

Riesgos, seguridad y gobernanza de la IA: revisión sistemática del marco legislativo argentino

Alejandro Covello

TecnocenoLab, Universidad de Buenos Aires (Argentina)

<https://orcid.org/0009-0008-2410-939X>

Flavia Costa

TecnocenoLab, Universidad de Buenos Aires (Argentina)

<https://orcid.org/0000-0001-8519-5860>

Iago Novidelsky

TecnocenoLab, Universidad de Buenos Aires (Argentina)

<https://orcid.org/0009-0009-6154-809X>

Julián Mónaco

TecnocenoLab, Universidad de Buenos Aires (Argentina)

<https://orcid.org/0000-0002-1918-2591>

Recibido: 20 de noviembre de 2025 / Aceptado: 20 de diciembre de 2025

DOI: <https://doi.org/10.62174/rs.11084>

Resumen¹

Las tecnologías de inteligencia artificial (IA) presentan riesgos capaces de afectar a personas, infraestructuras críticas, derechos fundamentales e incluso al funcionamiento democrático. En Argentina, múltiples proyectos de ley buscan establecer marcos regulatorios para su desarrollo, implementación y uso, pero lo hacen desde enfoques heterogéneos. En este trabajo, el equipo de investigación de TecnocenoLab analiza esos proyectos desde la perspectiva de la seguridad (*safety*) y establece un contrapunto con el enfoque de la protección (*security*).

Para ello, se examinaron los proyectos de ley presentados en ambas cámaras del Congreso Nacional argentino hasta diciembre de 2025. Se identificaron en total cincuenta y tres proyectos relacionados con IA, de los cuales once proponen una ley integral. Sobre estos últimos se realizó una revisión sistemática. El análisis se estructuró a partir de ocho elementos clave: definición de IA, establecimiento de principios de la IA, definición de ciclo de vida, definición inequívoca de seguridad y protección, identificación de categorías de riesgo, modelo de gobernanza de la IA, tratamiento de incidentes y/o accidentes de IA y perspectiva proactiva o reactiva con respecto a la seguridad y la protección. Esta matriz analítica desarrollada por el TecnocenoLab es original y consideramos que es el aporte más sustantivo de este artículo a futuras investigaciones y discusiones acerca de una política integral para la IA.

¹ Agradecemos especialmente a Sebastián Mateo por sus valiosos aportes en la lectura y edición de este artículo.



Los datos muestran la necesidad de adoptar un enfoque estatal centrado en la seguridad sistémica de la IA. Una regulación eficaz debe incorporar definiciones operativas, un ciclo de vida preciso, taxonomías de riesgo y sistemas de gobernanza robustos, incluyendo autoridades independientes de investigación, para mitigar peligros y fortalecer la protección de la sociedad y del orden democrático.

Palabras clave: inteligencia artificial; gestión de riesgos de la IA; ética de la IA seguridad sistémica; análisis normativo; leyes argentinas.

Abstract

Artificial intelligence (AI) technologies present risks that can affect people, critical infrastructure, fundamental rights, and even democratic functioning. In Argentina, multiple bills seek to establish regulatory frameworks for their development, implementation, and use, but they do so from heterogeneous approaches. In this work, the TecnocenoLab research team analyzes these bills from a safety perspective and contrasts them with a security approach.

To this end, the bills presented in both chambers of the Argentine National Congress up to December 2025 were examined. A total of fifty-three AI-related bills were identified, eleven of which propose a comprehensive law. A systematic review was conducted on these eleven bills. The analysis was structured around eight key elements: definition of AI, establishment of AI principles, definition of the AI lifecycle, definition of safety and security, identification of risk categories, AI governance model, handling of AI incidents and/or accidents, and proactive or reactive perspective regarding safety and security. This analytical matrix developed by TecnocenoLab is original, and we believe it is the most substantial contribution of this article to future research and discussions on a comprehensive AI policy.

The data demonstrates the need to adopt a state-led approach focused on the systemic security of AI. Effective regulation must incorporate operational definitions, a precise lifecycle, risk taxonomies, and robust governance systems, including independent investigative authorities, to mitigate risks and strengthen the protection of society and the democratic order.

Keywords: artificial intelligence; AI risk management; AI ethics; systemic security; normative analysis; Argentine laws.

Resumo

As tecnologias de inteligência artificial (IA) apresentam riscos que podem afetar pessoas, infraestruturas críticas, direitos fundamentais e até mesmo o funcionamento democrático. Na Argentina, diversos projetos de lei buscam estabelecer marcos regulatórios para seu desenvolvimento, implementação e uso, mas o fazem a partir de abordagens heterogêneas. Neste artigo, a equipe de pesquisa do TecnocenoLab analisa esses projetos de lei sob uma perspectiva de segurança (*safety*) e os contrasta com uma abordagem de proteção (*security*).



Para tanto, foram examinados os projetos de lei apresentados em ambas as casas do Congresso Nacional Argentino até dezembro de 2025. Um total de cinquenta e três projetos de lei relacionados à IA foram identificados, onze dos quais propõem uma lei abrangente. Uma revisão sistemática foi conduzida sobre esses onze projetos de lei. A análise foi estruturada em torno de oito elementos-chave: definição de IA, estabelecimento de princípios da IA, definição do ciclo de vida da IA, definição inequívoca de segurança e proteção, identificação de categorias de risco, modelo de governança da IA, tratamento de incidentes e/ou acidentes com IA e perspectiva proativa ou reativa em relação à segurança e proteção. Esta matriz analítica desenvolvida pelo TecnoCenoLab é original e acreditamos que seja a contribuição mais substancial deste artigo para futuras pesquisas e discussões sobre uma política abrangente de IA.

Os dados demonstram a necessidade de uma abordagem liderada pelo Estado, focada na segurança sistêmica da IA. Uma regulamentação eficaz deve incorporar definições operacionais, um ciclo de vida preciso, taxonomias de risco e sistemas de governança robustos, incluindo autoridades investigativas independentes, para mitigar os perigos e fortalecer a proteção da sociedade e da ordem democrática.

Palavras-chave: inteligência artificial; gestão de riscos da IA; ética da IA; segurança sistêmica; análise normativa; leis argentinas.

Introducción

El desarrollo y el uso de los sistemas de Inteligencia Artificial (IA) puede generar daños a las personas, afectar infraestructuras críticas, comprometer derechos fundamentales o producir impactos severos en el ambiente y en las instituciones democráticas. Partiendo de esta premisa, queda clara la importancia de un marco regulatorio que permita prevenir accidentes y mitigar los daños. Es por ello que el equipo de TecnoCenoLab, Observatorio de Tecnologías, Sociedad y Ecosistemas Digitales de la Facultad de Ciencias Sociales de la Universidad de Buenos Aires (UBA), se propuso examinar cómo cada proyecto de ley integral presentado en alguna de las dos cámaras del Congreso Nacional argentino define la IA, qué enfoque adopta frente a los riesgos y qué mecanismos propone para garantizar que estos sistemas operen dentro de niveles de riesgo aceptables para el Estado argentino. Se puso el foco en los siguientes factores de riesgo, que dialogan estrechamente con los peligros identificados en documentos internacionales de referencia, como la Ley de IA de la Unión Europea (UE) y la definición de incidentes y peligros de IA de la Organización para la Cooperación y el Desarrollo Económicos (OCDE):²

a) los daños a las personas;

² Es posible acceder a la Ley de la Unión Europea 2024/1689 aquí: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689>. Es posible acceder a la definición de incidentes y peligros de la OCDE aquí: https://www.oecd.org/en/publications/defining-ai-incidents-and-related-terms_d1a8d965-en.html

- b) la alteración grave del funcionamiento de infraestructuras críticas;³
- c) el incumplimiento de obligaciones constitucionales destinadas a proteger derechos fundamentales, y
- d) los daños graves a la propiedad o al medio ambiente.

1. Método

Inicialmente, se identificaron los proyectos de ley a nivel nacional vinculados con la IA. En un relevamiento que abarcó hasta diciembre de 2025, se encontraron cincuenta y tres proyectos que mencionan en su título la IA;⁴ no obstante, se seleccionaron para el análisis los once que proponen una ley integral para la regulación de la IA.

En segunda instancia, se analizó la muestra a través de ocho elementos clave estructurales para una gobernanza ética y segura de la IA. Definimos desde el TecnocenoLab estos criterios a partir de la integración de principios de seguridad aplicados históricamente en industrias de alto riesgo, estándares internacionales de gobernanza de IA y fundamentos de las teorías de seguridad sistémica. Consideramos que esta matriz constituye un insumo valioso para futuras investigaciones, así como para el análisis de ulteriores proyectos legislativos o normativos para instituciones en las escalas nacional, provincial o municipal.

Antes de definir cada elemento clave y el porqué de su importancia, es necesario aclarar el sentido que se da en este trabajo al concepto “seguridad”, para diferenciarlo del término “protección”. En inglés, el término *security* remite a la acción preventiva contra los delitos, sabotajes e infracciones, mientras que *safety* se refiere a la prevención de accidentes, enfermedades laborales, riesgos industriales, etcétera. El español no tiene una distinción equivalente, solo cuenta con el término seguridad. Ante esta dificultad, se consensuó en las industrias de alto riesgo utilizar dos conceptos inequívocos: protección para *security* y seguridad para *safety*. Es por

³ La Resolución 1523/2019 de la entonces Secretaría de Gobierno de Modernización, en el ámbito de la Jefatura de Gabinete de Ministros del Estado argentino, identifica como *infraestructuras críticas* a los siguientes espacios: energía, tecnologías de información y comunicaciones, transportes, hídrico, salud, alimentación, finanzas, nuclear, químico, espacio y Estado. Luego, las define como “aquellas que resultan indispensables para el adecuado funcionamiento de los servicios esenciales de la sociedad, la salud, la seguridad, la defensa, el bienestar social, la economía y el funcionamiento efectivo del Estado, cuya destrucción o perturbación, total o parcial, los afecte y/o impacte significativamente. Las *infraestructuras críticas de información* son las tecnologías de información, operación y comunicación, así como la información asociada, que resultan vitales para el funcionamiento o la seguridad de las infraestructuras críticas”. En Internet: <https://www.boletinoficial.gob.ar/detalleAviso/primera/216860/20190918>.

⁴ Un listado completo de estos proyectos se encuentra publicado en el sitio de TecnocenoLab (<https://tecnocenolab.ar/territorios/normativas-recomendaciones-y-etica-de-la-ia/>). La fuente de estos proyectos es la página del Congreso Nacional. A partir de enero de 2026, algunos de los proyectos analizados aquí perderán su estado parlamentario.

ello que, en la traducción castellana de los principios de la OCDE y de la UNESCO encontramos respectivamente “Robustez, seguridad y protección” (“*Robustness, safety and security*”)⁵ y “Seguridad y protección” (“*Safety and security*”)⁶. Adoptamos aquí esta distinción, si bien especificamos el término “seguridad” como “seguridad sistémica”, ya que el desafío de la gobernanza de la IA implica procurar que los riesgos asociados a un sistema de IA se encuentren en un nivel aceptable, y para esto es necesario abordar dichos riesgos desde el pensamiento sistémico

Los proyectos de ley analizados suelen identificar al derecho como el principal límite frente a eventos no deseados. Esto lo hacen mediante el establecimiento de leyes que regulan las prácticas. Si se infringe la legislación, se identifica al responsable y se impone un resarcimiento y/o se aplica una sanción. Este es un registro administrativo-judicial que actúa *a posteriori* del incidente; supone que, o bien hubo intención de daño (dolo), o bien negligencia (culpa); y pone el acento en la responsabilidad y en las amenazas externas al sistema. Sin embargo, no es el único registro o enfoque normativo para tratar la seguridad en relación con un sistema técnico. En el caso de un peligro inherente al funcionamiento del sistema, sin intención de daño, se deben implementar estrategias de prevención de accidentes y gestionar los riesgos que pueden advenir en el transcurso normal de la vida del sistema: este es el campo de la seguridad, o como diremos aquí, la seguridad sistémica.

En algunos casos, las tecnologías de IA pueden implicar un alto riesgo,⁷ por eso la estrategia reactiva (administrativo-judicial) es limitada. Se necesitan estrategias proactivas y predictivas que actúen desde el diseño del sistema de IA.

Explicación de los elementos clave

En este capítulo definiremos los ocho elementos clave o características que entendemos que constituyen requisitos básicos para una ley integral de inteligencia artificial, con acento en la gobernanza ética, segura y democrática dentro del

⁵ El cuarto de los cinco Principios de la IA de la OCDE es “Robustez, seguridad y protección”. En internet: <https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449>. Volveremos a esto en el apartado Principios de la IA.

⁶ El segundo de los diez Principios que establece la Unesco como base para un enfoque de la ética de la IA centrado en los derechos humanos es “Seguridad y protección”. En Internet: <https://www.unesco.org/es/artificial-intelligence/recommendation-ethics>.

⁷ Entre la abultada, aunque reciente, bibliografía sobre riesgos de la IA, sugerimos ver Hendrycks, Dan; Mazeika, Mantas y Woodside, Thomas (2023) y Bengio, Yoshua et al. (2024). Ya venimos señalando la importancia de esta dimensión desde el informe “Desafíos e impactos de la inteligencia artificial. Marcos normativos, riesgos y retos para la calidad democrática en la Argentina”, realizado por el equipo del TecnocenoLab en 2023 y que se encuentra en internet: <https://tecnocenolab.ar/proyecto/>. También puede encontrarse en: Costa, Flavia, Mónaco, Julián, Covello, Alejandro, Novidelsky, Iago et al (2023).

ecosistema digital: definición de Inteligencia Artificial, establecimiento de principios éticos de la IA, definición de ciclo de vida de los sistemas de IA, definición inequívoca de seguridad y protección, identificación de categorías de riesgo, modelo de gobernanza de la IA, tratamiento de incidentes o accidentes de IA y perspectiva proactiva o reactiva con respecto a la seguridad y la protección. Esta matriz analítica desarrollada por el TecnocenoLab es original y consideramos que es el aporte más sustantivo de este artículo a futuras investigaciones y discusiones sobre una política integral para la IA.

En cuanto al análisis de los proyectos, en la sección Desarrollo expondremos la síntesis del estudio comparativo realizado. La versión extendida se encuentra disponible en el sitio del TecnocenoLab.⁸

Definición de inteligencia artificial y/o sistema de inteligencia artificial

Kate Crawford, investigadora principal del Microsoft Research (MRS), afirma en su *Atlas de la inteligencia artificial*: “Cada forma de definir la IA cumple con un cometido y establece un marco de referencia para entenderla, medirla, valorarla y gobernarla” (2022, p.28). En efecto, hasta el momento no se ha llegado a una definición consensuada de IA. Se la describe como un sistema basado en máquinas, un sistema informático, un robot, una herramienta, un *software*, una rama de la informática y la digitalización. Tomaremos aquí como referencia cuatro definiciones que, si bien tienen diferencias, coinciden en que un sistema de IA tiene la capacidad de realizar un conjunto de tareas que solemos asimilar al comportamiento inteligente. Las definiciones de la OCDE y las dos que brinda la UNESCO son especialmente importantes porque se convirtieron en estándares o puntos de partida. Se suma la definición del Center of AI Safety (CAIS), una institución que es precursora en la investigación en seguridad y riesgos de la IA, liderada por Dan Hendrycks, ex investigador en DeepMind y ScaleAI.

Para la OCDE, la IA es “un sistema basado en máquinas que, para objetivos explícitos o implícitos, infiere, a partir de la información que recibe, cómo generar resultados como predicciones, contenido, recomendaciones o decisiones que pueden influir en entornos físicos o virtuales. Los diferentes sistemas de IA varían en sus niveles de autonomía y adaptabilidad tras su implementación” (Organización para la Cooperación y el Desarrollo Económicos [OCDE], 2024b, p.6, la traducción es nuestra).

Por su parte, la UNESCO, en su *Recomendación sobre Ética de la IA*, considera los sistemas de IA como “sistemas capaces de procesar datos e información de una

⁸ La versión ampliada de este trabajo se encuentra junto al listado de proyectos en el sitio mencionado en la nota 4.

manera que se asemeja a un comportamiento inteligente y que abarcan generalmente aspectos de razonamiento, aprendizaje, percepción, predicción, planificación o control” (Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura [UNESCO] 2021, p.10).

Esta misma organización, en su página de consulta de términos la define así: “La inteligencia artificial (IA) es un sistema basado en máquinas y la rama de la informática centrada en la creación de sistemas capaces de realizar algunas tareas que normalmente requieren inteligencia humana, como comprender el lenguaje, reconocer patrones y tomar decisiones”.⁹

Finalmente, para Dan Hendrycks, director del CAIS, las IA son “sistemas informáticos que realizan tareas típicamente asociadas con la inteligencia, como la resolución de problemas, la toma de decisiones y la previsión” (Hendrycks, 2023a: 4).

Por nuestra parte, en trabajos anteriores hemos definido las cinco características centrales de las IA generativas que es preciso tener en cuenta cuando se trata de diseñar estrategias de gobernanza:¹⁰

1. Las IA son *metatecnologías*. Operan y regulan otras tecnologías y sistemas técnicos. Son tecnologías de propósito general que automatizan procesos y facilitan la producción de conocimiento en muy diferentes campos.
2. Los ecosistemas digitales en los que se inscriben las IA generativas no constituyen solo una herramienta, sino que son también un *mundoambiente*, una infraestructura y a la vez un entorno de sentido indispensable para las actividades cotidianas de los usuarios.
3. Las IA no son solo “Inteligencia artificial” sino también “Sociedad artificial”. Operan con y sobre el mundo social. En particular los modelos de lenguaje grandes (*large language models*, LLM) aceleran y automatizan la (re)producción de lo social, (re)producen sociedad.
4. Las IA pueden ser tecnologías de alto riesgo. En ciertos usos, pueden producir riesgos no tolerables o de alto impacto. Requieren un tratamiento acorde desde una gestión de riesgos a lo largo de todo su ciclo de vida.
5. Para tratar esos riesgos, la ética de la IA debe complementarse con la ética organizacional y el pensamiento sistémico de la IA. El impacto social debe considerarse desde el momento mismo del diseño.

El primer elemento clave de análisis es, entonces, qué definición de IA y/o de sistema de IA enuncia cada proyecto de ley.

⁹ En Internet (y sólo en inglés): <https://www.unesco.org/en/query-list/a/ai>.

¹⁰ Ver Costa, Flavia Mónaco, Julián, Covello, Alejandro, Novidelsky, Iago et al (2023), op.cit. También Costa, Flavia y Mónaco, Julián (2024).

Principios de IA

Los principios definen un conjunto de valores que debe orientar a los sistemas de IA y a los distintos actores involucrados --desde diseñadores e implementadores hasta usuarios y diferentes tipos de autoridades--, a lo largo de todo el ciclo de vida. Es un marco de expectativas implícitas, actitudes y prácticas que pueden facilitar y promover decisiones y acciones beneficiosas en el uso de estos sistemas. Por la naturaleza de la IA, es necesario enunciar los principios rectores así como anticipar sus potenciales efectos disruptivos desde el momento mismo del diseño. Estas dos acciones constituyen el primer paso para un consenso orientado a la gobernanza de escala nacional, y a la creación de una gobernanza internacional.

En este sentido, la *Recomendación sobre ética de la IA* de la Unesco aborda la cuestión de los principios desde la “ética de la IA”, a la que define como

una reflexión normativa sistemática, basada en un marco integral, global, multicultural y evolutivo de valores, principios y acciones interdependientes, que puede guiar a las sociedades a la hora de afrontar de manera responsable los efectos conocidos y desconocidos de las tecnologías de la IA en los seres humanos, las sociedades y el medio ambiente y los ecosistemas. [...] Considera la ética como una base dinámica para la evaluación y la orientación normativas de las tecnologías de la IA, tomando como referencia la dignidad humana, el bienestar y la prevención de daños y apoyándose en la ética de la ciencia y la tecnología (Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura [UNESCO], 2021, p.10).

Por su parte, la OCDE ha desarrollado un estándar intergubernamental de la IA, originalmente redactado en 2019 y revisado en 2024, al que la Argentina adhirió en 2019, y en el que se definen cinco principios: Crecimiento inclusivo; Desarrollo sostenible y bienestar; Derechos humanos y valores democráticos, incluidas la equidad y la privacidad; Transparencia y explicabilidad; Robustez, seguridad y protección; y Responsabilidad (*Accountability*). A estos principios, se suman cinco *Recomendaciones para los responsables de las políticas*: Invertir en investigación y desarrollo de IA; Fomentar un ecosistema inclusivo que facilite la IA; Configurar un entorno de gobernanza y políticas interoperable y propicio para la IA; Desarrollar capacidades humanas y preparar a las personas para la transición del mercado laboral; Cooperar internacionalmente para una IA fiable (Organización para la Cooperación y el Desarrollo Económicos, 2024a, pp.4-10).

La existencia o ausencia de principios de la IA en cada proyecto es un indicio tanto de la reflexividad del legislativo respecto de los valores sociales que propone

custodiar cuanto de la disposición del proyecto a evaluar los sistemas de IA de alto riesgo.

Definición de ciclo de vida

Uno de los principios vitales para todo sistema de alto riesgo es la “seguridad inherente”, un enfoque que busca reducir o eliminar peligros desde el momento de la concepción del sistema hasta su salida del mercado. Esto evita tener que agregar otros sistemas o capas de protección en la etapa de la implementación, algo que puede introducir peligros residuales.

La gobernanza de la IA debe comenzar desde el inicio del diseño y estar presente en todo el ciclo de vida. A cada etapa le corresponden procesos distintos, de ahí que es importante delimitar estas etapas con precisión (identificar su inicio y su fin), como también los requisitos de supervisión y monitoreo.

Analizamos y comparamos los proyectos de ley según el modelo de ciclo de vida que elaboramos en TecnocenoLab (Figura 1), sobre la base de las etapas consignadas en las *Recomendaciones para una IA Fiable*, emitidas por la Subsecretaría de Tecnologías de la Información en junio de 2023 (Jefatura de Gabinete de Ministros, 2023).

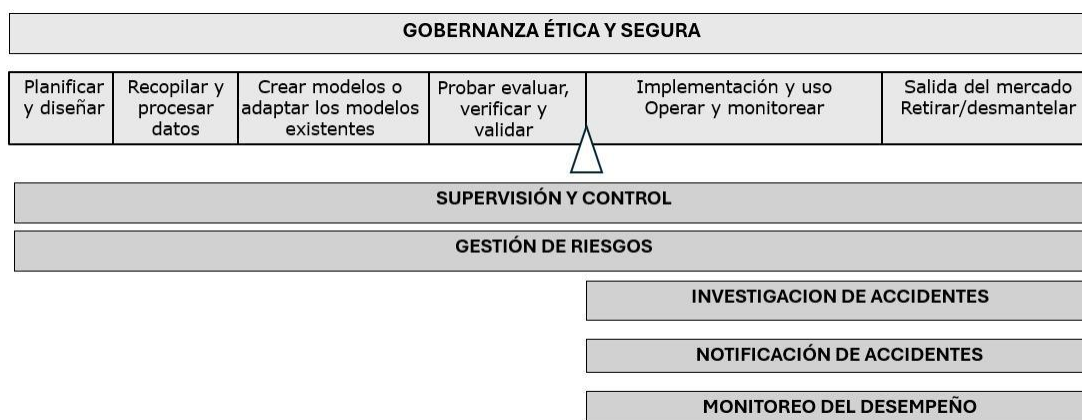


Figura 1. Ciclo de vida y gobernanza ética y segura. Fuente: Alejandro Covello para el TecnocenoLab (2025)

Distinción entre seguridad y protección

Al diseñar un sistema de IA, existen dos desafíos fundamentales y simultáneos: cumplir los objetivos de producción, por un lado, y proteger al sistema y a los usuarios de los peligros internos y externos durante todo el ciclo de vida, por otro lado.

Existen dos fuentes de peligros que amenazan al sistema de IA y de los cuales este debe protegerse: el uso malicioso, que implica la intención de causar un daño

(ciberataque, terrorismo, sabotaje, delitos, fraudes, etc.), y el accidente (un suceso adverso sin intención de daño provocado por una combinación de factores). En el primer caso, las áreas de competencia de la protección son la seguridad pública y la nacional, en el sentido del término inglés *security*. Es el ámbito donde actúan fuerzas de seguridad y defensa en conjunto con el poder judicial, con el fin de prevenir o castigar delitos, identificar responsabilidades, dictar penas, sanciones, resarcimientos, etc. El objetivo incluye compensar a las víctimas por los daños sufridos. En el segundo caso, el campo de aplicación refiere a la prevención de accidentes y a la reducción y control de los riesgos asociados a un sistema sociotécnico, hasta llevarlos a un nivel aceptable. Este análisis hace énfasis en un enfoque más amplio que el de la protección, que es el de la seguridad (*safety*), cuyas prioridades contemplan custodiar la robustez del sistema, proteger la salud de trabajadores y usuarios, limitar los daños al medioambiente. Busca asimismo garantizar modelos fiables que no avasallen valores y derechos fundamentales: evitar sesgos discriminatorios y prejuicios, minimizar la desigualdad, respetar la soberanía, promover la inclusión y preservar la diversidad.

Ante la perspectiva de que en el futuro las sociedades estarán cada vez más atravesadas por las tecnologías de IA, la seguridad debe ser un área fundamental de investigación y desarrollo. De allí la necesidad de definirla y diferenciarla de la protección en los proyectos de ley.

Cuando se considera que algunos usos de IA pueden ser de alto riesgo, se hace referencia a un tipo de riesgo especial que no todas las tecnologías presentan: el sistémico o catastrófico. Las llamadas industrias de alto riesgo, como la nuclear, la petroquímica o la aeronáutica, adoptaron sobre el final del siglo XX una nueva estrategia para abordarlo, llamada seguridad sistémica y, correlativamente, incorporaron la investigación sistémica de accidentes. La seguridad sistémica es un modelo de gestión de la seguridad basado en la teoría de sistemas cuyo objetivo es lograr un estado tal que los riesgos asociados a un sistema sociotécnico estén en un nivel aceptable durante todo su ciclo de vida. Por su parte, la investigación sistémica de accidentes no define la seguridad solo por los accidentes, incidentes o peligros que pueden surgir en la primera línea de trabajo (falla del trabajador, el llamado error humano) o en un fallo técnico, sino que pone el acento en los factores humanos y organizacionales (situación de trabajo, contexto en el que el operador toma una decisión, modos de gerenciamiento, entes normativos) y las defensas en profundidad (sistemas automáticos que “atrapan” y minimizan el error, y que están ya incorporados en el conjunto técnico, como un disyuntor eléctrico en un hogar o un antivirus en una PC; lo que denominamos “infraética”). Desde este marco analítico, el error humano o el fallo técnico son factores desencadenantes, pero no causas del accidente.

Volviendo a la especificidad de la IA, la normativa europea define riesgo sistémico del siguiente modo:



Un riesgo específico de las capacidades de gran impacto de los modelos de IA de uso general, que tienen unas repercusiones considerables en el mercado de la Unión debido a su alcance o a los efectos negativos reales o razonablemente previsibles en la salud pública, la seguridad, la seguridad pública, los derechos fundamentales o la sociedad en su conjunto, que puede propagarse a gran escala a lo largo de toda la cadena de valor (Ley de inteligencia artificial de la Unión Europea, 2024. Artículo 3, inc 65, p. 50).

Por su parte, Dan Hendrycks y otros han identificado a la seguridad sistémica como una de las cuatro áreas principales de investigación para una IA ética y segura (Hendrycks et al., 2021: 1-11). Por su parte, en *An overview on catastrophic AI Risk*, Hendrycks, Mantas Mazeika y Thomas Woodside definen cuatro tipos de riesgos catastróficos: uso malicioso, carrera de la IA, riesgos organizacionales e IA rebelde (Hendrycks, Mazeika y Woodside, 2023: 5). De un modo similar pero más exhaustivo, el equipo liderado por Uuk Risto, director de política e investigación de la UE en el Instituto Future of Life de Bruselas, construyó una taxonomía de doce tipos de riesgos sistémicos, ordenada por áreas clave que pueden verse afectadas severamente por la IA, entre las cuales destacan: cambios sociales, culturales o antropológicos irreversibles, defensa, democracia, derechos fundamentales, descontrol de la propia IA, discriminación, economía, información y medioambiente, entre otros (Risto Uuk et al., 2024: 2 y 3) (Ver Figura 2).



Figura 2: Fuente: Risto Uuk et al (2024).

La existencia o ausencia de definición de seguridad y protección, así como la consideración de los riesgos sistémicos en los diferentes proyectos de ley, es un indicador del amparo que ofrece la norma a los ciudadanos y a sus derechos fundamentales.

Identificación de categorías de riesgo de la IA

Algunos usos de los sistemas de IA son de alto riesgo y, por ello, se requiere un análisis detallado para definir qué nivel de riesgo corresponde a cada uso. Por lo general, las normativas proactivas definen una pirámide de riesgo. La pirámide clasifica los riesgos según su severidad: los “prohibidos” son los no tolerables; “alto riesgo”, los tolerables, pero mediante requisitos y medidas de mitigación y control; y “aceptables”, aquellos que permiten que el sistema se utilice tal como está, sin requisitos. El primer paso para un enfoque proactivo de seguridad es identificar en los proyectos de ley una pirámide de riesgos y la definición de requisitos como certificaciones, auditorías, monitoreos, notificaciones obligatorias y funciones de las autoridades reguladoras.

Sistema de gobernanza

Un sistema de gobernanza es el conjunto de estructuras, procesos, normas y prácticas que rigen la toma de decisiones, la gestión de recursos y la rendición de cuentas en una sociedad, organización o ente gubernamental. Abarca la interacción entre diferentes actores públicos y privados para asegurar el cumplimiento de objetivos y el interés general. Su propósito es establecer cómo se actúa, se decide y se rinde cuentas de manera efectiva y ética, buscando un orden social aceptable y la consecución de fines comunes.

Aquí analizaremos el mapa de actores claves (MAC) que define cada proyecto de ley. El MAC permite visualizar a quiénes se dirige la ley de manera específica; permite identificar la existencia o ausencia de autoridades de aplicación creadas por la propia ley, o de autoridades preexistentes a las que la ley otorga incumbencia en relación con las IA. Permite visualizar el rango de dichas autoridades y su ubicación dentro del organigrama del Estado. Permite asimismo observar si la ley identifica diferentes tipos de autoridad (por ej. de supervisión, de vigilancia de mercado, de notificación e investigación de accidentes, etc.). Permite, finalmente, visibilizar las diferentes posiciones en el ciclo de vida de la IA y sus incumbencias en dicho ciclo; por ejemplo: ente fiscalizador, prestadores de servicios, diseñadores, comercializadores, implementadores, usuarios. En suma, identifica la cadena de responsabilidades y la estructura de gobernanza de los sistemas de IA a nivel del país previstos por la ley. La definición de un MAC facilita la obtención de datos y el análisis ante eventuales incidentes, como también la gestión de riesgos.

Tomando como base la ley de IA de la UE, a la que incorporamos una autoridad de investigación de accidentes e incidentes, diseñamos un MAC de referencia (Figura 3) y sobre esta base analizamos los diferentes proyectos.



Figura 3. Mapa de actores claves, sobre la base del modelo de la Unión Europea. Fuente: TecnocenoLab (2025)

Notificación e investigación de accidentes e incidentes de IA

En este punto analizamos si los proyectos de ley contemplan la definición de accidente o incidente de IA, así como la obligatoriedad de notificación y de investigación de accidentes o incidentes por parte de una autoridad independiente. La existencia de estas dos instancias es indicio de la voluntad de implementar una política integral de seguridad sistémica.

Entendemos que, en el caso de establecerse por ley, la investigación de accidentes o incidentes no debería superponerse con la investigación administrativo-judicial; conviene que estos dos ámbitos sean independientes ya que sus objetivos, así como sus metodologías, difieren. La investigación sistémica busca identificar los factores relacionados con el incidente o accidente, con el fin de prevenir futuros eventos no deseados y lograr una arquitectura de seguridad más robusta. Para ello elabora informes y emite recomendaciones de seguridad a los distintos actores involucrados en el evento no deseado (operadores, empresas, instituciones de gobierno, organismos de control, etc.) con el fin de modificar las condiciones de posibilidad que provocaron el accidente, evitar su repetición, y contribuir así a un ecosistema de IA más seguro.

Perspectiva proactiva o reactiva de la ley

Las políticas de seguridad para una gobernanza de la IA pueden agruparse en dos grandes estrategias: reactivas o proactivas. Se las llama reactivas cuando responden y accionan luego de que el evento no deseado ha ocurrido. Ejemplo: una investigación de accidentes, delitos, etc. Se las llama proactivas cuando se anticipan a los eventos no deseados. Identifican fuentes de riesgos potenciales, miden probabilidades, buscan prevenir incidentes y crean medidas que minimicen los

potenciales accidentes. En la estrategia proactiva incluimos aquí la metodología predictiva que mide el monitoreo del desempeño. Al observar este elemento se busca identificar qué perspectiva prevalece en cada proyecto de ley analizado.

Corpus analizado

En este trabajo se analizaron seis proyectos presentados en la Cámara de Senadores y cinco presentados en la Cámara de Diputados. Los primeros son: “Proyecto de ley que dispone controlar el desarrollo, la implementación y utilización de sistemas basados en inteligencia artificial dentro del territorio argentino”, del Senador Juan Carlos Romero (Proyecto Romero); “Proyecto de ley que crea el Instituto Nacional de Inteligencia Artificial”, de la Senadora Silvia García Larraburu (Proyecto García Larraburu); “Proyecto de ley sobre sistemas de inteligencia artificial”, de las Senadoras Silvia Sapag y María Pilatti Vergara (Proyecto Sapag Pilatti); “Proyecto de ley de regulación de la inteligencia artificial”, del Senador Martín Doñate (Proyecto Doñate); “Proyecto de ley que establece un marco jurídico para el desarrollo, la comercialización, distribución y utilización de los sistemas de inteligencia artificial”, del Senador Flavio Fama (Proyecto Fama); “Proyecto de ley que establece el marco legal para la investigación, el desarrollo, uso y regulación de la inteligencia artificial”, de la Senadora Beatriz Ávila (Proyecto Ávila).

Por su parte, los proyectos presentados en la Cámara de Diputados son “Marco Normativo y de Desarrollo de los Sistemas de IA”, del Diputado Daniel Gollán (Proyecto Gollán); “Ley Nacional de Inteligencia Artificial”, del Diputado Gabriel Chumpitaz (Proyecto Chumpitaz); “Marco regulatorio de la inteligencia artificial. Creación del Registro Nacional de Sistemas de Inteligencia Artificial”, del Diputado Diego Giuliano y otros (Proyecto Giuliano); “Régimen jurídico aplicable para el uso responsable de la inteligencia artificial (IA), en la República Argentina”, del Diputado Juan Brügge (Proyecto Brügge); y “Responsabilidad algorítmica y promoción de la robótica, algoritmos verdes e inteligencia artificial en la República Argentina”, del Diputado Maximiliano Ferraro (Proyecto Ferraro).¹¹

2. Desarrollo sintético del análisis

Una vez definidos los ocho elementos clave se procedió al análisis y comparación de los once proyectos de ley de IA. Se estableció un sistema sintético de clasificación (Sí/No/Parcial) que rápidamente permite identificar en una tabla si cada elemento clave está o no contenido en cada proyecto de ley. El criterio de clasificación fue el siguiente:

¹¹ En la sección de Referencias y Bibliografía se encuentran los enlaces a cada proyecto.

Elemento 1, definición de IA: Si tiene al menos una definición de IA o de Sistema de IA: Sí. En caso contrario: No.

Elemento 2, principios de la IA: Si contiene principios de la IA: Sí/No

Elemento 3, definición de ciclo de vida: Sí/No.

Elemento 4, definición de seguridad y protección: Sí/No.

Elemento 5, identificación de categorías de riesgo: Sí/No/Parcial (es parcial cuando no explicita una pirámide de riesgo, pero admite la necesidad de elaborar alguna estrategia de riesgo)

Elemento 6, gobernanza de la IA: En este caso dividimos el análisis en 4 subelementos: Autoridad competente, sus funciones y jerarquía; Mapa de actores clave (MAC) sobre la base de la figura 3; Sistema de evaluación y registro por parte del Estado; y Sistema de monitoreo del desempeño a cargo de la autoridad competente. Estos se clasificaron como: Sí/No/Parcial.

Elemento 7, notificación e Investigación de incidentes y accidentes de IA: Se subdivide en dos elementos: Sistema de notificación de incidentes o fallas; e Investigación de seguridad sistémica. Ambos se clasificaron como: Sí/No

Elemento 8, perspectiva proactiva o reactiva: Se califica con Proactiva plena si la ley posee estos cuatro subelementos, y Proactiva parcial si falta alguno:

- a) Categorías de riesgo
- b) Los principios robustez, seguridad y protección
- c) Sistema de notificación o de investigación de accidentes/ incidentes
- d) En cuanto a gobernanza: al menos dos de los siguientes tres subelementos: Autoridad competente, Sistema de evaluación y registro, y Sistema de monitoreo.

El resultado de la clasificación realizada se presenta gráficamente en las dos tablas que siguen.

Elementos Clave	Definiciones IA	Principios de la IA	Definición Ciclo de vida	Definición Seg y Prot.	Categorías de riesgo	Sistema de Gobernanza: a) Autoridad. b) MAC c) Evaluación y registro d) Monitoreo	a) Notificación b) Investigación accidentes/inc.	Perspectiva Proactiva
Senadores								
Romero	Si	Si	No	No	Si	a) Si b) Parcial c) Si d) Si	a) No b) No	Parcial Proactiva
García Larraburu	No	No	No	No	No	a) Parcial b) No c) Si d) Parcial	a) No b) No	Reactiva
Sapag	Si	Si	No	No	Si	a) Si b) Si c) Si d) Si	a) Si b) No	Proactiva
Doñate	Si	Si	Si	No	Si	a) Si b) Si c) Si d) Si	a) Si b) No	Proactiva
Fama	Si	No	No	No	Si	a) Si b) Parcial c) Si d) Si	a) Si b) Si Parcial	Proactiva
Ávila	Si	Si	No	No	Si	a) Parcial b) No c) Si d) Parcial	a) No b) No	Parcial Proactiva

Tabla 1. Síntesis del análisis de los proyectos de Ley de la Cámara de Senadores (Congreso de la Nación Argentina). Fuente: elaboración propia, TecnocenoLab, 2025.

Elementos Clave	Definiciones IA	Principios de la IA	Definición Ciclo de vida	Definición Seg y Prot.	Categorías de riesgo	Sistema de Gobernanza: a) Autoridad. b) MAC c) Evaluación y registro d) Monitoreo	a) Notificación b) Investigación accidentes/inc.	Perspectiva Proactiva
Diputados								
Gollán	Si	Si	Si	No	Si	a) Si b) Si c) Si d) Parcial	a) No b) No	Parcial Proactiva
Chumpitaz	Si	Si	No	No	No	a) Si b) No c) Parcial d) No	a) Si b) No	Reactiva
Giuliano	Si	No	No	No	Si	a) No b) No c) Si d) Si	a) No b) No	Reactiva
Brügge	Si	Si	No	No	Si	a) Si b) Si c) Parcial d) Si	a) Si b) No	Parcial Proactiva
Ferraro	No	Si	No	No	Parcial	a) No b) Si c) Parcial d) No	a) No b) No	Reactiva

Tabla 2. Síntesis del análisis de los proyectos de Ley de la Cámara de Diputados (Congreso de la Nación Argentina). Fuente: elaboración propia, TecnocenoLab, 2025.

Análisis

En cuanto a las definiciones de IA, dos de los proyectos analizados no contienen definición de IA ni de sistema de IA. El resto de los proyectos sí definen ambos términos o uno de ellos. En general se considera a la IA como: un sistema basado en máquinas –definición similar a la OCDE– (tres proyectos), un asistente técnico (dos proyectos), un *software* (dos proyectos), una tecnología (un proyecto) y técnicas, algoritmos y procesos informáticos (un proyecto). Todas estas definiciones coinciden en dos aspectos: el primero es que mencionan la capacidad de las IA de realizar acciones complejas, de comportamiento inteligente similar al humano o acciones cognitivas humanas, en coincidencia con las definiciones de la OCDE, la UNESCO y el CAIS. El segundo aspecto es que consideran que la IA o los sistemas de IA tienen capacidad de actuar con niveles de autonomía, tal como afirma la OCDE.

En relación con los principios de la IA, ocho proyectos de ley mencionan estos principios. Un principio que consideramos clave frente a los riesgos de la IA, y por lo tanto para garantizar una IA ética y segura, es el de “Robustez, seguridad y protección”, debido a que contiene las salvaguardas para las dos principales amenazas a los sistemas de IA: externas (eventos no deseados con intención de daño) e internas (accidentes sin intención de daño). A los principios de “seguridad” y “protección” se les suma el principio de “robustez”, que designa la capacidad para mantener un desempeño preciso, estable y predecible bajo una amplia variedad de condiciones, incluyendo situaciones imprevistas o adversas. Se encontraron tres proyectos que contienen este principio; uno de ellos añade el principio de “Prevención de riesgos y posibles daños”. En otros proyectos aparecen principios similares como: “Robustez”, sumado a “privacidad y seguridad” (pero sin “protección”), “Solidez y seguridad técnica” (podemos suponer que el primer término refiere a robustez y el segundo a seguridad), “Seguridad y sostenibilidad” y “Garantía de seguridad y privacidad”.

En cuanto a la definición de ciclo de vida, encontramos solo dos definiciones a lo largo del análisis de los proyectos. En un caso se define en cuatro etapas: *creación, implementación, monitoreo y retiro* y en el segundo afirma que “el ciclo de vida de un sistema de IA suele implicar varias fases que incluyen: planificar y diseñar; recopilar y procesar datos; crear modelos o adaptar los modelos existentes a tareas específicas; probar, evaluar, verificar y validar; poner a disposición para su uso/implementación; operar y monitorear; y retirar/desmantelar. Estas fases suelen tener lugar de manera iterativa y no son necesariamente secuenciales” (Proyecto Gollán). Consideramos que esta última puede servir de referencia.

El cuarto elemento clave es la distinción entre seguridad y protección. A pesar de que todos los proyectos de una u otra forma refieren al término “seguridad”, ninguno lo define, como tampoco definen “protección”.

En cuanto a gobernanza de la IA, la exploración de los sistemas de gobernanza propuestos y el mapa de actores claves dio como resultado que la mayoría de los proyectos contemplan una autoridad de aplicación, aunque algunos no la definen específicamente, o no especifican dependencia. Por ello en las tablas 1 y 2 se califica a estas últimas como “Parcial”. Estas autoridades tienen diferentes niveles jerárquicos según el proyecto. El siguiente listado describe las propuestas:

1. Autoridades de nivel ministerial:

Un proyecto propone el Ministerio de Inteligencia Artificial de la Nación Argentina, como autoridad de regulación y aplicación de la ley, que también será encargado de la “estrategia nacional de IA”

2. Autoridades dependientes de ministerios y/o secretarías:

Un proyecto propone el Instituto Nacional de IA (INIA) dependiente del Ministerio de Ciencia y Tecnología e Innovación (actualmente, una secretaria dependiente de la Jefatura de Gabinete de Ministros). Se trata de un organismo de desarrollo pero no de gobernanza en relación con riesgos y seguridad, por lo que se lo consignó en la Tabla como Parcial.

Un proyecto propone la Comisión Nacional de Inteligencia Artificial, dependiente de la Secretaría de Innovación, Ciencia y Tecnología, bajo la órbita de la Jefatura de Gabinete de Ministros, cuyos integrantes serán nombrados por el Jefe de gabinete con recomendaciones recibidas por la comisión de Ciencia de la Cámara de Diputados de la Nación, por la Comisión de Ciencia y Tecnología del Senado de la Nación y por el Consejo Interuniversitario Nacional.

Un proyecto propone como autoridad de aplicación y de fiscalización de la ley al Instituto Nacional de Tecnología Industrial (INTI), que depende del Ministerio de economía, industria y comercio.

3. Autoridades independientes o autárquicas:

Un proyecto propone la Agencia Nacional de Supervisión de IA (ANSIA) organismo descentralizado, independiente.

Un proyecto propone la Agencia de Gestión y Conocimiento, entidad autárquica en el ámbito de la Secretaría de Innovación, Ciencia y Tecnología.

4. Otras

Un proyecto propone la Agencia de acceso a la información pública, sin especificar dependencia.

Un proyecto propone el Consejo Asesor de IA con carácter *ad honorem*, que será la autoridad de aplicación de la ley, sin especificación de dependencia ni fuerza de aplicación (se lo identificó como No).

Continuando con el análisis, se seleccionaron cuatro subelementos, que consideramos importantes para una gobernanza ética y segura de la IA. Uno de estos elementos, el referido a la autoridad de aplicación, fue recién descripto. Con respecto a los otros tres, son: Definición de actores claves; Evaluación de los sistemas de IA, certificación y registros; y Monitoreo periódico y continuo del desempeño del sistema de IA.

Con respecto a la Autoridad de aplicación, como dijimos, siete de los once proyectos la incluyen y describen su lugar con precisión; dos lo hacen de manera parcial, y dos no especifican claramente su lugar.

En cuanto a la Definición de actores clave, cinco proyectos tienen definiciones exhaustivas; dos proyectos, tienen definiciones parciales; y cuatro proyectos no delimitan actores clave.

En relación con la Evaluación de los sistemas y/o el registro de sistemas de IA, ocho proyectos incluyen mecanismos de Evaluación y/o registro; y tres los incluyen parcialmente.

Finalmente, en cuanto a la instancia de Monitoreo periódico o continuo, seis proyectos lo contemplan, tres lo hacen de manera parcial y dos no lo consideran.

A continuación, observaremos en conjunto los últimos tres elementos clave: Identificación de categorías de riesgo; Notificación e investigación de incidentes/accidentes de IA; y Perspectiva proactiva o reactiva, ya que consideramos que tienen una relación estrecha en la misión de crear un sistema de seguridad y gestión de riesgos de la IA.

Como ya dijimos, identificamos solo dos proyectos proactivos que contienen el principio “Robustez, seguridad y protección”. Luego encontramos que ocho de los once proyectos que contienen una pirámide de riesgo con identificación clara de categorías que van desde la prohibición hasta el uso sin requisitos, y uno parcial. En cuanto a la notificación de accidentes/incidentes y su investigación encontramos que si bien cinco proyectos contemplan la investigación, no dejan en claro el tipo de investigación al que se refieren: si se trata de una investigación administrativo-judicial, una investigación de seguridad sistémica o ambas, que sería lo recomendable. Es significativo también que ningún proyecto define a qué denomina incidente/accidente de IA. Por todo esto concluimos que, en resumen, siete proyectos de ley presentan una perspectiva proactiva de riesgos de la IA: tres de ellos de manera plena y cuatro de manera parcial.

Para cerrar esta sección, cabe subrayar que este es un documento vivo, que seguirá alimentándose de nuevos proyectos de ley, así como de las modificaciones que se realicen sobre los proyectos presentados hasta el momento.

3. Conclusión

El análisis sistemático de los proyectos de ley argentinos sobre inteligencia artificial evidencia un panorama normativo heterogéneo, en el que predominan los enfoques proactivos (algunos de manera plena, otros de manera parcial) frente a los riesgos, si bien muchas veces enfocados desde una perspectiva administrativo-jurídica, en detrimento de una concepción integral de la seguridad sistémica de la IA. La tendencia a concebir la regulación principalmente como un mecanismo sancionatorio *ex post* resulta insuficiente frente a tecnologías de alto riesgo cuya dinámica exige estrategias preventivas, predictivas y basadas en el pensamiento sistémico.

Si bien todas las iniciativas reconocen la necesidad de regular estas tecnologías, no todas contemplan principios éticos, que son la base para establecer un consenso orientado a la gobernanza nacional e internacional, así como un signo de reflexividad del Legislativo acerca de los valores sociales a custodiar. La mayoría contemplan autoridades de aplicación, pero no todas ellas establecen claramente funciones y jerarquía institucional.

La mayoría de los proyectos comprende que las IA pueden ser tecnologías de riesgo, si bien aún persisten vacíos en aspectos estructurales: definiciones precisas de IA, delimitación del ciclo de vida de los sistemas de IA, distinción inequívoca entre seguridad (*safety*) y protección (*security*), y mecanismos robustos de monitoreo, notificación de incidentes/accidentes, y una autoridad independiente de investigación sistémica de incidentes y accidentes. Estas ausencias limitan la capacidad del Estado para anticipar, comprender y mitigar riesgos complejos, incluidos aquellos de carácter sistémico o catastrófico. El fortalecimiento de estas instancias incentivaría, además, la formación de equipos técnicos especializados como parte de una necesaria estrategia de alfabetización y capacitación laboral de acuerdo con las nuevas exigencias que plantea la inteligencia artificial al mundo del trabajo.

La matriz analítica desarrollada por el TecnoCenoLab, aplicada en este trabajo comparativo, constituye un aporte original y replicable para el análisis normativo de la IA, tanto en el ámbito nacional como en otras escalas jurisdiccionales. Su valor reside en integrar aprendizajes provenientes de estándares internacionales de gobernanza y ética de la IA, lecciones de las industrias de alto riesgo y enfoques contemporáneos de seguridad sistémica, permitiendo evaluar de manera comparada la solidez de los marcos regulatorios propuestos.

De cara al futuro, una política pública eficaz en materia de inteligencia artificial en la Argentina debería avanzar hacia una ley integral que incorpore la promoción de la investigación y el desarrollo de la IA, la alfabetización digital interdisciplinaria y una política de seguridad exhaustiva, que incluya definiciones claras sobre sistemas de

IA, un ciclo de vida preciso, taxonomías de riesgo consistentes y sistemas de gobernanza robustos, conteniendo instancias independientes de investigación de accidentes/incidentes. Así será posible fortalecer tanto el despliegue de las potencialidades beneficiosas de los sistemas de IA para la sociedad y para las comunidades, como la seguridad y la protección de los ciudadanos, resguardando los derechos fundamentales, los valores democráticos y la integridad de las formas de vida, el medio ambiente y los ecosistemas.

Referencias bibliográficas

Bengio, Yoshua et al. (2024). Managing extreme AI risks amid rapid progress. *Science* 384, 842-845.

Costa, Flavia, Mónaco, Julián, Covello, Alejandro, Novidelsky, Iago et al (2023). Desafíos de la Inteligencia Artificial generativa. Tres escalas y dos enfoques transversales. *Question/Cuestión*. Nro.76, Vol.3.

Costa, Flavia y Mónaco, Julián (2024). ¿Pueden las humanidades coevolucionar con los humanos (y las máquinas)? Inteligencia artificial, sociedad artificial y los desafíos para las ciencias sociales. *Voces en el Fénix*, (93). Facultad de Ciencias Económicas, Universidad de Buenos Aires.

Crawford, Kate (2022). *Atlas de inteligencia artificial. Poder, política y costos planetarios*. Buenos Aires, Fondo de Cultura Económica.

Hendrycks, Dan, Mazeika, Mantas y Woodside, Thomas (2023). *An overview of catastrophic AI Risk* (CAIS).

Hendrycks, Dan (2023). *AI Safety, Ethics, and Society Course: Artificial Intelligence Fundamentals*. En Internet: bit.ly/3Nu1C6o.

Hendrycks, Dan, Nicholas Carlini, John Schulman and Jacob Steinhardt. "Unsolved Problems in ML Safety", ArXiv [abs/2109.13916](https://arxiv.org/abs/2109.13916) (2021): n. pp.1, 11.

Jefatura de Gabinete de Ministros, Secretaría de Gobierno de Modernización. (2019). *Resolución 1523/2019* (RESOL-2019-1523-APN-SGM#JGM). *Boletín Oficial de la República Argentina*, 18 de septiembre de 2019.
<https://www.boletinoficial.gob.ar/detalleAviso/primera/216860/20190918>

Jefatura de Gabinete de Ministros, Subsecretaría de Tecnologías de la Información (2023). *Disposición SSTI 2/2023* (DI-2023-2-APN-SSTI#JGM). *Boletín Oficial de la República Argentina*, 1 de junio de 2023.
<https://www.boletinoficial.gob.ar/detalleAviso/primera/287679/2023060>

Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura (UNESCO). (2021). *Recomendación sobre la ética de la inteligencia artificial*.



Organización para la Cooperación y el Desarrollo Económicos (OCDE). (2024a). *Principios de la OCDE sobre inteligencia artificial*. OECD.
<https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449>

Organización para la Cooperación y el Desarrollo Económicos (OCDE). (2024b). *Memorando explicativo sobre la definición actualizada de un sistema de inteligencia artificial (OECD Artificial Intelligence Papers, n.º 8)*. OECD Publishing.

Organización para la Cooperación y el Desarrollo Económicos (OCDE). (2024c). *Defining AI incidents and related terms* (OECD Artificial Intelligence Papers, n.º 16). OECD Publishing.

Parlamento Europeo y Consejo de la Unión Europea. (2024). *Reglamento (UE) 2024/1689 del Parlamento Europeo y del Consejo de 13 de junio de 2024*.
<https://eur-lex.europa.eu/legal-content/ES/TXT/?uri=CELEX:32024R1689>

TecnocenoLab. (2023). *Proyecto: Desafíos e impactos de la inteligencia artificial. Marcos normativos, riesgos y retos para la calidad democrática en la Argentina*. <https://tecnocenolab.ar/proyecto/>

Uuk, Risto et al. (2024). A Taxonomy of Systemic Risks from General-Purpose AI. *RAND Working Paper Series Forthcoming*. <http://dx.doi.org/10.2139/ssrn.5030173>

Proyectos analizados

- Proyecto de ley que dispone controlar el desarrollo, la implementación y utilización de sistemas basados en inteligencia artificial dentro del territorio argentino. En Internet: [Proyecto Romero](#)
- Proyecto de ley que crea el Instituto Nacional de Inteligencia Artificial. En Internet: [Proyecto García Larraburu](#)
- Proyecto de ley sobre sistemas de inteligencia artificial. En Internet: [Proyecto Sapag Pilatti](#)
- Proyecto de ley de regulación de la inteligencia artificial. En Internet: [Proyecto Doñate](#)
- Proyecto de ley que establece un marco jurídico para el desarrollo, la comercialización, distribución y utilización de los sistemas de inteligencia artificial. En Internet: [Proyecto Fama](#)
- Proyecto de ley que establece el marco legal para la investigación, el desarrollo, uso y regulación de la inteligencia artificial. En Internet: [Proyecto Ávila](#)
- Marco Normativo y de Desarrollo de los Sistemas de IA. En Internet: [Proyecto Gollán](#)
- Ley Nacional de Inteligencia Artificial. En Internet: [Proyecto Chumpitaz](#)
- Marco regulatorio de la inteligencia artificial. Creación del Registro Nacional de Sistemas de Inteligencia Artificial. En Internet: [Proyecto Giuliano](#)
- Régimen jurídico aplicable para el uso responsable de la inteligencia artificial (IA), en la República Argentina. En Internet: [Proyecto Brügge](#)
- Responsabilidad algorítmica y promoción de la robótica, algoritmos verdes e inteligencia artificial en la República Argentina. En Internet: [Proyecto Ferraro](#)