

Política de IA responsable de un proveedor: cómo debe ser

La política de IA responsable de un proveedor no es un código ético: es un documento operativo que nombra a un responsable, fija la regla de cambio de modelo (ningún cambio entra a producción sin re-medición contra la línea base con el mismo conjunto de evaluación y prueba estadística), publica la fiabilidad con intervalo de confianza y umbral contractual, obliga a la abstención, hace la supervisión humana no desactivable y se revisa cada año. Explicamos cada sección y cómo distinguir una política real de una de plantilla.

Por [David Pinto](#) — Fundador y arquitecto de la plataforma · [LinkedIn](#)

Revisión técnica: Innova y Cree

Publicado 17 ago 2026 · 18 min de lectura · [PDF](#)

La política de IA responsable de un proveedor debe ser un documento operativo, breve y verificable: nombra a un responsable, fija cómo y cuándo puede cambiar el modelo que sostiene el servicio, declara la fiabilidad medida con su intervalo de confianza, hace obligatoria la abstención y la supervisión humana, documenta guardarraíles y pruebas adversarias, y establece su propia revisión. No es una carta de valores. En Innova y Cree la redactamos así porque es lo que un comité de riesgo de una empresa regulada pide leer antes de firmar, y porque cada afirmación se puede contrastar contra el sistema. Este artículo describe sección por sección cómo debe ser esa política, la ancla en NIST AI RMF e ISO/IEC 42001, y cierra con las

señales que separan una política real de una de plantilla. La plantilla vacía se descarga al final; forma parte de la serie sobre [cómo evaluar a un proveedor de IA en Chile](#).

Qué es una política de IA responsable (y qué no es)

Una política de gestión de IA responsable es el documento con el que un proveedor de software con componentes de inteligencia artificial se obliga, ante sus clientes, a un conjunto de procedimientos sobre el ciclo de vida del sistema: selección y cambio de modelos, medición de fiabilidad, supervisión humana, guardarraíles, transparencia y revisión. Su unidad mínima es un procedimiento con responsable, no un principio.

Lo que no es: un código ético con palabras como «equidad», «beneficio» y «confianza» sin un solo verbo operativo. Ese texto puede ser sincero, pero un comprador no puede verificarlo, y por eso no le sirve. La prueba es simple: cada párrafo de la política debe responder a alguna de estas preguntas —quién aprueba, qué se mide, con qué frecuencia, qué se revierte, qué no se puede desactivar—. Si un párrafo no responde a ninguna, sobra. Una política real cabe en dos páginas; la brevedad no es una limitación, es una decisión: el evaluador lee una decena de documentos de respaldo, no cientos de páginas.

AVISO

Las normas citadas aquí son un ejemplo: Innova y Cree trabaja en toda Hispanoamérica y España, y el mismo requisito existe con otro nombre en cada país (México, Colombia, Perú, Argentina, España y la Unión Europea). El cuadro comparado, con la fuente primaria de cada jurisdicción, está en [Cómo evaluar a un proveedor de IA](#).

Alineación con NIST AI RMF e ISO/IEC 42001 sin certificación

El NIST AI Risk Management Framework 1.0 organiza la gestión del riesgo de IA en cuatro funciones: GOVERN (gobernanza y responsables), MAP (contexto y riesgos del caso de uso), MEASURE (medición de fiabilidad y de riesgos) y MANAGE (tratamiento, priorización y respuesta)¹. Es voluntario y no tiene traducción oficial al español; su perfil para IA generativa (NIST AI 600-1) añade doce riesgos específicos, entre ellos la confabulación —el nombre técnico de la alucinación—². ISO/IEC 42001:2023 es el primer sistema de gestión de IA certificable; en Chile existe como NCh-ISO IEC 42001:2024, oficializada por el INN el 29 de noviembre de 2024³⁴.

Una política seria se alinea con ambos y lo dice con precisión: «alineada con las funciones GOVERN, MAP, MEASURE y MANAGE del NIST AI RMF y con los principios de ISO/IEC 42001, en ausencia de certificación formal». Esa última cláusula es la que separa a un proveedor confiable de uno que insinúa un certificado que no tiene. La ausencia declarada de certificación, acompañada de evidencia verificable, es una respuesta legítima; la ambigüedad no lo es.

Gobernanza: un responsable con nombre, no un comité sin cara

La sección de gobernanza debe nombrar al responsable del sistema de IA —con cargo y canal de contacto— y describir qué aprueba: cambios de modelo, publicación de los informes de fiabilidad y respuesta ante el cliente. Debe además describir una función de revisión de seguridad y cumplimiento independiente del desarrollo, encargada de revisar los controles antes de cada entrega. Y debe declarar dónde quedan registradas las decisiones de diseño con impacto en el cliente: bajo control de versiones, con historial íntegro, de modo que cualquier auditoría pueda reconstruir qué se decidió y cuándo.

Lo que no corresponde en esta sección es declarar tamaños de equipo ni organigramas: el comprador no compra personas, compra procedimientos y

evidencia. La continuidad frente a la dependencia de personas clave se responde con arquitectura —formatos abiertos, exportación completa, depósito de código en custodia— y no con una lista de nombres. Un responsable nominado y una función de revisión separada bastan para que el evaluador sepa a quién dirigirse y quién controla a quién.

La regla central: ningún cambio de modelo entra sin re-medición

Ésta es la sección que justifica la existencia del documento. Debe decir, sin matices, que ningún cambio al modelo de lenguaje ni al motor de recuperación de información entra a producción sin re-medición de exactitud contra la línea base vigente, usando el mismo conjunto de evaluación, y que la comparación se valida con una prueba estadística pareada. Si el cambio degrada la exactitud, se revierte. La regla se aplica igual a una actualización del proveedor del modelo, a un cambio de proveedor y a un ajuste del propio motor de búsqueda que alimenta al modelo.

Por qué importa: los proveedores de modelos base retiran versiones y publican otras nuevas con calendario propio. Un producto que «usa el modelo más reciente» cambia de comportamiento sin que el contrato, el precio ni la interfaz cambien. La política obliga a que ese cambio pase por una compuerta medida. En nuestro benchmark público, la mejora entre dos versiones del motor se validó exactamente así: misma batería de preguntas, dos corridas, prueba de McNemar^{[12](#)}.

La compuerta de cambio de modelo

Ningún cambio de modelo ni de motor de recuperación entra a producción sin re-mediación contra la línea base.



Figura 1. La compuerta de cambio de modelo que una política de IA responsable debe describir: el candidato solo entra a producción si supera la re-mediación pareada; en caso contrario, se revierte.

McNemar explicado para un directorio

La prueba de McNemar compara dos versiones de un sistema sobre las mismas preguntas y responde una sola cosa: si la diferencia entre ellas puede explicarse por azar¹⁰. Funciona así. Se corre el mismo conjunto de evaluación con la versión actual y con la candidata, y para cada pregunta se anota si cada versión acertó o falló. Las preguntas que ambas aciertan o ambas fallan no informan sobre la diferencia; las que una acierta y la otra falla, sí. McNemar cuenta cuántas «mejoraron» y cuántas «empeoraron» y calcula la probabilidad de ver ese desbalance si en realidad las dos versiones fueran iguales. Si esa probabilidad (el valor p) es pequeña —por convención, menor de 0,05—, la diferencia es real; si no lo es, el cambio no demostró nada y no debe entrar. La ventaja para un directorio es que la decisión deja de depender de una impresión («se ve mejor») y pasa a depender de un número reproducible sobre las mismas preguntas.

DATO

En la medición pública de nuestro motor del 2026-07-10, la mejora entre versiones sobre las mismas 278 consultas se validó con McNemar pareado, $p = 0,0$. Fuente: innovaycree.com/empresas/benchmark/, consultado el 17 de agosto de 2026.

Versión fijada por configuración y aviso previo al cliente

La regla de cambio necesita dos condiciones para ser verificable. La primera: la versión del modelo base está fijada por configuración, usando el versionado oficial que publica el proveedor del modelo (un identificador de versión con fecha, no un alias comercial que apunta «a la última»). Un alias no es una versión: la configuración puede parecer estable y el modelo debajo cambiar igual. Sobre esta distinción y sobre cómo inventariar cada componente escribimos en [el inventario de componentes de IA \(AI-BOM\)](#).

La segunda: todo cambio, sustitución o deprecación del modelo se notifica previamente al cliente con el plazo pactado en el contrato, junto con el resultado de la re-medición. El aviso no es cortesía: permite al cliente decidir si acepta la nueva línea base, pedir una corrida adicional sobre sus propios casos o exigir la reversión. Una política que fija versión pero no avisa, o que avisa pero no fija versión, deja abierta la puerta al mismo riesgo por el otro lado.

Medición de fiabilidad publicada, con intervalo y umbral

La fiabilidad no se declara: se mide, se publica y se re-mide. La política debe describir el conjunto de evaluación (tamaño, cómo se construyó, quién lo construyó), la rúbrica con la que se califica cada respuesta, la frecuencia de re-medición y el umbral contractual bajo el cual se activa la remediación. Toda cifra va con su

intervalo de confianza y su fecha; una cifra sin intervalo ni fecha no es una medición, es un eslogan.

Las metodologías de referencia son públicas. El estudio de Stanford RegLab sobre herramientas jurídicas con IA definió una rúbrica exigente —una respuesta es alucinación si es incorrecta o si su cita real no sostiene lo afirmado— y separó la abstención como fallo seguro⁹. FACTS Grounding, de Google DeepMind, mide si la respuesta está fundamentada en el documento provisto y es completa⁸. Nuestro benchmark adopta ambas y separa los roles: quien genera las preguntas, el sistema evaluado y el juez son tres modelos distintos, y las categorías trampa se califican por código¹².

DAT0

Nuestra medición pública del 2026-07-10: 84,2 % de respuestas exactas y 2,5 % de alucinación bajo la rúbrica de Stanford, cada una con su intervalo de confianza de Wilson al 95 % publicado en la página del benchmark. Fuente: innovaycree.com/empresas/benchmark/, consultado el 17 de agosto de 2026.

El «cero» no se afirma: se acota

Una política que promete una tasa de error o de alucinación igual a cero falla la lectura de cualquier comité técnico en un minuto, porque ninguna medición finita demuestra un cero. Lo que sí se puede afirmar es una cota: cero fallos en N intentos permite acotar la tasa real con un nivel de confianza dado, y esa cota es la que se publica. La política debe usar ese lenguaje y obligar a usarlo en la interfaz, el material comercial y los informes.

Esto tiene una consecuencia práctica para el comprador: si un proveedor declara cero, la pregunta correcta no es «¿de verdad?» sino «¿en cuántos intentos, con qué

rúbrica y calificado por quién?». Si no hay respuesta, la afirmación no vale. Si la hay, se puede convertir en una cota y compararla con el umbral contractual.

DATO

En nuestra batería hostil pública (120 pruebas diseñadas para inducir la invención de normas y artículos, medición del 2026-07-10) se registraron 0 fabricaciones; la lectura correcta no es «cero», sino la cota superior con 95 % de confianza que publica la página. Fuente: innovaycree.com/empresas/benchmark/, consultado el 17 de agosto de 2026.

Abstención explícita: la salida que un sistema serio debe tener

La política debe establecer que, ante confianza insuficiente o ausencia de fundamento en el corpus autorizado, el sistema declara que no puede responder, en lugar de arriesgar una respuesta incorrecta. La abstención es una salida esperada y medida, no un defecto: la rúbrica de Stanford la trata como fallo seguro, y un buen benchmark mide dos cosas a la vez, la abstención innecesaria (el sistema calla cuando podía responder) y la fabricación (el sistema responde cuando debía callar)⁹.

Para el comprador, la señal es doble. Primero, que la política nombre la abstención y describa cómo se muestra al usuario. Segundo, que la medición publicada reporte la tasa de abstención junto con la de exactitud y la de alucinación; un sistema que nunca se abstiene está declarando, sin decirlo, que responde siempre, y ése es justamente el comportamiento que produce citas inventadas.

Supervisión humana por diseño, no por configuración

La política debe decir que la IA no toma decisiones: genera borradores y análisis con citas verificables que una persona revisa, corrige y aprueba; ninguna salida produce

efectos sin esa validación, y la bitácora registra quién validó y cuándo. Y debe añadir la frase que un comité busca: esta regla es de diseño y no es desactivable por configuración.

La diferencia es arquitectónica. «Configurable» significa que existe una casilla, un permiso o un parámetro que la apaga, y que por tanto un administrador apurado, un integrador o un incidente pueden apagarla. «Por diseño» significa que el flujo validar-y-registrar está en el código de la operación y no hay ruta alternativa. Bajo la Ley 21.719 —vigencia prevista para el 1 de diciembre de 2026, con una postergación en discusión legislativa al cierre de este artículo—, el titular de datos tiene derecho a no ser objeto de decisiones basadas únicamente en tratamiento automatizado con efectos jurídicos o significativos, con derecho a información, explicación, intervención humana y revisión¹¹. Un flujo con validación humana registrada es la forma más directa de que el cliente pueda demostrar que cumple.

AVISO

Este artículo describe buenas prácticas técnicas y de gobierno; no constituye asesoría legal. La aplicación de la Ley 21.719 —vigencia, excepciones del art. 8.º bis y obligaciones del responsable y del encargado— debe evaluarse con asesoría jurídica para cada caso de uso. Texto consultado en bcn.cl el 17 de agosto de 2026.

Guardarraíles: alcance restringido y verificación de citas

Los guardarraíles son las restricciones que el sistema aplica antes y después del modelo. La política debe listar los que operan en producción, sin revelar su implementación: alcance temático restringido al dominio contratado (el sistema rechaza consultas fuera de dominio); tratamiento de todo documento ingerido como dato y nunca como instrucción (la defensa estructural contra la inyección indirecta que OWASP lista como primer riesgo de las aplicaciones con modelos de

lenguaje)⁵⁶⁷; verificación automática de citas, contrastando cada referencia normativa contra el corpus antes de mostrarla; y ausencia de herramientas de ejecución, de modo que el sistema lee y redacta pero no escribe en sistemas del cliente, no envía ni borra.

Ese último punto merece una línea propia en la política, porque cierra de golpe la sección de «agencia excesiva» del cuestionario del comprador: si el sistema no puede actuar, una inyección exitosa no tiene nada que abusar y una interrupción no deja acciones a medias. En despliegues dedicados, los guardarraíles se ajustan al marco interno del cliente y esa adaptación queda documentada.

Red teaming repetido, no una vez

La política debe declarar que el sistema se somete a pruebas adversarias documentadas —manipulación de instrucciones, inducción a fabricar normas, elusión de restricciones, premisas falsas embebidas— y, sobre todo, que esas pruebas se repiten en cada corrida de medición, no una vez antes del lanzamiento. Un red teaming único describe el sistema del día en que se hizo; una batería hostil que corre en cada re-medición describe el sistema que el cliente usa hoy.

Conviene distinguir aquí dos cosas que los cuestionarios de compradores regulados separan con razón: las pruebas adversarias propias del proveedor, que la política documenta y repite, y la prueba de intrusión independiente sobre la capa de IA, que solo un tercero puede firmar. Ofrecer más de lo primero para cubrir lo segundo indica que no se entendió la diferencia. La política puede y debe declarar lo primero; sobre lo segundo, lo correcto es decir qué existe, qué no y con qué compromiso.

Transparencia: contenido identificado, confianza y citas verificables

La sección de transparencia debe comprometer tres cosas visibles en la interfaz: que el contenido generado por IA está identificado de forma expresa, que se muestra el nivel de confianza de cada respuesta y que las citas son verificables (enlazan o reproducen la fuente para que la persona que valida pueda contrastarla). Debe además comprometer que la cadena de modelos usada se declara al cliente y que sus cambios se notifican, lo que enlaza esta política con el inventario de componentes que la acompaña.

Los tres elementos tienen una razón operativa, no cosmética. La identificación del contenido generado permite al validador aplicar el nivel de escrutinio correcto. El nivel de confianza le dice dónde mirar primero. La cita verificable convierte la validación en un acto de un minuto en lugar de una investigación. Una política que promete «transparencia» sin estas tres piezas concretas promete un adjetivo.

Revisión anual y control de versiones de la propia política

La política debe llevar número de versión, fecha, responsable de aprobación y una regla de revisión: al menos una vez al año, y ante cambios relevantes del marco normativo aplicable —en Chile, la Ley 21.719 con vigencia prevista para diciembre de 2026, la Ley 21.663 de ciberseguridad y el proyecto de ley de IA en trámite— o del sistema. Cada versión conserva un registro de cambios legible: qué se modificó y por qué.

Esto no es formalismo. Una política sin fecha ni versión no permite saber si describe el sistema actual o el de hace dos años. Y una política que no menciona su propia revisión suele ser la que nadie ha vuelto a leer desde que se escribió. El mismo control de versiones que la política exige para las decisiones de diseño debe aplicarse a ella misma.

Cómo leerla como comprador: 8 señales de una política real vs una de plantilla

Cuando la política de IA responsable de un proveedor llega al escritorio del evaluador, estas ocho señales distinguen un documento operativo de uno decorativo.

#	SEÑAL DE UNA POLÍTICA REAL	LO QUE DELATA UNA DE PLANTILLA
1	Nombra a un responsable con cargo y correo, y una función de revisión separada del desarrollo.	«Un comité multidisciplinario vela por...» sin nombres ni canal.
2	Contiene una regla explícita de cambio de modelo: re-medición con el mismo conjunto, prueba estadística, reversión si degrada.	No menciona qué pasa cuando el proveedor del modelo publica una versión nueva.
3	Declara la versión del modelo fijada por configuración y el aviso previo al cliente con plazo.	«Usamos siempre la última tecnología disponible».
4	Publica cifras con intervalo de confianza, tamaño de muestra, fecha y umbral contractual.	Porcentajes redondos sin intervalo, sin fecha o con «cero».
5	Nombra la abstención como salida esperada y la mide.	El sistema «siempre responde».
6	Dice que la supervisión humana es de diseño y no desactivable, con bitácora de quién y cuándo.	«Recomendamos supervisión humana» o «configurable según el cliente».
7	Lista guardarraíles concretos y declara que el red teaming se repite en cada medición.	«Aplicamos las mejores prácticas de seguridad».
8	Se alinea con NIST AI RMF e ISO/IEC 42001 y declara sin ambigüedad si hay o no certificación.	Logos o menciones que sugieren certificación sin decirlo.

Ocho señales para leer la política de IA responsable de un proveedor como comprador.

Una nota final para el comprador: la política es la primera hoja del expediente, no la última. Su valor está en que cada afirmación remite a evidencia que se puede pedir —el informe de la última medición, la bitácora, el inventario de componentes, el registro de cambios— y en que las promesas que hoy son compromisos aparezcan como tales, con fecha, y no como hechos. Cuando la política se lee junto con [el inventario de componentes de IA](#) y con la [medición pública de fiabilidad](#), el evaluador tiene lo que necesita para la parte de IA de su cuestionario.

PLANTILLA DESCARGABLE

[Política de gestión de IA responsable — plantilla \(v1.0, Markdown\)](#)

Licencia CC BY-ND 4.0 · [descargar](#)

La plantilla reproduce las secciones descritas —objeto y alcance, gobernanza, regla de cambio de modelo, medición, abstención, supervisión humana, guardarraíles y pruebas adversarias, transparencia, revisión— con campos entre corchetes y sin cifras. Está pensada para que un proveedor la complete con datos verificables y para que un comprador la use como lista de contraste. Todas nuestras plantillas descargables se reúnen en [/recursos/plantillas/](#). Se publica bajo licencia CC BY-ND 4.0: se puede redistribuir citando la fuente, sin obras derivadas.

1. NIST, *AI Risk Management Framework 1.0* (AI 100-1), 26 de enero de 2023. Consultado el 17 de agosto de 2026. [↵](#)
2. NIST, *AI 600-1: Generative AI Profile*, julio de 2024. Consultado el 17 de agosto de 2026. [↵](#)
3. ISO/IEC 42001:2023, *Information technology — Artificial intelligence — Management system*. Ficha consultada el 17 de agosto de 2026. [↵](#)

4. INN, *NCh-ISO IEC 42001:2024*, oficializada el 29 de noviembre de 2024. Consultado el 17 de agosto de 2026. [↩](#)
5. OWASP, *Top 10 for LLM Applications 2025* (con traducción al español). Consultado el 17 de agosto de 2026. [↩](#)
6. OWASP, *Top 10 2025 de riesgos y mitigaciones para LLMs y aplicaciones de IA generativa*, traducción al español publicada el 12 de marzo de 2025. Consultado el 17 de agosto de 2026. [↩](#)
7. OWASP GenAI Security Project, *LLM Top 10 — 2026*, publicado en agosto de 2026. Consultado el 17 de agosto de 2026. [↩](#)
8. Google DeepMind, *FACTS Grounding*, 17 de diciembre de 2024. Consultado el 17 de agosto de 2026. [↩](#)
9. Magesh, Surani, Dahl, Suzgun, Manning y Ho, *Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools*, Stanford RegLab / Yale, publicado en *Journal of Empirical Legal Studies* (2025). Consultado el 17 de agosto de 2026. [↩↩](#)
10. McNemar, Q., *Note on the sampling error of the difference between correlated proportions or percentages*, *Psychometrika* 12 (1947). Consultado el 17 de agosto de 2026. [↩](#)
11. Ley 21.719, Biblioteca del Congreso Nacional de Chile. Consultado el 17 de agosto de 2026. [↩](#)
12. Innova y Cree, *Benchmark anti-alucinación*, medición del 2026-07-10 con re-mediación trimestral. Consultado el 17 de agosto de 2026. [↩↩](#)

Puntos clave

- 01 Una política de IA responsable es un documento operativo con responsable nominado, reglas de cambio, medición y revisión; no una declaración de valores.
- 02 La regla central: ningún cambio de modelo entra a producción sin re-mediación contra la línea base con el mismo conjunto de evaluación y una prueba estadística pareada; versión fijada por configuración y aviso previo al cliente.
- 03 La fiabilidad se publica con intervalo de confianza, fecha y umbral contractual; la abstención ante confianza insuficiente es una salida válida y esperada.

- 04 La supervisión humana debe estar en la arquitectura, con registro de quién validó y cuándo, y no ser desactivable por configuración.
- 05 Como comprador, hay ocho señales para distinguir una política real de una de plantilla; la primera es que nombre a un responsable con cargo y correo.

Preguntas frecuentes

¿Una política de IA responsable es lo mismo que un código de ética de IA?

¿Puede un proveedor sin certificación ISO/IEC 42001 tener una política válida?

¿Por qué la regla de cambio de modelo es la sección más importante?

¿Qué significa que la supervisión humana sea «por diseño y no por configuración»?

¿Es válido que la política prometa una tasa de error de cero?

Referencias

[1] NIST AI Risk Management Framework 1.0 (AI 100-1). <https://www.nist.gov/itl/ai-risk-management-framework>. Consultado el 17 de agosto de 2026. NORMA

[2] NIST AI 600-1 — Artificial Intelligence Risk Management Framework: Generative AI Profile. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>. Consultado el 17 de agosto de 2026.

NORMA

[3] ISO/IEC 42001:2023 — Information technology — Artificial intelligence — Management system. <https://www.iso.org/standard/42001>. Consultado el 17 de agosto de 2026. NORMA

[4] NCh-ISO IEC 42001:2024 (INN, versión chilena oficializada el 29 de noviembre de 2024). <https://ecommerce.inn.cl/nch-iso-iec-42001202489243>. Consultado el 17 de agosto de 2026.

NORMA

[5] OWASP Top 10 for LLM Applications (edición 2025 con traducción al español; edición 2026 vigente). <https://genai.owasp.org/llm-top-10/>. Consultado el 17 de agosto de 2026. NORMA

[6] OWASP — Top 10 2025 de riesgos y mitigaciones para LLMs y aplicaciones de IA generativa (traducción oficial al español, PDF). <https://genai.owasp.org/resource/top-10-2025-de-riesgos-y-mitigaciones-para-llms-y-aplicaciones-de-ia-generativa/>. Consultado el 17 de agosto de 2026.

NORMA

[7] OWASP GenAI LLM Top 10 — 2026. <https://genai.owasp.org/resource/owasp-genai-llm-top-10-2026/>. Consultado el 17 de agosto de 2026. NORMA

[8] FACTS Grounding: a new benchmark for evaluating the factuality of large language models (Google DeepMind). <https://deepmind.google/discover/blog/facts-grounding-a-new-benchmark-for-evaluating-the-factuality-of-large-language-models/>. Consultado el 17 de agosto de 2026.

PAPER

[9] Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools (Stanford RegLab / Yale; J. Empirical Legal Studies, 2025). <https://arxiv.org/abs/2405.20362>. Consultado el 17 de agosto de 2026. PAPER

[10] McNemar, Q. (1947). Note on the sampling error of the difference between correlated proportions or percentages. Psychometrika 12, 153–157. <https://doi.org/10.1007/BF02295996>. Consultado el 17 de agosto de 2026. PAPER

[11] Ley 21.719 (Chile) — regula la protección y el tratamiento de los datos personales y crea la Agencia de Protección de Datos Personales; art. 8.º bis sobre decisiones automatizadas. <https://www.bcn.cl/leychile/navegar?idNorma=1209272>. Consultado el 17 de agosto de 2026.

LEY

[12] Benchmark anti-alucinación de Innova y Cree — metodología, resultados e intervalos publicados. <https://innovaycree.com/empresas/benchmark/>. Consultado el 17 de agosto de 2026.

PROPIO

Citar este artículo

APA

Pinto, D. (2026, 17 de agosto). Política de IA responsable de un proveedor: cómo debe ser (v1.0). Blog técnico de Innova y Cree.
<https://innovaycree.com/blog/gobierno-ia/politica-ia-responsable-proveedor.html>

BIBTEX

```
@misc{pinto2026politica,  
  author      = {Pinto, David},  
  title       = {Política de IA responsable de un proveedor: cómo debe ser},  
  howpublished = {Blog técnico de Innova y Cree},  
  year        = {2026},  
  month       = {ago},  
  note        = {v1.0, revisado 2026-08-17},  
  url         = {https://innovaycree.com/blog/gobierno-ia/politica-ia-responsable-proveedor.html}  
}
```

v1.0 · 17 ago 2026

► HISTORIAL DE VERSIONES

Descargables de este artículo

- [Política de gestión de IA responsable — plantilla \(v1.0, Markdown\)](#) CC BY-ND 4.0
- [Política de gestión de IA responsable — plantilla \(v1.0, DOCX\)](#) CC BY-ND 4.0

SOBRE EL AUTOR

[David Pinto](#) · Fundador y arquitecto de la plataforma

Diseña la arquitectura de IA verificable de Innova y Cree: RAG con cita verificada contra la fuente oficial, benchmark anti-alucinación público y plataformas en producción para sectores regulados de habla hispana.

[LinkedIn](#) · [Todos sus artículos](#)