

La securitización de la inteligencia artificial: un análisis de sus impulsores y sus consecuencias*

Mónica A. Ulloa Ruiz y Guillem Bas Graells

Recibido: 29 de noviembre de 2024 | Aceptado: 5 de mayo de 2025 | Modificado: 2 de junio de 2025
<https://doi.org/10.7440/res93.2025.04>

Resumen | En este artículo se examina el proceso de securitización de la inteligencia artificial (IA) en Estados Unidos, entendida como la configuración discursiva de esta tecnología como una cuestión de seguridad que justifica un tratamiento político y regulatorio excepcional. En un contexto de creciente preocupación por los riesgos asociados al desarrollo de la IA, en el artículo se propone identificar los impulsores que motivan su securitización y evaluar las consecuencias de dicho enfoque en la formulación de políticas públicas. Para ello, se usa la metodología del análisis crítico del discurso a un corpus compuesto por veinticinco actos de habla provenientes de agencias gubernamentales, organizaciones técnico-científicas y medios de comunicación. El análisis revela la coexistencia de dos gramáticas discursivas predominantes: una basada en amenazas, que enfatiza la presencia de adversarios concretos y legítima medidas extraordinarias, y otra basada en riesgos, que favorece respuestas preventivas integradas en marcos regulatorios tradicionales. Además, se identifica una tensión entre enfoques de securitización nacional y procesos de macrosecuritización, en los que la IA se trata como un riesgo global que amenaza a la humanidad. Estas dinámicas, influenciadas por factores institucionales, ideológicos y geopolíticos, generan respuestas políticas dispares y, en ocasiones, contradictorias. La principal contribución del estudio radica en ofrecer un marco analítico que permita distinguir entre diferentes lógicas de securitización y comprender sus efectos en la gobernanza de la IA. Se concluye que una aproximación basada en la gestión de riesgos, más que en la identificación de amenazas, favorece respuestas sostenibles y cooperativas en el largo plazo. Asimismo, se argumenta que la integración de la IA en la política convencional, sin recurrir sistemáticamente a su securitización, es esencial para construir marcos regulatorios eficaces y compatibles con los desafíos globales actuales.

Palabras clave | inteligencia artificial; riesgo catastrófico global; riesgo existencial; securitización; seguridad nacional

* Esta investigación contó con el apoyo del Observatorio de Riesgos Catastróficos Globales (ORCG), Estados Unidos. Ambos autores participaron en todas las etapas del desarrollo del artículo: Guillem Bas Graells fue responsable principal de la construcción del marco teórico, la interpretación de los hallazgos y la redacción de las secciones de discusión; Mónica A. Ulloa Ruiz se encargó de la sistematización del análisis discursivo, el diseño metodológico, el procesamiento del corpus y la redacción de la sección de resultados. La sección de conclusiones se redactó de manera conjunta y refleja las contribuciones integradas de ambos autores. El diseño general del artículo y la revisión crítica del manuscrito se hicieron de forma colaborativa. A los autores les gustaría agradecer a Ivanna Alvarado, Roberto Tinoco, Jaime Sevilla y Gideon Futerman por sus útiles comentarios y por las discusiones que se dieron sobre varias versiones del artículo. Todos los errores restantes son responsabilidad de los autores.

The Securitization of Artificial Intelligence: An Analysis of its Drivers and Consequences

Abstract | This article examines the securitization of artificial intelligence (AI) in the United States, understood as the discursive framing of this technology as a security issue that justifies exceptional political and regulatory treatment. In a context of growing concern over the risks associated with the development of AI, the article identifies the drivers of securitization and evaluate its consequences for public policy-making. To this end, it applies critical discourse analysis to a corpus composed of twenty-five speech acts from government agencies, technical-scientific organizations, and the media. The analysis reveals the coexistence of two dominant discursive grammars: one threat-based, which emphasizes concrete adversaries and legitimizes extraordinary measures; and one risk-based, which promotes preventive responses embedded in traditional regulatory frameworks. The article also identifies a tension between national securitization approaches and macrosecuritization processes, where AI is portrayed as a global risk that threatens humanity. These dynamics—shaped by institutional, ideological, and geopolitical factors—lead to divergent and, at times, contradictory political responses. The study's main contribution lies in offering an analytical framework that distinguishes between different logics of securitization and clarifies their impact on AI governance. It concludes that an approach based on risk management, rather than threat identification, enables more sustainable and cooperative long-term responses. It further argues that integrating AI into conventional policymaking—without routinely resorting to securitization—is essential for building effective regulatory frameworks aligned with current global challenges.

Keywords | artificial intelligence; existential risk; global catastrophic risk; national security; securitization

A securitização da inteligência artificial: uma análise de seus vetores e consequências

Resumo | Neste artigo, analisa-se o processo de securitização da inteligência artificial (IA) nos Estados Unidos, entendida como a construção discursiva dessa tecnologia enquanto questão de segurança que justifica um tratamento político e regulatório excepcional. Em um contexto de crescente preocupação com os riscos associados ao desenvolvimento da IA, o artigo tem como objetivo identificar os fatores que motivam sua securitização e avaliar as consequências de tal abordagem na formulação de políticas públicas. Para tanto, adota-se a metodologia de análise crítica do discurso a um corpus composto por 25 atos de fala de órgãos governamentais, organizações técnico-científicas e meios de comunicação. A análise revela a coexistência de duas principais gramáticas discursivas: uma baseada em ameaças, que enfatiza a presença de adversários concretos e legitima medidas extraordinárias; e outra baseada em riscos, que favorece respostas preventivas integradas aos quadros regulatórios tradicionais. Além disso, observa-se uma tensão entre as abordagens nacionais de securitização e os processos de macrosecuritização, nos quais a IA é tratada como um risco global que ameaça a humanidade. Essas dinâmicas, influenciadas por fatores institucionais, ideológicos e geopolíticos, geram respostas políticas dispares e, em alguns momentos, contraditórias. A principal contribuição do estudo está em oferecer uma estrutura analítica para distinguir entre diferentes lógicas de securitização e compreender seus efeitos na governança da IA. Conclui-se que uma abordagem baseada na gestão de riscos, e não na identificação de ameaças, favorece respostas sustentáveis e cooperativas no longo prazo. Argumenta-se também que a integração da IA na política convencional, sem recorrer sistematicamente à sua securitização, é essencial para construir quadros regulatórios eficazes e compatíveis com os desafios globais atuais.

Palavras-chave | inteligência artificial; risco catastrófico global; risco existencial; securitização; segurança nacional

Introducción

La securitización es un proceso mediante el cual un tema se convierte en una cuestión de seguridad. Este puede desarrollarse a través de actos discursivos que enmarcan verbalmente un tema como una cuestión de seguridad, o mediante la práctica, al tratar un tema *de facto* como una amenaza (Williams 2003). En su forma más básica, la securitización es un acto discursivo conocido como *movimiento securitizador*, en el que un agente político construye una percepción colectiva dentro de una comunidad para presentar cierto fenómeno como una amenaza existencial hacia un objeto de referencia considerado valioso, y así justifica un llamado a medidas urgentes y excepcionales (Buzan y Wæver 2003).

La característica distintiva de este proceso es su estructura retórica. En el discurso de seguridad, un problema se amplifica y se presenta como de suprema prioridad: etiquetar un asunto como un tema de seguridad otorga a un actor el derecho de tratarlo con medios extraordinarios (Buzan, Wæver y de Wilde 1998). Así, para que una securitización sea exitosa, generalmente debe incluir tres elementos esenciales: la percepción de una amenaza existencial, la necesidad de acción urgente y la posibilidad de romper las normas políticas o sociales vigentes para enfrentar dicha amenaza (Buzan, Wæver y de Wilde 1998). Cuando un tema es securitizado, no significa necesariamente que represente una amenaza real para la supervivencia de un Estado, sino que ha tenido éxito en presentarse como tal. Esto puede conllevar riesgos, como permitir a los gobiernos tomar medidas autoritarias bajo la justificación de proteger la seguridad nacional.

El marco de la securitización se estructura en torno al objeto referente (la entidad que se busca proteger), la amenaza existencial (un peligro para la supervivencia de ese objeto), las medidas extraordinarias (acciones que exceden las reglas de la política convencional), el actor securitizante (quien enmarca un tema como amenaza) y la audiencia (quien acepta o rechaza la legitimidad de estas medidas). Si bien tradicionalmente el objeto referente ha sido la seguridad nacional, es decir, garantizar la supervivencia del Estado nación, el concepto ha evolucionado para abarcar desde individuos hasta la humanidad en su conjunto (Sears 2023).

En los últimos años, la teoría de la securitización ha considerado amenazas que trascienden el nivel nacional, incluyendo fenómenos globales como los ciberataques (Aguilar Antonio 2025), el cambio climático (Diez, von Lucke y Wellmann 2016) o la inteligencia artificial (IA) (Aguilar Antonio 2024; Zeng 2021). Este proceso marca una transición hacia lo que se denomina *macrosecuritización*, un enfoque en el que una amenaza de alcance global organiza y subsume una serie de securitizaciones de nivel medio o local en torno a un peligro central (Buzan y Wæver 2009). La IA, por su naturaleza de doble uso y su potencial disruptivo en múltiples sectores, ha pasado a enmarcarse discursivamente como una cuestión de seguridad nacional en Estados Unidos. Esto ha llevado a instituciones como el Departamento de Seguridad Nacional y la Agencia de Seguridad Nacional a involucrarse en su regulación. Al mismo tiempo, debido a ese potencial, se ha posicionado como un desafío transnacional que trasciende las preocupaciones del Estado nación, y que requiere una respuesta y coordinación multilateral.

Esta doble caracterización de la IA como amenaza nacional y global genera tensiones en su gobernanza. Si bien los procesos de securitización nacional y macrosecuritización no son necesariamente excluyentes y pueden amalgamarse a varios niveles, la priorización de la seguridad nacional frecuentemente entra en conflicto con los imperativos de la seguridad global (Buzan y Wæver 2009). Las macrosecuritizaciones pueden incorporar distintas securitizaciones nacionales unilaterales, pero estas dinámicas competitivas entre Estados pueden socavar los esfuerzos de coordinación internacional necesarios para abordar los riesgos globales de la IA.

En este artículo se examina cómo esta interacción y tensión entre discursos de securitización configura la percepción de la IA como amenaza y moldea su regulación. A través de un análisis crítico del discurso aplicado a veinticinco actos de habla de actores gubernamentales, técnicos y medios de comunicación, en el estudio se identifican los factores que impulsan esta securitización y se evalúan sus consecuencias bajo dos lógicas distintas: una orientada a la amenaza y otra al riesgo (Diez, von Lucke y Wellmann 2016).

En la siguiente sección se detalla la metodología adoptada para el análisis, basada en el enfoque de análisis crítico del discurso. A continuación, se exponen los resultados del estudio y se enfatiza en las lógicas de securitización identificadas y sus manifestaciones discursivas. Posteriormente, se analizan las tensiones entre la securitización nacional y la macrosecuritización, así como los factores que impulsan y desafían cada enfoque. Finalmente, se discuten las implicaciones de estos procesos para la gobernanza de la IA y se propone un posible enfoque para la formulación de políticas que integren estos riesgos en marcos regulatorios convencionales.

Metodología

La metodología utilizada fue la de análisis crítico del discurso (Van-Dijk 2016), que permite examinar cómo los discursos contribuyen a la construcción de realidades sociales, en este caso, la securitización de la IA. Para la selección de actos de habla se tuvieron en cuenta dos criterios: en primer lugar, el texto debía seguir las reglas formales para considerarse un acto legítimo dentro de un marco comunicativo (condición interna o lingüístico-gramatical). Esto incluye, por ejemplo, documentos oficiales o discursos emitidos a través de los canales institucionales establecidos por cada actor. En segundo lugar, el texto debía provenir de actores con capacidad de influencia, es decir, aquellos que ocupan posiciones de poder o autoridad en su contexto social (condición externa o contextual).

Con el fin de minimizar el sesgo en el análisis, se recurrió a diferentes fuentes. Los actos de habla seleccionados se enmarcaron en los discursos de tres tipos de actores. En primer lugar, se analizaron informes y declaraciones públicas de instituciones estatales estadounidenses, debido a su poder para imponer marcos de interpretación sobre la IA. En segundo lugar, se incluyeron declaraciones oficiales de grupos técnicos y científicos con sede en Estados Unidos, dada su capacidad de influir en las políticas públicas gubernamentales y de modelar el discurso en calidad de expertos. Finalmente, se incorporaron textos provenientes de la prensa escrita, considerada un actor funcional en la difusión y legitimación de discursos sobre seguridad, especialmente en un contexto en el que los medios digitales desempeñan un papel determinante en la formación de la opinión pública.

Para la selección de los informes gubernamentales, se realizó una búsqueda sistemática en las páginas oficiales de departamentos ejecutivos como el Departamento de Estado y el Departamento de Seguridad Nacional, identificando aquellos documentos que abordaran la IA en el marco de temas de seguridad nacional, riesgos tecnológicos o amenazas emergentes. La inclusión de cada documento se basó en la presencia explícita de estos temas, y no en una lectura interpretativa amplia. Para reducir la posibilidad de sesgo, se evitó seleccionar únicamente documentos que reforzaran una hipótesis específica, y se incluyeron también textos que ofrecieran posturas matizadas o incluso contradictorias sobre los riesgos de la IA.

En cuanto a las declaraciones de científicos y grupos técnicos, se consideraron informes y comunicados emitidos por organizaciones científicas y *think tanks* especializados en IA, como las declaraciones publicadas por el Centro para la Seguridad de la IA (CAIS, por sus siglas en inglés) durante los Diálogos internacionales sobre la seguridad de la IA (IDAIS, por sus siglas en inglés). También se incluyeron documentos de expertos independientes

con alto impacto, como el ensayo “Situational Awareness: The Decade Ahead” de Leopold Aschenbrenner (2024). Para el análisis de prensa escrita, se seleccionaron artículos de medios con amplia circulación, credibilidad internacional y posicionamiento relevante en la agenda informativa global; CNN y *The New York Times* fueron elegidos como casos representativos por su alcance, influencia y producción constante de contenido sobre tecnología y seguridad. La selección no se limitó a un canal específico dentro de CNN, sino que abarcó sus principales publicaciones escritas en línea, buscando mantener consistencia en el tipo de discurso analizado. Si bien esta elección puede excluir otras voces mediáticas, se privilegió la coherencia del corpus y la comparabilidad de los discursos.

El proceso de análisis textual incluyó la descomposición del discurso en categorías específicas y la reconstrucción del significado a través de la interpretación contextual (Santander 2011). El análisis de los textos se procesó con Atlas.ti y se usó una codificación que incluyó treinta códigos. Para determinar dichos códigos, se realizó un primer listado deductivo basado en el marco teórico revisado y una segunda iteración inductiva en la que se agregaron nuevos códigos al listado luego del primer análisis de los textos en cuestión. Posteriormente, se cuantificó la presencia de cada código en el corpus textual, generando métricas de enraizamiento que indicaron su frecuencia de aparición. Además, se adelantó un análisis de coocurrencias que permitió examinar las relaciones entre códigos e identificar patrones de asociación en el discurso. Los datos se procesaron a través de tablas de frecuencia, matrices de coocurrencia y visualizaciones de redes semánticas. Los datos cuantitativos se leyeron posteriormente mediante un análisis cualitativo que consideró el marco institucional, social y político de los textos.

Finalmente, los resultados se interpretaron a través del marco teórico de la securitización, vinculando los patrones observados con procesos de construcción social de amenazas y respuestas institucionales. Durante el análisis, se observó que, si bien algunos de los textos analizados están estructurados en una gramática de securitización basada en amenaza, como la propuesta por Buzan, Wæver y de Wilde (1998), otros hacen parte de una segunda categoría que no incluye amenazas existenciales ni medidas excepcionales, pero que trata la IA como un asunto de seguridad distinto de la política convencional (Corry 2012). Para acotar esta nueva categoría, se empleó un marco que diferencia entre securitización basada en amenaza, securitización basada en riesgos y politización tradicional (Rhinard *et al.* 2024). En este marco, la securitización construye escenarios de daño directo con amenazas existenciales que justifican medidas extraordinarias, mientras que la securitización basada en riesgos aborda condiciones que podrían generar daños futuros, lo que requiere medidas precautorias. Ambos enfoques se apartan de la política convencional, en la que los objetos son simplemente gobernables y sujetos a compensaciones políticas (ver [tabla 1](#)).

Tabla 1. Diferenciación entre marcos de securitización

Categoría	Securitización basada en amenaza	Securitización basada en riesgo	Politización (política normal)
Gramática	Construcción de un escenario de daño directo (una amenaza existencial) a un objeto de referencia valioso.	Construcción de condiciones que hacen posible el daño (un riesgo) a un objeto de gobernanza.	Construcción del objeto como gobernable (haciéndolo distinto, moldeable, medible).
Imperativo político	Plan de acción para defenderse contra una amenaza externa al objeto de referencia.	Plan de acción para aumentar la gobernanza y la resiliencia del objeto de referencia.	Plan de acción para maximizar la utilidad en compensaciones con otros bienes.
Efectos performativos	Legitimación de medidas excepcionales (secreto, acciones sin restricciones) dirigidas a la supervivencia.	Legitimación de medidas precautorias, es decir, inclusión de un margen de seguridad.	Legitimación de compensaciones en relación con otros bienes.

Fuente: elaboración propia a partir de una adaptación de Rhinard *et al.* (2024).

Para leer los datos a través de este marco, los indicadores discursivos de la lógica de amenaza incluyeron el uso claro del lenguaje de amenaza-seguridad y exhortaciones a medidas con un sentido de urgencia. En contraste, la lógica de riesgo se manifestó a través de referencias a causas constitutivas del daño y peligros cotidianos y omnipresentes, el uso del lenguaje de *riesgo-medida* y referencias a modelos y enfoques de regulaciones más tradicionales.

Resultados

Los resultados de la codificación revelan la distribución de los términos predominantes en el corpus¹. En la [tabla 2](#) se sintetiza el enraizamiento de los códigos y se muestran los conceptos más frecuentes, lo que evidencia no solo la frecuencia de aparición de cada término, sino también su importancia relativa en el contexto general del corpus.

Además del enraizamiento, se analizaron las coocurrencias entre los códigos para identificar relaciones entre conceptos dentro del discurso securitizador. Estas coocurrencias ayudaron a entender cómo los códigos interactúan entre sí en los textos analizados, lo que revela patrones discursivos dominantes y prioridades narrativas. En la [tabla 3](#) se detallan estas coocurrencias, mostrando la frecuencia con la que ciertos códigos aparecen juntos en el corpus.

Tabla 2. Enraizamiento de códigos en el corpus analizado

Código	Enraizamiento
Seguridad	1367
Riesgo	1135
Control	659
Militarización	575
Regulación	457
Protección	388
Amenaza	383
Impacto	293
Adversario	255

Fuente: elaboración propia a partir de los resultados obtenidos.

Tabla 3. Coocurrencia de códigos relevantes en el corpus analizado

	Riesgo	Amenaza	Catastrófico	Adversario	Peligro	Daño	Terrorismo
Seguridad	145	90	29	25	24	10	7
Control	106	14	41	13	14	9	2
Regulación	66	7	19	1	8	9	3
Militarización	27	11	0	18	2	2	3
Proteger	27	20	2	5	4	11	4
Impacto	45	4	12	9	4	6	1
Beneficio	40	6	5	0	3	7	0

Fuente: elaboración propia a partir de los resultados obtenidos.

1 Los conceptos y la sistematización de los resultados se hicieron en inglés. Para este artículo se tradujeron al español.

Con base en esto, se analizaron los fragmentos de texto asociados a las coocurrencias más relevantes, lo que permitió destacar dos hallazgos principales. En primer lugar, a nivel gubernamental, la securitización de la IA se materializa a través de dos lógicas discursivas distintas: una basada en amenazas y otra en riesgos. Esta diferenciación no es arbitraria, sino que responde a la estructura institucional y orientación ideológica del actor que articula el discurso. Los departamentos y agencias gubernamentales, según su naturaleza y mandato, adoptan patrones narrativos específicos que reflejan su aproximación particular a la seguridad.

En segundo lugar, el análisis del discurso técnico-científico de organizaciones con base en Estados Unidos revela una tendencia predominante, aunque no unánime, hacia la macro-securitización. Allí, la IA se construye como una amenaza potencial para la humanidad en su conjunto, trascendiendo las fronteras nacionales y demandando respuestas coordinadas a nivel global. Este enfoque contrasta con las narrativas más compartimentadas del sector gubernamental y genera tensiones específicas en la formulación de políticas de seguridad que se revisarán en secciones posteriores.

Securitización basada en riesgo vs. securitización basada en amenaza en el discurso gubernamental

Durante el análisis, se constató la presencia de dos lógicas de securitización. La primera, orientada a la amenaza, identifica un antagonista claro, enfatiza la urgencia de acción y justifica medidas extraordinarias. La segunda, orientada al riesgo, proyecta escenarios futuros de peligro con efectos multifacéticos y aborda amenazas menos directas sin un *otro* claramente identificable (Diez *et al.* 2016; Englund y Barquet 2023) y se enfoca en prevenir la materialización de riesgos mediante gestión continua.

El contraste entre estas dos lógicas se evidencia en la distinción de sus patrones narrativos. En el caso del patrón de amenaza, se observa una alta coocurrencia entre códigos como “seguridad”, “riesgo”, “amenaza” y otros como “adversario” o “militarización”. Este patrón identifica un *otro* amenazante y presenta llamados a la acción con un marcado sentido de urgencia. Por el contrario, el patrón de riesgo se caracteriza por el predominio de términos como “regulación” o “monitoreo”, asociados al código “riesgo”. Este patrón hace referencia a causas del daño y a peligros cotidianos y omnipresentes, y se asocia con un enfoque de gobernanza orientado a la prevención y al desarrollo sostenido de capacidades institucionales.

Es necesario aclarar que, además de la gramática riesgo-amenaza, dos de los actos de habla procesados estuvieron dentro del espectro de politización, es decir, no presentaron una gramática securitizadora. Entre ellos está la más reciente declaración del Departamento de Estado sobre la IA (U.S. Department of State 2024), en la que se prescinde por completo de un discurso de amenaza inminente o de riesgo futuro. Allí no existe ninguna mención a un *otro* amenazante o medidas asociadas a la militarización, ni se evidencia un objeto referente a proteger.

Ahora bien, a nivel gubernamental, se encontró que la gramática discursiva depende del actor que ejecuta el acto de habla. Las estructuras de gobernanza —ya sean departamentales o de agencias centralizadas o descentralizadas— parecen determinar el desarrollo y el tipo de securitización (Rhinard *et al.* 2024). Algunos estudios previos ya han mostrado la influencia de estas estructuras en los procesos de securitización de otros objetos de referencia, como la gestión de inundaciones y el cambio climático (Elander, Granberg y Montin 2022; Wesselink *et al.* 2013). Nuestro análisis evidenció que este factor también es determinante en la construcción del discurso sobre la seguridad de la IA.

Así, en algunos de los documentos asociados al poder ejecutivo, se presenta una securitización centrada en consecuencias de segundo orden². Por ejemplo, textos provenientes de la Casa Blanca, como la “Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence” (Biden 2023), evidencian un discurso enfocado principalmente en el riesgo. En dicho texto, los códigos con mayor enraizamiento y coocurrencias entre sí, “riesgo” (42,24 %) y “regulación” (20,32 %), parecen sugerir un enfoque preventivo centrado en gestionar riesgos futuros sin tratar a la IA como una amenaza inmediata. Al analizar los fragmentos codificados en estas categorías, se evidenció una narrativa consistente con la gestión tradicional, pero que se presenta como excepcional en la medida en que (i) no es gobernado de forma equiparable con otros riesgos de aparente menor escala y (ii) se menciona afectación sobre un objeto de referencia, en este caso la seguridad nacional:

La inteligencia artificial debe ser segura [...] También requiere abordar los riesgos de seguridad más apremiantes de los sistemas de IA —incluidos los relacionados con la biotecnología, la ciberseguridad, la infraestructura crítica y otros peligros para la seguridad nacional— al tiempo que se navega por la opacidad y complejidad de la IA³.
(Biden 2023, traducción propia)

Por su parte, en otros textos analizados, como el informe de la Comisión de Seguridad Nacional sobre Inteligencia Artificial (NSCAI, por sus siglas en inglés), comisión establecida para asesorar al presidente y al Congreso, se evidencia una securitización basada en amenaza (NSCAI 2019). Allí se construye explícitamente la figura del *otro* hostil, ilustrado por aspectos como la explícita referencia a China en 696 ocasiones. La dimensión securitizadora del discurso se evidencia también en la frecuencia relativa de códigos como “adversario” (5,5 %) que se encuentran menos presentes en los textos del poder ejecutivo y que sugieren la conformación de un marco en el que la IA se posiciona como un elemento crítico en la competencia geopolítica y la seguridad nacional.

A modo ilustrativo, consideremos los dos textos mencionados previamente: el de Biden (2023) y el de la NSCAI (2019). Ambos documentos pertenecen a instituciones con estructuras y propósitos diferentes, cuya tendencia securitizadora es claramente distinta. En la [tabla 4](#) se muestra la densidad relativa (porcentaje respecto al total de codificaciones de cada documento) de los códigos “militarización” y “regulación”, lo que ejemplifica la tendencia de estos indicadores distintivos de las lógicas de amenaza y riesgo. En el informe de la NSCAI, el código “militarización” representa el 10,67 % de las codificaciones, en contraste con el 1,60 % en la orden ejecutiva. Por el contrario, el código “regulación” abarca el 20,32 % de las codificaciones en la orden ejecutiva, mientras que en el documento de la NSCAI esta proporción se reduce al 3,34 %

Tabla 4. Contraste en la densidad de códigos indicadores

Código	Orden Ejecutiva (Biden 2023)	NSCAI (2019)
Militarización	1,60 %	10,67 %
Regulación	20,32 %	3,34 %

Fuente: elaboración propia a partir de los resultados obtenidos.

2 Aunque algunos autores, como Corry (2012), sostienen que la denominada *riskificación* es un proceso distinto de la securitización, otros, como Rhinard *et al.* (2024), consideran el enfoque basado en riesgos como parte de la securitización. Siguiendo a este último autor, en este artículo se abordan ambas lógicas (amenaza y riesgo) como parte del mismo proceso securitizador, aunque con implicaciones diferentes para la política y la regulación.

3 En el original: “Artificial Intelligence must be safe and secure. [...] It also requires addressing AI systems’ most pressing security risks—including with respect to biotechnology, cybersecurity, critical infrastructure, and other national security dangers—while navigating AI’s opacity and complexity”.

De esta manera, mientras que la NSCAI, como comisión de seguridad nacional, tiende a enfocarse en las amenazas y soluciones de carácter militar o de seguridad, la Casa Blanca, durante la administración Biden, al emitir una orden ejecutiva, tiende a un enfoque más moderado para abordar los desafíos de la IA desde una óptica de política pública y de gestión de riesgos. Esto podría ser reflejo de cómo sus estructuras institucionales y mandatos específicos modelan sus respectivos discursos y prioridades en torno a la seguridad de la IA.

Otros textos analizados, como el informe “Artificial Intelligence and National Security”, elaborado por el Servicio de Investigación del Congreso (Sayler 2020) y enfocado en seguridad nacional, también tiene un discurso más enfático en las categorías de “militarización” (43,7 %) y “adversario” (6,21 %), que enmarcan la IA en un contexto de rivalidad geopolítica considerablemente mayor a la evidenciada en los documentos anteriores. A diferencia de los textos previamente citados, aquí los términos “amenaza” y “seguridad” tienen una mayor coocurrencia con “adversario” y “militarización”, y se centran más en la amenaza inmediata.

De esta manera, la variación en estos patrones discursivos indica que la securitización de la IA en el contexto gubernamental de Estados Unidos no sigue un proceso uniforme, sino que se presenta de manera diferenciada según el actor. Mientras los textos de departamentos ejecutivos federales parecen privilegiar mayormente la construcción de marcos de gestión del riesgo excepcionales, pero basados en consecuencias de segundo orden, los órganos especializados en seguridad nacional y defensa adoptan posturas más explícitamente securitizantes y basadas en amenaza, alineadas con sus propósitos institucionales.

Una hipótesis adicional podría sugerir que la distinción entre discursos basados en el riesgo y la amenaza, además de estar influenciada por las estructuras institucionales de los actores, podría estar condicionada por la orientación ideológica y política de la administración en turno. En este sentido, los documentos de securitización analizados que responden a una lógica de amenaza se publicaron mayoritariamente durante la administración Trump (2017-2021), mientras que aquellos que adoptan un enfoque centrado en el riesgo han sido más comunes bajo la administración Biden (2021-2025).

En línea con este argumento, en la orden ejecutiva “Maintaining American Leadership in Artificial Intelligence” (Trump 2019) se subraya en la sección 8, “Action Plan for Protection of the United States Advantage in AI Technologies”, una preocupación notable por la seguridad nacional y la ventaja geopolítica asociada a la IA. No obstante, este hallazgo no es concluyente, ya que la extensión limitada de estas referencias dificulta realizar un análisis comparativo exhaustivo sobre la densidad de códigos en relación con otros textos. Incluso, otras órdenes ejecutivas de la primera administración Trump —como la orden ejecutiva 13960 “Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government” (Trump 2020)—, aunque mencionan el uso de la IA con fines de seguridad nacional o defensa, presentan un discurso más centrado en consecuencias de segundo orden.

Securitización nacional vs. macrosecuritización en el discurso técnico

Durante el análisis se evidenció que, en su mayoría, los discursos técnicos y científicos de organizaciones y expertos establecidos en Estados Unidos se enmarcan predominantemente en un proceso de macrosecuritización. Esto significa que posicionan a la humanidad como objeto principal de protección frente a la amenaza existencial de la IA, lo que demanda medidas extraordinarias que trascienden securitizaciones menores como la del Estado nación.

De acuerdo con esta perspectiva, declaraciones como las emitidas en Venecia, Pekín y Oxford por la organización de científicos en los IDAIS se sitúan dentro de este marco

de referencia macrosecuritizador (2023, 2024a, 2024b). En estos textos se presenta una alta frecuencia de términos como “humanidad”, “global”, “internacional” y “estados” (este último en plural), haciendo referencia a la necesidad de actores multilaterales para atender una amenaza que trasciende las fronteras de un Estado nación. Además, el uso del término “catastrófico”, como adjetivo recurrente del riesgo, sugiere un orden superior de prioridad que no se evidenció en documentos gubernamentales internos (ver figura 1).

Figura 1. Nube de palabras de las tres declaraciones públicas del IDAIS



Fuente: elaboración propia a partir de las declaraciones del IDAIS 2023, 2024a y 2024b.

Las declaraciones del IDAIS analizadas se enfocan en la gobernanza y tienden a una macrosecuritización basada en el riesgo; buscan incorporar la gestión de la IA en las agendas gubernamentales sin aludir a medidas extraordinarias, pero mantienen el sentido de urgencia y la identificación de un riesgo existencial para el objeto referente, es decir, la humanidad:

El desarrollo, despliegue o uso no seguro de sistemas de IA puede suponer riesgos catastróficos o incluso existenciales para la humanidad en el curso de nuestras vidas. [...] La combinación de esfuerzos técnicos conjuntos de investigación con un régimen prudente de gobernanza internacional podría mitigar la mayoría de los riesgos derivados de la IA, permitiendo aprovechar sus múltiples beneficios potenciales⁴. (IDAIS 2024a, traducción propia,)

Mientras tanto, otras declaraciones cortas como la del CAIS (2023), que ha alcanzado una amplia repercusión, incluyen factores suficientes para considerarlo un texto macrosecuritizador basado en amenaza. Allí se evidencia una amenaza existencial para la humanidad (extinción), la necesidad de un plan de acción y el llamado de urgencia que lo separa de medidas convencionales de gestión del riesgo, situándose junto a otros escenarios que han requerido medidas excepcionales de gestión, como la amenaza de guerra nuclear: “El riesgo de extinción asociado con la inteligencia artificial debería ser una prioridad global, al mismo nivel que otros riesgos a escala societal como las pandemias y la guerra nuclear”⁵ (CAIS 2023, traducción propia).

4 En el original: “Unsafe development, deployment, or use of AI systems may pose catastrophic or even existential risks to humanity within our lifetimes. [...] The combination of concerted technical research efforts with a prudent international governance regime could mitigate most of the risks from AI, enabling the many potential benefits”.

5 En el original: “Mitigating the risk of extinction from AI should be a global priority alongside other societal-scale risks such as pandemics and nuclear war”.

También, artículos de prensa opuestos al discurso que presenta la IA como un riesgo catastrófico o existencial reconocen que la IA presenta riesgos globales, aunque de menor escala, los cuales deben gestionarse con medidas imperativas, además de mostrar a la humanidad como su objeto de referencia amenazado (Science Media Centre 2023): “Por lo tanto, es imperativo que el público y los responsables políticos comprendan esta distinción para elaborar regulaciones apropiadas que aborden los riesgos reales de estas aplicaciones y no los ficticios”⁶ (Carusso 2024, traducción propia).

En contraste con la tendencia principal hacia la macrosecuritización, otros documentos provenientes del sector técnico-científico se apartan de los discursos expuestos y permanecen en el campo de la securitización nacional basada en amenaza. Entre ellos, el ensayo “Situational Awareness: The Decade Ahead” de Leopold Aschenbrenner (2024)⁷, quien, aunque identifica al “mundo libre” como el objeto de protección, sitúa a Estados Unidos (representante del mundo libre) como amenazado, en oposición a rivales como China. De manera similar, el Project 2025, una agenda política elaborada por más de cien *think tanks* conservadores, también se alinea con esta perspectiva al proponer “detener la contribución a los objetivos autoritarios de China en inteligencia artificial” (Dans y Groves 2023). Esta tensión entre securitización nacional y macrosecuritización, presente tanto en actores gubernamentales como científicos, tiene implicaciones para la gobernanza de la IA que se analizarán en secciones posteriores.

Discusión

¿Qué habilita los procesos de securitización?

Como se evidencia en los apartados previos, los procesos de securitización de la IA responden a diferentes impulsores cuya interacción determina tanto la forma como el alcance de las respuestas institucionales. A continuación, se muestran los patrones distintivos según el tipo de proceso, aunque algunos operan transversalmente.

Securitización basada en amenaza

La securitización tradicional refleja la persistencia de marcos en los que los Estados y sus instituciones de defensa y seguridad mantienen un rol central en la definición y respuesta a amenazas. Esta perspectiva se materializa a través de dos factores. Por un lado, la percepción de competencia geopolítica articula las preocupaciones sobre la IA dentro de narrativas de rivalidad estratégica y amenazas inmediatas, por ejemplo, el marco ideológico de la carrera armamentística que sitúa a la IA como un elemento decisivo en la distribución de poder global. Por otro lado, las estructuras institucionales militares y de defensa priorizan marcos de seguridad nacional sobre consideraciones globales.

Securitización basada en riesgo

El proceso de securitización basada en riesgos representa un giro hacia la gestión administrativa, en la que las estructuras burocráticas transforman amenazas excepcionales en objetos de regulación rutinaria. Este proceso se sustenta en tres impulsores principales: (i) las estructuras institucionales administrativas traducen preocupaciones excepcionales en problemas de gestión cotidiana; (ii) los marcos ideológicos de gestión

6 En el original: “It is, therefore, imperative that the public and policymakers understand this distinction to devise appropriate regulations that address the real risks of these applications and not fictional ones”.

7 Leopold Aschenbrenner, incluido en esta sección como experto independiente, es también fundador de una empresa de inversión en IA y un antiguo empleado de OpenAI, por lo que sus opiniones se encuentran potencialmente influenciadas por sus intereses profesionales y su trayectoria previa en el sector.

administrativa reconfiguran la IA como objeto de gobernanza antes que como amenaza existencial; (iii) los instrumentos previos de gobernanza catalizan la absorción de nuevas preocupaciones dentro de estructuras regulatorias existentes.

Macrosecuritización

La macrosecuritización emerge como respuesta a la naturaleza transnacional de los riesgos asociados a la IA, estableciendo marcos de evaluación y respuesta que trascienden las fronteras nacionales. Sus factores característicos incluyen la percepción de riesgo global, que fundamenta discursos técnico-científicos sobre amenazas multilaterales; la participación de actores securitizadores, como la comunidad técnico-científica, que emplea su experticia para movilizar a otros actores funcionales, como gobiernos nacionales y organismos multilaterales; y un marco ideológico que integra a la IA dentro de un conjunto más amplio de amenazas globales.

La coexistencia de distintos procesos de securitización genera dinámicas complejas, en las que ciertos factores operan simultáneamente en múltiples niveles. Los marcos institucionales preexistentes condicionan la forma en que las preocupaciones emergentes se traducen en políticas concretas y esta multiplicidad de impulsores da lugar a tensiones estructurales dentro del campo de la gobernanza de la IA.

Una primera tensión se presenta entre seguridad nacional y seguridad global, es decir, entre la securitización y la macrosecuritización. En este caso, los imperativos de soberanía tienden a entrar en conflicto con las necesidades de coordinación internacional frente a riesgos compartidos. Una segunda tensión aparece entre la lógica de la gestión administrativa y la lógica de la excepcionalidad, lo que se refleja en la competencia entre enfoques burocráticos que buscan institucionalizar la gobernanza del riesgo y estrategias orientadas a la acción urgente y extraordinaria.

Finalmente, existe una tensión entre el conocimiento técnico y el proceso político, expresada en disputas entre la comunidad científica y los actores políticos tradicionales por la autoridad epistémica. Esta disputa también se entrelaza con la tensión previa entre macrosecuritización y securitización nacional, ya que ambas perspectivas suelen estar respaldadas por distintos tipos de legitimidad. La configuración de estos impulsores no solo explica la coexistencia de múltiples procesos de securitización, sino que también anticipa los desafíos futuros en la gobernanza de la IA. Estos desafíos se revisarán en la siguiente sección.

Implicaciones para la gobernanza

Macrosecuritización y securitización nacional

Existe una tensión entre la macrosecuritización y la securitización nacional de la IA. Sears (2023) apunta que, cuando prevalece la narrativa de securitización de la humanidad, se puede abrir espacio a un consenso entre grandes potencias para la macrosecuritización. Sin embargo, cuando triunfa la securitización nacional, las rivalidades entre las grandes potencias las llevan a priorizar el poder y la seguridad nacional por encima de la seguridad y la supervivencia de la humanidad. La tendencia de los Estados hacia esto último es, según Sears (2023), el motivo principal de los fracasos históricos para macrosecuritizar desafíos globales. Dicha tendencia parece darse porque, como apuntan Buzan y Wæver (2009), las *colectividades limitadas*⁸ de escala media (Estados, civilizaciones) son

8 Son actores sociales o políticos que no representan a toda la humanidad ni a comunidades estrictamente locales, sino unidades intermedias como los Estados nación. Según Buzan y Wæver (2009), estas colectividades son los actores más comunes y duraderos en los procesos de securitización.

objetos de referencia⁹ más persistentes que colectividades más amplias (el conjunto de la humanidad), posiblemente porque se enfrentan a una alteridad más tangible que refuerza su sentido de pertenencia.

En los siguientes párrafos, evaluamos cómo distintas amenazas concretas de la IA dan pie a diferentes interacciones entre el nivel medio (securitización nacional) y el nivel sistema (macrosecuritización). En particular, y a partir de esta distinción propuesta por Buzan y Wæver (2009), argumentamos que distintos tipos de amenazas provocan que las diferentes securitizaciones se vinculen entre sí positivamente —cuando los actores comparten una definición de amenaza y un objeto de referencia— o negativamente —cuando los actores se construyen entre sí como amenazas—. En la práctica, esta distinción no es siempre clara: en algunos casos, los esfuerzos de un actor para protegerse de la percibida amenaza que constituye otro actor pueden ser valiosos para proteger a un conjunto más amplio de actores de una amenaza común.

En el caso de vinculaciones positivas, cuando un Estado nación adopta el marco de la securitización frente a amenazas comunes, las medidas implementadas pueden beneficiar indirectamente la seguridad de otros Estados. Por ejemplo, un Estado nación puede considerar que la posibilidad de que la IA facilite el desarrollo de armas biológicas es una amenaza para sí mismo. Una reacción lógica a esta preocupación podría ser requerir a sus empresas desarrolladoras de IA que adopten medidas para restringir las capacidades de sus modelos en biología o fortalecer los mecanismos de cribado en el suministro de secuencias de ADN con uso dual. En ambos casos, incluso aunque el actor esté eminentemente interesado en proteger la seguridad nacional, el resto de la humanidad también podría beneficiarse de la resultante mejora en la biocustodia global. De hecho, varios Estados podrían desarrollar sentimientos comunes, como la pulsión por defender su monopolio sobre la violencia (Bull 1977), ante la prospectiva de que la IA capacite a actores no estatales o se vuelva cada vez más difícil de controlar.

En ese sentido, la securitización nacional puede ser un instrumento para abordar riesgos globales o incluso un precursor de la macrosecuritización. Esto se debe a que activar y mantener activado un aparato estatal es un primer paso más factible que movilizar esfuerzos multilaterales, tanto a nivel logístico como a nivel discursivo. A modo de ejemplo, los Estados abandonaron sus programas de armas biológicas debido a cálculos de seguridad nacional y solamente más tarde negociaron la Convención sobre Armas Biológicas. De la misma manera, podría tener sentido que los Estados aborden los problemas comunes de manera autosuficiente antes de tratar de encontrar soluciones colaborativas a aspectos del riesgo para los que la unilateralidad sea insuficiente.

Por otro lado, cabe destacar que hay ciertas instancias en las que la securitización nacional y la macrosecuritización de la IA entran en conflicto directo. Por ejemplo, es probable que aquellos Estados que realicen un movimiento de securitización nacional tiendan a militarizar la tecnología, dado que una ventaja comparativa en el ámbito militar podría ser particularmente transformadora en el tablero internacional (Horowitz *et al.* 2020). Además, es posible que la securitización nacional exacerbe dinámicas competitivas peligrosas al incentivar a que los Estados inviertan cada vez más recursos en la carrera por el desarrollo de la IA. Estas presiones competitivas podrían incentivar a los Estados a ser menos cautelosos y dificultar la aprobación de acuerdos internacionales (Gruetzemacher *et al.* 2024). Llevado al extremo, esto también podría resultar en un conflicto entre grandes potencias. Por ejemplo, si un país se encuentra cerca de

9 Como fue mencionado previamente, en el enfoque de la Escuela de Copenhague, un objeto de referencia es aquello que se presenta como amenazado y cuya preservación justifica medidas extraordinarias. Los Estados o civilizaciones son más efectivos como objetos de referencia porque tienen fronteras más claras y una identidad colectiva más sólida frente a la alteridad.

desarrollar un sistema de IA transformativo, otros Estados podrían verse incentivados a perpetrar un ataque preventivo contra los centros de datos del país líder para evitar que este obtenga una ventaja inalcanzable.

En otros casos, aunque exista un juego de suma cero entre dos actores, los efectos de una medida resultante de la securitización nacional pueden ser positivos para el conjunto de la humanidad. Por ejemplo, el gobierno de los Estados Unidos podría requerir a las principales compañías desarrolladoras de IA que fortalezcan sus prácticas de ciberseguridad para proteger los pesos de los modelos de IA más avanzados con el objetivo de evitar que China copie dichos modelos. A pesar de que la política venga motivada por la competición con un adversario directo, esta medida tendría un efecto más amplio: evitar la proliferación de modelos de IA punteros para impedir que actores no estatales maliciosos obtengan acceso a capacidades potencialmente peligrosas o que el progreso en IA se acelere descontroladamente. Siguiendo con la analogía anterior de armas biológicas, Estados Unidos abandonó este programa para disminuir el riesgo de proliferación a sus adversarios, pero esto desencadenó en un mundo con menor riesgo de ataques biológicos en términos absolutos.

En definitiva, los efectos de distintos movimientos securitizadores no dependen tanto de cuál sea el objeto de referencia, sino de cómo se defina la amenaza y qué medidas de protección se adopten. La predominancia de una u otra amenaza dependerá de factores como la percepción del riesgo asociado a la IA o la animadversión hacia adversarios políticos. Transversalmente, ambos factores interactuarán con la ventaja competitiva que se espere que la IA transformativa ofrezca a aquel que la controle. A su vez, el peso de cada uno de estos factores determinará las políticas que implementen los Estados. Si los peligros de la IA se construyen como una amenaza grave para la seguridad nacional, los Estados podrían adoptar medidas con notables externalidades positivas y encontrar espacios de colaboración incluso partiendo de un interés propio limitado a sus confines. Esto se debe a que las amenazas globales más prominentes que se han identificado —incluyendo el uso de la IA para ejecutar ciberataques o para desarrollar armas biológicas o para la pérdida de control sobre sistemas de IA avanzados— también afectan necesariamente a los Estados de manera individual. En cambio, si los Estados otorgan más peso a vencer a sus rivales políticos en la carrera hacia la IA transformativa, la securitización nacional podría socavar la posibilidad de que se produzca una macrosecuritización que proteja a la humanidad.

En cualquier caso, los marcos existentes de seguridad humana y nacional muestran limitaciones significativas para abordar amenazas como la IA. La ausencia de una autoridad política mundial capaz de actuar en nombre de la humanidad hace que la macrosecuritización dependa del consenso entre las grandes potencias, lo cual se ve constantemente desafiado por narrativas en conflicto sobre securitización nacional versus securitización global. Este proceso debe navegar tensiones entre experticia técnica y participación democrática, consenso global y soberanía nacional, y entre la universalidad del riesgo y la particularidad de sus impactos.

Securitización y politización

Además de la tensión entre la securitización nacional y la macrosecuritización, también existen disputas alrededor del propio movimiento securitizador. Independientemente de cuál sea el objeto de referencia a proteger, la securitización implica trascender la lógica de la política “normal” y legitimar medidas extraordinarias. En ese sentido, la securitización probablemente moverá la ventana de Overton —el rango de políticas que son consideradas aceptables por la opinión pública— hacia la aceptación de medidas más radicales. Allen y Chan (2017) argumentan que las implicaciones de la IA para la seguridad nacional serán revolucionarias y, como consecuencia, los gobiernos considerarán medidas políticas extraordinarias como las que se consideraron durante las primeras décadas de

las armas nucleares. Asimismo, Leung (2019) predice que las empresas desarrolladoras de IA se enfrentarán cada vez más a restricciones legislativas impuestas por los Estados, en particular en forma de políticas como los controles de exportación, motivados por las preocupaciones de seguridad nacional.

En todo caso, es amplio el espectro de posibles políticas potencialmente legitimadas por la securitización. En Estados Unidos, la constatación de que la IA podría suponer riesgos graves para la seguridad ha incentivado una mayor participación del Departamento de Seguridad Nacional a través de subdivisiones como la Oficina para la Lucha contra las Armas de Destrucción Masiva o la Agencia de Seguridad de Infraestructura y Ciberseguridad. En un paso más allá, el informe de 2024 de la Comisión de Revisión Económica y de Seguridad de Estados Unidos-China, que asesora al Congreso, recomendó establecer un programa similar al Proyecto Manhattan para avanzar hacia el desarrollo de una IA general (U.S.-China Economic and Security Review Commission 2024). Por último, en un escenario hipotético, en el que la securitización estuviera ya muy arraigada, un gobierno podría establecer un sistema de vigilancia masivo para interceptar y prevenir desarrollos tecnológicos con altas capacidades de destrucción (Bostrom 2019). Así, al movilizar recursos en una dirección en particular, resulta importante que los tomadores de decisiones consideren las posibles implicaciones negativas de sus acciones. En este sentido, como apunta Wæver (1993), existe el riesgo de que los Estados instrumentalicen la securitización para incrementar abusivamente su poder.

Finalmente, cabe decir que muchas de las políticas que se han propuesto para gobernar la IA no requieren de un proceso de securitización. Por ejemplo, la Orden Ejecutiva emitida en el gobierno Biden (2023) obliga a los desarrolladores de modelos de IA por encima de un determinado umbral que reporten el entrenamiento de dichos modelos y los resultados de sus evaluaciones. Estos requerimientos pueden formar parte de un escenario de “política normal” y, de hecho, son habituales en otros sectores no securitizados como la industria farmacéutica o de la aviación. Sin embargo, la aprobación de un régimen en el que todos los sistemas de IA a partir de cierta capacidad deban ser registrados, monitoreados y controlados por el gobierno —o una agencia intergubernamental— sí podría estar reservado a un escenario en el que la IA haya sido securitizada *a priori*.

¿Securitización basada en amenaza o en riesgo?

El análisis invita a considerar las consecuencias de implementar una “política de seguridad de segundo orden”, es decir, una securitización orientada a la gestión de riesgos. A diferencia de una securitización estrictamente basada en amenazas concretas e inmediatas, la securitización basada en riesgos reconoce que estos no pueden eliminarse completamente, sino que requieren una gestión continua. Esta tendencia marca un cambio significativo desde una política de emergencia y excepcionalidad hacia una de permanencia y planificación a largo plazo, lo que genera tres modificaciones clave en la gobernanza: (i) la disociación entre la seguridad y la amenaza existencial, permitiendo una mayor flexibilidad en la gestión de los riesgos; (ii) el desplazamiento del tema de la seguridad como foco central, lo cual posibilita una aproximación más amplia e integradora que incorpora, por ejemplo, aspectos sociales y económicos; y (iii) el reemplazo de criterios de emergencia por políticas de ingeniería social sostenidas en el tiempo (Corry 2012).

Ambos enfoques de securitización, amenaza o riesgo, tienen desventajas. La securitización basada en amenazas, como se ha evidenciado, suele politizar la seguridad, permitiendo a los Estados justificar medidas extraordinarias y evadir procesos habituales de toma de decisiones (Buzan, Wæver y de Wilde 1998). Esto puede llevar a una asignación de recursos hacia intereses estratégicos específicos, afectando otras prioridades gubernamentales. Además, la lógica de amenazas fomenta una narrativa de crisis constante que legitima la intervención estatal y refuerza estructuras de vigilancia y control. Por

su parte, la securitización basada en riesgos genera una percepción de inseguridad difusa y constante que permea la vida pública y privada. Este enfoque puede normalizar el riesgo, haciendo que la inseguridad se perciba como una condición permanente, y promover prácticas de autovigilancia, como se ha visto en otros procesos de securitización (Aradau y Van Munster 2007).

Al considerar estas implicaciones, creemos que una gobernanza eficaz debería equilibrar los imperativos de seguridad nacional con la creación de marcos regulatorios armonizados con la política tradicional. Además, debería considerar los riesgos globales asociados a la IA, que, dada su naturaleza transnacional, exigen una respuesta multilateral y coordinada. Una posible solución sería incorporar estos riesgos en las políticas convencionales de gestión, impulsando la integración de la IA en las estructuras de gobernanza habituales; esto podría lograrse a largo plazo, pasando de un estado de securitización basada en riesgos a uno de politización.

Sin embargo, este enfoque enfrenta retos, como la implementación de un sistema de gobernanza de la IA en un contexto de competencia geopolítica y asimetrías de poder. Para que esta política de seguridad de segundo orden sea efectiva, será necesario fomentar la confianza entre actores internacionales, gestionar los riesgos de manera transparente y estandarizada, y adaptar los marcos regulatorios para que respondan efectivamente al rápido avance de esta tecnología sin requerir medidas excepcionales. De este modo, la gobernanza de la IA podría transformar los discursos securitizadores basados en amenaza en discursos basados en riesgo, más compatibles con los intereses multilaterales, que reconozcan la gravedad de los riesgos sin fomentar un estado de urgencia y excepción, y que, en última instancia, se integren en la política convencional y fomenten acuerdos a escala global.

Conclusiones

En este artículo se muestra cómo la securitización de la IA se está construyendo en Estados Unidos a partir de dos lógicas diferenciadas: una basada en amenazas y otra orientada a la gestión de riesgos que promueve la regulación precautoria y de largo plazo. Las estructuras institucionales y la orientación política son factores determinantes en la preferencia de uno u otro enfoque: mientras algunas agencias identifican una amenaza directa, a menudo incluyendo adversarios políticos, otras buscan integrar los riesgos de la IA en los marcos regulatorios convencionales, ampliando el espectro de la “política normal”.

La predominancia de la seguridad nacional como marco de referencia para la IA plantea riesgos específicos. Como se observa en la literatura sobre macrosecuritización, la competencia entre grandes potencias por el liderazgo en IA puede exacerbar tensiones internacionales y dificultar la cooperación en torno a la gobernanza de esta tecnología. Este contexto sugiere que, si bien la securitización nacional puede impulsar acciones rápidas y eficaces en contextos específicos, es posible que también limite la capacidad de establecer acuerdos globales y de abordar de manera coordinada los riesgos compartidos de la IA. En contraste, un enfoque de macrosecuritización basado en la seguridad de un objeto referente, como la humanidad, permitiría una respuesta colaborativa más amplia, aunque su efectividad depende del consenso y de la voluntad de cooperación entre las grandes potencias.

Este estudio sugiere la posibilidad de una aproximación híbrida en la gobernanza de la IA. Incorporar los riesgos de la IA a la política y la gestión de riesgo convencional podría facilitar una gestión sostenida sin recurrir necesariamente a la securitización basada en amenaza, evitando así las implicaciones negativas asociadas a la movilización de recursos en perjuicio de otras prioridades sociales. Sin embargo, este enfoque plantea un desafío

para los formuladores de políticas, quienes deben equilibrar la urgencia de acción que demanda esta tecnología con la construcción de un marco regulatorio que garantice la seguridad a largo plazo.

Referencias

1. Aguilar Antonio, Juan Manuel. 2024. "Trayectoria y modelo de gobernanza de las políticas de inteligencia artificial (IA) de los países de América del Norte". *Justicia* 29 (45). <https://doi.org/10.17081/just.29.45.7162>
2. Aguilar Antonio, Juan Manuel. 2025. *Tech Leap or Tech Lag: Latin America's Quest to Keep up with Emerging Technologies*. Miami: Steven J. Green School of International & Public Affairs / Florida International University.
3. Allen, Greg y Taniel Chan. 2017. *Artificial Intelligence and National Security*. Cambridge, MA: Belfer Center for Science and International Affairs / Harvard Kennedy School.
4. Aradau, Claudia y Rens Van Munster. 2007. "Governing Terrorism Through Risk: Taking Precautions, (Un)Knowing the Future". *European Journal of International Relations* 13 (1): 89-115. <https://doi.org/10.1177/1354066107074290>
5. Aschenbrenner, Leopold. 2024. "Situational Awareness: The Decade Ahead". *Situational Awareness*, junio. <https://situational-awareness.ai/wp-content/uploads/2024/06/situationalawareness.pdf>
6. Biden, Joseph R. 2023. "Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence". *The White House*, 30 de octubre. <https://bidenwhitehouse.archives.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>
7. Bostrom, Nick. 2019. "The Vulnerable World Hypothesis". *Global Policy* 10 (4): 455-476. <https://doi.org/10.1111/1758-5899.12718>
8. Bull, Hedley, ed. 1977. "How is Order Maintained in World Politics?" En *The Anarchical Society: A Study of Order in World Politics*, 51-73. Londres: Palgrave Macmillan.
9. Buzan, Barry y Ole Wæve. 2003. *Regions and Powers: The Structure of International Security*. Cambridge: Cambridge University Press.
10. Buzan, Barry y Ole Wæver. 2009. "Macrosecuritisation and Security Constellations: Reconsidering Scale in Securitisation Theory". *Review of International Studies* 35 (2): 253-276. <https://doi.org/10.1017/S0260210509008511>
11. Buzan, Barry, Ole Wæver y Jaap de Wilde. 1998. *Security: A New Framework for Analysis*. Boulder: Lynne Rienner Publishers.
12. CAIS (Centre for AI Safety). 2023. "Mitigating the Risk of Extinction from AI Should Be a Global Priority Alongside Other Societal-Scale Risks Such as Pandemics and Nuclear War". CAIS, s.f. <https://www.safe.ai/work/statement-on-ai-risk>
13. Carusso, Jeff. 2024. "Drink the Kool-Aid All You Want, but Don't Call AI an Existential Threat". *Bulletin of the Atomic Scientists*, 29 de abril. <https://thebulletin.org/2024/04/drink-the-kool-aid-all-you-want-but-dont-call-ai-an-existential-threat/>
14. Corry, Olaf. 2012. "Securitisation and 'Riskification': Second-order Security and the Politics of Climate Change". *Millennium* 40 (2): 235-258. <https://doi.org/10.1177/0305829811419444>
15. Dans, Paul y Steven Groves. 2023. *Mandate for Leadership: The Conservative Promise*. Nueva York: The Heritage Foundation.
16. Diez, Thomas, Franziskus von Lucke y Zehra Wellmann. 2016. *The Securitisation of Climate Change: Actors, Processes and Consequences*. Londres: Routledge.
17. Englund, Mathilda y Karina Barquet. 2023. "Threatification, Riskification, or Normal Politics? A Review of Swedish Climate Adaptation Policy 2005-2022". *Climate Risk Management* 40: 100492. <https://doi.org/10.1016/j.crm.2023.100492>
18. Elander, Ingemar, Mikael Granberg y Stig Montin. 2022. "Governance and Planning in a 'Perfect Storm': Securitising Climate Change, Migration and Covid-19 in Sweden". *Progress in Planning* 164: 100634. <https://doi.org/10.1016/j.progress.2021.100634>
19. Gruetzemacher, Ross, Shahar Avin, James Fox y Alexander K. Saeri. 2024. "Strategic Insights from Simulation Gaming of AI Race Dynamics". *Computers and Society*. <https://doi.org/10.48550/arXiv.2410.03092>
20. Horowitz, Michael, Lauren Kahn, Christian Ruhl, Mary Cummings, Erik Lin-Greenberg, Paul Scharre, et al. 2020. "Policy Roundtable: Artificial Intelligence and International Security". *Texas National Security Review*, 2 de junio. <https://tnsr.org/roundtable/policy-roundtable-artificial-intelligence-and-international-security/>
21. IDAIS (International Dialogues on AI Safety). 2023. "The International Dialogues on AI Safety-Oxford Statement". IDAIS, 31 de octubre. <https://idaais.ai/dialogue/idaais-oxford/>
22. IDAIS (International Dialogues on AI Safety). 2024a. "The International Dialogues on AI Safety-Beijing Statement". IDAIS, 11 de marzo. <https://idaais.ai/dialogue/idaais-beijing/>
23. IDAIS (International Dialogues on AI Safety). 2024b. "The International Dialogues on AI Safety-Venice Statement". IDAIS, 8 de septiembre. <https://idaais.ai/dialogue/idaais-venice/>

24. Leung, Jade. 2019. "Who Will Govern Artificial Intelligence? Learning From the History of Strategic Politics in Emerging Technologies". Disertación doctoral, University of Oxford. <https://ora.ox.ac.uk/objects/uuid:ea3c7cb8-2464-45f1-a47c-c7b568f27665>
25. NSCAI (National Security Commission on Artificial Intelligence). 2019. *National Security Commission on Artificial Intelligence: Interim Report*. November 2019. <https://digital.library.unt.edu/ark:/67531/metadci851191/>
26. Rhinard, Mark, Claudia Morsut, Elisabeth Angell, Simon Neby, Mathilda Englund, Karina Barquet *et al.* 2024. "Understanding Variation in National Climate Change Adaptation: Securitization in Focus". *Environment and Planning C: Politics and Space* 42 (4): 676-696. <https://doi.org/10.1177/23996544231212730>
27. Santander, Pedro. 2011. "Por qué y cómo hacer análisis de discurso". *Cinta de moebio* 41: 207-224. <https://doi.org/10.4067/S0717-554X2011000200006>
28. Sayler, Kelley M. 2020. "Artificial Intelligence and National Security". Congressional Research Service 45178, 10 de noviembre. <https://nsarchive.gwu.edu/document/27080-document-226-congressional-research-service-kelley-m-sayler-artificial-intelligence>
29. Science Media Centre. 2023. "Expert Reaction to a Statement on the Existential Threat of AI Published on the Centre for AI Safety Website". Science Media Centre, 30 de mayo. <https://www.sciencemediacentre.org/expert-reaction-to-a-statement-on-the-existential-threat-of-ai-published-on-the-centre-for-ai-safety-website/>
30. Sears, Nathan Alexander. 2023. "Great Power Rivalry and Macrosecuritization Failure: Why States Fail to 'Securitize' Existential Threats to Humanity". Disertación doctoral, University of Toronto. <http://hdl.handle.net/1807/126870>
31. Trump Donald. 2019. "Executive Order on Maintaining American Leadership in Artificial Intelligence". *The White House*, 11 de febrero. <https://trumpwhitehouse.archives.gov/presidential-actions/executive-order-maintaining-american-leadership-artificial-intelligence/>
32. Trump, Donald. 2020. "Executive Order 13960 of December 3, 2020. Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government". *Federal Register*, 8 de diciembre. <https://www.federalregister.gov/documents/2020/12/08/2020-27065/promoting-the-use-of-trustworthy-artificial-intelligence-in-the-federal-government>
33. U.S. Department of State. 2024. "Secretary Antony J. Blinken at the Advancing Sustainable Development Through Safe, Secure, and Trustworthy AI Event". *U.S. Department of State*, 23 de septiembre. <https://2021-2025.state.gov/secretary-antony-j-blinken-at-the-advancing-sustainable-development-through-safe-secure-and-trustworthy-ai-event/>
34. U.S.-China Economic and Security Review Commission. 2024. *2024 Annual Report to Congress. U.S.-China Economic and Security Review Commission*. <https://www.uscc.gov/annual-report/2024-annual-report-congress>
35. Van-Dijk, Teun A. 2016. "Análisis crítico del discurso". *Revista Austral de Ciencias Sociales* 30: 203-222. <https://doi.org/10.4206/rev.austral.cienc.soc.2016.n30-10>
36. Waever, Ole. Securitization and Desecuritization. Working Papers, vol. 5. Copenhagen: Centre for Peace and Conflict Research, 1993.
37. Wesselink, Anna, Karen S. Buchanan, Yola Georgiadou y Esther Turnhout. 2013. "Technical Knowledge, Discursive Spaces and Politics at the Science-policy Interface". *Environmental Science & Policy* 30: 1-9. <https://doi.org/10.1016/j.envsci.2012.12.008>
38. Williams, Michael C. 2003. "Words, Images, Enemies: Securitization and International Politics". *International Studies Quarterly* 47 (4): 511-531. <https://doi.org/10.1046/j.0020-8833.2003.00277.x>
39. Zeng, Jinghan. 2021. "Securitization of Artificial Intelligence in China". *The Chinese Journal of International Politics* 14 (3): 417-445. <https://doi.org/10.1093/cjip/poab005>

Mónica Ulloa Ruiz

Antropóloga y magíster en Estudios Sociales de la Ciencia por la Universidad Nacional de Colombia. Responsable de Transferencia de Políticas en el Observatorio de Riesgos Catastróficos Globales (ORCG), Baltimore, Estados Unidos. <https://orcid.org/0000-0002-7878-041X> | mulloar@orgc.info

Guillem Bas Graells

Especialista en Gobernanza de Inteligencia Artificial y magíster en Seguridad y Tecnología por la National School of Political and Administrative Studies (SNSPA), Rumania. Coordinador de Inteligencia Artificial en el Observatorio de Riesgos Catastróficos Globales (ORCG), Baltimore, Estados Unidos. <https://orcid.org/0009-0003-3541-2208>

