



PROPUESTA

Un modelo de tres capas

Un aporte a la discusión del
Proyecto de Ley que regula la IA en Chile

Jorge Tuñón S.
Centro de Estudios Democracia y Progreso
Julio 2026



ÍNDICE

Introducción	3
--------------------	---

PRIMER EJE

EL MODELO DE TRES CAPAS: PROHIBIR, PROTEGER, INCENTIVAR

Diagnóstico: El problema de la talla única	4
La posición de Chile en la región: fortalezas y brechas	5
La propuesta: un modelo de tres capas	6
3.1 Capa 1: prohibiciones absolutas	6
3.2 Capa 2: alto riesgo y sectores críticos	6
3.3 Capa 3: el resto del ecosistema: con incentivo a certificarse	10
3.4 La distinción entre Proveedor, Desplegador y Distribuidor	11
Modelos de propósito general: “no ser actor” no implica vacío total	12
Coherencia con el enfoque institucional	14

SEGUNDO EJE

SIMPLIFICACIÓN ADMINISTRATIVA: presunción de diligencia, ventanilla única y sandbox

Diagnóstico: dos instrumentos regulatorios sin arquitectura de ejecución	14
La presunción de diligencia como beneficio legal	15
La Ventanilla Única: registro y administración del sandbox	15
El sandbox como puerta de entrada por defecto	16
Coherencia con los ejes precedentes	16

TERCER EJE

ACTUALIZACIÓN PERIÓDICA DELEGADA

Diagnóstico: la ausencia de un mecanismo de actualización	17
La propuesta: revisión con plazo fijo y facultad de adelantamiento fundado	18
Límites a la delegación: la Administración actualiza dentro de los márgenes que fija la ley	19
Cierre de la arquitectura	19

CONCLUSIONES	20
--------------------	----

Un Modelo de Tres Capas

Jorge Tuñón Subercaseaux¹

Introducción

El Proyecto de Ley que regula los sistemas de inteligencia artificial, resultado de fundir los Boletines N° 15.869-19 y N° 16.821-19 y actualmente en segundo trámite constitucional ante el Senado, llega en un momento en que la regulación de la inteligencia artificial como fenómeno global está en plena revisión. La Unión Europea, cuyo Reglamento de IA constituyó el principal referente del proyecto aprobado por la Cámara, aprobó definitivamente, el 29 de junio de 2026, un paquete de simplificación que posterga las obligaciones de alto riesgo antes de que estas hubieran terminado de entrar en vigor. El estado de Colorado, en Estados Unidos, aprobó, suspendió y reemplazó su propio marco regulatorio dentro del mismo año. El debate ya no es si regular, sino cómo hacerlo sin frenar lo que se busca proteger.

El Centro de Estudios Democracia y Progreso participa en este debate desde una doble convicción: que la inteligencia artificial debe servir a la persona humana y que sus usos más lesivos no admiten tolerancia, y que la competitividad del ecosistema chileno de innovación es también un bien que la regulación debe cuidar, no sacrificar. Desde esa convicción, el Centro no propone rechazar el proyecto ni replicar mecánicamente el modelo europeo que lo inspira. Propone adaptarlo: calibrar su intensidad a la realidad de un país cuyo sector de inteligencia artificial está compuesto mayoritariamente por microempresas desplegadas de modelos de terceros, que lidera la región en talento y políticas públicas, pero exhibe una brecha crítica en adopción industrial, y que carece hoy de la infraestructura supervisora que haría viable el régimen de cumplimiento más exigente.

Este documento desarrolla esa adaptación en tres ejes. El primero propone un modelo de tres capas, prohibir, proteger e incentivar, que concentra las obligaciones más exigentes donde el riesgo de daño a las personas es real y verificable, y sustituye la carga preventiva universal por un incentivo legal afirmativo para el resto del ecosistema. El segundo eje diseña la arquitectura de ejecución: la presunción de diligencia como beneficio concreto de certificar, una Ventanilla Única que no crea burocracia nueva, y un sandbox con contenido real para los casos de clasificación incierta. El tercero asegura que esta arquitectura se actualice al ritmo de la tecnología que regula, mediante un mecanismo de revisión periódica con facultad de adelantamiento, sin crear para ello un órgano permanente adicional.

¹ Ingeniero Comercial, Universidad Diego Portales.

El Centro de estudios agradece la inestimable colaboración del Sr. Emilio Pfeffer Berger, Abogado de la PUC, MBA HEC de París y experto en Inteligencia Artificial.

PRIMER EJE

EL MODELO DE TRES CAPAS: PROHIBIR, PROTEGER, INCENTIVAR

DIAGNÓSTICO: EL PROBLEMA DE LA TALLA ÚNICA

El Proyecto de Ley, aprobado por la Cámara de Diputados y actualmente en tramitación en el Senado, adopta una arquitectura regulatoria basada en cuatro niveles de riesgo —inaceptable, alto, limitado y sin riesgo evidente— tomando como referencia principal el Reglamento (UE) 2024/1689, conocido como el AI Act europeo. Esta decisión tiene un mérito indiscutido: establece un marco estructurado, coherente con los estándares internacionales más avanzados, y consagra prohibiciones explícitas frente a los usos más lesivos de la inteligencia artificial.

Sin embargo, trasladar este modelo al contexto chileno sin las adaptaciones necesarias entraña un riesgo regulatorio que consideramos necesario abordar: el de aplicar estándares diseñados para un mercado de 450 millones de personas, con una institucionalidad supervisora robusta y empresas tecnológicas de gran escala, a un ecosistema radicalmente distinto en tamaño, capacidad y estructura.

Los datos disponibles sobre el ecosistema nacional de IA son elocuentes. Según la primera caracterización de la industria de inteligencia artificial en Chile, elaborada por la Cámara de Comercio de Santiago en julio de 2025, el sector está compuesto por aproximadamente 400 empresas que emplean a poco más de 14.500 personas, con ventas que alcanzaron los US\$ 755 millones en 2024 y una proyección de crecimiento del 33% para 2025. El promedio de trabajadores por empresa es de 25 personas, y el 46% de las compañías opera con entre 1 y 10 colaboradores, tramo en el que se ubica, en consecuencia, la mediana del sector.

Cabe precisar el alcance metodológico de estas cifras. El estudio de la CCS define el universo del sector a partir de las actividades que las empresas declaran como principales: desarrollo de aplicaciones y software de IA, implementación de IA, desarrollo de agentes inteligentes, consultoría sobre IA y desarrollo de modelos propios, entre otras. En consecuencia, la cifra de 400 empresas no incluye a empresas de otros rubros que utilizan IA como herramienta instrumental en sus procesos, sino únicamente a aquellas cuya actividad principal está vinculada al desarrollo o implementación de sistemas de IA. Esta distinción es relevante para la aplicación regulatoria del Proyecto pues alcanza tanto a proveedores como a implementadores de sistemas de IA, por lo que el universo de empresas sujetas a sus obligaciones es significativamente más amplio que las 400 aquí caracterizadas.

Este perfil—donde predominan la microempresas y pequeñas empresas— determina con precisión el problema regulatorio central: las cargas de cumplimiento propias de los sistemas de alto riesgo en el modelo europeo fueron concebidas para organizaciones con departamentos jurídicos, equipos técnicos especializados y capacidad financiera para absorber costos que se derivan de cumplir con estas exigencias. Aplicadas sin modulación al ecosistema chileno, estas exigencias no protegen mejor a las personas: simplemente excluyen a los actores más pequeños del mercado legal, empujándolos hacia la informalidad o inhibiendo directamente su desarrollo.

La Consulta Experta de la Biblioteca del Congreso Nacional, realizada en julio de 2024 sobre el proyecto de ley, recogió esta preocupación con claridad. Especialistas en sistemas de IA señalaron que la clasificación de riesgos, tal como está redactada, es "extremadamente punitiva y no da lugar a la autorregulación", advirtiendo que podría producir el efecto contrario al buscado: que ninguna organización declare utilizar o desarrollar tecnologías de IA para evitar las multas, generando una aplicación vacía de la ley.

LA POSICIÓN DE CHILE EN LA REGIÓN: FORTALEZAS Y BRECHAS

El Índice Latinoamericano de Inteligencia Artificial 2025 (ILIA 2025), elaborado por CENIA y la CEPAL, sitúa a Chile por tercer año consecutivo en el primer lugar regional, con 70,5 puntos sobre 100, destacando por su infraestructura tecnológica, gobernanza de datos, talento humano y desarrollo de políticas públicas. En el indicador de políticas públicas, Chile alcanza 83,2 puntos, y lidera en alfabetización en IA (84,7 puntos) y adopción de IA generativa (78,7 puntos).

No obstante, el mismo informe identifica una brecha crítica que define con precisión el nicho estratégico del país: en adopción industrial de IA, Chile apenas alcanza 39,3 puntos, ubicándose en la posición 12 de 19 países evaluados. Esta paradoja —liderazgo en talento y políticas, rezago en aplicación productiva— revela que el desafío de Chile no es frenar el desarrollo de grandes modelos de frontera, sino acelerar la adopción de IA especializada en los sectores donde el país tiene ventajas competitivas reales.

Chile no es ni puede ser, en el corto plazo, un actor relevante en el desarrollo de modelos de propósito general (GPAI; definido en el artículo 3°, numeral 2°, del texto aprobado por la Cámara como aquel sistema de IA “capaz de realizar múltiples funciones de aplicación general a la vez, tales como el reconocimiento de imágenes o voz, procesamiento de audio, generación de video, detección de patrones, respuesta a preguntas, traducción, entre otros, y que puede proporcionar resultados o salidas tanto previsibles como no previsibles”): no dispone de la infraestructura de cómputo, la masa crítica de datos, ni la inversión en I+D que ese desarrollo requiere. Esta afirmación admite hoy una excepción parcial: en febrero de 2026, el Centro Nacional de Inteligencia Artificial (CENIA), con apoyo del Ministerio de Ciencia, la Corporación Andina de Fomento (CAF) y Amazon Web Services (AWS), lanzó LatamGPT, un modelo de lenguaje de código abierto de más de 70.000 millones de parámetros entrenado con datos de la región. Se trata, con todo, de una iniciativa de investigación, no de una

industria privada de modelos de frontera, por lo que no altera la conclusión de que Chile no cuenta, de manera sostenida, con la escala de cómputo ni el financiamiento privado que ese desarrollo exige. Su ventaja competitiva real está en la aplicación de IA especializada en sectores donde tiene fortalezas. Los datos de Start-Up Chile confirman esta orientación: en la última selección del programa BIG 11 (febrero de 2026), el 45% de los 64 emprendimientos tecnológicos seleccionados se basa en IA, con foco principal en agricultura, recursos naturales y herramientas TI.

Esta posición de Chile, receptor y adoptante de la tecnología, es la base de la arquitectura que se propone a continuación: no tiene sentido regular con intensidad lo que Chile no hace (el entrenamiento de modelos de frontera), pero sí regular con precisión lo que Chile sí hace en volumen: el despliegue de esos modelos en decisiones que afectan a personas.

LA PROPUESTA: un modelo de tres capas

El Centro de Estudios Democracia y Progreso propone al Senado asumir una lógica de "prohibir, proteger, e incentivar", distribuida en tres capas. Este modelo asume las características del ecosistema chileno, que, en coherencia con el enfoque institucional del Centro, según el cual la primacía de la persona y sus derechos fundamentales opera como límite infranqueable, mientras que la intensidad de los mecanismos de cumplimiento se calibra a la capacidad real de los actores regulados y a dónde se concentra el riesgo real de daño.

3.1. Capa 1: prohibiciones absolutas

El Centro respalda íntegramente el listado de usos de riesgo inaceptable del artículo 7° del texto aprobado por la Cámara que incluye la manipulación subliminal, la explotación de vulnerabilidades de personas y grupos, la categorización biométrica discriminatoria, la calificación social genérica, la identificación biométrica remota en tiempo real en espacios públicos, el scraping facial y la evaluación de estados emocionales en entornos laborales y educativos. No se propone ninguna modificación a este artículo. Estas prohibiciones no requieren capacidad supervisora instalada para ser exigibles ni admiten gradualidad: operan como mandatos absolutos cuya infracción constituye una ilicitud en sí misma. A su vez estas prohibiciones, tienen vigencia inmediata desde la publicación de la ley.

3.2. Capa 2: alto riesgo y sectores críticos: el deber probado con una metodología común

Para los sistemas de alto riesgo en sectores críticos (salud, crédito, justicia, empleo, comunicaciones, servicios esenciales y gestión pública), se propone mantener una obligación legal exigible en tres elementos: transparencia frente a la persona afectada por una decisión asistida por IA, supervisión humana efectiva sobre la decisión final, y reporte obligatorio de incidentes graves a la autoridad sectorial.

La incorporación de las comunicaciones como sector de alto riesgo responde a la preocupación del uso de sistemas de IA para la generación o difusión de contenido engañoso (desinformación, contenido sintético no rotulado, publicidad política algorítmicamente dirigida) constituye uno de los riesgos más sensibles del ecosistema informativo, con efectos directos sobre la formación de la opinión pública y los procesos electorales. Este sector queda sujeto a las mismas tres obligaciones mínimas de la Capa 2, transparencia, supervisión humana efectiva y reporte de incidentes, sin perjuicio de que el desarrollo reglamentario de esta ley precise exigencias adicionales de rotulado que complementen la obligación general de identificación del contenido ya establecida en el artículo 5° del texto aprobado por la Cámara.

El reporte obligatorio de incidentes graves a que se refiere este párrafo requiere, además, precisar ante qué autoridad se cumple, cuestión que el texto actual del proyecto no resuelve. El Centro propone que la ley determine que dicho reporte se dirija a la autoridad sectorial correspondiente a cada uno de los sectores de alto riesgo aquí listados, sin que ello implique crear un organismo nuevo: el reglamento de ejecución de esta disposición establecerá los criterios que permitan al propio operador identificar, para cada sector, cuál es la autoridad sectorial ya existente competente para recibir su reporte, entre las que cuenten con competencia legal establecida sobre ese sector. Corresponde al operador, conforme a esos criterios, dirigir el reporte a la autoridad sectorial pertinente, la que lo recibirá y derivará dentro del ámbito de su propia competencia. Esta remisión al reglamento no compromete un tema de fondo, sólo se concreta un aspecto de organización administrativa. La Agencia Nacional de Ciberseguridad se encuentra en condiciones de recibir y derivar los reportes conforme a lo que el propio operador indique, sin que sea necesario desarrollar infraestructura de recepción adicional para este sector en particular.

La función de la autoridad sectorial así identificada se limita a la recepción, coordinación y derivación del reporte dentro del ámbito de su propia competencia, sin que se le atribuya potestad sancionatoria bajo la presente ley. La potestad de fiscalizar el cumplimiento de esta ley y de aplicar las sanciones administrativas que ella contempla permanece radicada, en todos los casos, de manera exclusiva en la Agencia de Protección de Datos Personales, en los mismos términos ya previstos en el texto aprobado por la Cámara. Esta exclusividad no afecta la potestad sancionatoria que, conforme a su propia ley, ejerce la Agencia Nacional de Ciberseguridad respecto de los incidentes que constituyan infracción a la Ley Nº 21.663, por tratarse de una infracción distinta a la que sanciona la presente ley.

Esta regla de exclusividad se ve además reforzada en el propio texto aprobado por la Cámara: el artículo 17º, inciso final, dispone que, cuando unos mismos hechos y fundamentos jurídicos puedan ser sancionados conforme a esta ley y a otra u otras leyes, se aplicará la sanción de mayor gravedad. Esta regla general es la misma que permite la coexistencia, sin doble sanción, entre la potestad de la Agencia de Protección de Datos Personales bajo esta ley y la potestad de la Agencia Nacional de Ciberseguridad bajo la Ley Nº 21.663, y resulta igualmente aplicable a cualquier otra autoridad sectorial que, conforme a su propia legislación, pudiera imponer una sanción por los mismos hechos.

Con todo, el Centro propone que la ley faculte a la Agencia de Protección de Datos Personales para solicitar a la autoridad sectorial correspondiente un informe técnico sobre la dimensión y naturaleza del daño ocasionado por el incidente, previo a resolver el procedimiento sancionatorio. Este informe no es vinculante para la Agencia, que conserva la facultad exclusiva de resolver y determinar la sanción aplicable, pero constituye un antecedente relevante para ponderar los criterios de proporcionalidad que el propio régimen sancionatorio del proyecto ya contempla, entre ellos el número de personas afectadas. La autoridad sectorial deberá evacuar este informe dentro de un plazo que fije el reglamento; transcurrido ese plazo sin respuesta, la Agencia podrá continuar y resolver el procedimiento sancionatorio sin que ello afecte su validez, en el entendido de que la exigencia recae sobre la solicitud del informe y no sobre su efectiva evacuación dentro de plazo.

Se excluyen deliberadamente, respecto del actual artículo 9° del proyecto, la documentación técnica exhaustiva, los sistemas de gestión de riesgos iterativos y la auditoría externa obligatoria del modelo europeo. Esas obligaciones fueron diseñadas para proveedores de gran escala con departamentos jurídicos dedicados, no para los desplegados chilenos que integran modelos de terceros vía API. Cabe notar que el propio artículo 9°, numeral 6, ya contempla una cláusula facultativa de proporcionalidad por tamaño de operador, lo que confirma que el problema no es la ausencia de sensibilidad al tamaño del ecosistema, sino que esa sensibilidad es facultativa ("se podrán establecer").

La exclusión señalada implica sustituir, para los sistemas de alto riesgo así clasificados, los numerales 1, 3 y 4 del artículo 9° del texto aprobado por la Cámara (relativos al sistema de gestión de riesgos iterativo continuo, la documentación técnica y el sistema de registros), reemplazándolos por el estándar de tres elementos aquí propuesto. La referencia a una auditoría externa obligatoria alude, en cambio, a una exigencia propia del Reglamento (UE) 2024/1689 que no tiene correlato en el texto actualmente aprobado por la Cámara.

El artículo 10° del texto aprobado por la Cámara impone, adicionalmente, un sistema de seguimiento posterior a la implementación de los sistemas de alto riesgo, de recolección y análisis continuo de datos durante toda su vida útil. El Centro advierte que esta exigencia puede superponerse, en la práctica, con el reporte de incidentes graves aquí propuesto para la Capa 2, y recomienda que el reglamento de ejecución precise que el cumplimiento del reporte de incidentes conforme a esta Capa 2 satisface, para los desplegados de menor escala, el estándar de seguimiento exigido por el artículo 10°, evitando así una duplicación de cargas sobre un mismo hecho.

Estas exigencias (transparencia, supervisión humana efectiva, reporte de incidentes), enfrentan, sin embargo, el mismo problema de fondo que afectaba al aparato completo del artículo 9°: la ley puede exigir que la supervisión sea "efectiva", pero si no existe una metodología común para probar esa efectividad, la exigencia queda abierta a interpretaciones heterogéneas entre operadores y a una discrecionalidad fiscalizadora sin criterio objetivo.

El Centro recomienda, en cambio, que el reglamento de ejecución de esta disposición adopte como contenido técnico la metodología TEVV —Test, Evaluation, Verification and Validation—, que estructura la evaluación de sistemas de IA en cuatro fases:

- prueba de funcionalidad,
- evaluación de desempeño bajo condiciones representativas,
- verificación de cumplimiento de especificaciones técnicas, y
- validación de que el sistema cumple su propósito real una vez desplegado.

Esta metodología se encuentra ya operacionalizada en el marco AI Verify, desarrollado por la Infocomm Media Development Authority (IMDA) de Singapur, que estructura once principios de gobernanza —entre ellos, equidad, robustez, transparencia en las decisiones y supervisión humana— bajo una arquitectura común de resultados esperados y evidencia técnica, incorporando herramientas de código abierto que permiten el testeado automatizado de sistemas de IA, evitando que Chile deba diseñar una metodología propia desde cero.

Cabe advertir, en todo caso, que a diferencia de un estándar de certificación como ISO/IEC 42001, TEVV no constituye un esquema con un proceso, tiempo y costo de implementación estandarizados a nivel internacional: su aplicación concreta dependerá del diseño específico que adopte el reglamento de ejecución de esta disposición, razón adicional por la cual se recomienda que dicho reglamento sea elaborado con participación de una institución técnica acreditada conforme a lo señalado más adelante en este mismo eje, y no de manera apresurada. Sugerimos que el reglamento debe ser claro para acreditar que la supervisión humana exigida por ley es efectiva y no meramente formal.

El reconocimiento y actualización periódica del contenido técnico de TEVV para estos efectos reglamentarios no debe radicarse, por ley, en una entidad privada nombrada específicamente. El Centro recomienda que la ley habilite al Ministerio de Ciencia, para reconocer, mediante un sistema de acreditación abierto y revisable, a instituciones académicas o técnicas idóneas para cumplir esta función, siguiendo el modelo ya empleado en Chile para organismos certificadores y laboratorios acreditados bajo el Sistema Nacional de Acreditación supervisado por el Instituto Nacional de Normalización.

Centros de investigación con trayectoria reconocida en inteligencia artificial en Chile, como el Centro Nacional de Inteligencia Artificial (CENIA), que participa en la elaboración del Índice Latinoamericano de Inteligencia Artificial junto a la CEPAL, constituyen ejemplos idóneos de la institución que podría cumplir este rol.

Esta concentración de funciones exige resguardos expresos de independencia, dado que la misma institución quedaría a cargo tanto de definir la metodología de evaluación como de recomendar qué sectores entran o salen del listado de alto riesgo. El Centro propone incorporar en la ley, al menos: (i) la prohibición de que quienes participen en la elaboración de la metodología de evaluación intervengan también en su administración o revisión en casos concretos; (ii) una inhabilidad temporal para quienes transiten desde roles de política pública hacia roles de certificación dentro de la misma institución acreditada; y (iii) la publicidad de

los criterios de reconocimiento de las instituciones acreditadas, de modo que el proceso de acreditación mismo sea auditable.

Mientras no se encuentre implementada la capa de acreditación nacional a través de TEVV, el cumplimiento de la norma ISO/IEC 42001 debiera constituir una presunción de diligencia equivalente para efectos de acreditar la supervisión humana efectiva exigida en esta Capa 2. Esta disposición cesará en su aplicación una vez que la institución acreditada conforme al párrafo anterior establezca el contenido técnico definitivo de TEVV.

El listado de sectores de alto riesgo no existe hoy en ninguna parte del proyecto: el artículo 8° solo describe en términos generales qué se entiende por alto riesgo, sin nombrar a la salud, la justicia, el crédito o el empleo, ni a las comunicaciones, como sectores específicos. Crear ese listado directamente en la ley permite que el propio Poder Legislativo que decida, con discusión pública y transparente, en qué áreas de la vida de las personas la inteligencia artificial requiere mayor protección, en lugar de dejar esa decisión a una definición posterior del Gobierno de turno, sujeto a su propia prioridad legislativa.

3.3. Capa 3: el resto del ecosistema: con incentivo a certificarse

Para el universo de sistemas de IA especializados que no califican como alto riesgo — la inmensa mayoría de los aproximadamente 400 operadores caracterizados por la CCS, predominantemente desplegados de modelos de terceros—, el Centro propone una certificación voluntaria bajo estándares reconocidos internacionalmente como la norma ISO/IEC 42001:2023 sobre sistemas de gestión de inteligencia artificial, o el AI Verify Testing Framework ya descrito. Esta Capa 3 corresponde específicamente a los usos sin riesgo evidente definidos en el artículo 6°, numeral 4°, del texto aprobado por la Cámara. Los usos de riesgo limitado, regulados en el artículo 11°, mantienen su propio régimen de transparencia obligatoria ya vigente, sin alteración por esta propuesta.

Para este universo, el centro propone que no se establezca obligación legal de partida que se deba cumplir. El valor jurídico de certificarse voluntariamente debiera constituirse en una “presunción legal de diligencia” por el operador certificado ante un eventual reclamo bajo el régimen de responsabilidad subjetiva del Código Civil chileno, mecanismo que se desarrolla en el Segundo Eje de este documento.

La certificación genera para un usuario o cliente, un estándar reconocido internacionalmente sin que el Estado deba fiscalizar individualmente cada uno de los cientos de sistemas que el mercado produce cada año.

Un punto a considerar, además, es el costo real de esta certificación al evaluar su carácter voluntario. La evidencia de mercado disponible sobre la norma ISO/IEC 42001 indica que, para una organización pequeña que certifica por primera vez, el proceso completo (incluyendo honorarios del organismo certificador, preparación documental y dedicación interna), demanda entre cuatro y doce meses y un costo que, según la fuente y el alcance de la certificación, se ubica entre los US\$ 4.000 y los US\$ 40.000, a lo que se suman auditorías de

mantención anuales y una recertificación completa cada tres años. Para una empresa del tramo de 1 a 10 trabajadores, que concentra al 46% del ecosistema chileno de IA, este costo representa una inversión considerable en relación con su tamaño. Esta evidencia refuerza, con datos de mercado y no solo con el argumento comparado frente al modelo europeo, por qué la certificación en esta capa debe ser voluntaria y generadora de un beneficio, en lugar de ser una condición de acceso al mercado o puerta de salida. Si la certificación fuera la única vía de operar formalmente, su costo funcionaría, en la práctica, como una barrera de entrada para buena parte del ecosistema incentivando la informalidad.

Esta arquitectura evita expresamente el riesgo, advertido por la propia evidencia comparada, de "crear primero una carga para luego eximirla", lo que generaría el incentivo perverso de que la certificación se convierta en una vía de cumplimiento cosmético, máxime cuando no existen aún estándares de acreditación para los propios organismos certificadores en Chile.

3.4. La distinción entre Proveedor, Desplegador y Distribuidor para modular la responsabilidad

Un elemento adicional de modulación, ausente con suficiente precisión en el texto actual del proyecto, consiste en establecer con claridad las obligaciones diferenciadas que corresponden a cada uno de los tres actores centrales de la cadena de valor de la IA: el proveedor, el desplegador —denominado implementador en el texto del proyecto— y el distribuidor.

El distribuidor es un actor de la cadena de suministro comercial cuya función es intermediar entre el proveedor y el mercado, sin tomar decisiones sobre el uso del sistema. Su responsabilidad es esencialmente comercial: no introducir en el mercado sistemas que no cumplan con la ley. El desplegador o implementador es quien determina el caso de uso concreto —el banco que evalúa solicitudes de crédito, el hospital que apoya diagnósticos clínicos—, y sobre él recaen las obligaciones de la Capa 2 cuando corresponda: transparencia, supervisión humana y reporte de incidentes.

El texto aprobado en primer trámite reconoce, además, otras dos categorías de operador el representante autorizado y el importador, cuyo rol específico la ley pondera expresamente al determinar la sanción aplicable. El representante autorizado, mandatario en Chile de un proveedor extranjero para el cumplimiento de las obligaciones de esta ley, y el importador, que introduce en el mercado nacional un sistema de IA bajo el nombre o marca de un tercero domiciliado en el extranjero, son categorías especialmente relevantes para el ecosistema chileno, atendida su posición de receptor y no de productor de tecnología de frontera ya descrita en este mismo eje. El Centro no advierte, en esta etapa, razones para atribuirles obligaciones sustantivas distintas de las que corresponden al distribuidor, esto es, una responsabilidad esencialmente de cumplimiento formal y de no introducir en el mercado sistemas que no satisfagan la ley, pero recomienda que el reglamento precise su posición dentro de la cadena de valor, en particular cuando el representante autorizado o el

importador son, en los hechos, quienes efectivamente controlan la relación comercial con el proveedor extranjero.

Esta distinción es especialmente relevante para el ecosistema chileno, donde la gran mayoría de las empresas son desplegadas de modelos desarrollados por terceros —OpenAI, Google, Anthropic— y no cuentan con control sobre la arquitectura del modelo subyacente. Asignarles obligaciones propias del proveedor del modelo resulta técnicamente imposible e institucionalmente inequitativo.

MODELOS DE PROPÓSITO GENERAL: “NO SER ACTOR” NO IMPLICA VACÍO TOTAL

Chile no debe intentar regular el entrenamiento de modelos de frontera con la misma intensidad que aplican jurisdicciones con una industria privada consolidada de esa naturaleza: salvo por la excepción de LatamGPT ya señalada, el país no cuenta con productores privados de modelos de frontera en escala relevante, ni con la capacidad técnica instalada para fiscalizar umbrales de cómputo como el que utiliza el AI Act europeo para presumir riesgo sistémico. Construir esa infraestructura de fiscalización general sería costoso y tendría, en la práctica, un universo de aplicación muy reducido dentro del territorio nacional.

Con todo, el desarrollo de sistemas de inteligencia artificial de código abierto entrenados con datos de la región, como el señalado en el párrafo anterior, constituye precisamente el tipo de activo que la ley debiera fomentar y no cargar con el mismo régimen aplicable a los grandes proveedores extranjeros de modelos cerrados. El Centro recomienda que el artículo 14° del texto aprobado por la Cámara, que habilita al Ministerio de Ciencia, para solicitar la colaboración de entidades del sector privado y universidades en la promoción de la innovación en inteligencia artificial, incorpore una mención expresa al fomento del desarrollo de sistemas de inteligencia artificial de código abierto entrenados con datos de la región latinoamericana, incluyendo modelos lingüísticos que reflejen la diversidad cultural e idiomática de América Latina y el Caribe.

Existe, sin embargo, un piso de bajo costo que sí tiene sentido exigir: que los proveedores extranjeros de GPAI que operan en Chile entreguen a los usuarios y a la autoridad chilena la misma información de transparencia sobre capacidades y limitaciones que ya producen para cumplir con sus obligaciones de transparencia en otros mercados, particularmente la Unión Europea. El beneficio para Chile es un piso de información sin necesidad de construir burocracia nueva ni de perseguir a actores sobre los que no tiene jurisdicción efectiva de fiscalización técnica.

Este piso de transparencia debe exigirse por igual a todo proveedor de GPAI que opere en Chile, con independencia de si comercializa o no sus modelos en la Unión Europea. Un diseño que se apoye exclusivamente en la información que el proveedor ya divulga para el mercado europeo dejaría sin ninguna obligación de transparencia a los proveedores que no venden en la UE, lo que introduciría una asimetría de trato entre proveedores extranjeros sin

justificación para el usuario chileno, que enfrenta el mismo riesgo de opacidad independientemente del origen del modelo. El Centro propone que la exigencia de transparencia se defina de manera autónoma en la ley chilena de modo que todo proveedor de GPAI que opere en Chile quede sujeto a la misma exigencia, provenga o no de una jurisdicción que ya lo regula.

El piso de transparencia descrito en los párrafos anteriores no debiera, sin embargo, aplicarse de manera idéntica a todo proveedor de GPAI con independencia de su modelo de negocio. El Centro recomienda que el reglamento de ejecución incorpore un criterio de diferenciación análogo al que contempla el Reglamento (UE) 2024/1689 para los modelos de código abierto: los proveedores de GPAI que publiquen abiertamente los pesos, la arquitectura y los datos de entrenamiento de sus sistemas quedan sujetos únicamente al piso de transparencia y documentación ya descrito, sin las obligaciones adicionales de evaluación adversarial que esta propuesta reserva a los modelos cerrados calificados de riesgo sistémico. Esta diferenciación reduce la asimetría regulatoria entre iniciativas como LatamGPT y los grandes proveedores cerrados de modelos de frontera, sin relajar el estándar aplicable a estos últimos.

Sobre los modelos de lenguaje de gran escala (LLM), el Centro adopta la misma posición que el Reglamento (UE) 2024/1689: los LLM se tratan como una subcategoría de los modelos de propósito general, sujeta al piso de transparencia general descrito en este eje, y a exigencias reforzadas únicamente cuando el modelo respectivo sea calificado como de riesgo sistémico conforme a criterios técnicos objetivos. El Centro no recomienda crear, en esta etapa, una categoría normativa separada y más exigente exclusivamente para LLM: hacerlo introduciría una distinción tecnológica difícil de sostener en el tiempo frente a la convergencia de arquitecturas de IA generativa, y se apartaría del criterio, ya adoptado en este documento, de regular según el riesgo de la aplicación y no según la tecnología subyacente.

El artículo 4°, numeral 6°, del texto aprobado por la Cámara impone a todo operador, esto es, a todas las categorías de la cadena de valor: proveedor, implementador, representante autorizado, importador y distribuidor, la obligación de publicar informes anuales sobre el impacto medioambiental de sus sistemas de IA, sin distinguir según el tipo de sistema. El Centro considera poco razonable imponer esta carga de manera indiscriminada a toda la cadena de valor y respecto de cualquier sistema de IA, y propone restringirla a los sistemas de inteligencia artificial de uso general (GPAI) cuya operación efectiva tenga lugar en el territorio nacional, en coherencia con el ámbito de aplicación territorial ya definido en el artículo 2° del texto aprobado por la Cámara.

El Centro recomienda, además, incorporar en el artículo de objeto o principios de la ley la noción de soberanía tecnológica en inteligencia artificial como principio interpretativo, no como obligación exigible. Se trata de un criterio que orienta al reglamento y a la autoridad de aplicación a preferir, entre dos interpretaciones igualmente válidas de las obligaciones de esta ley, aquella que favorezca el desarrollo de proveedores locales o regionales de sistemas de inteligencia artificial, sin imponer mandatos de localización de datos ni preferencias

obligatorias por proveedores chilenos en las compras públicas, que resultarían incompatibles con los compromisos vigentes de Chile bajo el Tratado de Libre Comercio con los Estados Unidos y el Tratado Integral y Progresista de Asociación Transpacífico (CPTPP). La conveniencia de este principio se ve reforzada por episodios recientes en que decisiones regulatorias de un Estado extranjero han afectado, sin aviso previo, la disponibilidad de modelos de propósito general para usuarios fuera de su jurisdicción, incluidos usuarios chilenos, lo que ilustra el valor de contar con alternativas de respaldo desarrolladas o alojadas regionalmente. El fomento a la industria nacional de inteligencia artificial debe canalizarse, en todo caso, por la vía del subsidio y la priorización presupuestaria ya descritas, y no por la vía de la restricción de mercado.

COHERENCIA CON EL ENFOQUE INSTITUCIONAL

El modelo de tres capas aquí propuesto mantiene intacta, en su Capa 1, la primacía de la persona como límite infranqueable. En su Capa 2, conserva los tres elementos que más directamente protegen a la persona afectada por una decisión automatizada en un sector crítico, probados mediante una metodología técnica común que evita tanto la laxitud de una exigencia sin estándar como la sobrecarga de un aparato documental desproporcionado. En su Capa 3, no renuncia a la protección de las personas: la traslada de la prevención administrativa a la reparación civil, en el entendido de que el Estado chileno no cuenta hoy con la capacidad fiscalizadora para sostener un régimen preventivo sobre cientos de operadores de pequeña escala sin desnaturalizar su propio propósito.

SEGUNDO EJE

SIMPLIFICACIÓN ADMINISTRATIVA: presunción de diligencia, ventanilla única y Sandbox

DIAGNÓSTICO: DOS INSTRUMENTOS REGULATORIOS SIN ARQUITECTURA DE EJECUCIÓN

El texto del Proyecto de Ley contempla espacios controlados de pruebas, sandbox, regulados en los artículos 12° y 13°. El artículo 12° establece que los operadores en espacios controlados de pruebas responderán civilmente por los perjuicios causados a terceros, pero quedarán exentos del pago de multas administrativas si respetan la ley y las orientaciones del órgano habilitante. Esta fórmula es razonable: protege a las personas que puedan sufrir daño sin sacrificar el incentivo a innovar bajo supervisión.

El defecto está en lo que rodea esta fórmula: el artículo 13° remite la implementación de estos espacios a "la disponibilidad presupuestaria existente" de los ministerios de Ciencia y de Economía, sin definir qué organismo administra el sandbox, bajo qué criterios se admite a un operador, por cuánto tiempo, ni con qué obligaciones de reporte durante la experimentación. El análisis comparado de marcos regulatorios internacionales es categórico sobre esta omisión: sin estas definiciones, el sandbox es decorativo.

LA PRESUNCIÓN DE DILIGENCIA COMO BENEFICIO LEGAL

Conforme a la arquitectura de tres capas del Primer Eje, el Centro propone que la certificación voluntaria de operadores de la Capa 3 bajo estándares internacionales de gestión de IA (ISO/IEC 42001 u otros equivalentes reconocidos) no opere como exención de una carga administrativa preexistente, sino como generadora de una presunción legal de diligencia frente a responsabilidad por daños derivados del uso de un sistema de IA.

Para que esta presunción sea jurídicamente aplicable, el Centro recomienda que la ley la defina en términos generales así, quien implemente un sistema de gestión de riesgos de inteligencia artificial reconocido por el órgano nacional competente conforme al sistema de acreditación que se propone gozará de una presunción simplemente legal de diligencia. Esta fórmula permite que estándares como ISO/IEC 42001, el AI Verify Testing Framework u otros que se desarrollen en el futuro, sean reconocidos por la institución acreditada sin necesidad de modificar la ley cada vez que el estado de la técnica certificadora cambie. La presunción debe, en todo caso, ser simplemente legal, esto es, admitir prueba en contrario— y no de derecho.

Esta arquitectura tiene una ventaja: no exige que el Estado defina primero una carga de cumplimiento para luego eximirla, lo que habría requerido capacidad fiscalizadora adicional sobre un universo de cientos de operadores de pequeña escala. En su lugar, sitúa el incentivo a certificar directamente en el interés del propio operador de reducir su riesgo, sin necesidad de que la Agencia de Protección de Datos Personales deba auditar individualmente a cada certificado.

LA VENTANILLA ÚNICA: REGISTRO Y ADMINISTRACIÓN DEL SANDBOX

El segundo componente de este eje resuelve el problema de la dispersión institucional que el proyecto distribuye entre el Ministerio de Ciencia, la Agencia de Protección de Datos Personales, la Agencia Nacional de Ciberseguridad, la Secretaría de Gobierno Digital y el Consejo para la Transparencia, sin un punto de coordinación único frente al operador regulado.

El Centro propone que el proyecto establezca, por ley, un sistema de Ventanilla Única que concentre tres funciones específicas. La primera es la orientación a los operadores sobre las obligaciones aplicables según la categoría de riesgo de sus sistemas de IA conforme a la arquitectura de tres capas. La segunda es el registro público de las certificaciones voluntarias

obtenidas bajo estándares reconocidos, de modo que la presunción pueda ser verificada de manera expedita. La tercera, es la administración efectiva de los espacios controlados de pruebas regulados en los artículos 12° y 13° del proyecto, incluyendo la determinación de los criterios objetivos de admisión, la duración máxima de la experimentación y las obligaciones de reporte periódico de los operadores admitidos.

Esta Ventanilla Única no debe constituir un nuevo organismo paralelo. El Centro propone que sus tres funciones se distribuyan entre organismos ya existentes, cada uno dentro de su ámbito de competencia. La orientación a los operadores y el punto único de ingreso de solicitudes se radican en la Secretaría de Gobierno Digital, aprovechando la plataforma que ya opera para otros trámites del Estado (ClaveÚnica, Casilla Única, SIMPLE). El registro público de las certificaciones voluntarias y la acreditación de los organismos certificadores se radican en el Instituto Nacional de Normalización, bajo el mismo Sistema Nacional de Acreditación que ya opera en Chile para organismos certificadores y laboratorios acreditados. La administración efectiva de los espacios controlados de pruebas permanece en el Ministerio de Ciencia, en los mismos términos ya previstos en el artículo 13° del texto aprobado por la Cámara. La Secretaría de Gobierno Digital opera como punto único de entrada frente al operador regulado, derivando internamente hacia el Instituto Nacional de Normalización o el Ministerio de Ciencia según corresponda, de modo que se preserve el principio de ventanilla única sin concentrar en un solo organismo funciones ajenas a su mandato legal. Esta distribución no requiere crear ningún organismo nuevo, ni afecta la potestad exclusiva de fiscalización y sanción que esta propuesta mantiene radicada en la Agencia de Protección de Datos Personales.

Conviene precisar que la Agencia de Protección de Datos Personales no interviene en la administración de la Ventanilla Única. Las funciones que aquí se encomiendan a la Secretaría de Gobierno Digital, al Instituto Nacional de Normalización y al Ministerio de Ciencia son de naturaleza administrativa —plataforma de ingreso, registro de certificaciones, acreditación de organismos certificadores y administración del sandbox, distintas de la fiscalización técnica y de la potestad sancionatoria que esta propuesta mantiene radicada, en todos los casos y de manera exclusiva, en la Agencia de Protección de Datos Personales. La determinación técnica de qué constituye supervisión humana efectiva, y el reconocimiento de la metodología TEVV, corresponden a la institución acreditada descrita en el Primer Eje, no a la Agencia ni a los organismos administradores de la Ventanilla Única. Esta separación de funciones evita que un mismo organismo concentre tareas administrativas ajenas a su mandato junto con su potestad fiscalizadora, y no exige a ninguno de los tres organismos administradores dotación técnica adicional a la que ya poseen para cumplir sus funciones actuales.

EL SANDBOX COMO PUERTA DE ENTRADA POR DEFECTO

El Centro propone ampliar la función del sandbox regulado en los artículos 12° y 13° del proyecto, de modo que todo caso de uso de inteligencia artificial que no encuadre con claridad en ninguna de las tres capas definidas en el Primer Eje pueda operar, mientras se determina su clasificación definitiva, bajo un régimen de exenciones temporales y reporte que

ya contempla el artículo 12°, administrado por el Ministerio de Ciencia, a través de la Ventanilla Única.

Esta ampliación institucionaliza la prudencia regulatoria frente a casos nuevos, en lugar de dejarla a la inacción o a una clasificación apresurada y potencialmente errónea, y sigue el precedente de Singapur y Japón, que han optado deliberadamente por esquemas de gobernanza voluntaria y testeo técnico para no frenar ecosistemas emergentes mientras la autoridad determina el tratamiento regulatorio definitivo de una tecnología nueva.

COHERENCIA CON LOS EJES PRECEDENTES

La presunción de diligencia, la Ventanilla Única y el sandbox por defecto completan, en conjunto, la arquitectura de incentivos de la Capa 3 del Primer Eje: ofrece una razón concreta para certificar voluntariamente, la reducción de riesgos, un mecanismo para que esa certificación sea verificable con el registro público, y una vía para que la incertidumbre regulatoria frente a casos nuevos no se traduzca en parálisis ni en clasificación apresurada, esto es el sandbox por defecto.

Ninguno de estos tres elementos exige al Estado chileno construir una capacidad fiscalizadora nueva y costosa; todos ellos aprovechan capacidad institucional ya existente o trasladan parte de la función protectora al sistema judicial civil, en coherencia con la posición de Chile como adoptante y no productor de la tecnología que se regula.

TERCER EJE ACTUALIZACIÓN PERIÓDICA DELEGADA

DIAGNÓSTICO: LA AUSENCIA DE UN MECANISMO DE ACTUALIZACIÓN ES, EN SÍ MISMA, UN RIESGO REGULATORIO

El texto aprobado por la Cámara no determina en su articulado ningún parámetro técnico cuantitativo para definir categorías de riesgo sistémico, a diferencia del AI Act europeo, que fija un umbral de capacidad de cómputo directamente en su texto. El proyecto chileno tampoco contempla, en su versión actual, ningún órgano ni mecanismo institucional de revisión periódica de los criterios de clasificación de riesgo. Sólo existe una mención a una colaboración (Art. 14), que permite al Ministerio de Ciencia solicitar la colaboración de entidades del sector privado y universidades para promover la innovación, sin que esta facultad esté vinculada a ninguna función de actualización normativa.

Esta ausencia no es, en sí misma, una deficiencia: evita el error que la experiencia europea ilustra con claridad, donde un umbral técnico fijado directamente en el reglamento fue ajustado a instancias de los propios actores regulados, reduciendo el número de modelos

efectivamente sujetos a la norma. Sin embargo, la ausencia de cualquier mecanismo de actualización, deja al proyecto expuesto al riesgo de que las categorías de riesgo se desactualicen frente al ritmo de cambio tecnológico, o que su revisión quede sujeta exclusivamente a los tiempos políticos del Ejecutivo.

La evidencia comparada más reciente confirma la urgencia de este punto. Durante 2025 y 2026, múltiples jurisdicciones han debido modificar sus propios marcos regulatorios de IA en plazos considerablemente más breves que los ciclos legislativos ordinarios: la Unión Europea aprobó definitivamente, el 29 de junio de 2026, un paquete de simplificación que retrasa y ajusta la aplicación de las reglas de alto riesgo de su propio Reglamento, antes de que estas hubieran terminado siquiera de entrar en vigencia, y el estado de Colorado, en Estados Unidos, aprobó, suspendió y reemplazó su ley de IA por una versión más estrecha dentro del mismo año, ante la dificultad de la industria para cumplir con el diseño original. Estos antecedentes respaldan la conclusión de que un mecanismo de actualización con plazo fijo, aunque necesario, puede no ser suficiente por sí solo si no admite ajustes ante cambios que ocurran antes de que se cumpla el plazo regular de revisión.

LA PROPUESTA: REVISIÓN CON PLAZO FIJO Y FACULTAD DE ADELANTAMIENTO FUNDADO

El Centro propone que el proyecto incorpore una cláusula de actualización periódica obligatoria cada 24 meses, mediante la cual un informe técnico evalúe la necesidad de ajustar el listado de sectores de alto riesgo fijado, así como los criterios técnicos de clasificación de riesgo.

El listado de sectores de alto riesgo y los criterios técnicos de clasificación de riesgo quedan fijados, en su versión inicial, directamente en la ley, conforme a lo propuesto en el Primer Eje de este documento. Su actualización posterior, sin embargo, no requiere de una nueva ley: el Centro propone que se habilite al Ministerio de Ciencia, para actualizar, ampliar o reducir tanto el listado como los criterios mediante decreto, precedido del informe técnico señalado. Este mecanismo replica, adaptándolo, el que ya opera en el ordenamiento chileno bajo la Ley N° 21.663, que crea la Agencia Nacional de Ciberseguridad: su artículo 4°, incisos segundo y tercero, fija en la propia ley el listado inicial de servicios esenciales sujetos a esa institucionalidad, y habilita a la Agencia para incorporar nuevos servicios mediante resolución fundada, sin necesidad de una reforma legal posterior. La presente propuesta traslada ese mismo principio - el legislador fija en la ley el criterio y el listado inicial; el órgano competente los actualiza después, dentro de los márgenes que la propia ley define - sustituyendo la resolución fundada de la Agencia Nacional de Ciberseguridad por un decreto del Ministerio de Ciencia, en atención a que la elaboración del informe técnico que antecede a la actualización, al igual que el acto mismo que la dispone, constituyen actos de contenido eminentemente ejecutivo, cuya responsabilidad no resulta delegable en una institución técnica externa a la Administración del Estado.

El decreto respectivo deberá ir precedido, en todo caso, de: (i) el informe técnico elaborado por el Ministerio de Ciencia, que deberá fundar expresamente la necesidad de la

actualización propuesta; (ii) un período de consulta pública de a lo menos treinta días hábiles; y (iii) una exposición de motivos que acredite que los sectores o criterios que se incorporan, modifican o excluyen presentan, para los derechos fundamentales de las personas, un nivel de riesgo análogo al de los sectores y criterios ya contemplados en el artículo 8°. El decreto se someterá, además, a toma de razón de la Contraloría General de la República, conforme a las reglas generales aplicables a los decretos supremos.

Se agrega una facultad de adelantamiento: el Ministerio de Ciencia elaborar el informe técnico y dictar el decreto correspondiente antes de cumplirse los 24 meses del ciclo ordinario, mediante una decisión fundada que acredite que la velocidad de la evolución tecnológica, o de los riesgos observados en la operación de sistemas de IA en Chile, hace inconveniente esperar al plazo ordinario. Esta facultad de adelantamiento no sustituye el plazo máximo de 24 meses, que opera en todo caso como límite superior obligatorio, sino que permite acortarlo cuando las circunstancias lo justifiquen, sin perjuicio del cumplimiento de los resguardos procesales señalados en el párrafo anterior.

La elaboración de este informe técnico corresponde al Ministerio ya mencionado, y es distinta de la función de reconocimiento de la metodología técnica de evaluación TEVV descrita en el Primer Eje, que sigue radicada en la institución acreditada conforme al sistema de acreditación abierta allí propuesto. El Ministerio podrá, para la elaboración de este informe, solicitar antecedentes técnicos a dicha institución acreditada y a otras entidades públicas o privadas con conocimiento especializado, sin que ello sustituya su responsabilidad de suscribir el informe y el decreto correspondiente.

La separación entre ambas funciones responde a la distinta naturaleza de cada una: el reconocimiento de la metodología TEVV es una función de estandarización técnica que no produce, por sí sola, efectos normativos generales sobre el listado de sectores o los criterios de clasificación de riesgo, y por ello puede radicarse en una institución técnica acreditada. La actualización del listado y de los criterios mediante decreto supremo, en cambio, sí produce un efecto normativo general sobre las obligaciones que esta ley impone a los operadores regulados, razón por la cual se reserva a la Administración del Estado.

LÍMITES A LA DELEGACIÓN: LA ADMINISTRACIÓN ACTUALIZA DENTRO DE LOS MÁRGENES QUE FIJA LA LEY

El decreto supremo que actualiza el listado de sectores de alto riesgo o los criterios de clasificación de riesgo no sustituye la decisión legislativa que fija, en el artículo 8°, el concepto mismo de alto riesgo y el listado y los criterios iniciales. Su función se limita a incorporar, ajustar o excluir sectores y criterios dentro de ese mismo concepto de riesgo para los derechos fundamentales ya definido por la ley, sin que pueda, por esta vía, alterar las tres obligaciones mínimas de la Capa 2 (transparencia, supervisión humana efectiva y reporte de incidentes), ni la potestad sancionatoria que esta propuesta mantiene radicada, en todos los casos y de manera exclusiva, en la Agencia de Protección de Datos Personales, materias que permanecen reservadas a la ley. El decreto queda, además, sujeto al control de legalidad que ejerce la

Contraloría General de la República mediante la toma de razón, y a la revisión jurisdiccional que corresponda conforme a las reglas generales sobre actos de la Administración.

CIERRE DE LA ARQUITECTURA

Con este tercer eje se completa la arquitectura que propone el Centro como aporte al Proyecto de Ley. El primer eje, define la arquitectura de tres capas, resuelve quien está sujeto y a qué intensidad de regulación, y aloja en su desarrollo la metodología técnica que acredita la supervisión humana exigida en la capa 2. El Segundo eje define el incentivo legal afirmativo de la Capa 3, su mecanismo de verificación, la ventanilla única y la vía para tratar casos de uso no clasificado mediante el sandbox. Este Tercer Eje, asegura que la arquitectura completa no se fosilice frente al cambio tecnológico, mediante un plazo fijo de revisión con una válvula de adelantamiento fundado para los casos en que la evolución del sector lo exija antes de tiempo, delegando la función de actualización en el Ministerio de Ciencia mediante decreto supremo sujeto a los resguardos procesales descritos, y no en un órgano nuevo y permanente, en plena coherencia con el principio que atraviesa toda la propuesta: regular con firmeza donde el riesgo de daño a las personas es real, y no construir institucionalidad nueva donde el riesgo es bajo o donde el costo de esa institucionalidad recaería desproporcionadamente sobre un ecosistema de escala reducida.

Conclusiones

UNA ARQUITECTURA INTEGRADA

Los ejes desarrollados en este documento proponen una arquitectura regulatoria integrada, en la que cada componente resuelve una pregunta distinta y necesaria para que el conjunto funcione en la práctica. El Primer eje, establece quién está sujeto a qué intensidad de regulación, mediante un modelo de tres capas que prohíbe de manera absoluta los usos incompatibles con la dignidad humana, exige un deber mínimo y técnicamente verificable en los sectores donde el riesgo de daño es real. El Segundo Eje desarrolla el incentivo de la presunción de diligencia derivada de la certificación voluntaria, su mecanismo de verificación mediante una Ventanilla Única que no requiere crear un nuevo organismo, y una vía para tratar con prudencia los casos de uso que aún no encuadran en ninguna categoría definida. El Tercer Eje, finalmente, asegura que esta arquitectura no se fosilice frente al ritmo de cambio de la tecnología que regula, mediante un plazo de revisión periódica que admite adelantarse cuando la evolución del sector lo exija, sin crear para ello un órgano nuevo y permanente.

COHERENCIA CON LOS BORDES DECLARADOS

El enfoque institucional que guía este trabajo establece que la inteligencia artificial debe servir a la persona humana, que la competitividad y la protección de derechos son

objetivos compatibles, y que, en caso de tensión irreductible entre ambos, la primacía corresponde siempre a la persona y sus derechos fundamentales. La arquitectura desarrollada es consistente con este principio en tres planos.

En cuanto a la primacía de la persona, el Primer Eje mantiene íntegras las prohibiciones absolutas ya aprobadas por la Cámara, y conserva en la Capa 2 los tres elementos que más directamente protegen a quien es afectado por una decisión automatizada en un sector crítico —transparencia, supervisión humana efectiva y reporte de incidentes—, probados con una metodología técnica común. En la Capa 3, la protección no desaparece: se traslada de la prevención administrativa a la reparación civil, ante la falta de capacidad del Estado para fiscalizar preventivamente a cientos de operadores de pequeña escala.

Respecto de la viabilidad práctica, cada eje calibra sus exigencias a un ecosistema compuesto mayoritariamente por microempresas: se adopta una metodología de prueba ya disponible en lugar de diseñar una nueva, se delega su reconocimiento en instituciones académicas existentes, y el incentivo de cumplimiento se basa en un beneficio civil que no exige fiscalización estatal extendida sobre todo el universo regulado.

Sobre la certeza jurídica, la arquitectura permite que el operador chileno, especialmente el de menor tamaño, sepa de antemano qué se le exige, cómo probarlo, qué beneficio obtiene al certificar y ante quién dirigirse. Esta previsibilidad no es una concesión a la industria: es la condición misma de que la ley se cumpla, en lugar de ser evadida por inviabilidad práctica, como ya advirtieron los especialistas consultados por la Biblioteca del Congreso Nacional.

VALOR ESTRATÉGICO DE LA PROPUESTA

La arquitectura desarrollada ofrece una alternativa completa, técnicamente fundada en el derecho comparado de múltiples jurisdicciones y empíricamente respaldada en datos verificables sobre el ecosistema chileno de inteligencia artificial y sobre el costo real de los mecanismos de certificación que propone. Este trabajo entrega al Senado un conjunto de soluciones articuladas entre sí, coherentes con los compromisos internacionales que Chile ya ha asumido como miembro de la OCDE, y alineadas con la evidencia más reciente sobre el ritmo de ajuste que incluso las jurisdicciones más avanzadas en regulación de IA han debido adoptar frente a una tecnología en cambio constante.

La arquitectura propuesta permite una adaptación responsable de la experiencia europea en términos de proteger los derechos fundamentales de las personas frente a la inteligencia artificial, sin replicar mecánicamente una arquitectura institucional diseñada para una realidad de mercado, capacidad fiscal e infraestructura técnica que no es la suya.



Democracia Y Progreso

Centro de Estudios
Guillermo Le Fort Varela

www.democraciayprogreso.org

