



Vigilancia algorítmica y derechos fundamentales en la seguridad pública latinoamericana

Algorithmic Surveillance and Fundamental Rights in Latin American Public Security

FAUSTO ANDRÉS BERRAZUETA CADENA

Investigador independiente, Quito, Ecuador

Correo electrónico: andresberrazueta@gmail.com

ORCID: <https://orcid.org/0000-0002-4145-3634>

Resumen

En este artículo se analiza la incorporación de sistemas de inteligencia artificial (IA) en las políticas de seguridad pública en América Latina, con énfasis en Ecuador, Colombia y Perú, y sus implicaciones para los derechos fundamentales. La pregunta de investigación que guía el estudio es: ¿en qué medida la implementación de estas tecnologías se desarrolla sin mecanismos adecuados de control jurídico y supervisión institucional, generando riesgos para los derechos fundamentales?

Desde un enfoque jurídico comparado y dogmático, el estudio examina tres tipos de tecnologías: videovigilancia inteligente, policía predictiva y reconocimiento facial, a partir de dimensiones analíticas como privacidad, transparencia algorítmica, sesgo y rendición de cuentas.

Los resultados evidencian patrones comunes de opacidad en la adopción de estos sistemas, ausencia de auditorías independientes y una débil articulación entre innovación tecnológica y garantías constitucionales. Asimismo, se identifican riesgos estructurales asociados a la vulneración del derecho a la privacidad, la presunción de inocencia y la no discriminación.

En este contexto, el artículo propone el fortalecimiento de los marcos regulatorios mediante la incorporación de evaluaciones de impacto en derechos humanos, estándares de transparencia algorítmica y mecanismos de control judicial, con el fin de garantizar un uso legítimo de la IA en la seguridad pública.

Palabras clave

Inteligencia artificial, Seguridad nacional, Derechos fundamentales, Ecuador, América Latina, Videovigilancia.

Abstract

This article analyzes the incorporation of artificial intelligence (AI) systems into public security policies in Latin America, with a focus on Ecuador, Colombia, and Peru, and their implications for fundamental rights. The research problem examines the extent to which these technologies, implemented under efficiency-driven approaches, lack adequate mechanisms of legal control and institutional oversight.

Using a comparative and doctrinal legal methodology, the study evaluates three types of technologies—intelligent video surveillance, predictive policing, and facial recognition—based on analytical dimensions such as privacy, algorithmic transparency, bias, and accountability. The findings reveal common patterns of opacity in the adoption of these systems, a lack of independent audits, and weak integration between technological innovation and constitutional guarantees. These conditions generate structural risks affecting fundamental rights, particularly privacy, due process, and non-discrimination.



The article proposes the strengthening of regulatory frameworks through the incorporation of human rights impact assessments, algorithmic transparency standards, and effective judicial oversight mechanisms to ensure the lawful use of AI in public security.

Keywords

Artificial Intelligence, Public Security, Fundamental Rights, Ecuador, Latin America, Video Surveillance.

1. Introducción

América Latina atraviesa un proceso de incorporación de tecnologías de IA en sus sistemas de seguridad pública, lo que plantea desafíos en materia de regulación, control institucional y protección de derechos fundamentales. En países como Ecuador, Colombia y Perú esta adopción se ha desarrollado en el contexto de políticas de modernización orientadas a la eficiencia y al combate de la criminalidad, pero también ha implicado una transformación en las formas de vigilancia estatal, ahora mediadas por sistemas algorítmicos capaces de procesar y analizar grandes volúmenes de datos.

En este contexto, la investigación se estructura a partir de la siguiente pregunta: ¿en qué medida la incorporación de sistemas de IA en las políticas de seguridad pública en la región andina se desarrolla sin mecanismos adecuados de control jurídico, generando riesgos para la protección de los derechos fundamentales?

El objetivo general del artículo es analizar, desde un enfoque jurídico comparado, el impacto de estas tecnologías en Ecuador, Colombia y Perú. Como objetivos específicos, se busca:

- (i) Identificar los principales sistemas implementados.
- (ii) Examinar su base normativa y los mecanismos de control institucional.
- (iii) Evaluar los riesgos que plantean para derechos como la privacidad, la presunción de inocencia y la no discriminación.

La hipótesis sostiene que la adopción de sistemas de IA en la seguridad pública en América Latina se ha producido de manera acelerada y con niveles insuficientes de transparencia, supervisión y rendición de cuentas, lo que genera riesgos estructurales para los derechos fundamentales.

El problema técnico-jurídico central radica en la ausencia de mecanismos efectivos de interpretabilidad y auditoría de los sistemas algorítmicos utilizados por el Estado, particularmente en el marco de su contratación pública. Estas tecnologías suelen ser adquiridas a proveedores privados bajo condiciones de opacidad, lo que impide conocer su funcionamiento, criterios de decisión y posibles sesgos, generando asimetrías de información que limitan el control institucional y la supervisión judicial.

El análisis se delimita geográficamente en Ecuador, Colombia y Perú, temporalmente en el periodo 2018-2024 y materialmente en el estudio de sistemas de videovigilancia inteligente, policía predictiva y reconocimiento facial aplicados a la seguridad pública.

En términos comparados, se observa que, pese a los avances normativos como la Ley Orgánica de Protección de Datos Personales (2021) en Ecuador, persisten brechas entre regulación y práctica. En Colombia y Perú se identifican patrones similares de opacidad en la adquisición de tecnologías, ausencia de auditorías independientes y uso de estos sistemas en contextos sensibles, como la protesta social.

Desde una perspectiva jurídica, estos desarrollos plantean tensiones con derechos fundamentales como la privacidad, la presunción de inocencia y la libertad de expresión, particularmente

en contextos de vigilancia masiva y tratamiento intensivo de datos. La ausencia de controles efectivos puede derivar en escenarios de opacidad algorítmica y concentración de poder.

En este sentido, sostiene como tesis central que la utilización de sistemas de IA en la seguridad pública, carentes de mecanismos de transparencia, auditoría independiente y control judicial, resulta incompatible con los estándares de protección de derechos fundamentales, lo que exige su sometimiento a un régimen reforzado de control jurídico basado en principios de legalidad, proporcionalidad y debido proceso.

2. Metodología

Para el estudio se adoptó un enfoque jurídico cualitativo, con una metodología de carácter comparado y dogmático. El análisis se centra en Ecuador, Colombia y Perú, seleccionados por su relevancia en la incorporación de tecnologías de IA en la seguridad pública en la región andina.

La comparación se estructura a partir de tres ejes:

- (i) Los tipos de tecnologías implementadas: videovigilancia inteligente, policía predictiva y reconocimiento facial.
- (ii) La existencia de marcos normativos y mecanismos de control institucional.
- (iii) Los riesgos asociados a la posible afectación de derechos fundamentales.

Para el desarrollo del estudio se revisaron fuentes normativas, informes institucionales, literatura académica y reportes de organizaciones especializadas en derechos digitales. Como limitación, se reconoce la escasa disponibilidad de información pública sobre contratos tecnológicos y funcionamiento de los sistemas algorítmicos, lo que restringe el acceso a datos técnicos y forma parte del problema analizado.

3. Marco normativo y estándares jurídicos aplicables

La incorporación de sistemas de IA en la seguridad pública debe analizarse a la luz de los estándares jurídicos nacionales e internacionales en materia de derechos humanos, particularmente aquellos relacionados con la privacidad, la protección de datos personales, la igualdad y el debido proceso.

3.1. Estándares internacionales de derechos humanos

En el plano internacional, el derecho a la privacidad se encuentra reconocido en el artículo 12 de la Declaración Universal de Derechos Humanos y en el artículo 11 de la Convención Americana sobre Derechos Humanos, los cuales establecen la protección frente a injerencias arbitrarias o abusivas en la vida privada.

En desarrollo de estos principios, la jurisprudencia del Sistema Interamericano de Derechos Humanos ha establecido estándares estrictos en materia de vigilancia estatal y tratamiento de datos personales. La Corte Interamericana de Derechos Humanos (Corte IDH) ha señalado que toda injerencia en la vida privada, incluidas la recolección, el almacenamiento y el procesamiento de datos, debe cumplir con los principios de legalidad, finalidad legítima, necesidad y proporcionalidad, así como estar prevista en normas claras y ser sujetas a control judicial efectivo. En casos como *Escher y otros vs. Brasil* (Corte IDH, 2009) y *Tristán Donoso vs. Panamá* (2009), la Corte IDH ha determinado que la interceptación y recopilación de información sin garantías adecuadas vulnera el derecho a la privacidad.

Estos estándares son plenamente aplicables al análisis de sistemas de IA en seguridad pública, en la medida en que implican formas avanzadas de vigilancia y procesamiento automatizado de datos personales.

En el contexto de tecnologías emergentes, los organismos internacionales han advertido, además, sobre los riesgos asociados a la opacidad algorítmica. Subrayan la necesidad de garantizar transparencia, explicabilidad y mecanismos efectivos de rendición de cuentas en el uso de sistemas automatizados.

3.2. Regulación en América Latina

En el ámbito regional, los países analizados tienen niveles diferenciados de desarrollo normativo en materia de protección de datos y regulación de tecnologías de vigilancia.

En Ecuador, la Constitución (2008) reconoce el derecho a la protección de datos personales y a la privacidad, el cual está reforzado en la Ley Orgánica de Protección de Datos Personales (2021). No obstante, la implementación de sistemas de IA en seguridad pública ha avanzado más rápidamente que los mecanismos institucionales de control. En este contexto, la Corte Constitucional ha desarrollado criterios relevantes a través de la acción de hábeas data. La institución reconoce que la protección de datos implica no solo el acceso a la información, sino también el control sobre su uso, finalidad y circulación.

En Colombia existe un marco normativo consolidado en materia de protección de datos personales, encabezado por la Ley Estatutaria 1581 (2012), que establece principios como legalidad, finalidad, transparencia, acceso restringido y seguridad. Sin embargo, esta normativa contempla excepciones en materia de seguridad y defensa nacional, lo que introduce zonas de menor control jurídico. En este escenario, la incorporación de sistemas de analítica predictiva plantea desafíos en términos de transparencia, auditabilidad y supervisión institucional.

En el caso peruano, la Ley N.º 29733 establece un marco de protección de datos basado en principios de legalidad, consentimiento y finalidad. Este régimen se complementa con la Ley N.º 31814 (2023), que incorpora principios de ética, transparencia y responsabilidad en el uso de IA. No obstante, persisten limitaciones en su aplicación en el ámbito de la seguridad pública, especialmente frente al uso de tecnologías biométricas como el reconocimiento facial.

4. Inteligencia artificial y seguridad pública: experiencias y desafíos en América Latina

La incorporación de la IA en las estrategias de seguridad es una tendencia creciente en América Latina. Varios gobiernos de la región están explorando el potencial de estas herramientas para optimizar la prevención del delito y la gestión policial. No obstante, esta adopción tecnológica conlleva un intenso debate sobre sus implicaciones éticas, particularmente en lo que respecta a la protección de datos personales, la equidad social y los límites de la vigilancia estatal. Un análisis de las experiencias en Ecuador, Colombia y Perú revela un panorama común de oportunidades y riesgos, que se replica en otros países de la región como México.

El caso de Ecuador es emblemático porque implementó a nivel nacional del Sistema de Servicio Integrado de Seguridad ECU-911. Este sistema, que centraliza la respuesta a emergencias, incorporó a su red de más de 6000 cámaras de videovigilancia desplegadas en Quito, Guayaquil y Cuenca, aumentando las capacidades de IA para el análisis de video en tiempo real. La implementación técnica, desarrollada con empresas como Huawei durante su fase de expansión, incluye algoritmos de reconocimiento facial y análisis de comportamiento. En la práctica, estos sistemas no solo identifican rostros en listas de personas buscadas, sino que están programados para detectar incidentes específicos mediante la interpretación de movimientos. Por ejemplo, los algoritmos pueden generar alertas automáticas ante comportamientos como caídas de personas en espacios públicos —para asistencia médica—.

Desde una perspectiva técnico-jurídica, resulta relevante evaluar si estos sistemas incorporan mecanismos de privacidad desde el diseño (*Privacy by Design*), especialmente en aquellos

supuestos en los que la detección se basa en patrones conductuales y no en la identificación individual. En estos casos, la utilización de técnicas de anonimización o minimización de datos, como el desenfoque de rostros o el procesamiento de información no identificable, permitiría cumplir con el principio de proporcionalidad, con el fin de evitar el tratamiento innecesario de datos biométricos.

La ausencia de estos mecanismos implica que los sistemas diseñados para finalidades asistenciales o preventivas puedan derivar en formas de vigilancia intensiva, lo que plantea tensiones con el derecho a la privacidad y con los estándares internacionales de protección de datos, como el principio de minimización establecido en el Reglamento General de Protección de Datos (GDPR) y las recomendaciones sobre IA de organismos como la OCDE (2019) o la Unión Europea (2016).

En el caso de Ecuador, estas alertas se canalizan directamente a los monitores del ECU-911, lo que da lugar a una verificación humana y el despliegue coordinado de la Policía Nacional, ambulancias o bomberos en un promedio de 6 a 8 minutos en zonas urbanas cubiertas, según datos oficiales del sistema.

Sin embargo, la Superintendencia de Protección de Datos Personales del Ecuador ([SPDP], 2025) ha reconocido en el Plan Nacional de Protección de Datos Personales que la implementación del reconocimiento facial en el ECU-911 se hizo sin un marco normativo específico que regulara su uso. Tal hecho generó un riesgo latente de vigilancia masiva y perfilamiento de la población. Esta capacidad de monitorización, sin los debidos controles, podría socavar libertades civiles fundamentales, como el derecho a la privacidad y la libertad de reunión pacífica.

Colombia ha implementado de forma pionera en la región el modelo de policía predictiva, materializado en herramientas como la Plataforma CAIVAS (Centro de Análisis e Inteligencia de Datos para la Seguridad), . La Policía Nacional gestiona esta plataforma con el apoyo del Ministerio de Defensa.

Este sistema integra y analiza masivos volúmenes de datos en tiempo real. El núcleo de su funcionamiento son algoritmos de *machine learning* que cruzan datos históricos de delincuencia –hurto a personas, homicidios, hurto a celulares–, e incorporan más de 40 variables en tiempo real, que incluyen lo siguiente:

- Reportes del cuadrante (SIEC) y llamadas al 123
- Flujo de personas a partir de datos de telefonía móvil –convenios con operadores–
- Información socioeconómica del DANE por manzana
- Geolocalización de establecimientos comerciales y cajeros automáticos,
- Patrones de movilidad y reportes de ruido –aglomeraciones–.

El resultado no es que predice un delito concreto, sino que genera *cuadrantes de alto riesgo*. El sistema produce mapas de calor dinámicos que se actualizan cada seis horas, asignando probabilidades de incidencia delictiva a cuadrantes específicos –áreas de aproximadamente 300x300 m–. Esto permite a los mandos tácticos desplegar efectivos de *policía por cuadrante* de manera preventiva y con una asignación de recursos basada en evidencia sustentada en datos (*data-driven*). No obstante, la eficacia y ética de este modelo son objeto de escrutinio. La Corte Constitucional de Colombia (2024a) señala que la fiabilidad es cuestionable, ya que los algoritmos pueden crear “profecías autocumplidas”. Al predecir mayor delincuencia en zonas históricamente marginalizadas y con mayor presencia policial, como en algunas comunas de Medellín o localidades de Bogotá, se incrementa la tasa de detección en esas áreas, lo que retroalimenta el ciclo de datos sesgados y perpetúa la estigmatización territorial. La falta de transparencia en los algoritmos y la auditoría pública de sus resultados amplifica este riesgo.

En Perú, la implementación de IA en seguridad se concretó con la puesta en marcha del Sistema de Reconocimiento Facial en el Sistema Metropolitano de Transporte de Lima, a cargo

de la Autoridad de Transporte Urbano (ATU). Este proyecto, lanzado a fines de 2023 e implementado por la empresa española Facephi, despliega más de 1200 cámaras de videovigilancia en las estaciones y dentro de los autobuses articulados de la ruta troncal.

El sistema funciona cotejando en tiempo real los rostros de los pasajeros con una *lista de alertas* o *lista caliente* proporcionada por la Policía Nacional del Perú (PNP). La lista incluye a personas buscadas por delitos como secuestro, homicidio, tráfico de drogas y miembros de organizaciones criminales. Cuando se produce una coincidencia con un alto grado de confianza, el Centro de Control del Metropolitano recibe una alerta y notifica de inmediato a la PNP para una intervención coordinada en la siguiente estación.

Sin embargo, la eficacia de esta tecnología ha sido cuestionada. La Autoridad Nacional de Protección de Datos Personales del Perú ([ANPD], 2025) ha advertido, mediante Opinión Consultiva, que en entornos de alta densidad como la Estación Central del Metropolitano (Naranjal), donde transitan más de 700 000 pasajeros al día, factores como el ángulo de las cámaras, la velocidad del flujo peatonal y las condiciones de iluminación pueden degradar la precisión del algoritmo. Aunque Facephi reporta una precisión del 99.3 %, la ANPD subraya que un falso positivo en este contexto podría llevar a la detención injustificada de un ciudadano inocente. Esta preocupación se agrava porque, si bien Perú cuenta con las Leyes N.º 29733 (2011) y N.º 31814 (2023), la propia autoridad de control ha identificado una insuficiencia de controles específicos para el uso de tecnologías biométricas en espacios públicos (ANPD, 2025).

5. Implicaciones jurídicas de la vigilancia algorítmica para los derechos fundamentales

A partir de los casos previamente sistematizados, este apartado se centra exclusivamente en su análisis jurídico. El fin es analizar sus implicaciones jurídicas y sociales desde una perspectiva de derechos fundamentales.

La incorporación de la IA en las estrategias de seguridad pública constituye una de las transformaciones más significativas del siglo XXI en América Latina. Gobiernos como los de Ecuador, Colombia y Perú han apostado por el uso de tecnologías predictivas, reconocimiento facial y análisis de redes criminales para optimizar la prevención del delito y la gestión policial. No obstante, estas innovaciones, al tiempo que prometen eficiencia y modernización institucional, plantean profundos dilemas éticos y jurídicos relacionados con la protección de los derechos fundamentales. La tensión entre seguridad y libertad se reconfigura bajo nuevas lógicas de control algorítmico que, lejos de ser neutras, pueden reproducir o incluso amplificar las desigualdades estructurales.

El discurso de la eficiencia tecnológica ha legitimado el uso de la IA en el ámbito de la seguridad como una herramienta “objetiva” y “basada en datos”. Sin embargo, la literatura crítica advierte que los algoritmos de predicción del delito, como los utilizados en México o Colombia, no son neutrales, sino que reflejan los sesgos que existían en los datos con los que fueron entrenados (Eubanks, 2018, p. 34; Crawford, 2021, p. 8).

En particular, es posible distinguir entre distintos tipos de sesgo relevantes para el análisis. Por un lado, el sesgo de medición, que se origina en la forma en que los datos policiales son recolectados, puede reflejar prácticas de sobrerrepresentación territorial o focalización selectiva de la vigilancia. Por otro lado, el sesgo de agregación, que se produce cuando los algoritmos ponderan variables, como condiciones socioeconómicas o localización geográfica, refuerza correlaciones estructurales sin considerar sus causas subyacentes (Barocas et al., 2019).

Estos sesgos distorsionan la capacidad predictiva del sistema y reproducen dinámicas de criminalización selectiva, porque concentran la vigilancia en barrios populares y comunidades históricamente marginadas. Como ha señalado la Corte Constitucional de Colombia (2024a), esta tendencia genera profecías autocumplidas: al aumentar la presencia policial en zonas consideradas

de riesgo, se incrementa la detección de delitos menores, lo que refuerza estadísticamente la idea de que dichos territorios son más peligrosos.

En el contexto latinoamericano, el uso de algoritmos de *machine learning* para predecir el delito se inscribe en una lógica de control territorial que frecuentemente vincula indicadores socioeconómicos con riesgos de criminalidad. Ello puede implicar una clasificación automatizada de poblaciones, donde la pobreza se convierte en variable predictiva del delito (Ferguson, 2017, p. 54). De este modo, la IA contribuye a consolidar una forma de “gobernanza algorítmica” (Rouvroy & Berns, 2013 p. 168), en la que las decisiones estatales se justifican bajo la autoridad del dato y la opacidad del cálculo matemático.

Desde una perspectiva jurídica, esta opacidad plantea tensiones directas con el principio de motivación de los actos administrativos, en la medida en que las decisiones de seguridad pública, como el despliegue de fuerzas policiales o militares, pueden fundamentarse en sistemas algorítmicos cuyo funcionamiento no es accesible ni comprensible para los ciudadanos. Esta situación compromete el derecho a conocer las razones que sustentan la actuación estatal, así como la posibilidad de ejercer mecanismos efectivos de control y defensa.

En este contexto, adoptar herramientas como las *algorithmic impact assessments* (evaluaciones de impacto algorítmico) se presenta como un estándar necesario, en tanto permiten evaluar *ex ante* los riesgos asociados al uso de estos sistemas, particularmente en términos de transparencia, proporcionalidad y afectación a derechos fundamentales. Así, estas herramientas contribuyen a reforzar la legitimidad de las decisiones públicas en entornos de gobernanza algorítmica.

6. Impactos de la vigilancia algorítmica en los derechos fundamentales

6.1. Privacidad y vigilancia masiva

El caso del Sistema ECU-911, cuyas características técnicas y riesgos ya se expusieron, ilustra la expansión de un modelo de vigilancia masiva bajo la retórica de seguridad ciudadana. Esta situación plantea interrogantes sobre la proporcionalidad y legitimidad de la vigilancia estatal, especialmente en consideración de que la Ley de Protección de Datos Personales entró en vigencia en 2021 y aún carece de una autoridad de control plenamente operativa (SPDP, 2025).

La recopilación y procesamiento automatizado de rostros y movimientos en espacios públicos afecta directamente el derecho a la privacidad, reconocido en el artículo 12 de la Declaración Universal de Derechos Humanos y el artículo 11 de la Convención Americana sobre Derechos Humanos. Según Lyon (2018, p. 109), la vigilancia digital contemporánea transforma a los ciudadanos en “sujetos de datos”, reducidos a flujos de información susceptibles de clasificación, predicción y control. Cuando tales mecanismos operan sin transparencia ni supervisión independiente, el riesgo de abuso y de uso político de la información se vuelve inminente.

6.2. Estigmatización y sesgo territorial

Una de las consecuencias más problemáticas de la vigilancia algorítmica es su contribución a la estigmatización territorial y social. Los sistemas predictivos de criminalidad y las herramientas de análisis espacial, como las implementadas en Colombia y otros países de la región, tienden a asociar la pobreza con el riesgo delictivo. Los estudios sobre *poverty mapping* demuestran que la IA aplicada a datos satelitales y censales puede generar representaciones distorsionadas de los espacios urbanos e invisibilizar dinámicas sociales complejas (Aiken, et al., 2023, p. 1).

Esta relación entre espacio, tecnología y estigma tiene raíces en la tradición latinoamericana de control de la marginalidad. Wacquant (2008) conceptualiza la “estigmatización territorial” como una forma de dominación simbólica que refuerza la segregación y legitima la intervención coercitiva del Estado en los barrios populares. Cuando la IA refuerza estas representaciones, lo

que ocurre es una “digitalización del prejuicio” (Benjamin, 2019, p. 4), en la que la pobreza deja de ser entendida como fenómeno estructural y se convierte en marcador algorítmico de sospecha.

6.3. Riesgos de discriminación

El uso de tecnologías de reconocimiento facial y análisis de datos biométricos, como en el caso peruano del Metropolitano, evidencia un déficit estructural en la gobernanza de la IA. La ausencia de auditorías públicas de los algoritmos, la opacidad de los contratos con proveedores privados y la falta de marcos de rendición de cuentas constituyen violaciones indirectas al principio de legalidad y al derecho a la no discriminación.

Diversos estudios han demostrado que los sistemas de reconocimiento facial presentan tasas de error significativamente más altas al identificar rostros de mujeres y personas racializadas (Buolamwini & Gebru, 2018, p. 5). En contextos de alta densidad y precariedad urbana, como Lima o Guayaquil, estos sesgos pueden derivar en detenciones arbitrarias o prácticas de perfilamiento basadas en rasgos fenotípicos o indicadores socioeconómicos.

6.4. Implicaciones éticas y gobernanza

La expansión de la vigilancia algorítmica plantea dilemas éticos que trascienden la eficiencia técnica. Según Zuboff (2019, p. 11), el capitalismo de la vigilancia convierte la vida humana en materia prima para acumular datos, lo cual difumina los límites entre el interés público y la explotación comercial de la información. En América Latina, donde la infraestructura digital depende en gran medida de proveedores extranjeros, se agrega el riesgo de dependencia tecnológica y de afectación a la soberanía informacional (Deibert, 2020).

Desde una perspectiva de ciberseguridad estratégica, este escenario implica riesgos adicionales vinculados a la posible existencia de vulnerabilidades ocultas (*backdoors*) y a la falta de acceso al código fuente de los sistemas implementados. Cuando el Estado no controla el ciclo de vida del dato, particularmente en el caso de datos biométricos, y este se encuentra bajo la lógica propietaria de proveedores privados, se debilita su capacidad de supervisión, auditoría y respuesta ante incidentes.

En este contexto, resulta pertinente considerar mecanismos como el depósito de código fuente (*source code escrow*) y la exigencia de auditorías de caja blanca (*white-box auditing*) en los procesos de contratación pública, con el fin de garantizar niveles adecuados de transparencia, seguridad y control estatal sobre infraestructuras críticas de información.

Asimismo, es imprescindible adoptar marcos normativos que garanticen la transparencia algorítmica, la supervisión judicial de los sistemas predictivos y la evaluación de impacto en derechos humanos antes de su implementación. La Unión Europea ha adoptado el Reglamento (UE) 2024/1689 (Parlamento Europeo y Consejo de la Unión Europea, 2024), conocido como AI Act, que establece un marco normativo armonizado en materia de IA, y que entró en vigor el 1 de agosto de 2024. Esta norma clasifica los sistemas de IA según su nivel de riesgo y prohíbe, desde el 2 de febrero de 2025, aquellas prácticas que impliquen vigilancia masiva indiscriminada o manipulación subliminal (Parlamento Europeo y Consejo de la Unión Europea, 2024). América Latina podría inspirarse en ese modelo, adaptándolo a sus contextos institucionales y sociales.

6.5. Presunción de inocencia y sesgo algorítmico

La policía predictiva introduce un riesgo estructural: la denominada “culpabilidad algorítmica”. Modelos similares al *software* COMPAS han demostrado que reproducen sesgos raciales y socioeconómicos, lo cual erosiona la presunción de inocencia y el principio de igualdad ante la ley.

El análisis de Propublica, un estudio empírico ampliamente citado, evidenció las siguientes diferencias significativas en las tasas de error del sistema COMPAS: los acusados

afroamericanos fueron clasificados erróneamente como de “alto riesgo” en un 44.9 % de los casos –falsos positivos–, frente a un 23.5 % de acusados blancos. A la inversa, los acusados blancos presentaron mayores tasas de falsos negativos, es decir, un 47.7 %, lo que implica una subestimación de su riesgo (Angwin et al., 2016, párr. 2-4). Estos resultados muestran que los sistemas algorítmicos pueden generar errores sistemáticos con impacto diferenciado según variables estructurales.

Si bien estos datos corresponden al contexto estadounidense, su relevancia para América Latina radica en la advertencia metodológica que plantean: en ausencia de auditorías independientes y métricas públicas de evaluación, los sistemas de predicción del delito implementados en la región podrían estar reproduciendo patrones similares de error y discriminación sin que existan mecanismos institucionales para detectarlos o corregirlos. Esta opacidad refuerza los riesgos de afectar a la presunción de inocencia, al desplazar el análisis individual del caso por inferencias probabilísticas de carácter colectivo.

7. Análisis comparado de sistemas de IA en seguridad pública

Con el fin de sistematizar los hallazgos del estudio y fortalecer el análisis comparado, se presenta a continuación una matriz (tabla 1) que integra los principales elementos jurídicos, tecnológicos y de riesgo identificados en los casos de Ecuador, Colombia y Perú. Esta herramienta permite visualizar de manera estructurada las convergencias y divergencias en la implementación de sistemas de IA en seguridad pública, así como evaluar su compatibilidad con los estándares de protección de derechos fundamentales.

País	Ecuador	Colombia	Perú
Tipo de sistema	Videovigilancia inteligente (ECU-911)	Policía predictiva (CAIVAS)	Reconocimiento facial (transporte público)
Base legal	Constitución (art. 66), LOPDP (2021)	Ley N. 1581 (2012), normativa policial	Ley N.º 29733 (2011), Ley N.º 31814 ([2023], IA)
Salvaguardas	Protección formal de datos personales	Protección de datos con excepciones por seguridad	Principios de protección de datos
Riesgos identificados	Vigilancia masiva, reconocimiento facial sin regulación específica	Sesgo algorítmico, opacidad, profecías autocumplidas	Falsos positivos, detenciones arbitrarias, baja supervisión
Casos documentados	ECU-911, Red Nacional de Cámaras	Plataforma CAIVAS	Metropolitano de Lima (Facephi)
Estándares aplicables	CADH art. 11, <i>Escher vs. Brasil</i> (2009a)	CADH, <i>Tristán Donoso vs. Panamá</i> (2009b)	CADH, principios de proporcionalidad
Propuesta normativa	Regulación biométrica, auditorías algorítmicas	Transparencia algorítmica, control judicial	DPIA, límites al uso biométrico

Tabla 1. Matriz comparativa: IA en la seguridad pública en la región andina
Nota: Elaboración propia.

El análisis comparado muestra la existencia de patrones comunes en la región andina. En primer lugar, se observa una adopción acelerada de tecnologías de IA en el ámbito de la seguridad pública, sin que exista una regulación específica que aborde sus particularidades técnicas y jurídicas. En segundo lugar, si bien todos los países analizados cuentan con marcos normativos en materia de protección de datos personales, estos son insuficientes para regular fenómenos como la opacidad algorítmica, la toma automatizada de decisiones y el uso de sistemas biométricos en contextos de alta sensibilidad.

Asimismo, se identifican riesgos estructurales compartidos, tales como la posibilidad de vigilancia masiva, la reproducción de sesgos algorítmicos y la afectación a derechos fundamentales como la privacidad, la presunción de inocencia y la no discriminación. Estos riesgos se ven agravados por la limitada transparencia en la adquisición de tecnologías, la ausencia de auditorías independientes y la falta de mecanismos efectivos de rendición de cuentas.

Finalmente, el análisis revela una brecha significativa entre el reconocimiento normativo de derechos fundamentales y su protección efectiva frente al uso de sistemas de IA, lo que plantea la necesidad de desarrollar marcos regulatorios más robustos y adaptados a los desafíos de la gobernanza algorítmica.

A partir de estos hallazgos, la principal contribución de este artículo consiste en proponer un enfoque integrado de análisis jurídico de la IA en seguridad pública, basado en la articulación entre estándares de derechos humanos y mecanismos de control estatal. En particular, se plantea la necesidad de incorporar tres elementos fundamentales en la regulación de estos sistemas:

- (i) La exigencia de evaluaciones de impacto en derechos humanos previas a su implementación.
- (ii) La obligación de transparencia y explicabilidad de los algoritmos utilizados.
- (iii) El establecimiento de mecanismos de supervisión independiente y control judicial efectivo.

Este enfoque permite superar aproximaciones meramente descriptivas y avanzar hacia una propuesta normativa orientada a garantizar un uso legítimo y compatible con el Estado de derecho de las tecnologías de IA.

8. Propuesta de gobernanza y auditoría algorítmica en el sector público

A partir de los hallazgos comparados, se propone avanzar de la fase diagnóstica hacia una propuesta normativa concreta. En este sentido, organismos internacionales y expertos recomiendan establecer mecanismos obligatorios de control y transparencia algorítmica que salvaguarden los derechos fundamentales. En particular, la Unesco (UNESCO, 2022, p. 20) enfatiza que los sistemas de IA “deben ser auditables y trazables” mediante “mecanismos de supervisión, evaluación de impacto, auditoría y diligencia debida”, lo que guía los lineamientos mínimos aquí planteados. A continuación se resumen las principales medidas:

8.1. Evaluaciones de impacto en derechos humanos (EIDH) previas

Toda implementación de IA gubernamental de alto riesgo debe acompañarse de una evaluación anticipada sobre su impacto en derechos. La EIDH es un instrumento inspirado en los principios rectores de la ONU, que tiene como fin analizar cómo la intervención de la IA afectará derechos como la privacidad, la igualdad y el debido proceso, considerando la existencia de incorporando principios como el de transparencia, rendición de cuentas y no discriminación. La propia Unesco reconoce la importancia de estas evaluaciones para anticipar riesgos éticos y sociales en la IA. Se sugiere utilizar marcos metodológicos formales, por ejemplo, *algorithmic impact assessments*, que permitan cuantificar efectos adversos y proponer medidas mitigadoras antes de contratar o desplegar sistemas algorítmicos.

8.2. Transparencia algorítmica y apertura de información

Es necesario exigir que los organismos públicos divulguen información relevante sobre los sistemas de IA que utilizan. La transparencia algorítmica se define como la capacidad de actores internos o externos de examinar la lógica, funcionamiento y criterios de un algoritmo. Prácticamente,

esto implica publicar detalles técnicos –modelo, datos de entrenamiento, parámetros– y posibilitar auditorías externas independientes. Países como Uruguay ya han elaborado recomendaciones para promover la transparencia y explicabilidad de los sistemas estatales. A su vez, normativas como el Reglamento General de Protección de Datos - RGPD europeo exigen que cualquier información sobre decisiones automatizadas sea accesible y comprensible para los interesados. Estos principios facilitan la supervisión ciudadana y fortalecen la legitimidad del uso de la IA.

8.3. Auditorías independientes y periódicas

Deben establecerse mecanismos de auditoría externa obligatoria para los sistemas de IA gubernamentales. Una auditoría es un proceso “sistemático, independiente y documentado” para recoger evidencias y evaluar el cumplimiento de criterios predefinidos. Aplicada a la IA, la auditoría algorítmica revisa el código, los datos de entrenamiento y los efectos en precisión, equidad, privacidad, entre otros. Dicha auditoría puede tener enfoque técnico –rendimiento y sesgos del algoritmo– y jurídico –compatibilidad con la Constitución y las leyes–. Es fundamental que sean periódicas y vinculantes; por ejemplo, la guía del BID recomienda realizarlas con cierta frecuencia y ajustarlas según el nivel de riesgo del sistema. Las recomendaciones resultantes deben poder obligar a corregir deficiencias o detener el sistema si viola estándares.

8.4. Derecho a la explicación y a la impugnación

Se debe garantizar que las personas afectadas por decisiones automatizadas conozcan sus fundamentos y puedan reclamarlas. La Unesco destaca que el despliegue ético de la IA depende de su “transparencia y explicabilidad”. Sin embargo, en los casos de algoritmos “caja negra”, debe cumplirse con trazabilidad y auditabilidad alternativa, de modo que la opacidad no sea excusa para evadir derechos. En consecuencia, el marco normativo debe reconocer expresamente el derecho a recibir explicaciones comprensibles y acceso a procesos de revisión de decisiones algorítmicas –administrativos o judiciales–, de forma similar a las garantías de debido proceso en entornos tradicionales.

8.5. Fortalecimiento del control judicial y capacidad institucional

Finalmente, se requiere reforzar los mecanismos de supervisión judicial y administrativa. La capacitación de jueces y funcionarios en IA es indispensable: la Unesco (2022, p. 14, párr. 44) subraya la necesidad de promover la sensibilización pública y la “comprensión del público respecto de la IA” mediante educación y formación. De igual modo, deben elaborarse protocolos específicos para la revisión judicial de tecnologías algorítmicas y criterios sobre estándares mínimos –por ejemplo, test de proporcionalidad y no discriminación–. Este enfoque asegura que, ante una posible vulneración de derechos, existan recursos efectivos y especializados.

Estos lineamientos integran un modelo mínimo de gobernanza algorítmica para el sector público, alineado con estándares internacionales –Unesco, OCDE, ISO– y regiones. De este modo, se trasciende la mera denuncia diagnóstica para ofrecer una agenda práctica: la adopción de EIDH, la transparencia proactiva, auditorías técnicas y jurídicas, el derecho a la explicación y un control institucional robusto constituyen herramientas concretas para garantizar que la IA en seguridad pública sea legítima, ética y compatible con los derechos fundamentales.

9. Conclusiones

En esta investigación se propuso analizar en qué medida la incorporación de sistemas de IA en las políticas de seguridad pública de Ecuador, Colombia y Perú se ha desarrollado sin mecanismos adecuados de control jurídico, lo que genera riesgos para los derechos fundamentales. Los hallazgos confirman la hipótesis planteada: la adopción de estas tecnologías se

ha producido de manera acelerada y con niveles insuficientes de transparencia, supervisión y rendición de cuentas.

La incorporación de la IA en las políticas de seguridad de América Latina configura un nuevo modelo de gestión de la seguridad basado en análisis algorítmico. Si bien estas tecnologías ofrecen potencial para mejorar la gestión pública y la respuesta ante el delito, su uso sin salvaguardas adecuadas amenaza con profundizar desigualdades, vulnerar derechos y reforzar estigmas históricos.

La pretensión de neutralidad algorítmica resulta insostenible frente a la evidencia de que los sistemas aprenden de datos socialmente sesgados. En consecuencia, la vigilancia predictiva tiende a criminalizar la pobreza y a generar una diferenciación estructural en los niveles de exposición a sistemas de vigilancia.

Garantizar que la IA opere al servicio de la dignidad humana y no como instrumento de discriminación exige diseñar políticas públicas basadas en principios de transparencia, rendición de cuentas y justicia social. Solo bajo esos presupuestos la tecnología podrá ser una aliada, y no una amenaza, de los derechos fundamentales.

En conjunto, estos casos permiten identificar patrones comunes en la implementación de sistemas de IA en la región, lo que sirve de base para un análisis comparado de sus implicaciones jurídicas y su impacto en los derechos fundamentales.

Referencias bibliográficas

Libros

- Aiken, E., Rolf, E., et al. (2023). Fairness and representation in satellite-based poverty maps: Evidence of urban-rural disparities and their impacts on downstream policy. *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence (IJCAI-23) Special Track on AI for Good*, pp. 5888-5896. <https://www.ijcai.org/proceedings/2023/0653.pdf>
- Barocas, S., Hardt, M., et al. (2019). *Fairness and machine learning: Limitations and opportunities*. Cambridge MA: MIT Press.
- Benjamin, R. (2019). *Race after technology: Abolitionist tools for the New Jim Code*. Cambridge, UK: Polity.
- Buolamwini, J., & Gebru, T. (2018). Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification. *Proceedings of Machine Learning Research*, 81, 1-15. <http://proceedings.mlr.press/v81/buolamwini18a/buolamwini18a.pdf>
- Crawford, K. (2021). *Atlas of AI: Power, politics, and the planetary costs of artificial intelligence*. New Heaven CT: Yale University Press.
- Deibert, R. (2020). *Reset: Reclaiming the internet for civil society*. Toronto: House of Anansi Press.
- Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police, and punish the poor*. Nueva York, NY: St. Martin's Press.
- Ferguson, A. G. (2017). *The rise of big data policing: Surveillance, race, and the future of law enforcement*. Nueva York: NYU Press.
- Lyon, D. (2018). *The culture of surveillance: Watching as a way of life*. Cambridge, UK: Polity.
- OCDE. (2019). *OECD Principles on artificial intelligence*. OECD Publishing. <https://doi.org/10.1787/9789264316453-en>
- SPDP. (2025, mayo). Plan Nacional de Protección de Datos Personales. PLAN NRO. PLN-SPDP-PGE-2025-0003

- Wacquant, L. (2008). *Urban outcasts: A comparative sociology of advanced marginality*. Cambridge, UK: Polity.
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. Nueva York: PublicAffairs.

Artículos de publicaciones periódicas

- Rouvroy, A., & Berns, T. (2013). Gouvernamentalité algorithmique et perspectives d'émancipation. *Réseaux*, 177 (1), 163-196. <https://doi.org/10.3917/res.177.0163>

Publicaciones web

- Angwin, J., Larson, J., et al. (2016, mayo 23). Machine bias. *ProPublica*. <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>

Leyes

- Constitución de la República del Ecuador. (2008). R. O. 449, 20 de octubre.
- Ley Estatutaria 1581 (2012). Por la cual se dictan disposiciones generales para la protección de datos personales. Congreso de la República de Colombia.
- Ley N.º 29733. (2011). Ley de Protección de Datos Personales. Congreso de la República del Perú, 31 de julio.
- Ley N.º 31814. (2023). Ley que promueve el uso de la inteligencia artificial en favor del desarrollo económico y social del país. Congreso de la República del Perú, 5 de julio.
- Ley Orgánica de Protección de Datos Personales (2021). R. O., Suplemento 459, 26 de mayo.
- Unión Europea. (2016). Regulation (EU) 2016/679 (General Data Protection Regulation). *Official Journal of the European Union*.
- Parlamento Europeo y Consejo de la Unión Europea. (2024). Reglamento (UE) 2024/1689 del Parlamento Europeo y del Consejo, de 13 de junio de 2024, por el que se establecen normas armonizadas en materia de inteligencia artificial y por el que se modifican determinados actos legislativos de la Unión (Reglamento de Inteligencia artificial). *Diario Oficial de la Unión Europea*, L 2024/1689, 12 de julio.

Sentencias

- Corte Constitucional de Colombia. (2024a). Sentencia sobre inteligencia artificial, vigilancia y sesgos algorítmicos. Sentencia T-323/24, 12 de agosto.
- Corte IDH. (2009). *Caso Escher y otros vs. Brasil. Fondo, Reparaciones y Costas*. Sentencia de 6 de julio de 2009.
- Corte IDH. (2009). *Caso Tristán Donoso vs. Panamá. Fondo, Reparaciones y Costas*. Sentencia de 27 de enero de 2009.

Observaciones, opiniones, recomendaciones e informes

- ANPD. (2025). Sobre ámbito de aplicación de la LPDP y tratamiento de datos a través de videovigilancia con reconocimiento facial por las municipalidades [Opinión consultiva]. OC N.º 049-2025-JUS/DGTAIPD, Perú, 5 de diciembre. <https://www.gob.pe/institucion/anpd/informes-publicaciones/7479404-oc-n-049-2025-jus-dgtaipd-sobre-ambito-de-aplicacion-de-la-lpdp-y-tratamiento-de-datos-a-traves-de-videovigilancia-con-reconocimiento-facial-por-las-municipalidades>
- Unesco. (2022). Recommendation on the Ethics of Artificial Intelligence. Adoptada el 23 de noviembre de 2021. <https://unesdoc.unesco.org/ark:/48223/pf0000381137>