

CONSTITUCIONALISTA RESEARCH

International Center for Constitutional Research

Inteligencia Artificial y Constitución

¿Quién responde cuando el algoritmo toma una decisión sobre un ciudadano?

Serie	Research Paper 001
Autor institucional	Constitucionalista Research
Página oficial	https://constitucionalista.com/inteligencia-artificial-y-constitucion/
Versión	1.0.1
DOI	10.5281/zenodo.22141803
Fecha de publicación	28 de agosto de 2026
Fecha de corte de la investigación	28 de agosto de 2026
Licencia	Creative Commons Attribution 4.0 International (CC BY 4.0)

Citación recomendada

Constitucionalista Research. *Inteligencia Artificial y Constitución: ¿Quién responde cuando el algoritmo toma una decisión sobre un ciudadano?* Research Paper 001, versión 1.0.1, 2026. DOI: [10.5281/zenodo.22141803](https://doi.org/10.5281/zenodo.22141803). Disponible en: <https://constitucionalista.com/inteligencia-artificial-y-constitucion/>

Contenido

Resumen

Nota metodológica y alcance

1. Cuando el algoritmo se convierte en quien decide
2. Responsabilidad constitucional frente a la delegación algorítmica
3. Unión Europea: de la protección de datos al debido proceso algorítmico
4. Estados Unidos: protección constitucional sin un código unificado de derechos algorítmicos
5. América Latina: hacia el control constitucional del poder algorítmico
6. Seis garantías constitucionales para decisiones automatizadas
7. ¿Quién responde? Distribución de responsabilidad en la cadena algorítmica
8. Propuesta doctrinal: Principio de Responsabilidad Constitucional No Delegable
9. Conclusiones

Referencias normativas y jurisprudenciales principales

Resumen

La expansión de sistemas automatizados e inteligencia artificial hacia ámbitos como crédito, empleo, prestaciones públicas, salud, seguridad y justicia obliga a replantear una pregunta constitucional clásica: quién ejerce el poder y quién responde por sus consecuencias. Este trabajo sostiene que el problema no se resuelve verificando si existe una persona que formalmente firma la decisión. Cuando una recomendación, puntuación o clasificación automatizada condiciona de forma material el resultado, la intervención tecnológica puede adquirir relevancia decisoria aunque exista una ratificación humana posterior.

A partir de un análisis comparado de la Unión Europea, Estados Unidos y América Latina, el estudio identifica seis garantías mínimas para decisiones automatizadas de alto impacto: legalidad; transparencia y explicación; igualdad y no discriminación; intervención humana significativa; impugnación y revisión efectiva; y responsabilidad acompañada de remedio. La investigación examina el RGPD, la jurisprudencia SCHUFA y Dun & Bradstreet, el Reglamento Europeo de Inteligencia Artificial, el debido proceso estadounidense, las políticas federales sobre IA de alto impacto, reglas sectoriales y estatales, la LGPD brasileña, la jurisprudencia constitucional colombiana, las reformas chilenas, la experiencia uruguaya, el marco peruano y los estándares interamericanos.

El trabajo propone como contribución doctrinal un Principio de Responsabilidad Constitucional No Delegable: cuando un sistema automatizado participa materialmente en una decisión que afecta derechos fundamentales, intereses jurídicamente protegidos o condiciones esenciales de una persona, debe permanecer identificable una persona física o institución jurídica con autoridad efectiva para explicar, revisar, corregir y responder por esa decisión. La tecnología puede redistribuir tareas; no puede convertirse en el lugar donde termina la responsabilidad.

Palabras clave

Inteligencia artificial; derecho constitucional; decisiones automatizadas; responsabilidad algorítmica; debido proceso; supervisión humana; derechos fundamentales; derecho constitucional comparado; Unión Europea; Estados Unidos; América Latina.

Nota metodológica y alcance

Este Research Paper adopta un enfoque jurídico-comparado y se concentra en decisiones automatizadas capaces de producir efectos jurídicos o consecuencias materialmente significativas sobre personas. No pretende catalogar toda la regulación mundial de inteligencia artificial ni ofrecer una teoría general de responsabilidad civil por productos de IA. El foco es más limitado: la relación entre automatización decisoria y garantías constitucionales o funcionalmente equivalentes.

La fecha de corte de la investigación es el 28 de agosto de 2026. Esta precisión es esencial porque varias normas analizadas se encuentran en transición. Entre ellas figuran las obligaciones europeas para determinadas categorías de sistemas de alto riesgo, las nuevas garantías chilenas sobre decisiones automatizadas, requisitos californianos sobre automated decisionmaking technology y reformas estatales estadounidenses con entrada en vigor posterior. El texto distingue entre normas vigentes, normas ya aprobadas con aplicación futura, políticas administrativas y proyectos legislativos.

La selección de jurisdicciones responde a tres razones. La Unión Europea ofrece una arquitectura consolidada de protección de datos y una regulación horizontal de IA. Estados Unidos combina garantías constitucionales frente al poder público con regulación sectorial, administrativa y estatal. América Latina permite observar la interacción entre protección de datos, justicia constitucional, habeas data, tutela o amparo y el sistema interamericano de derechos humanos. Las diferencias entre estos modelos son parte del objeto comparativo y no deben interpretarse como equivalencias normativas.

1. Cuando el algoritmo se convierte en quien decide

Una persona solicita un crédito, un beneficio público, un empleo, una prestación sanitaria o una autorización administrativa. Entrega sus datos y recibe una respuesta: denegado. Puede que ningún funcionario haya examinado personalmente su situación. Puede que el resultado provenga de una puntuación, una clasificación predictiva, un modelo automatizado o una herramienta de inteligencia artificial que asignó un nivel de riesgo o recomendó una decisión.

Desde la perspectiva del ciudadano, la diferencia entre una decisión tomada automáticamente y una decisión humana basada casi por completo en una recomendación automatizada puede ser mínima. En ambos casos, el resultado modifica oportunidades, derechos o condiciones materiales de vida. Lo constitucionalmente relevante no es únicamente si existe una persona que pulsa el último botón, sino quién ejerce realmente el poder decisorio, quién puede explicar su ejercicio y quién responde cuando el resultado vulnera un derecho.

No toda utilización de IA plantea el mismo problema. Un sistema que ordena expedientes, corrige errores formales o resume documentos no ocupa la misma posición que un sistema cuyo resultado determina, condiciona o pesa decisivamente en una resolución sobre una persona. La frontera aparece cuando la tecnología deja de limitarse a asistir una actividad y comienza a participar materialmente en el ejercicio de poder.

El artículo 22 del Reglamento General de Protección de Datos (RGPD) constituye uno de los puntos de partida más claros. Reconoce, con las excepciones y garantías previstas por el propio Reglamento, protección frente a decisiones basadas únicamente en tratamiento automatizado que produzcan efectos jurídicos o afecten de modo similar y significativo a la persona. La jurisprudencia del Tribunal de Justicia ha mostrado además que no basta con observar la forma del procedimiento. En *SCHUFA Holding (Scoring)*, C-634/21, el Tribunal consideró que una puntuación automatizada puede constituir una decisión relevante a efectos del artículo 22 cuando el tercero que la recibe le atribuye un papel determinante en el resultado final.

La opacidad algorítmica se convierte en un problema constitucional cuando impide ejercer derechos. Una persona que desconoce qué sistema intervino, qué información fue decisiva, qué criterio condujo al resultado o qué autoridad asumió la decisión puede quedar materialmente imposibilitada para cuestionarla. En *Dun & Bradstreet Austria*, C-203/22, el Tribunal de Justicia vinculó la información sobre la lógica de la decisión con la capacidad práctica de comprenderla y controvertirla.

La experiencia estadounidense ilustra otro aspecto del problema. El Memorando OMB M-25-21 define como high-impact AI ciertos usos federales en los que la salida de la IA sirve como base principal para decisiones con efectos legales, materiales, vinculantes o significativos sobre derechos o seguridad. El memorando aclara que un uso puede ser de alto impacto exista o no supervisión humana. Para esos usos exige, entre otras medidas, supervisión y responsabilidad humana apropiadas y, cuando corresponda, acceso a revisión humana o a mecanismos de apelación.

En América Latina, la Sentencia T-323 de 2024 de la Corte Constitucional de Colombia ofrece una formulación especialmente potente. La Corte admitió usos auxiliares de IA en la justicia, pero sostuvo que la tecnología no puede sustituir al juez en funciones que exigen interpretar hechos y pruebas, motivar y adoptar la decisión. El caso muestra que ciertas competencias jurídicas no son simples tareas susceptibles de optimización: están vinculadas a un titular que debe conservar la responsabilidad.

Este trabajo parte, por tanto, de una pregunta: ¿quién responde cuando el algoritmo toma una decisión sobre un ciudadano? ¿El desarrollador, el proveedor, la institución que adquirió la tecnología, el funcionario que aceptó su recomendación o todos ellos en ámbitos distintos? La fragmentación de la cadena tecnológica no puede producir un vacío de responsabilidad frente a la persona afectada.

La hipótesis central es que la responsabilidad constitucional no puede delegarse a una máquina. La IA puede asistir, recomendar, clasificar y automatizar análisis. Pero cuando una decisión afecta derechos fundamentales o intereses jurídicamente protegidos, debe permanecer identificable una persona o institución con capacidad de explicar, revisar, corregir y responder por el resultado.

2. Responsabilidad constitucional frente a la delegación algorítmica

2.1. ¿Qué significa realmente decidir?

Para este estudio conviene distinguir tres niveles. En la asistencia algorítmica, la IA organiza información o produce apoyo y la persona realiza una valoración independiente. En la recomendación algorítmica, el sistema propone una puntuación, clasificación o resultado que influye en la decisión humana. En la determinación algorítmica, la salida del sistema produce directamente el resultado o tiene un peso tan dominante que la intervención humana se reduce a una ratificación.

La diferencia no es terminológica. A mayor influencia material del sistema, mayor necesidad de garantías. El caso SCHUFA es importante precisamente porque impide que una organización convierta una decisión materialmente automatizada en formalmente humana por el simple hecho de insertar un tercero al final de la cadena. La pregunta jurídicamente relevante es cuánto poder real conserva la persona que decide.

2.2. La firma humana no elimina una decisión algorítmica.

Uno de los riesgos más frecuentes es el rubber-stamping: un funcionario, empleado o profesional recibe una recomendación automatizada y la acepta de forma prácticamente automática. La firma existe, pero el juicio independiente puede no existir. Una intervención humana constitucionalmente significativa requiere comprensión suficiente del resultado, autoridad efectiva para apartarse de él y responsabilidad por la decisión final.

La mera presencia de una persona no es control. La supervisión es significativa únicamente cuando el humano puede discrepar de la máquina. Si rechazar la recomendación exige autorizaciones extraordinarias, genera sanciones internas, carece de información suficiente o nunca ocurre en la práctica, la arquitectura decisoria puede seguir siendo algorítmica en lo sustancial.

2.3. ¿Quién es responsable: desarrollador, proveedor, institución o funcionario?

En la práctica existe una cadena: desarrollador, proveedor, deployer o entidad usuaria, funcionario decisor y persona afectada. La solución no consiste en atribuir todo a un solo actor, sino en impedir que la distribución de funciones produzca responsabilidad inexistente.

El desarrollador puede tener deberes vinculados al diseño, evaluación y documentación. El proveedor puede responder por obligaciones de puesta en el mercado y soporte. La organización que despliega el sistema tiene deberes propios de selección, gobernanza y supervisión. El funcionario o profesional conserva las obligaciones inherentes a la competencia que jurídicamente ejerce. La externalización tecnológica no convierte automáticamente en externalizable la responsabilidad constitucional.

El AI Act europeo refleja una distribución de obligaciones entre providers y deployers. Para los sistemas de alto riesgo, la regulación exige medidas de supervisión humana y, cuando las disposiciones correspondientes sean aplicables, impone al deployer deberes de gobernanza y uso. Del mismo modo, OMB M-25-21 atribuye responsabilidades organizativas a las agencias federales y exige que funcionarios responsables acepten y gestionen riesgos de usos de alto impacto.

De ello se desprende un criterio de proximidad al poder decisorio: cuanto más cerca se encuentra un actor del ejercicio efectivo de poder sobre la persona, mayor debe ser su deber de asegurar que el resultado sea legal,

revisable y atribuible. El ciudadano no debería tener que reconstruir toda una cadena contractual y técnica para descubrir quién responde por la decisión que lo afectó.

2.4. Principio de Responsabilidad Constitucional No Delegable.

Este trabajo propone la siguiente formulación: cuando un sistema automatizado o de inteligencia artificial participa materialmente en una decisión que afecta derechos fundamentales, intereses jurídicamente protegidos o condiciones esenciales de una persona, debe permanecer identificable una persona física o institución jurídica con autoridad efectiva para explicar, revisar, corregir y responder por esa decisión.

El principio exige cuatro capacidades mínimas: explicar de forma jurídicamente útil la intervención del sistema; revisar el caso individual; corregir o dejar sin efecto un resultado inadecuado; y responder ante un mecanismo administrativo o judicial de remedio. No exige que una sola entidad soporte todas las responsabilidades técnicas o contractuales. Exige que, frente al ciudadano, la arquitectura no termine en una máquina.

Una prueba práctica ayuda a identificar el vacío de responsabilidad. Ante una decisión automatizada de alto impacto deben poder responderse cinco preguntas: quién decidió utilizar el sistema; quién puede explicar el papel de la IA; quién puede apartarse de su resultado; quién puede revisar y corregir la decisión; y quién responde jurídicamente frente a la persona afectada. Si esas respuestas son indeterminadas, el problema deja de ser puramente tecnológico y se convierte en un problema constitucional.

3. Unión Europea: de la protección de datos al debido proceso algorítmico

3.1. El artículo 22 del RGPD como punto de partida.

El derecho europeo reconoce una protección específica frente a determinadas decisiones basadas únicamente en tratamiento automatizado que producen efectos jurídicos o afectan de forma similar y significativa a una persona. El artículo 22 no configura una prohibición absoluta: admite supuestos vinculados a contratos, autorización legal o consentimiento explícito, pero exige salvaguardias en los casos previstos.

Entre esas salvaguardias figuran, en determinados supuestos, la posibilidad de obtener intervención humana, expresar el propio punto de vista e impugnar la decisión. El artículo 15(1)(h), junto con las obligaciones de información del RGPD, incorpora además el acceso a información significativa sobre la lógica implicada, así como la importancia y consecuencias previstas del tratamiento automatizado.

3.2. SCHUFA: mirar la realidad funcional.

En C-634/21, el Tribunal de Justicia examinó una puntuación de solvencia generada automáticamente y utilizada por terceros. La importancia del caso reside en que el Tribunal no aceptó una separación puramente formal entre quien genera el score y quien adopta la decisión contractual. Cuando la puntuación desempeña un papel determinante, puede quedar comprendida en el régimen de decisiones automatizadas.

Este razonamiento permite identificar una regla comparativa: no basta con preguntar quién firmó. Debe identificarse qué elemento determinó realmente el resultado. La existencia de una interfaz humana o de una aprobación final no convierte por sí sola el proceso en genuinamente humano.

3.3. Dun & Bradstreet: explicar para poder impugnar.

En C-203/22, decidido el 27 de febrero de 2025, el Tribunal de Justicia profundizó en el contenido de la información significativa sobre decisiones automatizadas. La explicación debe ser comprensible y permitir que la persona conozca de forma útil cómo se llegó al resultado y pueda cuestionarlo. La complejidad técnica o la invocación de secretos empresariales no pueden vaciar automáticamente el derecho.

La relevancia constitucional de esta doctrina es directa. Un recurso sin información suficiente puede ser nominal pero no efectivo. La explicación no exige necesariamente revelar el código fuente ni convertir al ciudadano en auditor técnico; exige proporcionar razones que conecten el resultado con factores relevantes del caso y permitan detectar errores, datos incorrectos o criterios jurídicamente cuestionables.

3.4. El AI Act y la regulación del sistema.

El Reglamento (UE) 2024/1689 añade una capa de gobernanza centrada en los propios sistemas de IA. Su arquitectura combina prácticas prohibidas, obligaciones de transparencia, reglas para modelos de propósito general y un régimen especial para sistemas de alto riesgo. El Reglamento (UE) 2026/1744 modificó el calendario de aplicación de las Secciones 1, 2 y 3 del Capítulo III para sistemas de alto riesgo: 2 de diciembre de 2027 para los clasificados por el artículo 6(2) y Anexo III, y 2 de agosto de 2028 para los vinculados al artículo 6(1) y Anexo I.

La modificación temporal es importante para evitar una afirmación incorrecta: a la fecha de corte de este estudio, gran parte del AI Act es aplicable, pero varias obligaciones estructurales sobre sistemas de alto riesgo

todavía se encuentran en transición. La existencia de una norma vigente no significa que todas sus obligaciones sean ya exigibles a todas las categorías de sistemas.

El artículo 14 formula un concepto fuerte de supervisión humana para sistemas de alto riesgo: el diseño debe permitir que personas físicas los supervisen efectivamente, conozcan sus capacidades y limitaciones, eviten la confianza automática en la salida y puedan, según corresponda, ignorar, sustituir o revertir el resultado. La lógica coincide con la idea desarrollada en este trabajo: la supervisión requiere capacidad real de desacuerdo.

El artículo 26 atribuye obligaciones a deployers de sistemas de alto riesgo, entre ellas la asignación de supervisión a personas con competencia, formación y autoridad adecuadas cuando las disposiciones correspondientes resulten aplicables. La palabra autoridad es esencial: conocimiento sin capacidad institucional para apartarse del resultado no constituye control efectivo.

El artículo 27 incorpora evaluaciones de impacto sobre derechos fundamentales para determinados deployers y usos de alto riesgo. El artículo 86 reconoce un derecho a obtener explicaciones claras y significativas sobre el papel del sistema y los principales elementos de ciertas decisiones individuales basadas en la salida de sistemas de alto riesgo del Anexo III. El artículo 85, por su parte, prevé la posibilidad de presentar reclamaciones ante autoridades de vigilancia del mercado por infracciones del Reglamento.

3.5. Hacia un debido proceso algorítmico europeo.

El conjunto formado por RGPD, SCHUFA, Dun & Bradstreet, Carta de Derechos Fundamentales y AI Act permite identificar elementos que funcionalmente se aproximan a un debido proceso algorítmico: conocimiento de la automatización; información suficiente sobre su papel; capacidad de impugnación; intervención humana significativa cuando corresponda; control de riesgos sobre derechos fundamentales; y acceso a una autoridad capaz de proporcionar remedio.

No se trata de afirmar que el derecho de la Unión haya codificado una doctrina constitucional unitaria con ese nombre. Se trata de observar una convergencia normativa. A medida que aumenta la capacidad del sistema para producir consecuencias significativas, aumenta también la exigencia de trazabilidad, explicación, supervisión y revisión.

3.6. El límite europeo: regulación compleja y aplicación escalonada.

El modelo europeo es avanzado, pero no elimina los problemas. La coexistencia de RGPD, AI Act, legislación sectorial, secreto empresarial y calendarios transitorios puede dificultar a la persona identificar el derecho exacto aplicable. Además, no toda recomendación algorítmica cae dentro del artículo 22 del RGPD ni toda herramienta automatizada es un sistema de alto riesgo del AI Act. Por ello, el valor comparativo del modelo europeo no reside en ofrecer una respuesta automática para cada caso, sino en haber normalizado jurídicamente la idea de que la automatización de alto impacto necesita garantías adicionales.

4. Estados Unidos: protección constitucional sin un código unificado de derechos algorítmicos

4.1. Un modelo fragmentado.

Estados Unidos no dispone de un equivalente federal general al AI Act europeo que establezca un régimen único de derechos individuales frente a decisiones automatizadas. La protección surge de varias fuentes: las cláusulas constitucionales de debido proceso y equal protection cuando existe acción estatal; legislación sectorial; normas de derechos civiles y protección del consumidor; políticas administrativas federales; y regulación estatal y local.

Esta fragmentación tiene una consecuencia metodológica. No puede afirmarse que toda decisión automatizada privada esté sujeta directamente a la Quinta o Decimocuarta Enmienda. El análisis constitucional se activa principalmente frente a decisiones del gobierno federal, estatal o local y otros supuestos de state action. En el sector privado, la protección depende del ámbito material: crédito, empleo, vivienda, privacidad, salud u otras normas aplicables.

4.2. Debido proceso y sistemas opacos.

Los casos estadounidenses muestran que el problema algorítmico puede presentarse aunque la legislación no utilice la palabra inteligencia artificial. En *K.W. v. Armstrong*, 789 F.3d 962 (9th Cir. 2015), un litigio sobre presupuestos de servicios de Medicaid calculados mediante una metodología automatizada, el tribunal destacó la necesidad de notificaciones suficientes para permitir que los beneficiarios comprendieran la base de la reducción y pudieran impugnarla. El caso es útil como precedente funcional sobre razones, posibilidad de controversia y beneficios públicos.

En *Houston Federation of Teachers, Local 2415 v. Houston Independent School District*, 251 F. Supp. 3d 1168 (S.D. Tex. 2017), docentes cuestionaron un sistema propietario de evaluación de valor añadido utilizado en decisiones laborales. El tribunal consideró viable la reclamación de debido proceso cuando las personas afectadas no podían verificar o reproducir adecuadamente la puntuación que influía en consecuencias laborales. Aunque el litigio terminó mediante acuerdo y su alcance debe leerse con cautela, el caso muestra el conflicto entre decisiones públicas de alto impacto y modelos cuyo funcionamiento no puede ser efectivamente controvertido por quien sufre el resultado.

State v. Loomis, 2016 WI 68, 881 N.W.2d 749, ofrece otro ángulo. La Corte Suprema de Wisconsin permitió el uso de una evaluación COMPAS en la sentencia penal bajo limitaciones y advertencias, rechazando la pretensión de que ese uso específico violara por sí mismo el debido proceso. El caso no legitima una delegación automática a software propietario: por el contrario, muestra la necesidad de delimitar el peso del instrumento y reconocer limitaciones antes de integrarlo en una decisión sobre libertad.

Estos casos no crean un derecho federal general a explicación algorítmica. Pero juntos muestran un principio procesal reconocible: cuando el Estado utiliza una herramienta opaca como fundamento importante de una decisión adversa, la imposibilidad de comprender o controvertir el fundamento puede adquirir relevancia de debido proceso.

4.3. El crédito como laboratorio de explicabilidad.

La Equal Credit Opportunity Act (ECOA), 15 U.S.C. § 1691(d), exige comunicar razones específicas por determinadas acciones adversas. Regulation B, 12 C.F.R. § 1002.9, rechaza explicaciones genéricas que se limiten a afirmar que el solicitante no alcanzó estándares internos o una puntuación.

El CFPB ha aplicado este principio a modelos complejos. En Circular 2022-03 sostuvo que la complejidad de un algoritmo no excusa el cumplimiento de las obligaciones de notificación de acciones adversas. El mensaje regulatorio es especialmente útil para la tesis de este trabajo: la caja negra puede ser un problema técnico, pero no debe convertirse en una defensa jurídica para omitir razones cuando la ley exige razones.

4.4. M-25-21 y el concepto federal de high-impact AI.

El Memorando OMB M-25-21, de 3 de abril de 2025, reemplazó M-24-10 y estableció un marco de gobernanza para usos de IA por agencias del Poder Ejecutivo. Define high-impact AI cuando la salida del sistema sirve como base principal para decisiones o acciones con efectos legales, materiales, vinculantes o significativos sobre derechos o seguridad. El memorando afirma expresamente que un uso puede ser high-impact exista o no supervisión humana.

Para usos de alto impacto, las agencias deben realizar pruebas previas, evaluaciones de impacto, seguimiento, capacitación y medidas de gobernanza. El memorando exige supervisión, intervención y responsabilidad humana apropiadas y, cuando corresponda, acceso a revisión humana o a una oportunidad de apelación. También exige considerar impactos sobre privacidad, derechos civiles y libertades civiles.

M-25-21 no es una enmienda constitucional ni crea por sí solo un derecho privado general equivalente al artículo 22 del RGPD. Su importancia comparativa es institucional: el propio gobierno federal reconoce que los usos de IA que sirven de base principal para decisiones significativas requieren una gobernanza reforzada y no dejan de ser de alto impacto porque haya un humano en el circuito.

4.5. Estados y ciudades como laboratorios regulatorios.

Nueva York City aplica Local Law 144 a determinadas automated employment decision tools, exigiendo bias audit, publicación de resultados y avisos a candidatos o empleados. Es una protección específica y limitada al ámbito cubierto, pero demuestra la traducción de preocupaciones sobre discriminación algorítmica en obligaciones ex ante.

California adoptó regulaciones de privacidad sobre automated decisionmaking technology. Las regulaciones entraron en vigor el 1 de enero de 2026, pero las obligaciones específicas relacionadas con decisiones significativas mediante ADMT comienzan el 1 de enero de 2027. La California Privacy Protection Agency describe derechos de aviso, opt-out donde proceda y acceso a información significativa sobre el funcionamiento y efecto del ADMT. La precisión temporal es esencial: a la fecha de este estudio, parte relevante del régimen todavía es futura.

Colorado aprobó SB26-189, con protecciones para determinadas consequential decisions y un derecho a solicitar meaningful human review y reconsideración tras resultados adversos de sistemas cubiertos, con entrada en vigor principal en 2027. Texas, mediante HB 149, estableció desde el 1 de enero de 2026 un marco estatal de gobernanza de IA que incluye prohibiciones sobre ciertos usos discriminatorios o dirigidos a infringir derechos protegidos.

4.6. La conclusión estadounidense.

El derecho de Estados Unidos no converge en una única cláusula de derechos algorítmicos. Sin embargo, varias líneas apuntan en la misma dirección: el gobierno sigue siendo responsable por sus decisiones aunque utilice software; determinados sectores exigen razones específicas; la opacidad puede agravar problemas de debido proceso; y las políticas de IA de alto impacto exigen supervisión, evaluación y vías de revisión. El modelo estadounidense es menos uniforme que el europeo, pero no acepta necesariamente la idea de que el algoritmo rompa la cadena jurídica de responsabilidad.

5. América Latina: hacia el control constitucional del poder algorítmico

5.1. Una región heterogénea con herramientas constitucionales potentes.

América Latina no presenta un único modelo regulatorio. Conviven leyes de protección de datos inspiradas en tradiciones europeas, acciones constitucionales rápidas, tribunales constitucionales con fuerte intervención en derechos fundamentales y reformas recientes sobre IA. Esta combinación convierte a la región en un espacio especialmente relevante para construir garantías frente a decisiones automatizadas.

5.2. Brasil: revisión y transparencia bajo la LGPD.

El artículo 20 de la Lei Geral de Proteção de Dados Pessoais reconoce el derecho del titular a solicitar revisión de decisiones tomadas únicamente con base en tratamiento automatizado de datos personales que afecten sus intereses, incluidas decisiones de perfil personal, profesional, de consumo y crédito. El §1 obliga al controlador, cuando sea solicitado, a proporcionar información clara y adecuada sobre criterios y procedimientos utilizados, respetando secretos comerciales e industriales.

El texto vigente merece una precisión importante. Una versión temprana del artículo 20 contenía referencia expresa a revisión por persona natural, pero la redacción actual, dada por la Ley 13.853/2019, ya no incluye esa exigencia literal. Por ello, sería incorrecto presentar la LGPD vigente como si garantizara siempre revisión humana. Sí garantiza solicitud de revisión y transparencia sobre criterios y procedimientos, y el §2 permite a la autoridad nacional realizar auditorías para verificar aspectos discriminatorios cuando la información se niega por secreto comercial o industrial.

Brasil discute además un marco general de IA mediante PL 2338/2023. El Senado aprobó un texto sustitutivo en diciembre de 2024 y lo remitió a la Cámara de Diputados. En agosto de 2026 continuaba pendiente de votación en la Cámara y era objeto de trabajo en comisión especial. Por tanto, debe tratarse como proyecto legislativo, no como derecho vigente.

5.3. Colombia: la no sustitución del juez.

La Sentencia T-323/24 de la Corte Constitucional es uno de los pronunciamientos más directamente vinculados con la tesis de este estudio. Al revisar el uso de ChatGPT por un juez de tutela, la Corte distinguió entre asistencia tecnológica y sustitución del razonamiento jurisdiccional. La IA puede apoyar ciertas tareas, pero no reemplazar al juez en la valoración de hechos y pruebas, la motivación ni la adopción de la decisión.

La Corte articuló principios de transparencia, responsabilidad, privacidad, no sustitución, prevención de riesgos, igualdad y control humano. Su importancia excede el uso de un chatbot concreto. El fallo expresa una intuición constitucional generalizable: una competencia atribuida jurídicamente a un órgano humano no puede vaciarse mediante una delegación tecnológica que haga desaparecer el razonamiento y la responsabilidad del titular.

En 2026, la propia jurisprudencia constitucional continuó debatiendo el alcance de T-323/24 en otros contextos de uso de IA. Esa evolución confirma que el precedente debe leerse con precisión según el ámbito del caso, pero también que el problema de la automatización de funciones públicas ya forma parte del discurso constitucional colombiano.

5.4. Chile: un derecho de nueva generación, todavía en transición.

La Ley 21.719 reformó la legislación chilena sobre datos personales e incorporó un nuevo artículo 8 bis relativo a decisiones individuales automatizadas, incluida la elaboración de perfiles. La reforma reconoce garantías de información y transparencia y prevé mecanismos vinculados a explicación, intervención humana, expresión del punto de vista y revisión en los supuestos previstos.

Sin embargo, las reformas relevantes entran en vigor el 1 de diciembre de 2026. En la fecha de corte de esta investigación han sido promulgadas pero todavía no son operativas. Chile constituye así un ejemplo de cómo los derechos asociados a decisiones automatizadas están migrando desde principios generales de protección de datos hacia reglas más explícitas de contestabilidad y control humano.

5.5. Uruguay: un precedente anterior a la actual ola de IA.

El artículo 16 de la Ley 18.331 protege frente a decisiones con efectos jurídicos o efectos significativamente perjudiciales basadas en tratamiento automatizado destinado a evaluar aspectos de la personalidad, como rendimiento laboral, crédito, fiabilidad o conducta. La persona afectada puede impugnar actos administrativos o decisiones privadas y obtener información sobre criterios de valoración y el programa utilizado.

La importancia uruguaya es histórica y doctrinal. Muestra que el problema de la decisión automatizada no nace con los modelos generativos. Antes del actual auge de IA ya existía una preocupación jurídica por sistemas de evaluación automatizada capaces de afectar de forma significativa a personas.

5.6. Perú: principios de IA responsable y desarrollo reglamentario.

La Ley 31814 promueve el uso de la inteligencia artificial en favor del desarrollo económico y social bajo principios centrados en la persona, ética, transparencia, responsabilidad y respeto por derechos. El Decreto Supremo 115-2025-PCM aprobó su reglamento y desarrolla un marco de gestión de riesgos, transparencia y supervisión humana para determinados usos.

El modelo peruano es relevante como arquitectura de gobernanza, pero no debe equipararse automáticamente al régimen individual del artículo 22 del RGPD. Su aportación al análisis está en mostrar la progresiva incorporación de lenguaje de riesgos, responsabilidad y control humano al derecho público digital de la región.

5.7. El sistema interamericano.

La Convención Americana sobre Derechos Humanos proporciona anclajes generales: artículo 8 sobre garantías judiciales, artículo 24 sobre igualdad ante la ley y artículo 25 sobre protección judicial efectiva. Estas disposiciones no fueron redactadas pensando en algoritmos, pero sus funciones permanecen relevantes cuando una autoridad utiliza tecnología para adoptar decisiones sobre derechos y obligaciones.

En marzo de 2026, la Comisión Interamericana de Derechos Humanos llamó a los Estados a prevenir discriminación algorítmica y destacó la necesidad de marcos basados en derechos humanos, transparencia, explicabilidad, acceso a información sobre decisiones automatizadas, mecanismos de impugnación con revisión humana significativa, remedios, monitoreo y evaluaciones de impacto. Al tratarse de orientación regional y no de una sentencia que cree por sí sola obligaciones idénticas para todos los Estados, debe presentarse como evidencia de una dirección interpretativa emergente.

5.8. Una posible convergencia regional.

Brasil aporta revisión y transparencia; Colombia, no sustitución y control humano en la función judicial; Chile, garantías explícitas de nueva generación; Uruguay, contestabilidad desde una legislación más antigua; Perú, gobernanza de IA responsable; y el sistema interamericano, un lenguaje regional de explicación, no discriminación y remedio. No existe todavía un código latinoamericano uniforme, pero sí elementos suficientes para identificar un movimiento hacia el control jurídico del poder algorítmico.

6. Seis garantías constitucionales para decisiones automatizadas

La comparación permite construir un estándar funcional compuesto por seis garantías. No todas aparecen con igual fuerza ni bajo el mismo nombre en cada jurisdicción. Su valor está en identificar qué condiciones debe cumplir una decisión automatizada de alto impacto para permanecer dentro de una estructura de Estado de Derecho.

6.1. Legalidad.

La primera pregunta es si la autoridad o entidad tiene base jurídica para utilizar el sistema en ese contexto y para procesar los datos o producir el tipo de resultado utilizado. En el sector público, la tecnología no crea una competencia nueva. La administración debe actuar dentro de las atribuciones legales existentes y respetar los límites constitucionales aplicables.

La legalidad incluye también la temporalidad normativa. Un régimen aprobado pero todavía no aplicable no puede ser presentado como obligación vigente. El calendario del AI Act europeo, la reforma chilena, California y Colorado muestra que la gobernanza algorítmica exige precisión no solo sobre qué norma existe, sino sobre cuándo resulta exigible.

6.2. Transparencia y explicación.

La persona necesita saber que la automatización intervino de manera material y debe recibir una explicación jurídicamente útil. Explicar no significa necesariamente revelar código fuente o secretos comerciales. Significa ofrecer razones suficientes para comprender el papel de la tecnología, los factores relevantes y la conexión entre esos factores y el resultado.

Dun & Bradstreet, ECOA/Regulation B, la LGPD brasileña, la ley uruguaya y las reformas chilenas muestran distintas expresiones de esta exigencia. La regla funcional puede formularse así: si una persona no puede entender suficientemente por qué recibió un resultado, su capacidad de impugnarlo se reduce.

6.3. Igualdad y no discriminación.

Los sistemas automatizados pueden reproducir desigualdades mediante datos históricos, variables proxy, decisiones de diseño o contextos de despliegue. La igualdad no se satisface afirmando que el algoritmo trata a todos de manera matemática. Debe analizarse si el sistema produce distinciones jurídicamente prohibidas, impactos desproporcionados o errores concentrados sobre grupos protegidos.

La Carta de Derechos Fundamentales de la UE, la Equal Protection Clause cuando existe acción estatal, las normas de derechos civiles, los bias audits de NYC, la LGPD y la orientación interamericana ofrecen instrumentos diferentes para el mismo problema: la automatización no neutraliza la prohibición de discriminación.

6.4. Intervención humana significativa.

La supervisión debe ser real. El revisor necesita conocimiento, tiempo, acceso a información suficiente y autoridad para apartarse del resultado. Un humano que solo confirma lo que la interfaz muestra no agrega una garantía sustantiva.

El AI Act europeo desarrolla esta idea mediante supervisión que permita ignorar, sustituir o revertir resultados. OMB M-25-21 exige supervisión e intervención apropiadas para high-impact AI. Colorado define una revisión humana significativa en la que el revisor tiene autoridad para aprobar, modificar o anular la decisión automatizada. Colombia insiste en la no sustitución del juez.

6.5. Impugnación y revisión efectiva.

Una persona afectada debe disponer de un canal para cuestionar datos, inferencias, procedimiento y resultado. El canal debe ser accesible y producir una revisión que pueda modificar la decisión. Un formulario que solo reenvía el caso al mismo algoritmo no constituye revisión efectiva.

El RGPD contempla contestación e intervención humana en supuestos del artículo 22; OMB M-25-21 prevé revisión y apelación cuando corresponda; la legislación uruguaya permite impugnar; la Convención Americana protege el acceso a recursos efectivos. La forma institucional puede variar, pero la función es común: convertir la decisión en algo controvertible.

6.6. Responsabilidad y remedio.

Finalmente debe existir alguien que responda. La persona necesita saber contra qué institución presentar una reclamación, quién puede corregir el daño y qué autoridad puede revisar la actuación. Sin atribución de responsabilidad, transparencia y supervisión pierden parte de su eficacia.

Las seis garantías se refuerzan mutuamente. La explicación permite impugnar; la impugnación requiere un revisor con autoridad; el revisor necesita trazabilidad; la igualdad necesita auditoría y datos adecuados; el remedio necesita una institución responsable. El debido proceso algorítmico no es una única regla, sino una arquitectura de garantías conectadas.

7. ¿Quién responde? Distribución de responsabilidad en la cadena algorítmica

7.1. La cadena de valor no puede convertirse en una cadena de excusas.

Los sistemas contemporáneos pueden involucrar desarrolladores de modelos, proveedores de infraestructura, integradores, consultores, entidades que despliegan la herramienta, administradores, funcionarios y profesionales que adoptan la resolución. Cada actor controla una parte diferente del riesgo.

La responsabilidad técnica puede surgir cuando el sistema presenta defectos previsibles de diseño, entrenamiento, validación o documentación. La responsabilidad del proveedor puede relacionarse con información, soporte, limitaciones conocidas o incumplimiento de obligaciones regulatorias. El deployer controla el contexto de uso, los datos operativos, la finalidad concreta y el grado de dependencia institucional. El funcionario o profesional controla, al menos en teoría, la decisión que jurídicamente se le atribuye.

7.2. El proveedor no reemplaza al titular del poder.

Cuando una administración compra un sistema privado, la contratación puede externalizar ingeniería, pero no transforma a la empresa tecnológica en titular de la potestad pública. Si la ley atribuye a una agencia la decisión sobre una prestación, licencia o sanción, esa agencia sigue siendo responsable de que el procedimiento respete los derechos del ciudadano.

Este punto es especialmente importante para sistemas propietarios. La institución no debería poder responder a una impugnación diciendo que desconoce la lógica porque el proveedor protege el modelo como secreto. Si el sistema no puede ser utilizado de forma compatible con las obligaciones jurídicas de la institución, el problema se encuentra en la decisión de adquirirlo o desplegarlo, no en el ciudadano que reclama.

7.3. El funcionario no puede ser un sello.

A la inversa, tampoco resulta suficiente mantener un nombre humano en el expediente si el diseño organizativo convierte a esa persona en un ratificador automático. La responsabilidad individual requiere autoridad efectiva y condiciones para ejercer juicio independiente.

En contextos como la justicia, la exigencia puede ser especialmente intensa. T-323/24 muestra que el deber de motivación y la función decisoria del juez no pueden trasladarse a una herramienta. En otros ámbitos la intensidad puede variar, pero el principio es similar: una competencia jurídicamente personal o institucional no desaparece porque una recomendación automatizada parezca objetiva.

7.4. Responsabilidad compartida, punto de contacto único.

Este trabajo propone separar dos planos. Internamente, varias entidades pueden compartir o distribuir responsabilidad contractual, técnica, administrativa, disciplinaria o civil. Externamente, la persona afectada debe disponer de al menos una institución claramente identificable que reciba la reclamación y garantice explicación, revisión y remedio.

Esta solución evita imponer al ciudadano una investigación forense sobre contratos de software antes de poder ejercer sus derechos. La institución decisora puede después repetir contra el proveedor, exigir indemnización, iniciar procedimientos contractuales o distribuir responsabilidades según corresponda. Pero la arquitectura interna no debe impedir el acceso inicial a la justicia o a la revisión administrativa.

7.5. Una matriz de responsabilidad.

Puede formularse una distribución orientativa: el desarrollador responde principalmente por diseño y documentación bajo las normas aplicables; el proveedor, por conformidad, información y soporte; el deployer, por selección, finalidad, datos de uso, evaluación de impacto y gobernanza; la institución decisora, por legalidad, debido proceso, igualdad, explicación, revisión y remedio; y el funcionario, por los deberes propios de su función cuando conserva competencia decisoria.

Esta matriz no pretende sustituir los regímenes nacionales de responsabilidad. Su función es mostrar que la pregunta '¿quién responde?' admite respuestas múltiples sin aceptar una respuesta nula. La existencia de varios responsables posibles nunca debe terminar en ausencia de responsable frente a la persona.

8. Propuesta doctrinal: Principio de Responsabilidad Constitucional No Delegable

8.1. Naturaleza de la propuesta.

El Principio de Responsabilidad Constitucional No Delegable no se presenta como una regla positiva ya reconocida con ese nombre por todas las jurisdicciones estudiadas. Es una propuesta doctrinal construida a partir de elementos convergentes de derecho vigente, jurisprudencia, regulación y políticas públicas.

Su propósito es ofrecer una herramienta común para evaluar decisiones automatizadas sin depender de una tecnología concreta. La doctrina debe funcionar tanto para un sistema de scoring tradicional como para un modelo de aprendizaje automático, un sistema generativo integrado en un flujo administrativo o una herramienta futura que cumpla una función equivalente.

8.2. Formulación.

Cuando un sistema automatizado o de inteligencia artificial participa materialmente en una decisión que afecta derechos fundamentales, intereses jurídicamente protegidos o condiciones esenciales de una persona, debe permanecer identificable una persona física o institución jurídica con autoridad efectiva para explicar, revisar, corregir y responder por esa decisión.

La palabra materialmente es importante. El principio no pretende imponer el mismo nivel de control a una calculadora que a un sistema que recomienda cancelar una prestación social. Debe examinarse la influencia real del sistema sobre el resultado, la gravedad de las consecuencias, la reversibilidad del daño, la vulnerabilidad de la persona y la disponibilidad de mecanismos alternativos.

8.3. Test de cinco pasos.

Primer paso: identificar el efecto. ¿La decisión produce consecuencias jurídicas o significativas sobre derechos, oportunidades, servicios esenciales, libertad, empleo, crédito, salud, educación u otros intereses protegidos? Segundo paso: identificar la influencia de la IA. ¿El sistema solo asiste o sirve como base principal, determinante o prácticamente vinculante?

Tercer paso: identificar control humano real. ¿Existe una persona que comprende suficientemente el resultado y tiene autoridad práctica para apartarse de él? Cuarto paso: identificar contestabilidad. ¿La persona afectada puede conocer la intervención tecnológica, obtener razones, corregir datos y solicitar una revisión capaz de modificar el resultado? Quinto paso: identificar responsabilidad. ¿Existe una institución o persona jurídica claramente responsable ante la cual puede reclamarse y obtenerse remedio?

Si las cinco preguntas reciben respuestas satisfactorias, la automatización tiene mayores posibilidades de integrarse en un procedimiento compatible con garantías constitucionales. Si el sistema tiene alta influencia, el revisor carece de autoridad, la explicación es imposible y ninguna entidad asume responsabilidad, existe un fuerte indicio de delegación constitucionalmente problemática.

8.4. Regla de intensidad.

Las garantías deben aumentar con el impacto. Un sistema utilizado para ordenar correos no requiere el mismo régimen que uno que influye en detención, sentencia, elegibilidad para prestaciones, contratación, crédito o acceso a tratamiento médico. La proporcionalidad permite evitar dos errores opuestos: regular toda

automatización como si fuera una decisión soberana, o tratar decisiones de alto impacto como simples productos tecnológicos.

8.5. Contratación pública y deber de explicabilidad.

El principio tiene una consecuencia práctica para adquisiciones. Antes de contratar una herramienta destinada a decisiones de alto impacto, una institución debería verificar si puede cumplir sus propias obligaciones jurídicas con ese producto. Debe conocer al menos sus limitaciones, mecanismos de auditoría, posibilidades de revisión, trazabilidad y condiciones para proporcionar razones al afectado.

Una cláusula de secreto empresarial no debería ser suficiente para anular garantías constitucionales. Si una institución no puede explicar o revisar una decisión porque el contrato tecnológico se lo impide, existe un defecto de gobernanza anterior al resultado individual.

8.6. Diseño institucional.

Para decisiones de alto impacto, la institución debería mantener un registro del sistema utilizado, finalidad, responsable, categorías de datos relevantes, medidas de supervisión, evaluación de riesgos, vías de reclamación y procedimientos de corrección. Estos elementos se aproximan a obligaciones ya presentes, con diferentes alcances, en el AI Act, M-25-21 y diversas reglas de protección de datos.

8.7. Ventaja comparativa del principio.

La propuesta no depende de afirmar que todos los ordenamientos reconocen un derecho idéntico a explicación. En cambio, se apoya en una pregunta estructural que puede trasladarse entre sistemas jurídicos: ¿puede el ordenamiento permitir que una decisión que afecta derechos exista sin un sujeto capaz de justificarla y responder por ella? La respuesta de un Estado constitucional debería ser negativa.

9. Conclusiones

La inteligencia artificial modifica la manera en que se produce el poder decisorio. Lo hace no solo cuando una máquina adopta automáticamente un resultado, sino también cuando una puntuación o recomendación se convierte en una referencia tan dominante que la intervención humana pierde contenido.

La comparación muestra caminos distintos. Europa combina protección de datos, jurisprudencia sobre decisiones automatizadas y una regulación de IA basada en riesgos. Estados Unidos ofrece garantías constitucionales frente al gobierno, deberes sectoriales de explicación y un ecosistema creciente de políticas federales y normas estatales. América Latina combina protección de datos, justicia constitucional y nuevos marcos de IA con una capacidad especialmente relevante de traducir el problema algorítmico al lenguaje de derechos fundamentales.

Ningún sistema estudiado ofrece una solución perfecta. Existen excepciones, calendarios transitorios, fragmentación institucional, secretos empresariales y áreas todavía sin regulación específica. Pero la ausencia de uniformidad no impide observar una convergencia funcional: las decisiones automatizadas de alto impacto requieren mayor transparencia, control humano, contestabilidad y atribución de responsabilidad.

La principal conclusión de este trabajo es sencilla: la existencia de un algoritmo no rompe la cadena constitucional. Un Estado no puede dejar de responder porque compró un modelo a un proveedor. Un banco no puede reemplazar una razón legalmente exigida por la frase 'el modelo lo determinó'. Un funcionario no ejerce verdadero control si carece de autoridad para contradecir la salida. Un juez no conserva su función si delega a una máquina el razonamiento que la Constitución le exige realizar.

Por ello se propone el Principio de Responsabilidad Constitucional No Delegable. Su núcleo no es anti-tecnológico. La IA puede aumentar capacidad administrativa, detectar patrones, reducir tareas repetitivas y mejorar servicios. El límite aparece cuando la eficiencia tecnológica se utiliza para hacer desaparecer a quien debe justificar el ejercicio del poder.

En una democracia constitucional, toda decisión significativa debe poder regresar desde la arquitectura tecnológica hacia un sujeto jurídicamente responsable. Debe existir alguien que explique. Alguien que revise. Alguien que corrija. Y alguien que responda.

La inteligencia artificial puede automatizar el análisis. Puede automatizar recomendaciones. Puede incluso automatizar partes enteras de un procedimiento. Lo que no debería poder automatizar es la desaparición de la responsabilidad.

Referencias normativas y jurisprudenciales principales

UNIÓN EUROPEA Y EUROPA. Reglamento (UE) 2016/679 (RGPD), especialmente arts. 15 y 22. Carta de los Derechos Fundamentales de la Unión Europea, especialmente arts. 21 y 47. Tribunal de Justicia de la Unión Europea, SCHUFA Holding (Scoring), C-634/21, ECLI:EU:C:2023:957, sentencia de 7 de diciembre de 2023. Tribunal de Justicia de la Unión Europea, Dun & Bradstreet Austria, C-203/22, ECLI:EU:C:2025:117, sentencia de 27 de febrero de 2025. Reglamento (UE) 2024/1689 (Artificial Intelligence Act), especialmente arts. 14, 26, 27, 85 y 86. Reglamento (UE) 2026/1744, Digital Omnibus on AI. Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, CETS No. 225.

ESTADOS UNIDOS. U.S. Constitution, Amendments V and XIV. K.W. v. Armstrong, 789 F.3d 962 (9th Cir. 2015). State v. Loomis, 2016 WI 68, 881 N.W.2d 749. Houston Federation of Teachers, Local 2415 v. Houston Independent School District, 251 F. Supp. 3d 1168 (S.D. Tex. 2017). Equal Credit Opportunity Act, 15 U.S.C. § 1691(d). Regulation B, 12 C.F.R. § 1002.9. CFPB Circular 2022-03. OMB Memorandum M-25-21 (3 April 2025). New York City Local Law 144 (Automated Employment Decision Tools). California Privacy Protection Agency regulations concerning Automated Decisionmaking Technology. Colorado SB26-189. Texas HB 149, Texas Responsible Artificial Intelligence Governance Act.

AMÉRICA LATINA Y SISTEMA INTERAMERICANO. Brasil, Lei 13.709/2018 (LGPD), art. 20. Brasil, PL 2338/2023 (proyecto legislativo; no ley vigente a la fecha de corte). Corte Constitucional de Colombia, Sentencia T-323/24. Chile, Ley 21.719 y nuevo art. 8 bis de la Ley 19.628, con aplicación desde 1 de diciembre de 2026. Uruguay, Ley 18.331, art. 16. Perú, Ley 31814 y Decreto Supremo 115-2025-PCM. Convención Americana sobre Derechos Humanos, arts. 8, 24 y 25. Comisión Interamericana de Derechos Humanos, Comunicado de Prensa 047/26, 20 de marzo de 2026.

Fuentes primarias en línea

RGPD: <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>

AI Act consolidado: <https://eur-lex.europa.eu/eli/reg/2024/1689/2026-07-27/eng>

Reglamento (UE) 2026/1744: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32026R1744>

SCHUFA C-634/21: <https://infocuria.curia.europa.eu/tabs/redirect/juris/liste.jsf?language=en&num=C-634%2F21>

Dun & Bradstreet C-203/22: <https://infocuria.curia.europa.eu/tabs/redirect/juris/liste.jsf?num=C-203%2F22>

Carta de Derechos Fundamentales:

<https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:12012P/TXT>

OMB M-25-21: <https://www.whitehouse.gov/wp-content/uploads/2025/02/M-25-21-Accelerating-Federal-Use-of-AI-through-Innovation-Governance-and-Public-Trust.pdf>

ECOA: <https://uscode.house.gov/view.xhtml?req=granuleid:USC-prelim-title15-section1691>

Regulation B: <https://www.ecfr.gov/current/title-12/chapter-X/part-1002/subpart-A/section-1002.9>

CFPB Circular 2022-03: <https://www.consumerfinance.gov/compliance/circulars/circular-2022-03-adverse-action-notification-requirements-in-connection-with-credit-decisions-based-on-complex-algorithms/>

NYC AEDT: <https://www.nyc.gov/site/dca/about/automated-employment-decision-tools.page>

California ADMT: https://coppa.ca.gov/regulations/ccpa_updates.html

Colorado SB26-189: <https://leg.colorado.gov/bills/sb26-189>

Texas HB149: <https://capitol.texas.gov/billlookup/BillSummary.aspx?Bill=HB149&LegSess=89R>

Brasil LGPD: https://www.planalto.gov.br/ccivil_03/ato2015-2018/2018/lei/L13709compilado.htm

Brasil PL 2338/2023: <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233>

Colombia T-323/24: <https://www.corteconstitucional.gov.co/relatoria/2024/t-323-24.htm>

Chile Ley 21.719: <https://www.bcn.cl/leychile/Navegar?idNorma=1209272>

Uruguay Ley 18.331, art. 16: <https://www.impo.com.uy/bases/leyes/18331-2008/16>

Perú D.S. 115-2025-PCM: <https://www.gob.pe/institucion/pcm/normas-legales/7133522-115-2025-pcm>

Convención Americana: <https://www.oas.org/en/iachr/mandate/basics/convention.asp>

CIDH 047/26: https://www.oas.org/en/iachr/jsForm/?File=%2Fen%2Fiachr%2Fmedia_center%2Fpressreleases%2F2026%2F047.asp

Nota de publicación

Este documento corresponde a la versión 1.0.1 de Research Paper 001 de Constitucionalista Research. Se publica bajo licencia Creative Commons Attribution 4.0 International (CC BY 4.0). DOI: [10.5281/zenodo.22141803](https://doi.org/10.5281/zenodo.22141803). Página oficial: <https://constitucionalista.com/inteligencia-artificial-y-constitucion/>

Este trabajo es investigación académica y no constituye asesoramiento jurídico. Las referencias y estados normativos reflejan la fecha de corte indicada en la portada.