los derechos laborales, y el derecho a un recurso efectivo.

³⁸ Respecto a la cuestión de estudio encontramos desde la constitución de un Relator Especial de Naciones Unidas sobre el derecho a la privacidad, cuyo mandato incluye examinar los desafíos planteados por las nuevas tecnologías, la publicación del informe “El derecho a la privacidad en la era digital” (2018), que analiza las implicaciones de la IA para la privacidad y la protección de datos o la adopción de la Recomendación sobre la Ética de la Inteligencia Artificial (2021), el primer instrumento normativo global sobre la ética de la IA.

³⁹ Constituyendo un Comité *ad hoc* sobre IA (CAHAI), publicando el Convenio 108+ sobre Modernización del Convenio para la protección de las personas con respecto al tratamiento automatizado de datos de carácter personal, abordando los desafíos de la IA, o con el Plan para la Eficacia de la Justicia (CEPEJ) en 2018, establece principios éticos para el uso de la IA en el ámbito judicial.

⁴⁰ Con los Principios de la OCDE sobre IA (2019), el primer conjunto de principios intergubernamentales sobre IA y la creación del Observatorio de Políticas de IA (OECD. AI) para monitorear y analizar el desarrollo de la IA a nivel global.

⁴¹ Desarrollando una Estrategia de Transformación Digital para África, que incluye consideraciones sobre el uso ético de la IA.

⁴² Véase: https://unesdoc.unesco.org/ark:/48223/pf0000381137_spa

⁴³ Méndez (2024, 81) abunda en la importancia de la supervisión humana y el mantener una rienda de contención permanente en la IA, dado que sus progresos “surgen a partir de las decisiones humanas involucradas en el proceso”.

A nivel español existen normas aplicables a la IA con previsión constitucional con anclajes en el citado precepto anteriormente junto con otras mutaciones constitucionales posibles con normas como la Ley Orgánica 7/2021, de 26 de mayo, de protección de datos personales tratados para fines de prevención, detección, investigación y enjuiciamiento de infracciones penales y de ejecución de sanciones penales⁴⁴ actualmente, pero donde existe un “Anteproyecto de Ley para el buen uso y la gobernanza de la inteligencia artificial”⁴⁵ que desarrolla el régimen sancionador y de gobernanza previsto en el Reglamento (UE) 2024/1689 del Parlamento Europeo y del Consejo, de 13 de junio de 2024, por el que se establecen normas armonizadas en materia de inteligencia artificial —primera norma vinculante de la historia en materia de IA⁴⁶— que coincide en su art. 1 en el concepto de “fiabilidad” con la resolución de la ONU.

Este Reglamento Europeo remanece de una propuesta de la Comisión de la UE de 2021 titulado *Artificial Intelligence Act* o *AI Act*, que proponía *ab initio* un marco legal integral para el desarrollo y uso de la IA, y posteriormente se ha mantenido, como lo es el basarse en tres pilares principales: clasificación de riesgos en diferentes categorías (inaceptable, alto, limitado y mínimo)⁴⁷; los requisitos para los sistemas de alto riesgo donde los desarrolladores y usuarios de sistemas de IA de alto riesgo deben cumplir con una serie de requisitos, incluidos la gestión de riesgos, la transparencia, la supervisión humana y la garantía de seguridad; y el de la prohibición de prácticas de IA inaceptables como los sistemas de puntuación social utilizados por gobiernos.

A pesar de llevar meses aprobado el Reglamento podemos analizar que la clasificación de riesgos es un componente crucial, ya que determina el nivel de supervisión y los requisitos aplicables a los diferentes sistemas de IA donde los sistemas de riesgo inaceptable —como los que manipulan el comportamiento humano de manera

⁴⁴ BOE n. 126, de 27/05/2021.

⁴⁵ Véase: https://avance.digital.gob.es/_layouts/15/HttpHandlerParticipacionPublicaAnexos.ashx?k=19128

⁴⁶ DOUE n. 1689, de 12 de julio de 2024.

⁴⁷ Antecedente del reglamento en el sistema comunitario se aprobó su *soft law* con el Libro Blanco de la UE para la IA de 2020 que de la misma forma ya identificaba dos niveles de riesgo en las aplicaciones de IA (que no de las decisiones automatizadas en general) que luego han sido más, algo que se está desarrollando al alza a la vista del caso alemán en donde su Comisión Ética de Datos lo amplió a 5 categorías (Grigore 2022, 173).

perjudicial— están prohibidos⁴⁸ y los sistemas de alto riesgo deben cumplir con requisitos estrictos, lo que incluye la realización de evaluaciones de conformidad antes de su despliegue.

A continuación, podemos destacar desde su estudio esa visión preservadora o limitadora de la IA a la vista de una serie de elementos que incorpora como que los sistemas de alto riesgo están sujetos a una serie de requisitos diseñados quizás para mitigar los riesgos asociados desde su gestión —pues ahora los desarrolladores deben implementar medidas para identificar y mitigar los riesgos potenciales asociados—, primando la transparencia (los usuarios deben ser informados de que están interactuando con un sistema de IA, y se debe proporcionar información clara sobre su funcionamiento)⁴⁹, la supervisión humana (deben ser supervisados por humanos para asegurar que las decisiones tomadas por la IA sean apropiadas y justas), y la seguridad (los desarrolladores deben garantizar que los sistemas de IA sean seguros y funcionen según lo previsto).⁵⁰

Por último, es necesario hacer referencia al nivel autonómico dentro del sistema de tutela legislativa multinivel que resulta aplicable en materia de IA. En este ámbito, destaca la iniciativa de la Comunidad Autónoma de Galicia, cuyo Parlamento tramitó el Proyecto de Ley para el Desarrollo e Impulso de la Inteligencia Artificial en Galicia (12/PL-000007)⁵¹. Esta iniciativa culminó con la aprobación de la ley en marzo de 2025.

La norma tiene como finalidad promover el desarrollo, la innovación y la implantación de la inteligencia artificial en la comunidad autónoma. Entre sus principales medidas se encuentran la creación de órganos y entidades de supervisión y gobernanza de la inteligencia artificial, así como el impulso de un ecosistema tecnológico mediante la creación de las denominadas «factorías de inteligencia artificial», concebidas como polos de investigación, desarrollo e innovación vinculados al Centro de Supercomputación de

⁴⁸ El reglamento prohíbe prácticas de IA que se consideran inaceptables. Esto incluye sistemas de vigilancia masiva indiscriminada y sistemas que explotan vulnerabilidades de grupos específicos, como menores de edad o personas con discapacidades.

⁴⁹ Esto implica la divulgación completa de los algoritmos y los datos utilizados (Pasquale 2015).

⁵⁰ Estos mecanismos son esenciales para prevenir abusos y garantizar que la tecnología beneficie a la sociedad en su conjunto (European Commission 2021).

⁵¹ Véase: <https://www.es.parlamentodegalicia.es/ParticipacionCidada/LexislaConNos/3432/proxecto-de-lei-para-o-desenvolvemento-e-impulso-da-intelixencia-artificial-en-galicia-12pl-000007>.

Galicia (CESGA). Con ello, Galicia se posiciona como una de las primeras comunidades autónomas españolas en dotarse de un marco normativo específico para fomentar el desarrollo responsable de esta tecnología.

4.2. Desde la jurisprudencia y la doctrina

A pesar de su consideración como pseudofuentes, la jurisprudencia y la doctrina son aquí los gurús y guardianes del orden al no existir una normativa consolidada, por lo que han generado la reacción *ex novo* en la materia. Podemos destacar como referentes jurisprudenciales aperturistas y cercanos el *Caso Google España SL, Google Inc. contra Agencia Española de Protección de Datos (AEPD) y M. Costeja González* (2014)⁵², el *Caso Facebook Ireland y Schrems* (2020)⁵³ o incluso la STJUE ECLI:EU:C:2023:957 de 2023 sobre las afectaciones de la IA al art. 22 del Reglamento General de Protección de Datos (RGPD-6637/2016) relativo a las decisiones automatizadas, a las que han seguido el resto de pronunciamientos.

Si continuamos con respuestas judiciales en asuntos de IA y sus derivadas, hemos de destacar el *Caso Cambridge Analytica*⁵⁴ que ilustra los riesgos asociados con el uso indebido de datos personales y algoritmos para manipular la opinión pública y socavar procesos democráticos que se relaciona con un caso nacional donde la AEPD sancionó al BBVA por el uso de sistemas de IA para evaluar la solvencia crediticia de sus clientes sin el consentimiento explícito de estos. La Resolución PS/00070/2019 de 15 de diciembre de 2020 de la AEPD destacó la importancia del consentimiento informado y la transparencia en el uso de sistemas de IA que procesan datos personales.

⁵² En aquella sentencia, la ECLI: ECLI:EU:C:2014:317, se parte el conocido como “derecho al olvido”, que no trata directamente sobre IA, sí tiene implicaciones significativas para el uso de datos personales en sistemas automatizados pues el TJUE falló a favor del derecho de los individuos a solicitar la eliminación de sus datos personales de los motores de búsqueda, estableciendo un precedente importante para la protección de la privacidad en la era digital.

⁵³ En la ECLI: ECLI:EU:C:2020:559 el TJUE invalidó el acuerdo de transferencia de datos entre la UE y Estados Unidos, conocido como *Privacy Shield*. La decisión subrayó la importancia de proteger los datos personales de los ciudadanos europeos y la necesidad de establecer salvaguardias adecuadas cuando se transfieren datos a países fuera de la UE lo que tiene implicaciones directas para las empresas de IA que manejan grandes volúmenes de datos personales.

⁵⁴ Véase: <https://www.bbc.com/mundo/noticias-43472797>.

En el ámbito europeo, pero del Tribunal Europeo de Derechos Humanos (TEDH), destaca el caso *López Ribalda y otros c. España (Demandas n. 1874/13 y 8567/13)*⁵⁵ donde establece —*ratio decidendi*— que el uso de tecnologías digitales de vigilancia —incluyendo los que procedan de IA— deben ajustarse a los principios de proporcionalidad y transparencia, garantizando la protección de la vida privada⁵⁶, que era la cuestión central.

Al final el TEDH, en su análisis, estableció que toda medida tecnológica de vigilancia debe superar un test de proporcionalidad estricto, en particular, en los párrafos 116–122, el Tribunal desarrolla criterios que deben ponderarse: la existencia de una sospecha razonable, el alcance de la medida, su duración y el grado de intrusión en la privacidad del trabajador. Asimismo, en el párrafo 134, se concluye que, en este caso concreto, no se produjo vulneración del art. 8 CEDH debido a la gravedad de las sospechas y la limitación temporal de la medida; razonamiento esencial para la IA, en tanto que los sistemas automatizados de monitorización deben respetar estos mismos estándares de proporcionalidad, transparencia y necesidad (*principium proportionalitatis*).

Otra sentencia clave es *S. y Marper c. Reino Unido (Gran Sala, 2008)*⁵⁷, que, si bien no trata directamente la IA, aborda el almacenamiento masivo de datos biométricos —huellas dactilares y perfiles de ADN— por parte del Estado, y es que unos demandantes que habían sido absueltos supieron que sus datos permanecieron en bases policiales por lo que el TEDH consideró que esta retención indefinida vulneraba el art. 8 CEDH. En los párrafos 103–105, se subraya que la mera conservación de datos personales sensibles constituye una injerencia en la vida privada, incluso sin uso activo. Más relevante aún, en los párrafos 112–125, se enfatiza la necesidad de límites claros, previsibilidad normativa y garantías contra el uso arbitrario de datos, lo que destaca en materia de IA, pues los sistemas dependen de grandes volúmenes *big data*, estableciendo que la acumulación y tratamiento automatizado deben regirse por criterios estrictos de legalidad y proporcionalidad (*nulla ingerentia sine lege*).

Una tercera resolución de gran impacto es el caso *Big Brother Watch y otros c. Reino Unido (Gran Sala, 2021)*, relativa a los

⁵⁵ Los hechos del caso se sitúan en un supermercado donde el empleador instaló cámaras ocultas para detectar irregularidades en caja, lo que derivó en despidos.

⁵⁶ Art. 8 del Convenio Europeo de Derechos Humanos de 1950 (CEDH).

⁵⁷ Véase: <https://www.cepc.gob.es/publicaciones/revistas/revista-de-derecho-comunitario-europeo/numero-33-mayoagosto-2009/>

programas de vigilancia masiva de comunicaciones electrónicas por parte de los servicios de inteligencia británicos donde unos demandantes alegaron como en los casos anteriores, que estos programas vulneraban los arts. 8 y 10 CEDH y el TEDH reconoce que los sistemas de interceptación masiva implican el uso de tecnologías avanzadas de filtrado y análisis automatizado de datos, lo que los aproxima funcionalmente a sistemas de IA.

En los párrafos 347–350, se establece que los regímenes de vigilancia deben contar con “garantías de extremo a extremo” o *end-to-end safeguards*, incluyendo control independiente, criterios claros de selección de datos y mecanismos de supervisión efectiva y el párrafo 387, concluye que el régimen británico vulneraba el CEDH por la insuficiencia de estas garantías, lo que resulta fundamental para la IA, ya que subraya la necesidad de transparencia, *accountability* y control institucional en sistemas automatizados de gran escala.

De la revisión conjunta de estas sentencias se desprenden varios principios estructurales aplicables al uso de IA en contextos jurídicos y administrativos como el principio de proporcionalidad que actúa como eje central en la evaluación de cualquier tecnología que afecte derechos fundamentales. En segundo lugar, la exigencia de previsibilidad normativa (*lex certa*) implica que los ciudadanos deben conocer las condiciones bajo las cuales sus datos pueden ser tratados por sistemas automatizados y, en tercer lugar, la necesidad de supervisión humana efectiva se configura como garantía indispensable frente a la opacidad algorítmica.

El TEDH insiste en que el uso de tecnologías avanzadas no puede erosionar el núcleo esencial de los derechos fundamentales; aunque la IA ofrece importantes ventajas en términos de eficiencia y capacidad analítica, su implementación debe someterse a un marco jurídico riguroso que garantice el respeto de los principios del Estado de Derecho.

Finalmente, y abordando las primeras reacciones de la literatura científica especializada, destacar otro aspecto crucial que también constituye una reacción del ordenamiento como es la concienciación sobre los riesgos y beneficios de la IA tanto para los desarrolladores como para que los usuarios estén informados sobre las implicaciones éticas y legales de la IA, ahí sería conveniente la implementación de programas de formación para jueces, abogados y otros profesionales del derecho, principales guardadores ante desafíos legales de la IA o asegurar el consentimiento informado en materia de entrenamiento de IA en RRSS promoviendo la alfabetización digital entre la población (Chesney y Citron 2019).

Y es que esta reconfiguración social mediada por la tecnología plantea inquietudes fundamentales sobre la distribución del poder, el acceso al conocimiento y la formación de nuevas desigualdades (Goñi 2019, 64) pero también ciertas oportunidades apetitosas (De Asís y Ansuátegui 2020, 13) pues como advierte Zuboff (2019), estamos ante el surgimiento de un “capitalismo de vigilancia”, *ergo* supervisión, donde el imperativo de predicción declara su derecho a modificar el comportamiento humano con el fin de producir resultados comerciales y gubernamentales negociados por actores complejos (Foucault 1976).⁵⁸ Si aparte de las aplicaciones comerciales y privadas de la IA, pasamos a su afectación al sector público, ya tenemos reacciones doctrinales por la tan temida aplicación al estrado con ese “juez artificial” con el que se es tan escéptico a la vista de Pineros (2024, 76-77):

la utilización de sistemas de IA en el ámbito de las decisiones judiciales optimizaría sin duda, el funcionamiento de la Justicia, garantizando una tutela judicial más rápida, accesible. Pero, creemos, que su implementación en este momento, con la tecnología actual, y la configuración constitucional de la función jurisdiccional, no debe permitirse, simplemente no es posible. Podríamos ser optimistas, y afirmar que se trata de un futuro cercano, en el que seremos capaces de garantizar una tutela judicial verdaderamente efectiva, sin fisuras. ¿Quién sabe? Los técnicos lo ven posible, los juristas no tanto.

Paradójicamente⁵⁹, los primeros en incorporar la IA han sido letrados a los que la jurisprudencia también ha tenido que reaccionar, pues ya encontramos abogados que incorporan esta herramienta, no siempre con fines probos, pues como aquellos alumnos que ya torticeramente plagian sus TFT⁶⁰ con IA, ha habido abogados

⁵⁸ El reto, por tanto, no es simplemente crear reglas para controlar la IA, sino desarrollar un ecosistema regulatorio que sea tan adaptativo y reflexivo como la tecnología que busca gobernar un ecosistema que debe fundamentarse en un diálogo continuo entre disciplinas, desde informática, derecho, hasta filosofía y sociología, reconociendo que la regulación de la IA es, en última instancia, un ejercicio de definición y redefinición de lo que significa ser humano en la era digital.

⁵⁹ Un caso paradigmático se produjo en el Tribunal Superior de la provincia de Chubut (Argentina), donde la Cámara Penal de Esquel anuló una sentencia al constatar que el juez de primera instancia había utilizado un sistema de IA para redactar parcialmente el fallo sin supervisión humana suficiente ni transparencia. Véase: <https://www.palabrasdelderecho.com.ar/articulo/6613/Revocaron-la-sentencia-que-habia-anulado-una-condena-penal-por-uso-indebido-de-inteligencia-artificial>.

⁶⁰ Detectamos el problema de que los “antídotos” para detectar el uso de IA en textos son aún débiles tanto en lo académico como en lo profesional, dado que incluso la

amonestados por su uso –concretamente, el Auto TSJNA n. 2/2024, de fecha 4 de septiembre de 2024 (Rec. n. 17/2024), que aborda un caso en el que un abogado utilizó información errónea proporcionada por *ChatGPT* en la redacción de una querrela.⁶¹ El mismo mes el TC publicó la Nota Informativa n. 90/2024 exponiendo que la Sala Primera del TC por unanimidad sanciona a un abogado por la falta del debido respeto al tribunal, al incluir reiteradas citas de doctrina que se entrecomillaban como si fueran reales cuando en realidad no existían.⁶²

Por tanto, ofrecida ya una aproximación a la panorámica de la tutela multinivel dispositiva, jurisprudencial y doctrinal en torno a la IA, planteamos futuras líneas de investigación que analicen individualmente el Reglamento Europeo, las antinomias entre la normativa nacional y autonómica al respecto, un estudio comparado comunitario para conocer las velocidades a las que se despliega el futuro derecho de la IA y si está generando un derecho digital de segunda generación en la jurisprudencia que se vaya publicando.

Conclusiones

Una vez realizada la investigación propuesta, habiendo preparado la argumentación que sostiene las inquietudes respondidas a las cuestiones inicialmente propuestas podemos concluir:

Primera.- Que existe una reacción jurídica clara, aunque todavía en fase de consolidación, frente al impacto de la IA en los derechos

sofisticada plataforma *Turnitin* para detectar el plagio es vulnerable al uso de IA, en donde las propias plataformas *chatbot* están preparando sus propios detectores para proteger la honestidad, la propiedad intelectual, el mérito y la capacidad en los trabajos universitarios, si bien debiera construirse una herramienta gratuita de detección de la IA o una suerte de aquella “fortaleza digital” anti *translatr* que escribía Dan Brawn en aquella novela.

⁶¹ “TERCERO.- Utilización de tecnologías emergentes y buena fe procesal. El uso de las tecnologías emergentes y de los materiales generados por inteligencia artificial en los procedimientos judiciales no está exento de importantes consideraciones éticas y legales para garantizar un uso responsable. Lo que impone una verificación adicional, puesto que la revisión y validación de los documentos legales seguirá siendo responsabilidad de los abogados para garantizar la precisión y el cumplimiento normativo. [...] Ante la evidencia del error cometido, el querellante se ha apresurado para presentar sus “más sinceras excusas” tras haber detectado la Sala -como expusimos en el auto de inadmisión de querrela- que la cita textual del precepto penal que incluía en el texto de la querrela no correspondía en realidad a nuestro Código Penal sino al Código Penal de la República de Colombia”. Véase: <https://www.poderjudicial.es/search/AN/openDocumento/1f95323f61e2b4cfa0a8778d75e36f0d/20241108>

⁶² Véase: https://www.tribunalconstitucional.es/NotasDePrensaDocumentos/NP_2024_090/NOTA%20INFORMATIVA%20N%C2%BA%2090-2024.pdf

humanos. Esta respuesta se articula en un sistema multinivel —internacional, europeo y nacional— que opera bajo una lógica predominantemente preventiva *ex ante* pues instrumentos como la resolución de Naciones Unidas o el Reglamento europeo de IA reflejan una tendencia hacia la gestión del riesgo, la transparencia y la supervisión humana. No obstante, este modelo presenta un sesgo conservador, priorizando la contención frente a la promoción de la innovación lo que evidencia una tensión estructural entre el desarrollo tecnológico y la garantía de derechos fundamentales, que obliga a diseñar marcos regulatorios dinámicos capaces de evolucionar al ritmo de la tecnología (*lex adaptativa*).

Segunda.- Que la jurisprudencia —especialmente del TEDH y del TJUE— está desempeñando un papel estructural como fuente material del Derecho en materia de IA ante la insuficiencia normativa (*vacatio legis tecnológica*) a la vista de que los tribunales han desarrollado principios clave como la proporcionalidad, la transparencia y la protección de datos. Casos como *López Ribalda, S.* y *Marper o Big Brother Watch* constituyen precedentes que, aunque no aborden directamente la IA, establecen estándares plenamente aplicables a sistemas algorítmicos. Esta función jurisprudencial no solo es interpretativa, sino también creadora (*ius praetorium digital*), configurando un marco garantista que condiciona la actuación de legisladores y operadores jurídicos.

Tercera.- Que, desde una perspectiva de derecho comparado, se observa una notable heterogeneidad en los modelos regulatorios de la IA a nivel global. Mientras la UE adopta un enfoque garantista basado en la clasificación de riesgos, Estados Unidos opta por un modelo descentralizado y orientado a la innovación, y China desarrolla un sistema de control estatal intensivo. América Latina, por su parte, se encuentra en una fase de recepción y adaptación de estándares internacionales. Esta diversidad evidencia la ausencia de una *lex universalis* en materia de IA, lo que genera riesgos de fragmentación normativa (*fragmentatio legis digitalis*) y asimetrías regulatorias, por lo que resulta imprescindible avanzar hacia mecanismos de armonización internacional, al menos en estándares mínimos de protección de derechos humanos, para evitar vacíos legales y conflictos de jurisdicción.

Cuarta.- Que la incorporación de la IA en la Administración de Justicia plantea un riesgo significativo de erosión de garantías procesales, pues casos como el de la Cámara Penal de Esquel evidencian que el uso no supervisado de sistemas automatizados puede vulnerar principios esenciales como la motivación de las resoluciones, la independencia

judicial y el derecho de defensa donde la automatización acrítica de decisiones compromete la esencia de la función jurisdiccional (*jurisdictio humana*), que exige un juicio valorativo y contextual imposible de replicar plenamente por sistemas algorítmicos. En este sentido, la IA debe concebirse como herramienta auxiliar, nunca sustitutiva, reforzando la necesidad de mantener el principio de supervisión humana efectiva. Por ello, se pone de relieve que la regulación de la IA no puede abordarse exclusivamente desde el Derecho, sino que requiere un enfoque interdisciplinar que integre perspectivas tecnológicas, éticas y sociales. La complejidad de los sistemas de IA —caracterizados por su opacidad (*black box*) y capacidad de aprendizaje autónomo— exige mecanismos preventivos como evaluaciones de impacto, auditorías algorítmicas y formación especializada de operadores jurídicos. Asimismo, resulta esencial promover la alfabetización digital y el consentimiento informado, especialmente en contextos de uso masivo de datos. En última instancia, la regulación de la IA constituye un proceso de redefinición del propio concepto de dignidad humana en la era digital (*dignitas digitalis*), lo que exige un equilibrio constante entre innovación y protección de derechos.

Referencias

- Asís, Rafael de y Francisco J. Ansuátegui. 2020. «Inteligencia artificial y derechos humanos». *Materiales de Filosofía del Derecho* 4: 1-20. Acceso el 3 de julio 2026: https://www.academia.edu/65116371/Inteligencia_artificial_y_derechos_humanos
- Castellanos, Jorge y M^a Dolores Montero. 2020. *Perspectiva constitucional de las garantías de aplicación de la inteligencia artificial: La ineludible protección de los derechos fundamentales*. Córdoba: Universidad de Córdoba. Acceso el 3 de julio 2026: <https://helvia.uco.es/handle/10396/29229>
- Chesney, Robert y Danielle K. Citron. 2019. «Deepfakes and the new disinformation war». *Foreign Affairs* 98(1): 147–155. Acceso el 3 de julio 2026: <https://dialnet.unirioja.es/servlet/articulo?codigo=7172827>
- Comisión Europea. 2021. *Fostering a European approach to artificial intelligence*. Bruselas: Comisión Europea. Acceso el 3 de julio 2026: <https://digital-strategy.ec.europa.eu/es/policies/european-approach-artificial-intelligence>
- Corvalán, Juan G. 2017. «La primera inteligencia artificial predictiva al servicio de la Justicia: Prometea». *La Ley* 81(186): 1–8. Acceso el 3 de julio 2026: <https://www.pensamientopenal.com.ar/doctrina/47747-primera-inteligencia-artificial-predictiva-al-servicio-justicia-prometea>

- Crawford, Kate y Ryan Calo. 2016. «There is a blind spot in AI research». *Nature* 538(7625): 311–313. Doi: 10.1038/538311a.
- Cumplido, Francisco. 2024. «Derechos humanos en la Constitución: Pasado y futuro». *Revista de Derecho Público* 76: 63–84. Acceso el 3 de julio 2026: <https://revistaderechopublico.uchile.cl/index.php/RDPU/article/view/35367>
- European Commission. 2021. *Proposal for a regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union Legislative Acts (COM(2021) 206 final)*. Bruselas: European Commission. Acceso el 3 de julio 2026: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52021PC0206>
- Floridi, Luciano, Josh Cows, Monica Beltrametti, Raja Chatila, Patrice Chazerand, Virginia Dignum y Christoph Luetge. 2018. «AI4People—An ethical framework for a good AI society: opportunities, risks, principles, and recommendations». *Minds and Machines* 28(4): 689–707. Doi: 10.1007/s11023-018-9482-5
- Foucault, Michel. 1976. *Histoire de la sexualité I: La volonté de savoir*. París: Gallimard.
- Goñi, José L. 2019. «Innovaciones tecnológicas, inteligencia artificial y derechos humanos en el trabajo». *Documentación Laboral* 117: 57–72.
- Holmes, Wayne, Maya Bialik y Charles Fadel. 2019. *Artificial Intelligence in education: Promises and implications for teaching and learning*. Boston: Center for Curriculum Redesign.
- Janurário, Tiago F.X. 2023. «Inteligencia artificial y responsabilidad penal de personas jurídicas: Un análisis de sus aspectos materiales y procesales». *Estudios Penales y Criminológicos* 44 (extra): 1–30. Doi: 10.15304/epc.44.8910
- Larrondo, Manuel E. y Nicolás M. Grandi. 2021. «Inteligencia artificial, algoritmos y libertad de expresión». *Universitas-XXI* 34: 177–194. Doi: 10.17163/uni.n34.2021.08
- Lima, Guilherme, Meeyoung Cha, Chihyung Jeon y Kyomin Park. 2020. «The conflict between people's urge to punish AI and legal systems». *arXiv*. Doi: 10.48550/arXiv.2003.06507
- Mantelero, Alessandro. 2018. «AI and big data: A blueprint for a human rights, social and ethical impact assessment». *Computer Law & Security Review* 34(4): 754–772. Doi: 10.1016/j.clsr.2018.05.017
- Méndez, Mª del Mar. 2024. «Derechos fundamentales y personalidad jurídica de los robots». *Derecho Privado y Constitución* 44: 21–59. Doi: 10.18042/cepc/dpc.44.01
- Noble, Safiya U. 2018. *Algorithms of oppression: How search engines reinforce racism*. New York: New York University Press. Doi: 10.18574/9781479837243
- O'Neil, Cathy. 2016. *Weapons of math destruction: How big data increases inequality and threatens democracy*. New York: Crown.
- Palomares-Herrera, Manuel. 2017. *Estado de la transparencia: El emergente derecho fundamental de acceso a la información pública*. Tesis doctoral,

- Universidad de Jaén. Acceso el 3 de julio 2026: <https://hdl.handle.net/10953/860>
- Palomares-Herrera, Manuel. 2022. *Fundamentos de la transparencia y el acceso a la información*. Madrid: Dykinson.
- Palomares-Herrera, Manuel. 2024. *Constitución Española comentada: Especial 45º aniversario*. Madrid: Dykinson.
- Píneros, Elena. 2024. «El juez-robot y su encaje en la Constitución Española». *Estudios de Deusto* 72(1): 53–78. Doi: 10.18543/ed.3100
- Quesada, Jesús. 2025. «El plan de China para liderar la regulación mundial sobre IA». *National Geographic España*. Acceso el 3 de julio 2026: https://www.nationalgeographic.com.es/tecnologia/china-quiere-liderar-regulacion-ia-tendra-exito-su-objetivo_26892
- Tapia, Alberto J. 2025. «Inteligencia artificial: Orden ejecutiva de 11 de diciembre de 2025 del Presidente de los Estados Unidos sobre la Inteligencia Artificial. Por qué consideramos más eficiente la Ley Europea de Inteligencia Artificial que la regulación presente y la proyectada en los Estados Unidos». *AJ Tapia*. Acceso el 3 de julio de 2026: <https://ajtapia.com/author/admin/page/3/>
- Topol, Eric. 2019. *Deep medicine: how artificial intelligence can make healthcare human again*. New York: Basic Books.
- Valls, Javier. 2021. *Inteligencia artificial, derechos humanos y bienes jurídicos*. Granada: Universidad de Granada.
- Xu, Senxiao. 2025. «Polarización estructural en grafos y redes sociales: Un estudio temporal sobre OpenAI y DeepSeek en X». Trabajo académico, Universidad Complutense de Madrid. Acceso el 3 de julio de 2026: <https://docta.ucm.es/entities/publication/a4e8858d-925e-4619-8c9f-50f7d97d14c1>
- Zuboff, Shoshana. 2019. *The age of surveillance capitalism: the fight for a human future at the new frontier of power*. New York: PublicAffairs.