

TRAZABILIDAD, VALIDACIÓN HUMANA Y USO DEFENDIBLE DE INTELIGENCIA ARTIFICIAL Y AGENTES EN EL TRABAJO JURÍDICO: HACIA UN ESTÁNDAR MÍNIMO DE CONTROL PROFESIONAL

Jorge Eduardo Peralta¹

Consultor independiente en IA aplicada al derecho (Argentina)

consultoria@tracewarden.com.ar

<https://orcid.org/0009-0008-1772-8969>

Recibido: 23/03/2026

Aceptado: 18/05/2026

<https://doi.org/10.26422/RJA.2026.0701.per>

Resumen

La incorporación de inteligencia artificial generativa al trabajo jurídico ha dejado de ser una práctica excepcional para transformarse en una realidad creciente en estudios jurídicos, asesorías internas, consultoría especializada y, de manera más controvertida, en ciertos ámbitos de producción jurisdiccional. Sin embargo, el debate todavía suele quedar reducido a parámetros de rapidez, eficiencia, automatización o ahorro de costos. Este trabajo sostiene que esa formulación es jurídicamente insuficiente. Cuando la inteligencia artificial interviene de modo material en tareas de investigación, clasificación, síntesis, redacción o apoyo decisorio, la cuestión relevante ya no es únicamente la corrección aparente del resultado final, sino la posibilidad de reconstruir, controlar y justificar el proceso por el cual dicho resultado fue generado, revisado y adoptado por un operador humano responsable. La hipótesis central es que el uso profesio-

-
- 1 Consultor en inteligencia artificial aplicada al derecho. Desarrolla su trabajo especializado en gobernanza de IA, trazabilidad, validación humana y responsabilidad en el uso profesional de sistemas inteligentes en entornos jurídicos y judiciales con enfoque en la construcción de estándares mínimos de control, documentación y defensabilidad del uso de inteligencia artificial en el trabajo jurídico. Investiga sobre problemas vinculados con supervisión humana, el riesgo profesional y la atribución de responsabilidad en procesos asistidos por IA.

nalmente legítimo de la inteligencia artificial en el trabajo jurídico exige un estándar mínimo de trazabilidad, validación humana sustantiva y carácter defendible del proceso. Ese estándar se intensifica cuando la IA opera bajo arquitecturas agentivas o conectadas a fuentes externas, repositorios, herramientas o bases documentales, debido a riesgos específicos como la inyección de *prompts*, la contaminación de bases de conocimiento, la exfiltración de información sensible y la ejecución de acciones no previstas. Tales riesgos no deben ser tratados únicamente como problemas de ciberseguridad, sino también como cuestiones de diligencia profesional, competencia tecnológica, confidencialidad, integridad argumentativa y atribución de responsabilidad. Sobre esa base, el trabajo propone un criterio mínimo de control profesional orientado a distinguir el uso asistido legítimo de la delegación impropia del juicio jurídico.

Palabras clave: inteligencia artificial, agentes, trabajo jurídico, trazabilidad, validación humana, inyección de *prompts*, responsabilidad profesional.

Traceability, Human Validation, and Defensible Use of Artificial Intelligence and Agents in Legal Work: Toward a Minimum Standard of Professional Control

Abstract

The incorporation of generative artificial intelligence into legal work is no longer exceptional. It is increasingly present in law firms, in-house legal departments, specialized consultancy and, more controversially, in judicial settings. Yet the debate is still often framed in terms of efficiency, speed, or automation. This article argues that such framing is legally insufficient. Whenever artificial intelligence materially intervenes in legal research, classification, synthesis, drafting, or decision-support tasks, the relevant issue is not only the apparent correctness of the final output, but also whether the process through which that output was generated, reviewed, and adopted by a human actor can be reconstructed, controlled, and justified. The article's central hypothesis is that the professionally legitimate use of artificial intelligence in legal work requires a minimum standard of traceability, substantive human validation, and defensible process. That standard becomes more demanding when AI is deployed through agentic architectures or systems connected to external sources, repositories, tools, or knowledge bases, due to specific risks such as prompt injection, knowledge-base contamination, exfiltration of sensitive information, and unauthorized action execution. These risks should not be treated solely as cybersecurity problems, but also as matters of professional diligence, technological competence, confidentiality, argumentative integrity, and responsibility attribution. On that basis, the article proposes a minimum standard of professional control aimed at distinguishing legitimate assisted use from improper delegation of legal judgment.

Key words: artificial intelligence, agents, legal work, traceability, human validation, prompt injection, professional responsibility.

1. Introducción

La incorporación de inteligencia artificial al trabajo jurídico ya no puede ser pensada como una posibilidad futura ni como una cuestión marginal reservada a laboratorios de innovación. Se ha convertido en una práctica creciente y, en muchos casos, silenciosa. Herramientas capaces de resumir documentos, estructurar borradores, proponer argumentos, clasificar evidencia, recuperar normativa o asistir procesos internos ingresaron en estudios jurídicos, asesorías corporativas, consultorías especializadas y, de modo más discutido, en espacios vinculados a la actividad jurisdiccional. Sin embargo, esa incorporación no vino acompañada, en la misma medida, por una reflexión jurídica suficientemente rigurosa sobre las condiciones de legitimidad de su uso.

La conversación pública suele formular el problema en términos operativos: rapidez, ahorro de tiempo, automatización, escalabilidad, eficiencia. Esa forma de plantearlo puede tener utilidad empresarial, pero resulta insuficiente para el derecho. El trabajo jurídico no consiste únicamente en producir textos plausibles ni en acelerar flujos documentales. Supone juicio, responsabilidad, control de fuentes, prudencia, deber de fundamentación, posibilidad de revisión posterior y, en muchos casos, tutela de derechos y garantías. Desde esa perspectiva, el problema no puede quedar agotado en la calidad aparente del resultado final.

Una parte importante del debate reciente se concentró en las llamadas “alucinaciones” de la inteligencia artificial: referencias inexistentes, citas inventadas, hechos deformados o razonamientos plausibles pero falsos. Se trata, sin duda, de un problema real. Pero reducir el análisis a ese fenómeno supone mirar solo la manifestación más visible de una cuestión más profunda. El problema central no es únicamente que la IA pueda equivocarse, sino que el proceso mediante el cual se produce el resultado jurídico se vuelva opaco, difícilmente reconstruible y, en consecuencia, difícilmente atribuible a un humano que pueda explicarlo y defenderlo.

La relevancia de ese desplazamiento del foco —del mero error al proceso— se advierte con especial claridad en el derecho. A diferencia de otros campos en los que puede bastar un criterio funcional de desempeño, la práctica jurídica valora también la forma en la que se obtiene el resultado. Importan la reconstrucción del razonamiento, la verificabilidad de las fuentes, la racionalidad de la motivación, la posibilidad de impugnación y la asignación de responsabilidad. Dicho en términos simples: en el trabajo jurídico no basta con llegar a una conclusión, importa cómo se llega a ella y quién puede responder por ese itinerario.

Esa observación se vuelve todavía más importante cuando la inteligencia artificial ya no opera solo como asistente textual, sino como sistema agéntivo. Cuando un modelo puede consultar repositorios, recuperar documentos, interactuar con herramientas, navegar fuentes externas y encadenar acciones, el riesgo deja de estar concentrado en la incorrección de una frase. El problema pasa a involucrar la integridad del flujo de trabajo, la seguridad del entorno documental, la gobernanza de fuentes y permisos y la persistencia real del control humano. El National Institute of Standards and Technology (NIST, por sus siglas en inglés), en su taxonomía adversarial más reciente, trata expresamente estos riesgos como especialmente relevantes para asistentes conversacionales, sistemas con recuperación aumentada y sistemas agéntivos (Vassilev et al., 2025). De manera convergente, OpenAI advierte que el acceso de agentes a internet o a herramientas incrementa la exposición a inyección de *prompts*, exfiltración de secretos y acciones no deseadas (OpenAI Developers, s.f.-a, s.f.-b).

La hipótesis central de este artículo es la siguiente: cuando la inteligencia artificial interviene materialmente en el trabajo jurídico, y con mayor razón cuando lo hace a través de arquitecturas agéntivas o conectadas a fuentes y herramientas externas, la legitimidad profesional del resultado no puede descansar exclusivamente en su apariencia de corrección, sino en la existencia de un estándar mínimo de trazabilidad, validación humana sustantiva y uso defendible del proceso. Ese estándar no debe entenderse como una exigencia técnica absoluta ni como una transparencia total del sistema, sino como la posibilidad jurídicamente relevante de reconstruir, controlar y justificar cómo fue utilizada, qué lugar ocupó en el proceso de trabajo y por qué el resultado final sigue siendo atribuible a un operador humano responsable.

El objetivo del trabajo es desarrollar esa tesis y proponer un estándar mínimo de control profesional aplicable al uso de inteligencia artificial en el trabajo jurídico. Para ello, primero se examinan las razones por las cuales el problema no puede ser reducido al error visible; luego, se define el papel de la trazabilidad, la validación humana y el uso defendible; posteriormente, se analiza el salto cualitativo que representan los agentes; se diferencia la inyección de *prompts* del envenenamiento de datos o modelos; y, finalmente, se propone una formulación mínima de control profesional destinada a distinguir el uso asistido legítimo de la delegación impropia del juicio jurídico.

2. Metodología y delimitación del problema

Este trabajo adopta una metodología dogmática y analítica. No pretende ofrecer una explicación exhaustiva de la arquitectura técnica de los modelos de lenguaje ni una taxonomía completa de ataques adversariales. Su propósito es más acotado: identificar, desde una perspectiva jurídica, cuáles son las condiciones mínimas para que el uso de inteligencia artificial en tareas propias del trabajo jurídico siga siendo atribuible, controlable y defendible.

La investigación se estructura como una reconstrucción dogmática de deberes profesionales clásicos —diligencia, competencia, supervisión, confidencialidad y deber de explicación— frente a riesgos técnicos identificados por la literatura reciente sobre IA generativa y sistemas agentivos. El análisis no busca describir empíricamente usos de IA, sino elaborar un estándar normativo mínimo a partir de la integración entre ética profesional, gobernanza tecnológica y teoría jurídica contemporánea.

Para ello se toman en consideración cuatro tipos de materiales. En primer lugar, normas e instrumentos regulatorios recientes en materia de inteligencia artificial, transparencia, supervisión humana y gestión del riesgo. En segundo lugar, documentos de gobernanza y marcos de gestión que, aun cuando no constituyen derecho positivo en sentido estricto, ofrecen criterios valiosos para traducir riesgos técnicos en obligaciones organizacionales y profesionales. En tercer lugar, guías éticas y profesionales provenientes del ámbito jurídico, relevantes en materia de diligencia, competencia, confidencialidad y veracidad. En cuarto lugar, algunos antecedentes jurisprudenciales particularmente ilustrativos del problema.

La delimitación del objeto es importante. El artículo se centra en el uso de inteligencia artificial generativa y sistemas agentivos en tareas vinculadas con el trabajo jurídico profesional: investigación, recuperación de información, clasificación documental, síntesis, estructuración argumental, redacción de borradores y apoyo a la toma de decisiones. No se aborda aquí, salvo de modo tangencial, el análisis de sistemas biométricos, analítica predictiva penal, vigilancia masiva u otros supuestos regulatorios que plantean problemas adicionales. Tampoco se desarrolla una teoría general de la responsabilidad civil, penal o disciplinaria por daños derivados de sistemas de IA. El foco está puesto en el estándar mínimo de control exigible para que el uso cotidiano de estas herramientas no degrade el carácter justificable del trabajo jurídico.

El trabajo tampoco adopta una posición prohibicionista. No parte de la idea de que la inteligencia artificial sea incompatible, por definición, con el

ámbito jurídico. Parte, por el contrario, de un dato evidente: la herramienta ya ingresó. La cuestión relevante es bajo qué condiciones ese ingreso puede ser profesionalmente aceptable. Dicho de otro modo, no se trata de discutir si la IA debe entrar o salir del trabajo jurídico, sino de establecer qué exigencias mínimas deben cumplirse para que su utilización no erosione la responsabilidad humana ni vacíe de contenido la obligación de dar razones.

Finalmente, el trabajo adopta una pretensión deliberadamente modesta. No se propone formular una teoría total de la inteligencia artificial aplicada al derecho ni clausurar debates regulatorios abiertos. Su objetivo es más práctico y, precisamente por eso, más exigente: contribuir a construir un lenguaje jurídico capaz de evaluar el uso profesional de sistemas de IA allí donde lo decisivo no es solo el rendimiento funcional, sino también la posibilidad posterior de explicar, justificar y sostener el proceso de producción jurídica.

3. Marco normativo y profesional contemporáneo: de la IA confiable a la exigencia de control

La tesis de este trabajo no surge en el vacío. Aunque el derecho positivo no ha desarrollado todavía un régimen unificado y exhaustivo para todos los usos de inteligencia artificial en la práctica jurídica, sí existe un entramado normativo y profesional convergente que permite reconstruir exigencias mínimas de control, transparencia, revisión y responsabilidad.

Esta reconstrucción exige, sin embargo, evitar un riesgo metodológico: presentar la gobernanza de la IA como si se agotara en estándares técnicos o recomendaciones institucionales. El punto de partida debe seguir siendo jurídico. En términos próximos a los desarrollados por Ober y Tasioulas (2024), la ética de la inteligencia artificial no comienza desde una tabla rasa ni exige abandonar las categorías normativas tradicionales; los sistemas de IA deben ser comprendidos como herramientas inteligentes orientadas a fines humanos y no como sujetos capaces de desplazar la responsabilidad moral o jurídica de las personas que los diseñan, despliegan o utilizan.

Desde otra perspectiva, De Gregorio (2023) ha mostrado que la inteligencia artificial no opera siempre como herramienta neutra, sino que puede producir una capa técnica de normatividad capaz de condicionar prácticas sociales, decisiones públicas y formas de aplicación, generando tensiones con el Estado de derecho. Esta advertencia es especialmente relevante para el trabajo jurídico: si un sistema organiza fuentes, jerarquiza información, sugiere inferencias o

estructura borradores, su influencia no se limita al resultado visible, sino que puede condicionar el entorno decisonal antes de que el profesional adopte formalmente una posición.

En el plano general, la Recomendación de la UNESCO sobre la ética de la inteligencia artificial (2021) asentó una de las bases más influyentes del debate contemporáneo: la necesidad de que el desarrollo y uso de IA se mantenga compatible con la dignidad humana, los derechos fundamentales, la diversidad, la transparencia y la rendición de cuentas. Esa perspectiva no está pensada específicamente para la abogacía, pero sí aporta una idea decisiva para este trabajo: la evaluación de la IA no puede agotarse en su utilidad, sino que debe atender los efectos que produce sobre la autonomía humana, la equidad y la posibilidad de atribuir responsabilidades.

En una línea afín, el NIST ha desarrollado un enfoque de gestión de riesgos que resulta particularmente fértil para traducir problemas técnicos a categorías organizacionales. El AI RMF 1.0 parte de la idea de que comprender y gestionar los riesgos de los sistemas de IA es condición para potenciar su confiabilidad y sostener la confianza pública (NIST, 2023). El Perfil para IA generativa profundiza ese enfoque y enumera riesgos propios de estos sistemas, entre ellos, la confabulación, los sesgos dañinos, la sobredependencia humana y los problemas de seguridad de la información, advirtiendo además que las organizaciones pueden necesitar revisión humana adicional, seguimiento, documentación y mayor supervisión gerencial (NIST, 2024). Esa observación es central: los sistemas generativos no exigen menos control que otras herramientas, sino, en muchos casos, más.

En el ámbito europeo, el Reglamento europeo de IA tampoco regula específicamente el trabajo cotidiano de los abogados, pero sí ofrece una pauta doctrinal y regulatoria muy relevante. La idea de supervisión humana efectiva forma parte de la arquitectura del reglamento para sistemas de alto riesgo, y la doctrina especializada ha mostrado que esa exigencia no puede ser entendida como un mero rito formal, sino como una obligación orientada a prevenir o minimizar riesgos sobre salud, seguridad y derechos fundamentales (Enqvist, 2023; Lazcoz Moratinos, 2025). La relevancia de esta línea para el trabajo jurídico es evidente: cuando una herramienta incide materialmente en actos de relevancia jurídica, la pretensión de “humano en el circuito” solo es seria si ese humano conserva comprensión mínima de capacidades y límites, poder de intervención y responsabilidad efectiva sobre la decisión o adopción del resultado.

En el ámbito de la justicia, la Carta Ética de la Comisión Europea para la

Eficiencia de la Justicia (CEPEJ, por sus siglas en francés) agregó una formulación especialmente útil. Allí se afirma que la IA puede contribuir a la eficiencia y calidad de la justicia, pero solo si se implementa de manera responsable y compatible con los derechos fundamentales; además, identifica principios rectores como respeto de derechos fundamentales, no discriminación, calidad y seguridad, transparencia y control por el usuario (CEPEJ, 2018). Lo valioso de este documento es que evita un falso dilema. No parte de una lógica antitecnológica, pero tampoco admite que la mera mejora de eficiencia pueda justificar la desarticulación de garantías.

Más recientemente, las Directrices de la UNESCO para el uso de sistemas de IA en cortes y tribunales ubican la supervisión humana significativa, la transparencia, la rendición de cuentas y la protección de derechos humanos como principios rectores para la adopción judicial de IA (UNESCO, 2025). Aunque dirigidas primariamente a cortes y tribunales, esas directrices son útiles para el trabajo jurídico en general, porque confirman que el uso institucionalmente legítimo de IA no se mide solo por eficiencia, sino por la preservación de control humano, la responsabilidad y la confianza pública.

Ese entramado normativo general se complementa con directrices profesionales específicas. La American Bar Association (ABA), a través de la Formal Opinion 512 (2024a), sostuvo que los abogados y estudios que utilicen herramientas de IA generativa deben considerar plenamente sus obligaciones éticas, incluyendo competencia, información al cliente, confidencialidad y razonabilidad de honorarios. A nivel europeo, el Council of Bars and Law Societies of Europe (CCBE) (2025) subrayó que el uso de IA generativa por parte de abogados debe mantenerse dentro de las obligaciones profesionales y regulatorias de la práctica jurídica, prestando atención a los beneficios, pero también a los problemas urgentes para la profesión. En ambos casos, el mensaje es claro: la incorporación de IA no crea una zona inmune a las categorías clásicas del ejercicio profesional; por el contrario, reactiva y complejiza deberes preexistentes.

La doctrina iberoamericana también permite reforzar este punto sin depender exclusivamente de guías institucionales. Seleme (2021), al analizar la relación entre ética legal, inteligencia artificial y rol profesional, sostiene que la utilización de herramientas de IA no inaugura un vacío normativo, sino que genera nuevos modos de transgredir deberes profesionales ya existentes: diligencia y competencia, supervisión, información al cliente, confidencialidad, no facilitación del ejercicio ilegal de la profesión y no participación en conductas discriminatorias. Esta lectura resulta especialmente útil para el presente trabajo

porque desplaza el análisis desde la fascinación tecnológica hacia el rol profesional: la IA no suspende las obligaciones del abogado, sino que modifica los escenarios en los que esas obligaciones deben ser interpretadas y aplicadas.

También en el ámbito argentino comienzan a aparecer materiales institucionales y doctrinales que van en esa dirección. La guía del Colegio Público de la Abogacía de la Capital Federal (CPACF) (2025) para el uso de inteligencia artificial por abogados parte de una premisa realista: la IA generativa “irrumpió en el mundo jurídico y llegó para quedarse”, pero la clave está en usarla con “criterio, responsabilidad y conciencia plena de sus límites” (p. 1). La importancia de este documento no radica en ofrecer un marco teórico exhaustivo, sino en registrar institucionalmente que, en el foro argentino, el problema ya se percibe como un asunto de buen uso profesional y no solo de adopción tecnológica.

En suma, el marco normativo y profesional contemporáneo no habilita una lectura ingenua del uso de IA en derecho. Aunque sus instrumentos provengan de ámbitos distintos —ética global, gestión de riesgos, regulación europea, guías para justicia y deontología profesional—, todos convergen en una idea común: la legitimidad del uso de IA depende de la permanencia de estructuras de control humano, comprensión suficiente, responsabilidad y trazabilidad. Esa convergencia constituye el presupuesto normativo de la tesis que aquí se sostiene.

4. De las alucinaciones a la opacidad del proceso

La crítica más inmediata al uso de IA en derecho suele concentrarse en los errores visibles. El ejemplo paradigmático es la cita inexistente, el fallo inventado, la norma deformada, la fecha equivocada o el resumen que omite un dato decisivo. Esa reacción es comprensible: los errores ostensibles son fáciles de identificar y tienen alto poder pedagógico. Sin embargo, si el análisis se detiene allí, se corre el riesgo de trivializar el verdadero problema jurídico.

En el campo jurídico, un producto puede no contener errores groseros y, aun así, haber sido elaborado mediante un proceso defectuoso. Un borrador puede sonar razonable, un informe puede estar bien escrito y una línea argumental puede parecer consistente, pero todo ello no alcanza si nadie puede explicar con precisión qué función cumplió la herramienta, qué insumos recibió, qué fuentes consultó, qué materiales recuperó y qué revisión sustantiva realizó el profesional antes de adoptar el resultado. La opacidad del proceso puede coexistir con una salida aparentemente correcta.

Esa observación es particularmente importante porque el derecho no valora

solo resultados, sino también condiciones de producción. La legitimidad de una actuación jurídica no se mide únicamente por su eficacia pragmática, sino también por su racionalidad, verificabilidad, controlabilidad y atribuibilidad. Por eso, la pérdida de inteligibilidad del proceso no es un problema lateral. Afecta la posibilidad de dar razones, controlar el itinerario decisional y sostener la responsabilidad del profesional o del magistrado.

El caso *Mata v. Avianca Inc.* (2023) se ha convertido, por esa razón, en una referencia inevitable. Allí, el tribunal sancionó a los abogados que habían presentado precedentes inexistentes obtenidos mediante ChatGPT y destacó que la producción había sido trasladada al expediente sin verificación suficiente, imponiendo incluso una multa económica. Lo más importante del caso no es solo que la herramienta “alucinó”, sino que los profesionales rompieron el deber mínimo de control. El escándalo no fue tecnológico en sentido estricto, fue profesional y procesal.

La experiencia argentina reciente confirma que el problema no es excepcional ni meramente comparado. En *Giacomino c/ Monserrat* (2025), la Cámara de Apelación en lo Civil y Comercial de Rosario advirtió que el letrado había incorporado jurisprudencia inexistente obtenida mediante IA generativa y remarcó que la alegada buena fe no lo relevaba del deber de cotejar las fuentes invocadas. En *Acevedo c/ Cáceres* (2025), la Cámara de Apelación en lo Civil y Comercial de Morón declaró desierto un recurso y puso en conocimiento del letrado y del colegio profesional la utilización de precedentes jurisprudenciales inexistentes. Ambos casos muestran que el deber de verificación no es una carga accesorio, sino una condición mínima de probidad y responsabilidad profesional cuando se utilizan herramientas generativas en presentaciones judiciales.

Ese antecedente puso de manifiesto una cuestión que este artículo considera central: el defecto relevante comienza antes del error final. Comienza cuando la intervención humana se degrada a una revisión ritual o decorativa, cuando el operador no comprende suficientemente las limitaciones de la herramienta o cuando el proceso no deja un rastro inteligible de qué ocurrió. En esos casos, el problema ya no es únicamente la eventual incorrección del texto, sino la imposibilidad de defender el proceso de producción jurídica.

El NIST (2024) ofrece una formulación que ayuda a conceptualizar esta cuestión. El Perfil para IA generativa advierte que los humanos pueden sobreconfiar en estos sistemas o atribuirles una calidad superior a la de otras fuentes, fenómeno que vincula con sesgo de automatización y excesiva deferencia a sistemas automatizados. En el trabajo jurídico, esa observación es especialmente

grave. La sobreconfianza no solo incrementa la probabilidad de error, también reduce la calidad de la revisión humana y, con ello, deteriora las condiciones bajo las cuales el resultado puede seguir siendo considerado un producto profesional propio.

La advertencia puede extenderse más allá de las alucinaciones textuales. Matsumi y Solove (2025) sostienen que las predicciones algorítmicas presentan problemas propios porque operan como inferencias sobre el futuro y se ubican en una zona ambigua entre verdad y falsedad, difícilmente capturable por los deberes tradicionales de exactitud o corrección de datos. Aunque el presente trabajo no se centra en sistemas predictivos, esa tesis refuerza una idea aplicable al trabajo jurídico asistido por IA: no todo defecto se presenta como falsedad verificable; muchas veces aparece como inferencia, priorización o recomendación cuya debilidad solo puede detectarse si existe control humano sustantivo sobre el proceso.

La literatura reciente sobre ética profesional de los abogados refuerza esta lectura. Terzidou (2025), al estudiar el impacto de la IA generativa sobre las obligaciones de competencia y confidencialidad, muestra que la cuestión no puede analizarse desde el mero rendimiento funcional, ya que la calidad del servicio jurídico también depende de la conservación de la agencia profesional y de la protección efectiva de la información del cliente. En el mismo sentido, la doctrina sobre supervisión humana en el Reglamento europeo de IA ha insistido en que la supervisión no puede equivaler a un asentimiento pasivo, sino que debe articularse con comprensión, posibilidad de intervención y orientación preventiva del riesgo (Enqvist, 2023; Lazcoz Moratinos, 2025).

En síntesis, la crítica jurídicamente seria al uso de IA no puede quedar confinada a las “alucinaciones”. Lo que debe examinarse es si el proceso sigue siendo inteligible, si el profesional conserva un dominio suficiente de lo que ocurre y si el resultado puede ser posteriormente explicado y defendido sin fisuras. De allí surge la necesidad de reformular el problema en torno a tres categorías: trazabilidad, validación humana sustantiva y uso defendible.

5. Trazabilidad, validación humana y uso defendible

La primera categoría necesaria para ordenar este problema es la trazabilidad. En este trabajo no se la utiliza en un sentido técnico totalizante, como si exigiera una auditoría completa del funcionamiento interno del modelo. Se la utiliza en un sentido jurídicamente operativo: la capacidad de reconstruir, de manera

suficiente, la secuencia relevante de intervención de la inteligencia artificial dentro del flujo de trabajo jurídico.

Un proceso trazable debe permitir responder, al menos, seis preguntas: qué sistema fue utilizado, para qué tarea concreta, sobre qué insumos, mediante qué fuentes o repositorios, con qué límites o permisos y con qué intervención humana. La utilidad de esta reconstrucción es inmediata. Permite diferenciar el uso de una herramienta para tareas preliminares del uso en momentos cercanos a la decisión final, identificar si se expusieron datos sensibles y, sobre todo, sostener una explicación posterior del proceso cuando el resultado sea cuestionado por un cliente, un tribunal, una contraparte, un auditor o un órgano disciplinario.

La segunda categoría es la validación humana. Aquí conviene ser tajantes: no toda intervención humana equivale a una supervisión jurídicamente adecuada. Corregir estilo, mejorar redacción o leer superficialmente un borrador generado por IA no constituye validación suficiente. En el trabajo jurídico, la validación exigible debe ser sustantiva. Debe incluir comprobación de fuentes, contraste de hechos, revisión de adecuación normativa, control de coherencia argumentativa, detección de omisiones relevantes y evaluación del ajuste al caso concreto.

Esta exigencia se intensifica a medida que la herramienta se aproxima al núcleo del juicio jurídico. No es lo mismo pedir un esquema preliminar de un escrito que delegar la estructuración de la tesis central del caso, la selección de hechos jurídicamente relevantes, la ponderación de riesgos o la valoración preliminar de la prueba. Allí donde el error o la opacidad puedan afectar derechos, garantías o responsabilidades concretas, la revisión humana debe ser proporcionalmente más profunda. El Reglamento europeo de IA ofrece una analogía útil: la supervisión humana se concibe para prevenir o minimizar riesgos y para permitir que la persona comprenda, interprete y, si corresponde, ignore o corrija la salida del sistema (Lazcoz Moratinos, 2025; Unión Europea, 2024).

La tercera categoría, que articula las dos anteriores, es el uso defendible del proceso. Un proceso es defendible cuando puede ser explicado posteriormente de manera inteligible, razonable y documentada ante un tercero relevante. Ese tercero puede ser un cliente, un juez, una contraparte, un superior jerárquico, un auditor o una autoridad disciplinaria. El uso defendible no exige omnisciencia técnica ni acceso a cada parámetro del modelo. Exige algo más modesto y, para el derecho, mucho más importante: que el profesional pueda justificar por qué utilizó determinada herramienta, qué lugar ocupó en el flujo de trabajo, qué controles aplicó y por qué el producto final sigue siendo atribuible a él.

Conviene precisar el uso terminológico. En esta versión, se prefiere hablar de “uso defendible” y reservar “defensabilidad” como una categoría operativa de síntesis. No se la propone como una institución jurídica autónoma ni como un deber nuevo separado de la dogmática existente, sino como una forma abreviada de nombrar la convergencia entre diligencia, competencia, supervisión, confidencialidad, motivación y atribución de responsabilidad cuando esas exigencias se proyectan sobre procesos jurídicos asistidos por IA.

Hablar de uso defendible permite, además, evitar dos errores frecuentes. El primero es el entusiasmo ingenuo que equipara utilidad con legitimidad. Que una herramienta ahorre tiempo no demuestra, por sí misma, que su uso sea jurídicamente aceptable. El segundo error es el rechazo absoluto, que trata toda intervención de IA como incompatible con el trabajo jurídico serio. La categoría operativa de uso defendible permite una posición más sobria: el uso de IA puede ser legítimo, pero solo si permanece bajo condiciones de control y explicación compatibles con la responsabilidad profesional.

Desde un punto de vista dogmático, la utilidad de esta categoría es mayor de lo que parece. Traducir el problema al carácter defendible del proceso permite vincularlo con deberes jurídicos tradicionales: diligencia, veracidad, motivación, tutela judicial efectiva, deber de información, supervisión de auxiliares y responsabilidad por el producto profesional. No se trata, entonces, de imponer categorías “tecnológicas” extrañas al derecho, sino de releer exigencias clásicas a la luz de una nueva mediación técnica.

La literatura jurídica reciente se mueve cada vez más en esa dirección. Enqvist (2023) ha mostrado que la supervisión humana en el Reglamento europeo de IA no es una cláusula vacía, sino una exigencia estructural cuya eficacia depende de cómo se defina qué debe ser supervisado, cuándo y por quién. Lazcoz Moratinos (2025) profundiza esa idea al destacar que la supervisión debe ser efectiva por diseño y orientarse a prevenir o minimizar riesgos. En el ámbito de la ética profesional, la ABA (2024a) y la CCBE (2025) también convergen en la necesidad de que el abogado conserve comprensión razonable de la herramienta, proteja la información del cliente y no desplace sin más su responsabilidad sobre el sistema.

El aporte de Seleme (2021) permite precisar todavía más esta idea. Al examinar el deber de supervisión, sostiene que la aparición de herramientas de IA desplaza el foco desde el tipo de agente que presta asistencia hacia el tipo de servicio que no es realizado directamente por el abogado. En esa lógica, la herramienta puede asistir, pero no eximir del deber de control, y la transparencia

del sistema, aunque necesaria, no basta si el resultado no es explicable de modo interpretable para el usuario profesional. Esta formulación se corresponde con la noción aquí propuesta de validación humana sustantiva: no alcanza con que el abogado use o conozca superficialmente la herramienta, sino que debe poder vigilar razonablemente la calidad, los límites y la pertinencia del resultado que incorpora a su trabajo.

La noción de control humano significativo confirma esta exigencia. Robbins (2024) advierte que el concepto suele utilizarse de manera intuitiva y estrecha, sin aclarar qué humano controla, qué controla, cuándo controla y con qué capacidad efectiva de intervención. Esa observación es decisiva para el trabajo jurídico: no basta con afirmar que existe un profesional “en el circuito”, debe poder determinarse si ese profesional controla los insumos, las fuentes, las salidas, la adopción final del resultado o solo realiza una aprobación superficial posterior.

Por todo ello, trazabilidad, validación humana y uso defendible no deben entenderse como exigencias accesorias. Constituyen, conjuntamente, el mínimo vocabulario jurídico necesario para evaluar la legitimidad del uso de inteligencia artificial en el trabajo profesional del derecho.

6. Del uso asistivo al uso agentivo: ampliación de la superficie de riesgo

En una primera etapa, muchas herramientas de inteligencia artificial fueron utilizadas como asistentes textuales. Resumían, proponían estructuras de párrafos, sintetizaban materiales o respondían preguntas sobre un conjunto relativamente delimitado de datos. Aunque incluso en ese escenario ya existían riesgos claros, el modelo de uso seguía siendo relativamente acotado: el profesional consultaba, recibía una respuesta y decidía, al menos en teoría, qué hacer con ella.

La incorporación de agentes modifica ese escenario de manera cualitativa. Un agente no necesariamente se limita a responder una pregunta. Puede consultar recursos externos, recuperar documentos desde repositorios, usar herramientas conectadas, interactuar con bases de conocimiento, navegar contenidos y encadenar acciones. Esa ampliación no representa solo un incremento de capacidad. Representa también una ampliación de la superficie de riesgo.

El NIST lo expresa con claridad al señalar que varios ataques y mitigaciones deben analizarse en el contexto de sistemas de IA generativa, tales como sistemas con recuperación aumentada, asistentes conversacionales y sistemas agentivos (Vassilev et al., 2025). La importancia de esa precisión es decisiva. El

sistema ya no opera en un entorno acotado y estático, sino en un ecosistema dinámico, compuesto por recursos externos, conectores, archivos, sitios web, herramientas auxiliares y permisos de ejecución. En esa arquitectura, el riesgo no se limita a la precisión de la salida: alcanza el propio itinerario de producción.

Los materiales técnicos y de seguridad más recientes confirman esta ampliación. OpenAI Developers (s.f.a) advierte que habilitar acceso de agentes a internet incrementa riesgos como inyección de *prompts* desde contenido no confiable, exfiltración de código o secretos, descarga de contenido malicioso y acciones inseguras, recomendando restringir dominios y métodos permitidos, además de revisar la salida y el registro de trabajo del agente. Aunque esa advertencia proviene de documentación técnica, su traducción jurídica es evidente: cuando el sistema interactúa con el entorno, también interactúa con los riesgos del entorno.

En el trabajo jurídico, esta mutación es especialmente sensible. Estudios, asesorías y equipos jurídicos manejan documentos reservados, estrategias procesales, datos personales, comunicaciones sensibles y borradores cuyo valor depende precisamente de su control y confidencialidad. Si un agente funciona sobre ese ecosistema sin límites claros, sin segmentación de acceso, sin controles sobre fuentes y sin revisión suficiente, la cuestión deja de ser una mera innovación productiva y se transforma en un problema serio de diligencia profesional y de preservación de la integridad del trabajo jurídico.

Aquí aparece una primera consecuencia teórica relevante: la evaluación del uso de IA no puede seguir centrada exclusivamente en la interfaz conversacional. En arquitecturas agentivas, el objeto del análisis jurídico debe desplazarse hacia el flujo de trabajo completo: fuentes, permisos, herramientas, registros, recorridos, validaciones y puntos de intervención humana. En otras palabras, el centro de gravedad se traslada desde la respuesta aislada al sistema sociotécnico que la produce.

Este desplazamiento permite conectar el análisis con la tesis de De Gregorio (2023) sobre el poder normativo de la IA. En sistemas jurídicos mediados por herramientas que recuperan, priorizan y estructuran información, la normatividad no aparece únicamente en la regla jurídica formal, sino también en la arquitectura técnica que filtra qué materiales se vuelven visibles, qué relaciones se sugieren y qué cursos de acción parecen razonables. Por eso, en contextos agentivos, la trazabilidad no es un requisito meramente administrativo: es una condición de control del poder técnico que incide sobre el razonamiento jurídico.

La propia práctica jurídica sugiere por qué esto es importante. Un agente

conectado a bases documentales podría, por ejemplo, priorizar materiales no confiables; un agente con acceso a correos podría exponer información reservada; un agente que navegue fuentes externas podría recuperar instrucciones ocultas; y un agente que utilice herramientas automatizadas podría ejecutar acciones impropias o desalineadas con la finalidad jurídica original. El daño aquí no siempre se manifiesta en una frase falsa. A veces se evidencia en la contaminación del proceso, en la fuga de información o en la adopción de un curso de acción que el profesional no habría seguido si hubiera mantenido un control más fuerte.

La literatura doctrinal sobre la práctica profesional del derecho ayuda a comprender por qué este punto no es meramente técnico. Terzidou (2025) muestra que el uso de IA generativa por abogados impacta directamente en las obligaciones de competencia y confidencialidad, justamente porque la herramienta no es neutra respecto de la forma en que se produce y circula la información sensible del cliente. En el ámbito argentino, Santos (2025) advierte que la utilización de IA en dictámenes jurídicos administrativos puede aportar eficiencia, pero exige regulación adecuada para evitar falta de transparencia y afectaciones a la tutela administrativa efectiva. En ambos casos, el problema no es si la herramienta “sirve”, sino cómo reorganiza las condiciones de agencia y responsabilidad.

Por ello, el paso del uso asistivo al uso agentivo no debe ser tratado como un simple refinamiento tecnológico. Marca un cambio de escala del problema jurídico. La exigencia de trazabilidad, validación y uso defendible no desaparece, se vuelve más intensa. Y esa intensificación es, precisamente, una de las claves del estándar que aquí se propone.

7. Inyección de *prompts*, contaminación documental y distinción frente al envenenamiento

En este punto, es necesario ser conceptualmente rigurosos. No conviene utilizar como equivalentes categorías que, en términos técnicos, no lo son. La inyección de *prompts* y el envenenamiento de datos o modelos no constituyen el mismo fenómeno, aunque puedan converger funcionalmente en ciertos escenarios.

La inyección de *prompts* se refiere, en términos generales, a la introducción de instrucciones o contenidos diseñados para alterar el comportamiento esperado del sistema. Esa manipulación puede ser directa, cuando aparece en el propio insumo visible, o indirecta, cuando llega a través de documentos,

sitios, correos, archivos o recursos externos que el sistema recupera y procesa. En este segundo caso, el modelo o el agente puede tratar contenido aparentemente informativo como si fuera una instrucción operativa. El NIST describe los ataques de inyección indirecta de *prompts* como una vía para comprometer la privacidad e integridad de sistemas conectados, incluidos asistentes conversacionales de internet, sistemas con recuperación aumentada y otros sistemas generativos (Vassilev, 2025).

El envenenamiento, en cambio, se vincula con la contaminación de datos de entrenamiento, ajuste o componentes estructurales del sistema con el fin de modificar su comportamiento de manera más persistente o profunda. El NIST distingue expresamente entre ataques de inyección de *prompts*, *data poisoning* y *model poisoning*, y describe estos últimos como mecanismos capaces de inducir comportamientos alterados o *backdoors* mediante la inserción maliciosa de datos o actualizaciones comprometidas (Vassilev, 2025). No toda inyección es, por tanto, un envenenamiento. Esta distinción importa porque evita formular el problema de manera técnicamente débil.

Ahora bien, en entornos jurídicos con bases documentales conectadas, repositorios de conocimiento o sistemas con recuperación aumentada, sí puede aparecer un fenómeno que resulta jurídicamente muy relevante: la contaminación documental o contaminación de la base de conocimiento. Documentos, correos, páginas o archivos aparentemente neutros ingresan al circuito y contienen instrucciones, patrones o estructuras diseñadas para desviar el comportamiento del sistema, alterar prioridades, inducir respuestas impropias o facilitar exfiltración de información.

Desde la perspectiva del derecho, lo importante no es solo la taxonomía técnica, sino también la consecuencia jurídica. Si el sistema consulta una fuente contaminada y la trata como una directiva operativa, el proceso queda comprometido aunque el defecto no sea evidente a simple vista en el resultado final. Por eso, en arquitecturas agentivas o conectadas a repositorios, la exigencia de trazabilidad, separación entre fuentes confiables y no confiables, limitación de permisos y validación humana se vuelve todavía más intensa.

El Open Worldwide Application Security Project (OWASP) (2025), a través de su proyecto de seguridad en IA generativa, ha descrito la inyección de *prompts* como una vulnerabilidad estructural en aplicaciones basadas en modelos de lenguaje grandes (LLM, por sus siglas en inglés), asociada al hecho de que instrucciones y datos son procesados conjuntamente, lo que permite que contenidos maliciosos reorienten el comportamiento del modelo. Esa observación es

especialmente útil para el análisis jurídico, porque muestra que el problema no depende exclusivamente de mala fe del usuario final. Puede surgir de la propia arquitectura de integración cuando no existe una separación adecuada entre datos confiables y datos no confiables.

La discusión reciente sobre *legal prompting* también obliga a cierta cautela. Faliero (2025) advierte sobre la obsolescencia de enfoques oportunistas centrados en el *prompting legal/legal prompting* y en la IA generativa jurídica como si la formulación de instrucciones fuera, por sí sola, una solución suficiente para la práctica jurídica con IA generativa. Esa advertencia converge con la tesis de este trabajo: el riesgo no se resuelve mediante mejores *prompts* aislados, sino mediante gobernanza del proceso, control de fuentes, delimitación de tareas, validación humana y responsabilidad profesional.

La traducción dogmática de este punto es importante. Allí donde el sistema puede recuperar materiales desde múltiples fuentes, el deber de diligencia ya no se satisface solo con “leer la salida del sistema”. Exige diseñar o seleccionar arquitecturas en las que exista alguna gobernanza de fuentes, una política de permisos, un criterio de segmentación de datos y un nivel razonable de revisión de recorridos y salidas. De lo contrario, la revisión humana llegará demasiado tarde, cuando el proceso ya haya sido condicionado por insumos hostiles o no confiables.

En el terreno de la práctica jurídica, este riesgo se vuelve especialmente delicado por la naturaleza de los materiales involucrados. Un documento con instrucciones ocultas incrustadas, una página con contenido malicioso, una base de conocimiento no depurada o un correo aparentemente inocuo pueden alterar el comportamiento de un agente que trabaje sobre estrategia procesal, clasificación de evidencia, elaboración de informes o priorización de tareas. El problema, de nuevo, no es únicamente que una frase resulte falsa, es que el itinerario mismo de formación del resultado quede secuestrado o contaminado.

Por eso conviene insistir en una precisión conceptual adicional. En el discurso profesional sobre IA aplicada al derecho se observa a veces una tendencia a nombrar todo este conjunto de riesgos como “envenenamiento de datos”. Esa simplificación es comprensible, pero jurídicamente inconveniente. Lo correcto es distinguir entre inyecciones de instrucciones en tiempo de uso, contaminación de bases de conocimiento o repositorios y envenenamiento propiamente dicho de datos o modelos. Solo sobre esa base puede construirse una respuesta organizacional y jurídica proporcionada al riesgo.

8. Consecuencias jurídicas: diligencia, competencia, confidencialidad y deber de explicación

Los riesgos aquí examinados no deben leerse únicamente como problemas de seguridad informática. En el trabajo jurídico impactan directamente sobre deberes profesionales clásicos. La competencia profesional, la diligencia, la veracidad ante el tribunal, la confidencialidad, la prudencia en el manejo de información sensible y la capacidad de justificar el propio trabajo son dimensiones afectadas de manera directa por el uso no gobernado de herramientas de IA y, con mayor razón, de agentes.

La consecuencia más obvia es la que recae sobre el deber de diligencia. Si un profesional utiliza una herramienta cuyos límites no comprende mínimamente, cuyos resultados no verifica y cuyo proceso no puede reconstruir, su actuación difícilmente pueda considerarse diligente. La inteligencia artificial no desplaza la obligación de revisar, contrastar, verificar y asumir la responsabilidad por el producto final. En todo caso, la intensifica. La ABA (2024a) ha sido explícita al respecto: la competencia requiere comprensión razonable de capacidades y limitaciones de la herramienta, así como revisión y examen de la salida para asegurar su exactitud.

Esa idea se articula, además, con un deber más amplio de competencia tecnológica. Jara (2021) ya había advertido, antes del auge actual de la IA generativa, que la transformación digital modifica de manera profunda el ejercicio profesional y exige que el abogado conozca las nuevas tecnologías para estar preparado ante retos presentes y futuros. El punto, sin embargo, no debe ser leído de forma ingenuamente tecnocrática. Competencia tecnológica no significa que todo jurista deba convertirse en ingeniero. Significa, más modestamente, que no puede seguir utilizando herramientas cuyo impacto sobre la producción jurídica desconoce por completo.

También se ve comprometida la confidencialidad. La incorporación de sistemas conectados a fuentes, repositorios o herramientas externas obliga a una consideración mucho más cuidadosa de qué información ingresa al sistema, bajo qué condiciones, con qué permisos y con qué grado de exposición. Allí donde existen deberes de reserva o secreto profesional, una arquitectura técnicamente potente, pero mal gobernada, puede traducirse en una afectación grave del estándar profesional esperado. Terzidou (2025), al analizar la práctica de abogados que utilizan IA generativa, subraya justamente que la divulgación de información del cliente a terceros —incluidos proveedores del sistema— compromete el deber de confidencialidad si no existen salvaguardas adecuadas o consentimiento suficiente.

En el contexto argentino, los materiales recientes van en la misma dirección. La Guía del CPACF (2025) insiste en que la IA debe usarse con criterio, responsabilidad y conciencia de límites y presenta la adopción de herramientas de IA como un problema inseparable del uso seguro y ético. Santos (2025), a su vez, al estudiar el uso ético de IA en dictámenes jurídicos administrativos, pone el foco en riesgos como la falta de transparencia y la posible afectación de la tutela administrativa efectiva, mostrando que el problema no se agota en la privacidad, sino que se proyecta sobre la validez misma de la argumentación institucional.

Una dimensión adicional es la veracidad frente al tribunal o frente al destinatario del producto jurídico. No se trata solo de evitar falsedades deliberadas. Se trata de impedir que la mediación de IA permita introducir, por negligencia o por automatización acrítica, afirmaciones no verificadas en escritos, dictámenes, informes o decisiones. El caso *Mata* ilustra este punto de forma extrema, pero la lógica es más amplia: un profesional que no controla suficientemente el proceso de producción arriesga convertir al expediente o al cliente en el lugar donde se hace visible la falla que debió haber sido detectada antes.

Por último, aparece una dimensión que este trabajo considera especialmente importante: el deber de explicación. El trabajo jurídico no se legitima solo por producir una respuesta, sino por poder dar razones sobre cómo se la obtuvo. Esa exigencia no vale únicamente para la sentencia judicial. Vale, en otro nivel, para el abogado, el consultor, el equipo interno o la organización que adopta un producto elaborado con intervención de IA. La posibilidad de una defensa futura depende de esa capacidad de explicación. Sin ella, la atribución de responsabilidad se debilita y la confianza en el producto jurídico se vuelve precaria.

En este punto, la doctrina argentina aporta un matiz decisivo. Seleme (2021) organiza el impacto de la inteligencia artificial sobre el rol profesional alrededor de deberes éticos concretos y muestra que la aparición de herramientas de IA genera nuevos modos de incumplir exigencias tradicionales. En materia de diligencia y competencia, advierte que el abogado debe mantenerse actualizado respecto de los beneficios y riesgos asociados al uso de tecnología. En materia de supervisión, subraya que el deber de control ya no se limita a asistentes humanos, sino que se extiende a servicios prestados por herramientas de IA. En materia de confidencialidad, destaca que el uso de sistemas externos puede comprometer información protegida cuando los datos se almacenan o procesan fuera del ámbito de control del profesional. Esta reconstrucción permite reforzar la tesis de este artículo: el estándar de uso defendible no crea deberes

exóticos, sino que traduce obligaciones profesionales clásicas a un entorno técnico que altera sus formas de incumplimiento.

En definitiva, diligencia, competencia, confidencialidad y deber de explicación no son categorías superadas por la IA. Son precisamente las categorías a través de las cuales debe evaluarse si su uso profesional resulta legítimo.

9. Densidad argentina e iberoamericana: una lectura profesional del problema

Una crítica frecuente a los trabajos sobre inteligencia artificial en derecho elaborados desde América Latina consiste en que muchas veces importan, casi sin mediaciones, categorías regulatorias o deontológicas producidas en Europa o en Estados Unidos. Esa importación puede ser útil, pero resulta insuficiente si no se la contrasta con las condiciones institucionales, profesionales y culturales de nuestros sistemas jurídicos. El problema de la IA en el trabajo jurídico no se presenta del mismo modo en un gran estudio transnacional, en una asesoría jurídica pública, en un juzgado provincial o en un ejercicio liberal individual. Por eso, para que el estándar propuesto en este artículo sea útil, debe dialogar también con materiales argentinos e iberoamericanos.

En el plano argentino e iberoamericano, la literatura reciente todavía es incipiente, pero ya ofrece señales relevantes. Granero (2018) anticipó tempranamente la utilidad de programas de inteligencia artificial para asistir la labor de abogados, jueces y legisladores, aunque vinculando esa promesa con una reflexión prudente sobre responsabilidad. Seleme (2021), por su parte, situó el problema dentro de la ética legal y mostró que el uso de IA impacta sobre deberes profesionales concretos, especialmente diligencia, competencia, supervisión, información y confidencialidad. En ambos casos, el valor no está en la fascinación por la herramienta, sino en la insistencia sobre la responsabilidad humana que debe acompañar su adopción.

Jara (2021), por su parte, había mostrado tempranamente que la transformación digital del ejercicio profesional exige revisar incumbencias, hábitos y expectativas de formación del abogado contemporáneo. Aunque su trabajo antecede al auge de los agentes y de los modelos generativos actuales, conserva plena actualidad en un punto crucial: el profesional del derecho ya no puede refugiarse en una ignorancia tecnológica completa y pretender, aun así, conservar intacto su estándar de competencia. Esa observación dialoga directamente con la exigencia de comprensión razonable de capacidades y límites que luego formularían con mayor precisión la ABA y otras guías profesionales.

Más cerca del problema específico que aquí se analiza, Santos (2025) ha estudiado el uso ético de inteligencia artificial en la redacción de dictámenes jurídicos en el derecho administrativo argentino. Su trabajo destaca que la IA puede mejorar eficiencia y reducir costos, pero insiste en que debe estar adecuadamente regulada para evitar riesgos de discriminación algorítmica, falta de transparencia y vulneración de la tutela administrativa efectiva. La importancia de este enfoque es que muestra cómo el problema se desplaza rápidamente desde la productividad hacia categorías propiamente jurídicas: transparencia, razonabilidad, debido procedimiento y resguardo de derechos.

A nivel institucional, la CPACF (2025) para abogados constituye un indicador significativo del estado del debate local. El documento parte de una premisa pragmática: la IA generativa ya ingresó al mundo jurídico y no debe ser tratada como fenómeno pasajero, pero su uso requiere criterio, responsabilidad y plena conciencia de límites. Aunque se trata de un texto práctico y no de una teoría sistemática, tiene un valor normativo informal importante. Muestra que, en el foro argentino, el problema comienza a ser leído como un asunto de buena práctica profesional y no solo de curiosidad tecnológica o *marketing* legal.

En el espacio iberoamericano, la experiencia española resulta especialmente útil porque combina tradición jurídica cercana, regulación europea aplicable y desarrollo reciente de doctrina específica. Rayón Ballesteros (2025) propone para España una aproximación equilibrada que articula marco jurídico, fundamentos técnicos y buenas prácticas, subrayando que la adopción de IA generativa debe preservar confidencialidad, competencia, exactitud, independencia y responsabilidad. Su aporte es particularmente valioso porque traduce el problema a protocolos de gobernanza interna, formación, evaluación de riesgos y revisión humana, evitando tanto la fascinación por la eficiencia como el rechazo abstracto a la herramienta.

La doctrina procesal reciente también ofrece un punto de apoyo para comprender la dimensión estratégica del problema. Nieva Fenoll (2025) analiza cómo la inteligencia artificial generativa puede alterar la lógica tradicional de la defensa cuando la argumentación ya no se dirige exclusivamente a persuadir el marco mental de un juez humano, sino también a interactuar con sistemas capaces de procesar y valorar patrones textuales de otro modo. Sin asumir que la decisión jurisdiccional pueda ser sustituida por una máquina, esa reflexión refuerza la necesidad de conservar claridad sobre qué tareas son asistibles y cuáles permanecen indelegablemente vinculadas al juicio humano.

Esa producción iberoamericana muestra algo que conviene destacar. La

cuestión no se reduce a “regular IA”, sino a adaptar la estructura de obligaciones de la profesión jurídica a un entorno en el que la producción del trabajo se vuelve híbrida, mediada por sistemas técnicos que pueden asistir, expandir la capacidad, contaminar fuentes o debilitar controles. En América Latina, donde conviven altas exigencias formales, infraestructura desigual, prácticas forenses heterogéneas y una fuerte centralidad de la responsabilidad personal del abogado o funcionario, este punto es todavía más delicado.

De allí surge una consecuencia importante para el estándar que aquí se propone. No basta con trasladar sin más el lenguaje del *compliance* tecnológico o de la *AI governance* corporativa. Es necesario traducirlo a categorías operativas para despachos, organismos públicos y estudios jurídicos locales. En nuestro contexto, eso significa hablar de reserva profesional, debida verificación, control del expediente, reconstrucción del razonamiento, revisión de borradores, protección de datos personales, prudencia en el uso de proveedores externos y conservación de materiales de respaldo. El concepto de carácter defendible del proceso busca precisamente cumplir esa función de traducción.

También conviene advertir un riesgo específico del contexto regional. En ecosistemas jurídicos con fuerte presión por resultados, escasez de recursos y sobrecarga burocrática, la IA puede ser percibida como atajo irresistible. Ese incentivo es comprensible, pero peligroso. Cuanto mayor es la presión para acelerar la producción, más fácil resulta degradar la revisión humana a una mera firma o a un control superficial. La tentación de “dar por bueno” un resultado verosímil sin verificarlo es, en ese contexto, mayor. Por eso, la respuesta adecuada no consiste en negar la utilidad de la herramienta, sino en reforzar las condiciones de uso que impidan que la búsqueda de eficiencia vacíe de contenido el deber de control.

En definitiva, la producción argentina e iberoamericana disponible, aun siendo todavía parcial, confirma la intuición central de este artículo. La IA puede aportar valor al trabajo jurídico, pero su incorporación solo es defendible si se conserva una estructura visible de juicio humano, revisión sustantiva, protección de la confidencialidad y trazabilidad del proceso. Leído desde nuestra región, el problema no es únicamente técnico ni exclusivamente regulatorio. Es, ante todo, profesional e institucional.

10. Función jurisdiccional, nulidad automática y el criterio fijado en el caso *Payalef*

El problema se vuelve especialmente sensible cuando se proyecta al ámbito jurisdiccional. Allí, la tentación de resolver la cuestión mediante prohibiciones generales o, en sentido contrario, mediante permisividad acrítica es particularmente fuerte. Ninguno de esos extremos parece satisfactorio. Prohibir toda asistencia tecnológica desconoce la realidad y puede impedir mejoras razonables en tareas auxiliares; admitir cualquier uso sin mayores exigencias pone en riesgo garantías, motivación y atribuibilidad de la decisión.

El caso *Provincia del Chubut c/ Payalef, Raúl Amelio s/ recurso de queja* (2026) ofrece un punto de apoyo especialmente útil para pensar este problema. Según el texto del fallo y el Acuerdo Plenario 5435/2025 del Superior Tribunal de Justicia del Chubut, se reafirmó que la IA es una herramienta de apoyo, nunca una fuente autónoma de decisión, que el control humano debe ser real y previo a la firma del pronunciamiento y que la responsabilidad por el contenido del fallo recae íntegramente en el magistrado que lo suscribe.

La relevancia del precedente radica en que evita dos errores simétricos. Por un lado, rechaza la nulidad automática basada en la mera constatación o sospecha de uso de IA. El tribunal no sostuvo que cualquier referencia residual o utilización de apoyo tecnológico convierta, por sí misma, al acto jurisdiccional en inválido. Por otro lado, tampoco trivializa el problema. La decisión insiste en la necesidad de control humano efectivo, anonimización de datos y revisión crítica del resultado, dentro del marco de las directivas éticas aprobadas por el propio Superior Tribunal a través del Acuerdo Plenario 5435/2025.

Desde el punto de vista dogmático, la lección del caso es doble. En primer lugar, la invalidez de una sentencia no puede deducirse simplemente del uso de una herramienta, sino de la acreditación de una afectación sustancial a garantías o de un déficit real de fundamentación y atribuibilidad. En segundo lugar, precisamente porque la nulidad no es automática, el control relevante debe desplazarse a las condiciones de uso: si hubo juicio humano real, si la motivación exteriorizada es verificable y si el acto sigue siendo plenamente atribuible al magistrado.

Este razonamiento resulta compatible con la lógica general del trabajo. La pregunta decisiva no es si existió o no asistencia de IA, sino si el proceso siguió siendo trazable, validado y defendible. Lo que el caso *Payalef* impide es una lectura fetichista de la herramienta, como si la sola presencia de IA agotara la cuestión jurídica. Al mismo tiempo, su exigencia de control humano real impide caer en una banalización según la cual todo uso de IA sería neutro mientras el juez o el abogado firme el resultado.

La importancia del precedente excede, además, el ámbito judicial estricto. El razonamiento puede trasladarse, con los ajustes necesarios, al trabajo profesional no jurisdiccional. También allí el control no debe recaer obsesivamente en la mera presencia de la herramienta, sino en la persistencia del juicio humano, la verificabilidad de las razones exteriorizadas y la posibilidad de atribuirle el producto a quien lo adopta y lo presenta como propio.

La resolución chubutense dialoga, de manera interesante, con los marcos internacionales antes examinados. La CEPEJ (2018) insiste en que la IA en justicia debe implementarse de modo compatible con los derechos fundamentales y con control por parte del usuario. El Reglamento europeo de IA ubica la supervisión humana como una exigencia orientada a prevenir o minimizar riesgos (Unión Europea, 2024). El NIST (2024), por su parte, advierte que la IA generativa puede exigir revisión humana adicional, seguimiento y documentación. *Payalef* no replica mecánicamente esos instrumentos, pero sí se mueve dentro de esa misma sensibilidad: lo decisivo no es la herramienta aislada, sino el mantenimiento de estructuras de control y responsabilidad.

Las Directrices de UNESCO (2025) para cortes y tribunales refuerzan esa misma lectura al presentar la IA como herramienta asistiva y no sustitutiva, sometida a supervisión humana significativa y a principios de transparencia, rendición de cuentas y protección de derechos. La convergencia entre ese marco internacional y el criterio de *Payalef* permite sostener una posición equilibrada: la IA puede ser empleada en tareas de apoyo, pero no puede vaciar de contenido la autoría, la motivación ni la responsabilidad humana de la decisión jurídica.

Por todo ello, el caso ofrece un insumo particularmente valioso para la construcción de un estándar argentino serio. No valida la desresponsabilización tecnológica. Tampoco consagra una alergia abstracta a la innovación. Propone, en cambio, un criterio jurídicamente más maduro: la función jurisdiccional es indelegable, la herramienta puede asistir, la decisión debe seguir siendo atribuible al humano responsable y la nulidad requiere más que la mera presencia de IA.

11. Hacia un estándar mínimo de control profesional

A partir de lo expuesto, es posible proponer un estándar mínimo de control profesional para el uso de inteligencia artificial en el trabajo jurídico. Se trata de un estándar deliberadamente sobrio. No pretende resolver todos los pro-

blemas de implementación, ni reemplazar regulaciones sectoriales futuras, ni convertir al profesional del derecho en especialista en seguridad ofensiva o arquitectura de modelos. Su función es otra: fijar un umbral mínimo por debajo del cual el uso de IA deja de ser jurídicamente defendible.

La objeción de sobreerregulación debe ser tomada en serio. Un estándar excesivamente pesado podría desalentar usos inocuos o convertir la innovación profesional en una práctica burocráticamente inviable. Por eso, el estándar que aquí se propone es gradual y proporcional al riesgo: cuanto más cercana esté la herramienta al núcleo del juicio jurídico, mayor debe ser la trazabilidad, la validación y la documentación; cuanto más auxiliar, reversible y de bajo impacto sea la tarea, menor será la intensidad exigible. Esta estructura de proporcionalidad permite evitar tanto la permisividad acrítica como el formalismo paralizante.

El primer componente del estándar es la delimitación funcional previa. No todas las tareas presentan el mismo riesgo. No equivale utilizar IA para ordenar materiales, sugerir una estructura preliminar o resumir documentos que hacerlo para estructurar la argumentación central, realizar valoraciones decisivas o intervenir en núcleos duros del juicio jurídico. La intensidad del control exigible debe variar según el impacto jurídico potencial de la tarea y la proximidad del sistema al momento de adopción final de la decisión o del producto.

El segundo componente es la trazabilidad mínima del proceso. Debe poder reconstruirse qué sistema se utilizó, para qué finalidad, sobre qué materiales, con qué fuentes, con qué permisos y con qué controles humanos. No se trata de producir una burocracia paralizante ni de registrar cada interacción irrelevante. Se trata de conservar el rastro suficiente para poder explicar, posteriormente, cómo intervino la herramienta y bajo qué condiciones. Sin ese umbral, la discusión sobre responsabilidad y revisión se vuelve difusa.

El tercer componente es la validación humana sustantiva. No basta una lectura superficial del resultado ni un chequeo puramente formal. La profundidad de la validación debe ser proporcional al impacto de la tarea y a la relevancia jurídica del producto final. Cuando la IA intervenga en tareas cercanas al núcleo del juicio jurídico, la revisión debe incluir control de fuentes, contraste de hechos, detección de sesgos o vacíos argumentativos, evaluación de pertinencia y una decisión consciente sobre la adopción o rechazo del resultado.

El cuarto componente es la gobernanza de insumos y fuentes. En arquitecturas agentivas o con recuperación, toda fuente externa debe ser tratada como potencialmente riesgosa. Deben existir límites de acceso, segmentación de permisos, distinción entre datos confiables y no confiables, y revisión de recorridos

o registros de trabajo cuando el sistema interactúe con herramientas o recursos externos. Aquí la lección de la seguridad técnica es clara, pero su traducción jurídica también lo es: si no se gobiernan las fuentes, no se gobierna el proceso.

El quinto componente es la protección reforzada de la confidencialidad y de la integridad documental. Los datos del cliente, los escritos en preparación, la estrategia procesal y los materiales de trabajo no pueden ser tratados como insumos indiferenciados. La selección de herramientas, la configuración de sus condiciones de uso y el alcance de sus accesos forman parte del deber de diligencia profesional. Esto exige, como mínimo, una evaluación *ex ante* de qué datos ingresan al sistema, bajo qué condiciones y con qué posibilidad de reutilización o exposición por parte de terceros.

El sexto componente es la documentación y conservación razonable. Debe preservarse un umbral mínimo de documentación suficiente para permitir explicación posterior en caso de controversia, auditoría o cuestionamiento. No se trata necesariamente de conservar cada *prompt*, pero sí de poder reconstruir la lógica de intervención de la herramienta, el contexto de uso y los controles realizados. En productos especialmente sensibles, esa documentación debería ser más robusta.

Un componente complementario, sugerido por la lectura de Seleme (2021), es el deber de informar adecuadamente cuando la utilización de IA incida de manera relevante en la prestación profesional. Si la herramienta modifica costos, riesgos, tiempos, fuentes utilizadas, exposición de información o grado de intervención humana, el cliente o destinatario institucional del trabajo jurídico no deberían quedar situados frente a una mediación técnica invisible. Esta exigencia no supone transformar todo uso auxiliar en una carga comunicacional formal, pero sí reconocer que, ante usos materialmente relevantes, la información sobre beneficios, límites y riesgos forma parte de la buena práctica profesional.

El séptimo componente es una regla de no delegación respecto de los núcleos duros del juicio jurídico. La interpretación decisiva, la ponderación final, la fundamentación que pueda afectar derechos o garantías, la estrategia jurídica adoptada como decisión final y los actos cuyo contenido no pueda ser razonablemente justificado por un humano no deben quedar entregados a la herramienta. La IA puede asistir, ordenar, sugerir e incluso ampliar la capacidad de procesamiento; no debe sustituir la responsabilidad intelectual por la que el profesional o magistrado responde.

Finalmente, el estándar exige una cultura de revisión continua. La evalua-

ción del uso de IA no es un acto único. Cambian los modelos, cambian las integraciones, cambian los riesgos, cambian las prácticas organizacionales. De allí que la gobernanza del uso de IA no pueda agotarse en una política inicial o en una capacitación aislada. Debe incorporar monitoreo, revisión periódica y aprendizaje institucional, en línea con el enfoque de gestión de riesgos que hoy domina los marcos más serios sobre la materia (NIST, 2023, 2024).

Este estándar no pretende ser maximalista. Es, en realidad, un mínimo. Pero es un mínimo exigente. Su función es impedir que la incorporación de inteligencia artificial degrade la responsabilidad profesional bajo una apariencia de eficiencia. Allí radica, precisamente, su sentido jurídico.

12. Conclusiones

La inteligencia artificial ya forma parte del trabajo jurídico. La pregunta relevante no es si esa realidad debe negarse, sino bajo qué condiciones su uso puede seguir siendo profesionalmente aceptable. El error más frecuente ha sido concentrarse únicamente en la calidad aparente de la salida. Sin embargo, el problema jurídicamente decisivo no está solo allí. Está en el proceso: quién decide, con qué fuentes, bajo qué límites, con qué revisión y con qué posibilidad posterior de explicación.

Ese problema se intensifica cuando la herramienta ya no funciona solo como asistente textual, sino como agente conectado a fuentes, repositorios y herramientas. Allí donde el sistema puede recuperar información, navegar entornos documentales, usar conectores o encadenar acciones, el derecho no puede conformarse con un humano meramente simbólico “en el circuito”. Debe exigir un humano con control real sobre la tarea, la fuente, la interpretación y la adopción del resultado.

La propuesta de este trabajo es deliberadamente sobria. No pretende formular una teoría total de la inteligencia artificial aplicada al derecho ni clausurar todos los debates regulatorios abiertos. Propone algo más acotado, pero más útil: un estándar mínimo de trazabilidad, validación humana sustantiva y uso defendible del proceso; no como refinamiento académico, sino como condición mínima para que el uso de IA en el trabajo jurídico no erosione la responsabilidad profesional ni vacíe de contenido la exigencia de justificación que define al derecho. En esa línea, la incorporación de Seleme (2021) permite expresar con mayor precisión el núcleo de la tesis: la IA no elimina los deberes clásicos de rol, sino que genera nuevas formas de incumplirlos cuando el profe-

sional no controla, no informa, no supervisa o no protege la confidencialidad en contextos tecnológicamente mediados.

En esa línea, la propuesta se apoya en dos premisas doctrinales complementarias. La primera, sugerida por Ober y Tasioulas (2024), es que la IA debe seguir siendo tratada como herramienta inteligente al servicio de fines humanos, no como sustituto de la agencia moral o jurídica. La segunda, desarrollada por Robbins (2024), es que el control humano solo es significativo si se especifica quién controla, qué controla y con qué capacidad real de intervención. Trasladadas al trabajo jurídico, ambas premisas conducen a una misma conclusión: la asistencia tecnológica solo es aceptable si preserva un proceso atribuible, explicable y profesionalmente defendible.

Si la inteligencia artificial ha entrado para quedarse, entonces la discusión relevante ya no puede girar exclusivamente en torno a cuán rápido escribe, resume o clasifica, sino a si el resultado puede seguir siendo atribuido, explicado y defendido por un humano responsable. En el derecho, esa es la frontera decisiva entre asistencia tecnológica y delegación impropia.

Bibliografía

- American Bar Association. (29 de julio de 2024a). *Formal Opinion 512: Generative artificial intelligence tools*. https://www.americanbar.org/content/dam/aba/administrative/professional_responsibility/ethics-opinions/aba-formal-opinion-512.pdf
- American Bar Association. (29 de julio de 2024b). *ABA issues first ethics guidance on a lawyer's use of AI tools*. <https://www.americanbar.org/news/abanews/aba-news-archives/2024/07/aba-issues-first-ethics-guidance-ai-tools/>
- Colegio Público de la Abogacía de la Capital Federal. (2025). *Guía para el uso de inteligencia artificial para abogados*. https://www.cpacf.org.ar/uploads/files/com/11072515_Gu%C3%A-DaparaelusodeInteligenciaArtificialparaAbogados.pdf
- Council of Bars and Law Societies of Europe. (2025). *CCBE guide on the use of generative AI by lawyers*. <https://www.ccbe.eu/>
- De Gregorio, G. (2023). The normative power of artificial intelligence. *Indiana Journal of Global Legal Studies*, 30(2), 55-80. <https://ssrn.com/abstract=4436287>
- Enqvist, L. (2023). "Human oversight" in the EU Artificial Intelligence Act: What, when and by whom? *Law, Innovation and Technology*, 15(2), 508-535. <https://doi.org/10.1080/17579961.2023.2245683>
- European Commission for the Efficiency of Justice. (2018). *European Ethical Charter on the use of artificial intelligence in judicial systems and their environment*. Council of Europe. <https://rm.coe.int/ethical-charter-en-for-publication-4-december-2018/16808f699c>

- Faliero, J. C. (2025). La obsolescencia del oportunismo del *Prompting Legal/Legal Prompting* e IA Generativa/IAGen Jurídica. *Revista Jurídica Rubinzal-Culzoni*, RC D 329/2025.
- Granero, H. R. (2018). La inteligencia artificial aplicada al Derecho: el cumplimiento del sueño de Hammurabi. *Informática y Derecho: Revista Iberoamericana de Derecho Informático* (segunda época), (5), 119-133.
- Jara, M. L. (2021). Abogacía 4.0: Transformación digital del ejercicio profesional. *Revista de Teoría y Práctica Jurídica*, 1(1), 1-14. <https://calz.org.ar/>
- Lazcoz Moratinos, G. (2025). Human oversight (Article 14). En Huergo Lora, A. y M. Díaz González, G. (Eds.), *The EU regulation on artificial intelligence: A commentary* (pp. 243-266). Wolters Kluwer/CEDAM. <https://dialnet.unirioja.es/descarga/articulo/10126859.pdf>
- Matsumi, H. y Solove, D. J. (2025). The prediction society: AI and the problems of forecasting the future. *University of Illinois Law Review*, 2025(1), 1-62. <https://illinoislawreview.org/print/vol-2025-no-1/the-prediction-society/>
- National Institute of Standards and Technology. (2023). *Artificial intelligence risk management framework* (AI RMF 1.0). <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>
- National Institute of Standards and Technology. (2024). *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile* (NIST AI 600-1). <https://doi.org/10.6028/NIST.AI.600-1>
- Nieva Fenoll, J. (2025). Inteligencia Artificial Generativa (IAGEN) y defensa: ¿cómo lograr la “convicción” de una máquina? *Revista Ítalo-española de Derecho Procesal*, 1-17. <https://doi.org/10.37417/rivitsproc/3204>
- Ober, J. y Tasioulas, J. (2024). *The Lyceum Project: AI ethics with Aristotle*. Human-Centered Artificial Intelligence, Stanford University/Institute for Ethics in AI, University of Oxford/. <https://doi.org/10.2139/ssrn.4879572>
- OpenAI Developers. (s.f.-a). *Agent internet access*. <https://developers.openai.com/codex/cloud/internet-access/>
- OpenAI Developers. (s.f.-b). *Safety in building agents*. <https://developers.openai.com/api/docs/guides/agent-builder-safety/>
- OWASP GenAI Security Project. (2025). *LLM01:2025. Prompt injection*. <https://genai.owasp.org/llmrisk/llm01-prompt-injection/>
- Rayón Ballesteros, M. C. (2025). Generative AI for lawyers in Spain: A balanced approach to the legal framework, technical foundations and best practices, combining technological innovation with professional responsibility. *Law and Business*, 5, 1-9. <https://doi.org/10.2478/law-2025-0003>
- Robbins, S. (2024). The many meanings of meaningful human control. *AI and Ethics*, 4, 1377-1388. <https://doi.org/10.1007/s43681-023-00320-6>
- Santos, A. M. (2025). Inteligencia artificial y dictámenes jurídicos en el derecho administrativo. Uso ético. *Aequitas Virtual*, 19(42), 4-34. <https://p3.usal.edu.ar/index.php/aequitasvirtual/article/view/7485>
- Seleme, H. O. (2021). Ética legal, inteligencia artificial y rol profesional. En Azuaje Pirela, M. y Contreras Vásquez, P. (Coords.), *Inteligencia artificial y derecho: desafíos y perspectivas* (pp. 481-500). Tirant lo Blanch.
- Terzidou, K. (2025). Generative AI systems in legal practice offering quality legal services while

- upholding legal ethics. *International Journal of Law in Context*, 21, 431-452. <https://doi.org/10.1017/S1744552325000047>
- UNESCO. (2021). *Recommendation on the ethics of artificial intelligence*. <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics>
- UNESCO. (2025). *Guidelines for the use of AI systems in courts and tribunals*. <https://doi.org/10.58338/LIEY8089>
- Unión Europea. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence. *Official Journal of the European Union*. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- Vassilev, A., Oprea, A., Fordyce, A., Anderson, H., Davies, X., y Hamin, M. (2025). *Adversarial machine learning: A taxonomy and terminology of attacks and mitigations* (NIST AI 100-2e2025). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.100-2e2025>

Legislación citada

- Superior Tribunal de Justicia del Chubut. Acuerdo Plenario 5435/2025, “Directivas para el uso ético y responsable de inteligencia artificial generativa en el Poder Judicial de Chubut”.
- Unión Europea. 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.

Jurisprudencia citada

- Acevedo Gerardo Gabriel c/ Cáceres Mareco Willian Arsenio y Agrosalta Cooperativa de Seguros Limitada s/ daños y perjuicios autom. c/ les. o muerte (exc. Estado), Cámara de Apelación en lo Civil y Comercial, Sala I, Morón, 15/9/2025.
- Giacomino, César Adrián y otros c/ Monserrat, Facundo Damían y otros s/ daños y perjuicios, CUIJ 21-11893083-2, Cámara de Apelación en lo Civil y Comercial de Rosario, Sala II, agosto de 2025.
- Mata v. Avianca, Inc., 678 F. Supp. 3d 443 (S.D.N.Y. 2023).
- Provincia del Chubut c/ Payalef, Raúl Amelio s/ recurso de queja, Superior Tribunal de Justicia del Chubut, Expte. 101215, Carpeta Judicial 6209, Oficina Judicial Esquel (sentencia de marzo de 2026).

Declaración sobre uso de inteligencia artificial

En la elaboración de este manuscrito se utilizó ChatGPT, herramienta de inteligencia artificial generativa desarrollada por OpenAI, exclusivamente como apoyo auxiliar para revisión de redacción, organización preliminar de ideas y contraste exploratorio de estructura expositiva. Todas las afirmaciones sustantivas, la selección y verificación de fuentes, la interpretación jurídica, las conclusiones y la versión final del texto fueron revisadas, corregidas y asumidas íntegramente por el autor humano, quien conserva plena responsabilidad por el contenido.

Roles de autoría y conflicto de interés

El autor manifiesta que cumplió todos los roles de autoría del presente artículo y declara no poseer conflicto de interés alguno.