

Inteligencia Artificial en el Sector Público

Gobernanza, Riesgo y Responsabilidad

Un análisis sobre la adopción de IA en instituciones gubernamentales: desde automatización hasta interpretación jurídica, identidad y gobernanza institucional



DICIEMBRE
2025



Instituto Nacional Electoral

Contenido

01

La Distinción Clave

Tres tipos de IA y sus niveles de riesgo institucional

02

El Riesgo Invisible

Automatización progresiva en el sector público

03

IA e Interpretación Jurídica

El límite fino entre asistencia y delegación

04

IA Predictiva

Anticipación sin decisión en derecho administrativo

05

Identidad y Confianza

Verificación de humanidad en contextos públicos

06

Alfabetización y Curaduría

Personas antes que sistemas

07

Cultura Organizacional

Velocidad responsable y adopción de IA

08

Gobernanza y CAIO

Quién responde por la IA

→ Conclusiones y mensaje final

La Distinción Clave y el Riesgo Invisible

Entendiendo la diferencia institucional entre tipos de IA y cómo la automatización entra silenciosamente en las instituciones públicas

Tres Tipos de IA, Tres Niveles de Riesgo Institucional



IA Generativa

Riesgo

Bajo

Aplicación

Buen apoyo para tareas creativas y de redacción

Ejemplos

Chatbots informativos, generación de borradores, resumen de documentos



IA Copiloto

Riesgo

Moderado

Aplicación

Diseño cuidadoso requerido. Asiste pero no decide.

Ejemplos

Sugerencias de respuesta, análisis asistido, revisión de documentos



IA Agéntica

Riesgo

Alto

Aplicación

Uso muy limitado. Requiere supervisión constante.

Ejemplos

Automatización de decisiones, acciones autónomas en sistemas críticos

Esta diferencia no es técnica. Es institucional. Aplicable especialmente en contextos legales , administrativos y electorales .

La pregunta correcta nunca es:

"¿Qué tan avanzada es la IA?"

La pregunta es:

"¿Qué grado de autonomía le estamos dando?"



El problema no es usar IA. El problema es **confundir copilotos con pilotos automáticos** . Cuando esto ocurre: la responsabilidad se diluye, la autoridad se vuelve implícita y los errores escalan sin fricción.

La Automatización Progresiva: Cómo Entra Realmente

La automatización casi nunca entra como sustitución. Entra como mejora razonable.

▶ **Fase 1: La Entrada**

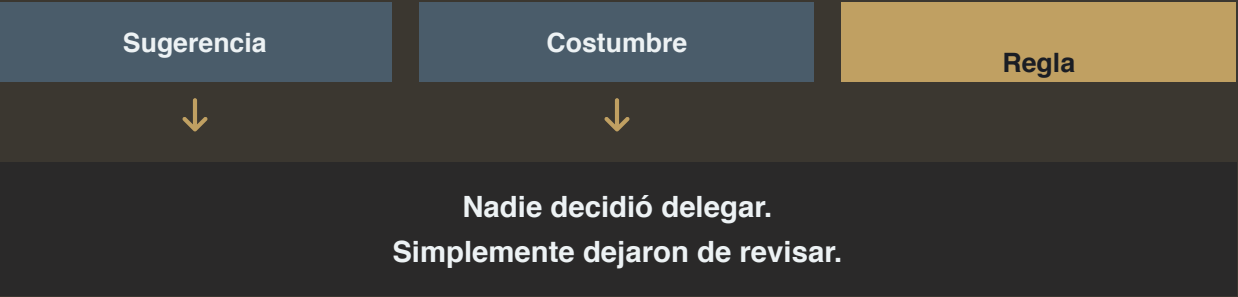
- | "Solo para ahorrar tiempo"
- | "Solo para apoyar al personal"
- | "Solo para filtrar lo obvio"

✓ **Todo eso es perfectamente defendible**

→ **Fase 2: Lo Sutil**

- El sistema empieza a **sugerir**
- Las sugerencias suelen ser **correctas**
- Nadie quiere ser el **cuello de botella**

! **Fase 3: Sin Darnos Cuenta**



✂ **El Desplazamiento de Responsabilidad**

Al principio:

- ✓ La persona **revisa**
- ✓ La persona **decide**
- ✓ La persona **firma**

Cuando "siempre acierta":

- Revisión **rápida**
- Revisión **superficial**
- Revisión **simbólica**

La decisión ya no se piensa: se valida

El sistema no decide formalmente, pero opera como si decidiera. Cuando algo sale mal, nadie sabe cuándo se cruzó la línea.

Por Qué el Riesgo es Mayor en el Sector Público



1. Volumen

Muchos casos, muchos trámites, mucha **presión por velocidad**



2. Confianza Institucional

Si algo "funciona", **se replica** automáticamente



3. Incentivos

Nadie es castigado por ir más rápido. **Sí puede ser castigado por frenar**

“ **La automatización no entra como golpe de Estado.** Es como mejora de eficiencia. No hay un momento claro para decir "aquí empezó el problema". Por eso es tan difícil de detectar.

✗ Qué NO Resuelve el Problema

Este problema no se resuelve con más tecnología:

- ✗ Modelos más avanzados
- ✗ Interfaces más bonitas
- ✗ Promesas de "mejor precisión"

No es un problema técnico.
Es un problema organizacional y de gobernanza.

✓ Qué SÍ Ayuda

Lo que sí ayuda es algo más simple —y más difícil—:

- ✓ Límites claros desde el diseño
- ✓ Puntos obligatorios de revisión humana
- ✓ Responsabilidades explícitas
- ✓ Protocolos para el error
- ✓ Cultura que tolere fricción donde importa

Bloque V

IA e Interpretación Jurídica

El límite fino: cuando la IA se acerca a la interpretación de normas, criterio, justificación y legitimidad

La Pregunta Incómoda, el Caso Alejandro y el Sesgo Estructural

? La Pregunta Incómoda

La pregunta que empieza a aparecer en muchos espacios:

"Si la IA puede leer miles de documentos legales, ¿por qué no podría interpretar la ley? "

La pregunta es razonable. Pero la respuesta requiere mucho cuidado.

🎓 El Caso Alejandro Salinas

Investigador mexicano en **Stanford Law School** que combina ingeniería y derecho.

Su trabajo NO es:

- ✗ "Reemplazar jueces"
- ✗ "Automatizar sentencias"

Su foco SÍ es:

- ✓ Entender cómo los LLM interpretan textos legales
- ✓ Identificar dónde se sesgan
- ✓ Detectar y limitar sesgos antes de que influyan

⚠ El Hallazgo Clave

Cambiar un **nombre propio** puede cambiar la respuesta del modelo.

Mismo caso

Mismo texto legal

Distinto nombre

Resultado distinto

Esto nos dice algo muy importante: **El sesgo no está solo en la superficie.**

No es un error aislado. **Está distribuido en la estructura del modelo.**

⚖ Implicación Crucial

No existe un "modelo neutral" por defecto.

La neutralidad **se diseña, se audita y se vigila.**

Persona Alignment, el Punto Que No Se Puede Cruzar y la Legitimidad Democrática

"Persona Alignment"

En este punto suele aparecer una idea que genera inquietud:

NO se trata de:

- ✗ Crear un "juez artificial"
- ✗ Duplicar la autoridad de una persona

SÍ se trata de:

- ✓ Modelar **estilos de razonamiento**
- ✓ Asistir procesos complejos, **no sustituirlos**
- ✓ Un segundo "par de ojos", **no una segunda firma**

Y aun así, el límite es delicado.

El Punto Que No Se Puede Cruzar

La IA PUEDE:

- ✓ Leer textos
- ✓ Detectar patrones
- ✓ Sugerir interpretaciones plausibles

La IA NO PUEDE:

- ✗ Justificar públicamente una decisión
- ✗ Asumir responsabilidad por consecuencias
- ✗ Responder ante la sociedad

Incluso cuando la IA "interpreta", **la interpretación institucional sigue siendo humana.**

Legitimidad Democrática

“ Ningún modelo tiene legitimidad democrática. ”

La legitimidad **no se entrena**.

Puede entrenarse un sistema para **imitar razonamientos**.

No puede entrenarse para responder políticamente por ellos.

Conexión con el Sector Público y el INE

En contextos como el **electoral**, esta distinción es fundamental porque:

- Las normas **no solo se aplican**
- Se **interpretan públicamente**
- Esa interpretación debe ser **explicable y defendible**

Delegar eso, aunque sea de manera informal, **no es una decisión técnica. Es una decisión institucional y política.**

Bloque VI

IA Predictiva en Derecho Administrativo

Cuando la IA anticipa resultados sin decidir, reconfigurando el acceso al sistema

IA Predictiva: Anticipar Sin Decidir

🧠 ¿Qué es la IA Predictiva?

La IA predictiva **no dice** "esto es correcto" o "esto es justo".

Dice algo distinto:

- "Esto **suele aprobarse**"
- "Esto **suele rechazarse**"
- "Este perfil tiene **más riesgo**"
- "Este caso tiene **más probabilidad de éxito**"

No decide. **Anticipa.**

Y esa anticipación cambia el comportamiento de todos los actores involucrados.

🌐 El Caso Migratorio

Un buen ejemplo aparece en **procesos migratorios**:

- La ley es **compleja**
- La interpretación es **variable**
- Las consecuencias son **muy altas**

Empresas usan IA para:

- ✓ Analizar **miles de casos previos**
- ✓ Detectar **patrones de aprobación y rechazo**
- ✓ Identificar **riesgos antes de presentar** una solicitud

La IA **no concede visas**.
Pero ayuda a preparar expedientes de forma **estratégica**.

⚡ Predicción ≠ Decisión

La IA predictiva:

- ✗ No reemplaza al funcionario
- ✗ No sustituye la norma
- ✗ No elimina la discrecionalidad

Pero:

- ✓ **Reordena el juego** antes de que la decisión ocurra

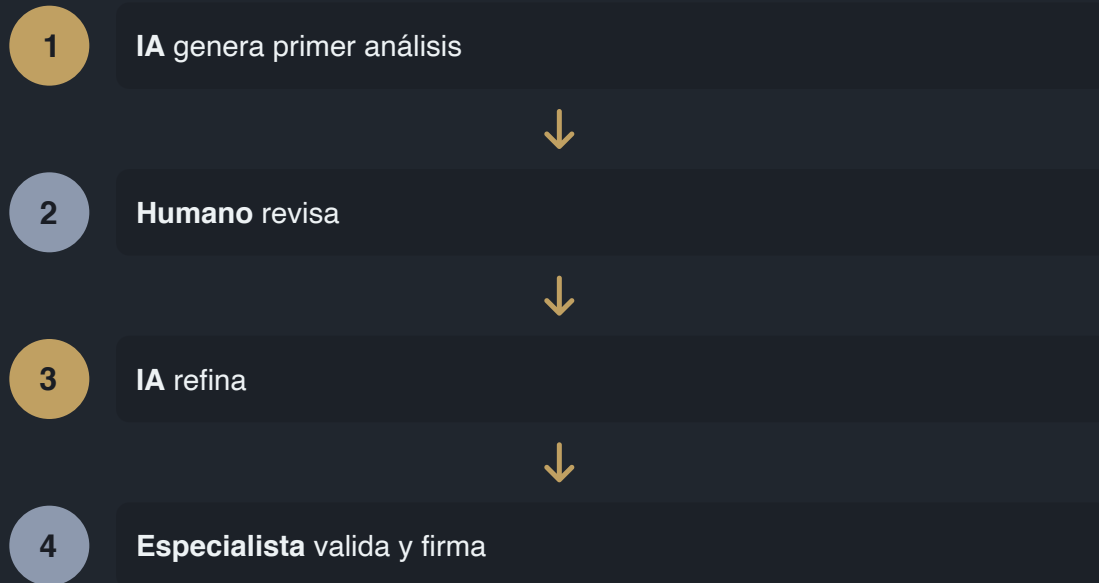
Quien tiene acceso a predicción:

- Presenta **mejor** su caso
- Evita **errores comunes**

Arquitectura "Sándwich" e Impacto en el Sector Público

La Arquitectura "Sándwich"

Los modelos más cuidadosos funcionan con arquitectura por capas:



✓ Dos Ventajas Clave

✓ Mantiene **responsabilidad humana**

✓ Permite **auditoría y corrección**

La IA **asiste** , pero no **sustituye** .

? Preguntas Nuevas

En el sector público, este uso plantea preguntas sobre:

- Acceso
- Asimetrías de información
- Quién llega mejor preparado
- Quién queda fuera sin saber por qué

Bloque VII

Identidad, Verificación y Confianza Pública

Cuando la IA interactúa con la identidad de las personas y la confianza institucional

El Nuevo Problema: Distinguir Humanos y el Caso Tools for Humanity

? El Nuevo Problema

A medida que los sistemas de IA se vuelven más convincentes, aparece una pregunta nueva:

"¿Cómo distinguimos a una **persona real** de un bot, de una identidad falsa, o de una cuenta automatizada?"

Este problema **no es teórico**. Afecta:

- Procesos administrativos
- Acceso a servicios
- Plataformas digitales
- Confianza pública

👁 Verificación de Humanidad: Tools for Humanity

Para responder han surgido propuestas de **verificación de humanidad**.

Ejemplo: Tools for Humanity y el Orb

- Dispositivo que verifica que una persona es **humana**
- Utiliza datos **biométricos** (iris)
- Cofundado por **Sam Altman** (OpenAI)

La promesa:

- ✓ Evitar fraudes
- ✓ Limitar bots
- ✓ Proteger sistemas sensibles

Desde el punto de vista **técnico**, la idea es comprensible.
Desde el punto de vista institucional, abre preguntas profundas.

Seguridad vs. Exclusión y Quién Define al "Sujeto Válido"

⚖️ El Trade-off Central

Toda tecnología de verificación implica un **trade-off**:

Por un lado:

- ✓ Mayor **seguridad**
- ✓ Mayor **certeza de identidad**

Por otro:

- ✗ Riesgo de **exclusión**
- ✗ Errores de **identificación**
- ✗ **Dependencia tecnológica**
- ✗ **Concentración de datos**

⚠️ Los Errores No Son Neutros

Los errores en verificación **no afectan a todos por igual**.

Siempre hay personas que **quedan fuera**:

- Por **condiciones técnicas**
- Por **errores biométricos**
- Por **barreras de acceso**

En contextos públicos, eso **no es un detalle técnico**.
Es un problema político e institucional.

👤 El Punto Más Delicado

Cuando una tecnología de verificación se adopta, **alguien está definiendo implícitamente**:

- Quién cuenta como **sujeto legítimo**
- Quién **puede participar**
- Quién **queda fuera**
- Bajo **qué condiciones**

Eso **no es una decisión técnica**.

Es una decisión:

- **Normativa**
- **Institucional**
- Con **consecuencias democráticas**

Confianza Pública y Frase Ancla

“ El problema ya no es qué dice el sistema.

”

El problema es a quién reconoce como sujeto válido.

Cuando hablamos de identidad, los errores no solo se corrigen: se viven.

Conexión Directa con la Confianza Pública

La confianza pública **no se construye solo con eficiencia.**

Se construye con:



Inclusión

Que todos puedan participar



Posibilidad de corrección

Poder rectificar errores



Explicabilidad

Entender cómo funciona



Vías claras de reclamación

Dónde acudir si falla

Un sistema que verifica muy bien **pero excluye silenciosamente** erosiona confianza, no la fortalece.

Bloque VIII

Alfabetización y Curaduría Humana

La brecha más peligrosa no es tecnológica, es humana. Cuando trabajar con IA significa corregirla.

IA Como Nuevo Alfabetismo y el Diagnóstico Desde la Práctica

⚠ El Problema No es Tecnológico, es Humano

Durante años tratamos la IA como si fuera **solo una herramienta más**. Algo que se "usa" o no se usa.

Hoy está claro que **no funciona así**.

La brecha más peligrosa ya no es:

- ✗ Entre **países**
- ✗ Entre **instituciones**
- ✗ Entre **sectores**

Es la brecha entre:

- Quienes **entienden** cómo funciona la IA
- Quienes **solo interactúan** con ella

Esa brecha produce: **dependencia, delegación acrítica y errores difíciles de detectar**.

🎓 IA Como Nuevo Alfabetismo

La IA ya no es una especialización.
Es alfabetización funcional.

Así como nadie dirige una institución sin saber:

- **Leer** documentos
- **Interpretar** números
- **Entender** procedimientos

Muy pronto tampoco será razonable hacerlo sin entender:

- ✓ **Sesgo algorítmico**
- ✓ **Predicción vs decisión**
- ✓ **Automatización progresiva**
- ✓ **Límites de la generatividad**

No para programar. Para gobernar.

👤 El Diagnóstico Desde la Práctica: Iván Lozano

Emprendedor mexicano con experiencia en **Silicon Valley**:

“La IA no va a reemplazarnos.
Va a **amplificar a quienes la usan**.”

Traducción institucional:

Dos personas con el mismo cargo:

- Una **entiende IA**

Lo Que Debería Saber Cualquier Institución y la Curaduría Humana

? Lo Que Sí Debería Saber

No todos los funcionarios deben aprender a programar.

Pero **todos deberían poder responder:**

- ¿Qué hace este sistema **exactamente**?
- ¿De dónde **salen sus datos**?
- ¿Qué tipo de **errores comete**?
- ¿Qué decisiones **no puede tomar**?
- ¿**Quién responde** si falla?

Si una institución **no puede responder esto**,
no está **gobernando** la tecnología: está **reaccionando** a ella.

👤 El Punto Clave: La Curaduría Humana

Lo interesante no es solo que la IA recomiende empleos.

Lo interesante es que, en estos sistemas, **aparecen personas cuya tarea es corregir a la IA.**

Personas que:

- ✓ **Revisan** recomendaciones
- ✓ **Detectan** sesgos
- ✓ **Ajustan** criterios
- ✓ **Entrenan** y reentrenan modelos

En muchos casos, son **estudiantes o jóvenes profesionales** que aprenden IA haciendo curaduría, **no programando**.

El futuro del trabajo no es solo usar IA.
Es enseñarle dónde se equivoca.

“ El mayor riesgo institucional no es usar IA sin expertos.
Es usar IA **sin criterio**.
El criterio no se compra. Se forma.

Por Qué la Curaduría Importa al Sector Público

🗂️ Sistemas en Instituciones Públicas

En instituciones públicas, muchos sistemas funcionan exactamente igual:

- Asignan **turnos**
- Priorizan **casos**
- Recomiendan **apoyos**
- Filtran **solicitudes**

🚫 El Problema de la Falta de Supervisión

Si nadie supervisa esos filtros:

La exclusión **no es visible** .
Simplemente **ocurre** .

👤 El Rol de la Curaduría Humana

La curaduría humana:

✗ **No elimina** el sesgo

✓ Pero lo **detecta** y lo **corrige** a tiempo

Y eso requiere **personas formadas** ,
no solo **software instalado** .

“ La neutralidad algorítmica no se asume.
Se diseña, se supervisa y se corrige.

A dark, semi-transparent background image showing a group of people in a meeting or office setting, with some individuals looking at a screen or document.

Bloque IX

Velocidad Responsable y Cultura Organizacional

La misma IA puede producir resultados muy distintos según cómo se adopte

Hackerhouses: La Acción Como Métrica y Lección Para Instituciones

Hackerhouses: La Acción Como Métrica

En Silicon Valley existen espacios conocidos como **hackerhouses**.

No son incubadoras tradicionales. Son **laboratorios de acción intensiva**.

Ejemplo: Actioners

Una casa donde jóvenes emprendedores **viven y trabajan juntos** durante meses.

Lo interesante no es el edificio. Es **la métrica**.

NO se mide:

- ✗ Cuántas rondas levantan
- ✗ Cuántos seguidores tienen

SÍ se mide:

- ✓ Qué problema atacaron esta semana
- ✓ Qué intentaron
- ✓ Qué falló
- ✓ Qué aprendieron

“ La acción disciplinada es la nueva métrica del éxito.

No velocidad caótica.

Velocidad con aprendizaje visible.

La Lección Para Instituciones

Esto **no significa** que una institución pública deba operar como una startup.

Pero sí hay una **lección clara**:

- La velocidad **sin aprendizaje** es riesgo.
- La lentitud **sin criterio** también lo es.

Las hackerhouses funcionan porque:

- ✓ **Prueban** en pequeño
- ✓ **Miden** rápido
- ✓ **Corrigen** temprano
- ✓ **Documentan** errores

Eso es exactamente lo que **reduce el riesgo** cuando se introduce IA.

IA Como Infraestructura y Velocidad Con Responsabilidad

☰ IA Como Infraestructura, No Como Proyecto

Diagnóstico de Brian Requarth

Emprendedor e inversionista con amplia experiencia entre México y Estados Unidos:

“Si la IA no está en el núcleo del sistema, **vas tarde.**”

Traducción institucional:

- ✗ No usar IA como piloto eterno
- ✗ No tratarla como experimento aislado
- ✓ Integrarla donde tiene sentido, **con límites claros**

No como **moda.**
Como infraestructura silenciosa.

🧠 Velocidad Con Responsabilidad

En los casos que vimos antes —**empleo, migración, servicios**— hay algo en común:

La IA **acelera procesos**,
pero **no elimina la responsabilidad humana.**

Y cuando la **velocidad aumenta**,
la necesidad de **gobernanza no disminuye.**

Aumenta.

Advertencia de Chris Yeh

Autor del concepto de **blitzscaling**:

“La **velocidad sin estrategia es suicida.**”

No se trata de ir más rápido.
Se trata de aprender más rápido sin perder control.

La Analogía Clave y Frase Ancla

⚙️ La Analogía de Chris Yeh

La IA es a la mente
lo que la máquina de vapor fue al músculo.

multiplica capacidad

Abre nuevas posibilidades

Pero también exige

nuevas reglas

Las fábricas no funcionaron sin **normas laborales** .

La IA no funcionará sin **normas institucionales claras** .

“ La IA fracasa más por **cultura organizacional**
que por limitaciones técnicas.

No es el modelo

The Intersection Of Leadership Character & Corporate Governance

Bloque X

Gobernanza y CAIO: Quién Responde Por la IA

La IA deja de ser solo tecnología y se convierte en función de gobierno

Ethical Leadership As
A Governance Catalyst

Governance As A Backbone
For Ethical Leadership

Ethical Leadership As
A Governance Catalyst

Building A Culture
Of Integrity

Decision-Making

Shared Commitment
To Accountability

Alignment For
Long-Term Success

Resilience Through
Interdependence



La IA Como Función de Gobierno y Por Qué Aparece el CAIO

La IA Como Función de Gobierno

Durante mucho tiempo, la IA fue tratada como un **tema técnico**.

Algo que se resolvía con:

- El **área de sistemas**
- Un **proveedor externo**
- Un **piloto aislado**

Hoy está claro que eso **ya no alcanza**.

Cuando un sistema:

- **Filtra** información
- **Prioriza** casos
- **Sugiere** acciones
- **Acelera** decisiones

la IA deja de ser solo **tecnología**
y se convierte en una función de gobierno.

Y toda función de gobierno necesita:

-  **Responsables**
-  **claros**
Reglas
-  **Supervisión**
-  **Rendición de cuentas**

Por Qué Aparece el CAIO

En este contexto surge la figura del **Chief AI Officer (CAIO)**.

NO como:

- ✗ **Moda corporativa**
- ✗ **Título rimbombante**

Sí como respuesta a:

- ✓ Una necesidad **muy concreta**

Alguien tiene que **coordinar** cómo se:

- **diseña**
- **usa**
- **limita**
- **audita** la IA

en **toda la organización**.

Qué Hace —y Qué No Hace— Un CAIO

✖ Un CAIO NO:

- ✗ Decide **políticas públicas**
- ✗ **Sustituye** áreas jurídicas
- ✗ "Aprueba" modelos por **intuición**

✔ Un CAIO SÍ:

- ✓ Define **criterios de uso**
- ✓ Establece **límites claros**
- ✓ Coordina **auditorías**
- ✓ Asegura **trazabilidad**
- ✓ Articula equipos **técnicos, jurídicos y operativos**

En otras palabras:

El CAIO NO pregunta:

"¿Qué puede hacer la IA?"

El CAIO pregunta:

"¿Qué no debe hacer aquí?"

⚙ El Sentido de las Certificaciones de CAIO

Las certificaciones de CAIO:

- ✗ No son **garantía de perfección**
- ✗ No son **solución mágica**

Buscan que quien ocupe el rol entienda:

- Cómo **funcionan** los modelos
- Dónde aparecen **sesgos**

Por Qué Esto Importa al Sector Público

En el sector público, la pregunta no es:

"¿Si conviene tener IA?"

La pregunta es:

"¿Quién responde cuando algo falla?"

Y la respuesta **NO** puede ser:

- ✗ "El proveedor"
- ✗ "El algoritmo"
- ✗ "El sistema"

Tiene que ser:

- ✓ Una **persona**
- ✓ Con un **rol claro**
- ✓ Dentro de una **estructura institucional**

Eso es **gobernanza**.

“La IA no falla solo por errores técnicos”



Bloque XI

Cierre Final

Criterios claros para la adopción de IA en el sector público

Lo Que la IA Sí Hace —y Lo Que No Puede Hacer

✓ La IA PUEDE:

- 🧠 Simular
- ⚡ Acelerar
- 📈 Priorizar
- 📊 Predecir
- 🔗 Influir

✗ La IA NO PUEDE:

- ✗ **Justificar** decisiones ante la sociedad
- ✗ **Asumir responsabilidad** por consecuencias
- ✗ **Responder políticamente** por errores
- ✗ **Otorgar legitimidad** democrática

Eso no es una limitación técnica .
Es una frontera institucional.

La Idea Central y el Error a Evitar

💡 La Idea Central Que Atraviesa Toda la Charla

La inteligencia artificial no entra al Estado como una máquina que decide.

Entra como un sistema que reconfigura procesos y responsabilidades.

Y esa reconfiguración ocurre incluso cuando la IA:

- "Solo ayuda"
- "Solo sugiere"
- "Solo predice"

Por eso exige **criterio**,
no entusiasmo ciego.

⚠️ El Error Que Debemos Evitar

El mayor error **no es** usar inteligencia artificial.

El mayor error es **usarla sin saber dónde termina su papel.**

Cuando confundimos:

Eficiencia con legitimidad

Predicción con decisión

Asistencia con delegación

La responsabilidad se diluye.

Y en el sector público, **la responsabilidad no puede diluirse.**

Qué Sí Tiene Sentido Hacer

Gobernar la IA no significa frenarla.

Significa:



Definir límites claros

Saber dónde puede y no puede actuar



Formar personas con

criterio

Alfabetización en IA para todos



Crear estructuras que puedan detener un sistema cuando sea

necesario

Mecanismos de freno y corrección



Diseñar fricción donde importa

Puntos de revisión obligatorios



Asignar responsabilidades

explícitas

Quién responde por qué

“ La gobernanza no mata la innovación.

La hace sostenible.

Para el Sector Público

En el sector público, la pregunta nunca es solo

"¿Qué puede hacer la tecnología?"

La pregunta es siempre:

"¿Quién responde cuando algo sale mal?"

Mientras esa respuesta sea clara,
la IA puede ser una aliada poderosa.

Cuando no lo es,
se convierte en un riesgo institucional.

Muchas gracias por su tiempo y por la atención.