



UNMSM
Fundación
SAN MARCOS
Por el desarrollo de la Ciencia y la Cultura



DIPLOMADO DE ESPECIALIZACIÓN

GESTIÓN DE LA CIBERSEGURIDAD Y PRIVACIDAD



RESUMEN EJECUTIVO

¿Quién le teme a la inteligencia?
Posibilidades y riesgos de la inteligencia artificial en el Estado digital

JULIO – OCTUBRE 2026



CENCAEP
CAPACITA

RESUMEN EJECUTIVO

¿Quién le teme a la inteligencia?

Posibilidades y riesgos de la inteligencia artificial en el Estado digital

Marcelo Cabrol y Roberto Sánchez A.
Banco Interamericano de Desarrollo (BID) · 2021

1. Introducción

La inteligencia artificial (IA) se ha convertido en uno de los principales motores de la transformación digital contemporánea. Su capacidad para procesar grandes cantidades de información, identificar patrones, automatizar actividades y generar predicciones abre importantes posibilidades para mejorar la gestión pública. Al mismo tiempo, su utilización plantea riesgos relacionados con la privacidad, la seguridad de la información, la discriminación algorítmica, la vigilancia masiva, la concentración de datos y la protección de los derechos fundamentales.

En ¿Quién le teme a la inteligencia? Posibilidades y riesgos de la inteligencia artificial en el Estado digital, Marcelo Cabrol y Roberto Sánchez analizan esta transformación desde la perspectiva de América Latina y el Caribe. Su planteamiento central es que la discusión sobre la IA no debería reducirse a decidir si esta tecnología es positiva o negativa. La cuestión verdaderamente importante es cómo incorporarla al Estado de manera responsable, segura, transparente y compatible con los derechos de las personas.

Los autores rechazan, por un lado, la visión de una inteligencia artificial autónoma que inevitablemente terminará sustituyendo o dominando al ser humano. Pero también cuestionan el extremo opuesto: asumir que cualquier innovación tecnológica constituye por sí misma una solución. La IA no es moralmente autónoma. Sus resultados dependen de los datos utilizados, de los objetivos establecidos y de las decisiones adoptadas por quienes diseñan, entrenan, implementan y supervisan los sistemas.

Desde esta perspectiva, el documento propone avanzar hacia un Estado de Bienestar Digital Inteligente (EBDI): un Estado capaz de utilizar datos y herramientas predictivas para anticiparse a los problemas sociales y mejorar sus decisiones, pero sometido simultáneamente a límites técnicos, éticos y jurídicos.

2. De un Estado reactivo a un Estado predictivo

Una de las ideas centrales de la publicación es que la transformación tecnológica debe analizarse como parte de un proceso histórico más amplio. Las responsabilidades del Estado han evolucionado conforme las sociedades enfrentaron nuevas crisis económicas, sanitarias, sociales y tecnológicas.

El Estado inicialmente reactivo, característico de las primeras etapas de industrialización, evolucionó progresivamente hacia modelos de bienestar y protección de derechos. En la actualidad, fenómenos como el cambio climático, la violencia, las crisis sanitarias, la automatización y la transformación del mercado laboral plantean un nuevo desafío: desarrollar instituciones capaces no solo de reaccionar frente a los problemas, sino también de anticiparlos.

La inteligencia artificial puede contribuir precisamente a esa transición. Los modelos predictivos permiten analizar información histórica y datos generados prácticamente en tiempo real para identificar riesgos, reconocer patrones y apoyar decisiones públicas antes de que determinadas situaciones alcancen niveles críticos.

Esto no significa delegar automáticamente las decisiones gubernamentales en algoritmos. La propuesta consiste en emplear la IA como instrumento de toma y soporte de decisiones, manteniendo la responsabilidad humana sobre las consecuencias de esas decisiones.

El concepto de EBDI resume esta transformación. Se trataría de un Estado digitalizado, interconectado y con capacidad analítica suficiente para desarrollar políticas públicas preventivas y adaptativas. Sin embargo, para que este modelo funcione se necesita infraestructura tecnológica, información confiable, interoperabilidad entre sistemas, capacidad institucional, personal especializado y un marco normativo que establezca límites claros al tratamiento de los datos.

La experiencia de la pandemia de COVID-19 permitió observar con especial claridad tanto las posibilidades como las limitaciones de esta transformación. El uso de datos en tiempo real, sistemas de rastreo, algoritmos predictivos, telemedicina y automatización demostró que la tecnología podía incrementar considerablemente la capacidad de respuesta del Estado. Al mismo tiempo, puso en evidencia los riesgos de utilizar tecnologías de vigilancia sin controles suficientes y reveló las importantes brechas de digitalización existentes entre países y grupos sociales.

3. Inteligencia artificial al servicio de las políticas públicas

La publicación presenta la IA, en términos funcionales, como software diseñado para realizar determinadas tareas consideradas inteligentes. Dentro de este concepto se incluyen sistemas basados en reglas, aprendizaje automático (machine learning) y aprendizaje profundo (deep learning), entre otras aproximaciones.

Para la administración pública, su importancia radica principalmente en la posibilidad de procesar información a una velocidad y escala difíciles de alcanzar exclusivamente mediante capacidades humanas. La IA puede utilizarse para automatizar procesos administrativos, clasificar información, detectar anomalías, estimar riesgos, apoyar diagnósticos, mejorar la focalización de programas sociales y proporcionar información para la toma de decisiones.

En este contexto, la automatización no debe entenderse únicamente como sustitución del trabajo humano. Su mayor potencial se encuentra en complementar las capacidades profesionales, reduciendo tareas repetitivas y permitiendo que los servidores públicos concentren sus esfuerzos en actividades que requieren análisis, criterio, interacción humana y responsabilidad institucional.

La obra analiza especialmente la utilización de estas tecnologías durante la emergencia sanitaria. Los sistemas digitales permitieron identificar patrones de propagación, procesar imágenes médicas, apoyar diagnósticos, realizar seguimiento epidemiológico y acelerar determinadas actividades de investigación científica.

También se describen experiencias en América Latina vinculadas con salud, educación, administración pública y prestación de servicios. Estas iniciativas demuestran que la región no parte de cero. Existen capacidades y proyectos innovadores que pueden servir como referencia para una transformación más amplia.

Sin embargo, los autores insisten en una condición esencial: la tecnología por sí sola no resuelve problemas institucionales. Un algoritmo avanzado aplicado sobre registros deficientes, bases de datos fragmentadas o instituciones sin capacidad de respuesta puede generar resultados poco útiles e incluso perjudiciales.

4. América Latina frente al desafío de la transformación digital

Uno de los principales obstáculos para el desarrollo de la inteligencia artificial en América Latina y el Caribe es la persistencia de importantes brechas de digitalización.

Los sistemas de IA necesitan datos para funcionar. Pero disponer de grandes cantidades de información no es suficiente. Los datos deben tener calidad, estar actualizados, estructurados adecuadamente y, cuando corresponda, ser interoperables entre instituciones.

En numerosos países latinoamericanos todavía existen procedimientos predominantemente presenciales, registros administrativos fragmentados, sistemas públicos que no intercambian información y niveles desiguales de conectividad.

Estas limitaciones quedaron especialmente expuestas durante la pandemia. La ausencia o insuficiencia de registros electrónicos de salud, padrones de beneficiarios y censos educativos dificultó incluso la implementación de soluciones digitales básicas. A ello se suman factores estructurales como la informalidad, la desigualdad y la brecha digital.

En consecuencia, la estrategia de IA de un Estado no debería comenzar simplemente adquiriendo algoritmos sofisticados. Primero es necesario fortalecer las bases de la transformación digital: infraestructura, interoperabilidad, identidad digital, gestión de datos, ciberseguridad, capacidades institucionales y formación de talento humano.

El riesgo de avanzar sin resolver estas condiciones es generar una nueva forma de desigualdad. Los grupos que ya tienen dificultades para acceder a servicios públicos podrían quedar nuevamente excluidos si los servicios migran al entorno digital sin garantizar conectividad, accesibilidad y competencias digitales.

5. El dato como activo y como superficie de riesgo

Desde una perspectiva de ciberseguridad, uno de los aportes más relevantes del documento es mostrar que el potencial de la inteligencia artificial está directamente relacionado con la disponibilidad de información, pero que esa misma dependencia crea riesgos significativos.

Los Estados administran grandes cantidades de datos laborales, educativos, financieros, sanitarios, tributarios y sociales. Cuando estas fuentes se integran, permiten identificar relaciones y patrones que serían prácticamente imposibles de detectar mediante procedimientos tradicionales.

Esta capacidad puede mejorar la focalización de políticas públicas y facilitar modelos predictivos. Sin embargo, la concentración de información también aumenta el impacto potencial de un acceso indebido, una filtración, un uso no autorizado o una vulneración de los sistemas que la almacenan.

Por ello, la protección de datos no debe considerarse un componente secundario de una estrategia de inteligencia artificial. Constituye una condición estructural para su implementación responsable.

Desde la óptica de la seguridad de la información, la transformación digital del Estado exige proteger la confidencialidad, integridad y disponibilidad de los datos y sistemas utilizados. La calidad de una decisión algorítmica también depende de que los datos permanezcan íntegros y no hayan sido manipulados.

El documento recuerda además las debilidades de ciberseguridad existentes en la región y la necesidad de fortalecer capacidades de respuesta frente a incidentes. La creciente digitalización amplía inevitablemente la superficie tecnológica expuesta a amenazas.

En consecuencia, un Estado que aumenta su dependencia de datos, servicios digitales e inteligencia artificial también debe incrementar proporcionalmente sus capacidades de prevención, detección, respuesta y recuperación frente a incidentes de seguridad.

6. Sesgo algorítmico y calidad de los datos

Otro riesgo central analizado en la obra es el sesgo algorítmico. Los sistemas de IA aprenden o construyen sus resultados a partir de datos, reglas y decisiones humanas. Por ello, no operan en un vacío social.

Si los datos utilizados reflejan desigualdades históricas, discriminación o representaciones incompletas de determinados grupos, el algoritmo puede reproducir esos patrones. Del mismo modo, las decisiones adoptadas durante el diseño, selección de variables, entrenamiento y evaluación del modelo pueden introducir nuevos sesgos.

Esto adquiere especial relevancia cuando la IA participa en decisiones relacionadas con derechos o servicios públicos. Un error en un sistema comercial puede afectar una recomendación; un error en un sistema público podría influir en el acceso a una prestación social, una evaluación de riesgo, una decisión sanitaria o una actuación de seguridad.

Por ello, los autores plantean que la mitigación de sesgos debe considerarse durante todo el ciclo de vida del sistema: desde la conceptualización y recopilación de datos hasta el entrenamiento, evaluación, implementación y monitoreo posterior.

Un sistema técnicamente preciso tampoco puede considerarse automáticamente justo. La precisión general puede ocultar tasas de error significativamente mayores para determinados grupos. La evaluación debe

incorporar criterios de equidad y analizar cómo se distribuyen los errores y beneficios entre diferentes segmentos de la población.

7. Explicabilidad, transparencia y rendición de cuentas

La expansión de algoritmos complejos introduce otro problema: la dificultad para comprender cómo se producen determinadas decisiones.

Algunos modelos funcionan como verdaderas “cajas negras”. Pueden generar resultados estadísticamente precisos sin que resulte sencillo explicar de manera comprensible por qué se produjo una predicción concreta.

En el sector público esta situación plantea un problema jurídico e institucional. Cuando una decisión afecta derechos, beneficios o servicios, el ciudadano debería poder conocer sus fundamentos y, cuando corresponda, cuestionarla.

La publicación diferencia en este contexto dos conceptos importantes. La explicabilidad se refiere a la capacidad de comunicar las razones que llevaron al sistema a producir determinado resultado. La interpretabilidad, por su parte, apunta a hacer comprensible el funcionamiento del algoritmo y sus resultados para las personas.

Por esta razón, transparencia, explicabilidad y auditabilidad constituyen requisitos esenciales de una IA responsable.

No todos los algoritmos necesitarán el mismo grado de explicación. El nivel requerido dependerá del riesgo y de las consecuencias de la decisión. Sin embargo, cuanto mayor sea el impacto sobre los derechos de las personas, mayor deberá ser la capacidad institucional para explicar, revisar y eventualmente impugnar el resultado.

8. Vigilancia, privacidad y concentración de poder

La misma infraestructura de datos que permite construir un Estado predictivo puede utilizarse para desarrollar mecanismos de vigilancia extremadamente sofisticados.

Reconocimiento facial, cámaras inteligentes, geolocalización, análisis de redes sociales, metadatos y sistemas de identificación automatizada permiten obtener información detallada sobre los movimientos, preferencias, relaciones y comportamientos de las personas.

Estas capacidades pueden utilizarse legítimamente en ámbitos como seguridad pública, investigación criminal, gestión sanitaria o prevención de riesgos. Sin embargo, sin límites claros pueden transformarse en mecanismos de vigilancia masiva.

La publicación analiza esta tensión mediante el concepto de Estado de vigilancia y advierte que el riesgo no procede únicamente de los gobiernos. Las grandes empresas tecnológicas acumulan enormes volúmenes de información sobre sus usuarios y poseen capacidades analíticas que también pueden influir en comportamientos individuales y colectivos.

De esta manera, aparecen dos posibles concentraciones de poder: el “súper Estado” y las “súper empresas”. Ambos modelos pueden aprovechar la IA y los datos para ampliar significativamente su capacidad de observación, predicción e influencia.

La experiencia internacional muestra que tecnologías creadas con una finalidad legítima pueden utilizarse posteriormente para objetivos distintos. Por ello, el análisis de riesgo debe contemplar no solo la finalidad inicial de un sistema, sino también sus posibles usos secundarios y las consecuencias no previstas de su implementación.

9. Ciberseguridad e inteligencia artificial

Aunque el libro aborda la IA principalmente desde la perspectiva de las políticas públicas, sus planteamientos tienen implicaciones directas para la ciberseguridad.

La digitalización creciente de los servicios públicos genera nuevas dependencias tecnológicas. Bases de datos centralizadas, plataformas digitales, servicios en la nube, sistemas automatizados y modelos de IA se convierten en activos críticos cuyo compromiso puede afectar tanto a las instituciones como a millones de ciudadanos.

La protección de estos ecosistemas requiere una visión integral. No basta con proteger el algoritmo. Deben asegurarse las fuentes de datos, las interfaces de acceso, la infraestructura tecnológica, los sistemas de almacenamiento, los mecanismos de autenticación y los canales utilizados para intercambiar información.

Además, la confiabilidad de un sistema de IA depende directamente de la confiabilidad de sus datos. Si la información de entrenamiento o los datos utilizados durante la operación son alterados, incompletos o manipulados, las decisiones resultantes pueden perder integridad.

Desde esta perspectiva, la gobernanza de la IA y la ciberseguridad deben desarrollarse de manera conjunta. Cuanto mayor sea la automatización de decisiones públicas, mayor será la necesidad de controles de seguridad, trazabilidad, auditoría, monitoreo continuo y capacidad de respuesta ante incidentes.

La publicación también pone de manifiesto una situación particularmente importante para América Latina: las capacidades regionales de ciberseguridad son heterogéneas y existen diferencias considerables entre países. La construcción de un Estado digital inteligente debe ir acompañada, por tanto, del fortalecimiento de capacidades institucionales de prevención y respuesta.

10. Gobernanza responsable: dimensiones técnica, ética y sociojurídica

Uno de los planteamientos más sólidos del documento es la necesidad de construir una gobernanza de la inteligencia artificial sobre tres dimensiones complementarias: técnica, ética y sociojurídica.

La dimensión técnica comprende los datos y algoritmos que sustentan los sistemas. Aquí resultan fundamentales la calidad de la información, la mitigación de sesgos, la explicabilidad, la transparencia, la auditabilidad y el monitoreo durante todo el ciclo de vida.

La dimensión ética establece los principios que determinan qué usos de la inteligencia artificial son aceptables y cuáles deberían limitarse o excluirse. La pregunta no debe ser únicamente si una aplicación es técnicamente posible, sino también si resulta socialmente deseable y compatible con la dignidad humana.

Finalmente, la dimensión sociojurídica convierte estos principios en instituciones, obligaciones, responsabilidades y mecanismos de control.

Esta estructura resulta especialmente relevante porque evita dos errores frecuentes: intentar solucionar mediante regulación problemas exclusivamente técnicos o, en sentido contrario, confiar únicamente en soluciones tecnológicas para problemas que tienen componentes políticos, sociales y jurídicos.

Una gobernanza efectiva requiere integrar las tres dimensiones desde el inicio del proyecto y mantenerlas durante todo el ciclo de vida del sistema.

11. Regulación sin paralizar la innovación

El documento sostiene que la inteligencia artificial requiere reglas claras. No resulta razonable trasladar al ciudadano la responsabilidad de comprender y evaluar todos los riesgos de sistemas algorítmicos altamente complejos.

El Estado tiene una doble responsabilidad: es simultáneamente usuario de inteligencia artificial y regulador de su utilización.

Como usuario, debe asegurarse de que los sistemas incorporados a la administración pública cumplan requisitos de seguridad, calidad, equidad, transparencia y respeto de los derechos fundamentales. Como regulador, debe establecer las condiciones bajo las cuales las empresas y otras

organizaciones pueden recopilar datos, desarrollar algoritmos y utilizar sistemas automatizados.

La regulación tampoco debería convertirse en un obstáculo innecesario para la innovación. El reto consiste en establecer mecanismos proporcionales al nivel de riesgo. Una aplicación administrativa de bajo impacto no requiere necesariamente los mismos controles que un algoritmo utilizado para identificar personas, asignar prestaciones sociales o apoyar decisiones judiciales.

La responsabilidad humana tampoco desaparece cuando una decisión es automatizada. Si un sistema provoca daños o decisiones injustas, debe existir una estructura que permita determinar responsabilidades, revisar resultados y corregir errores.

12. Conclusiones y consideraciones estratégicas

La principal conclusión de ¿Quién le teme a la inteligencia? es que América Latina y el Caribe no deberían observar la revolución de la inteligencia artificial únicamente como consumidores tardíos de tecnología. La región tiene la oportunidad de participar en su desarrollo y, especialmente, en la construcción de modelos propios de gobernanza.

La pandemia puso en evidencia tanto el enorme potencial de los sistemas digitales como las debilidades estructurales de la región: registros incompletos, baja interoperabilidad, brechas de conectividad, insuficientes capacidades técnicas e importantes desigualdades sociales.

Superar estas limitaciones exige considerar la inteligencia artificial como una política de Estado, y no como la adquisición aislada de herramientas tecnológicas.

Desde una perspectiva de ciberseguridad y gobernanza digital, las prioridades que se desprenden de la obra son: fortalecer la transformación digital del Estado; construir una gobernanza sólida de datos; integrar la ciberseguridad desde el diseño; evaluar y mitigar sesgos algorítmicos; garantizar transparencia, explicabilidad y auditabilidad; mantener supervisión y responsabilidad humana; proteger la privacidad y los derechos fundamentales; fortalecer capacidades institucionales; promover cooperación entre Estado, academia, empresas y sociedad civil; y adoptar un enfoque basado en riesgos.

La visión final de los autores no es contraria a la inteligencia artificial. Por el contrario, plantea acelerar su aprovechamiento, pero hacerlo con responsabilidad. La región necesita modernizar sus Estados y desarrollar capacidades predictivas que permitan anticipar crisis y prestar mejores

servicios. No obstante, esta modernización solo será sostenible si está acompañada de seguridad, transparencia, protección de datos, control democrático y respeto de los derechos humanos.

El desafío no consiste, por tanto, en elegir entre innovación y protección. Consiste en construir instituciones capaces de alcanzar ambas.

Un Estado de Bienestar Digital Inteligente puede utilizar algoritmos para anticipar riesgos, optimizar recursos y mejorar los servicios públicos, pero su verdadera inteligencia no se medirá únicamente por la sofisticación de sus modelos. Se medirá por su capacidad para utilizar la tecnología sin comprometer la seguridad, la privacidad, la equidad y la confianza de la ciudadanía.

En ese sentido, el aporte más relevante de la publicación es recordar que el futuro de la inteligencia artificial no depende exclusivamente de las máquinas. Depende de las decisiones humanas e institucionales que determinan qué datos se utilizan, con qué finalidad se procesan, quién controla los sistemas, quién responde por sus resultados y qué límites estamos dispuestos a establecer.

Para América Latina y el Caribe, esta discusión es particularmente urgente. El rezago tecnológico constituye un riesgo, pero avanzar sin controles también lo es. La estrategia adecuada consiste en aumentar el ritmo de la transformación digital y, al mismo tiempo, elevar el nivel de responsabilidad con el que se diseña, implementa y gobierna la inteligencia artificial.

Fuente

Cabrol, M. y Sánchez A., R. (2021). ¿Quién le teme a la inteligencia? Posibilidades y riesgos de la inteligencia artificial en el Estado digital. Banco Interamericano de Desarrollo (BID).