

NEURODERECHOS E INTELIGENCIA ARTIFICIAL EN LA EDUCACIÓN SUPERIOR LATINOAMERICANA: PRIVACIDAD MENTAL, AUTONOMÍA COGNITIVA E INFERENCIAS MÁS ALLÁ DE LOS NEURODATOS

NEURORIGHTS AND ARTIFICIAL INTELLIGENCE IN LATIN AMERICAN HIGHER EDUCATION: MENTAL PRIVACY, COGNITIVE AUTONOMY, AND INFERENCES BEYOND NEURODATA

Deisi Yunga-Godoy¹
Fredy Paredes-Cuenca²

RESUMEN: Este artículo analiza la relación entre neuroderechos, inteligencia artificial y educación superior en América Latina, con énfasis en los riesgos que las inferencias producidas por sistemas automatizados pueden generar sobre la privacidad mental, la autonomía cognitiva y la justicia algorítmica. La investigación adopta un enfoque cualitativo, doctrinal-conceptual y comparativo, sustentado en literatura especializada sobre neuroderechos, gobernanza de la IA, protección de datos personales y educación superior, así como en estándares internacionales y marcos normativos latinoamericanos, con especial atención a Chile y Ecuador. El estudio examina el uso de plataformas de analítica de aprendizaje, evaluación automatizada, asistentes generativos, sistemas predictivos de permanencia, herramientas de integridad académica y tecnologías de acompañamiento estudiantil. Los resultados muestran que, aunque las inferencias generadas por estos sistemas no constituyen neurodatos en sentido estricto, pueden incidir en dimensiones vinculadas con la vida mental, la autonomía intelectual, la trayectoria académica y las oportunidades educativas de estudiantes y docentes. Se concluye que la incorporación de IA en la educación superior exige garantías institucionales orientadas a prevenir la vigilancia cognitiva, el etiquetamiento algorítmico y la automatización acrítica de decisiones pedagógicas o administrativas.

Palabras clave: neuroderechos; inteligencia artificial; educación superior; privacidad mental; autonomía cognitiva; inferencias cognitivas; justicia algorítmica; gobernanza educativa

ABSTRACT: This article analyzes the relationship between neurorights, artificial intelligence, and higher education in Latin America, with particular attention to the risks that inferences produced by automated systems may pose to mental privacy, cognitive autonomy, and algorithmic justice. The study follows a qualitative, doctrinal-conceptual, and comparative approach, drawing on specialized literature on neurorights, AI governance, personal data protection, and higher education, as well as on international standards and Latin American regulatory frameworks, with special attention to Chile and Ecuador. It examines the use of learning analytics platforms, automated assessment, generative assistants, predictive retention systems, academic integrity tools, and student support technologies. The findings show that, although the inferences generated by these systems do not constitute neurodata in the strict sense, they may affect dimensions related to mental life, intellectual autonomy, academic trajectories, and educational opportunities for students and faculty. The article concludes that the use of AI in higher education requires institutional safeguards aimed at preventing cognitive surveillance, algorithmic labeling, and the

¹ Deisi Cecibel Yunga-Godoy. Ph.D. in Teacher Education por la Eötvös Loránd University, Budapest, Hungria

² Fredy Paredes-Cuenca, Mgtr. nm Investigación en Educación por la Universidad Técnica Particular de Loja

uncritical automation of pedagogical or administrative decisions.

Keywords: neurorights; artificial intelligence; higher education; mental privacy; cognitive autonomy; cognitive inferences; algorithmic justice; educational governance.

INTRODUCCION

La convergencia entre inteligencia artificial, plataformas educativas, analítica de aprendizaje y sistemas automatizados de decisión ha abierto un problema jurídico, ético y pedagógico de creciente relevancia: cómo proteger dimensiones asociadas con la vida mental cuando los sistemas tecnológicos no acceden directamente al cerebro, pero producen perfiles, predicciones o valoraciones sobre conducta, rendimiento, atención, vulnerabilidad, escritura, riesgo académico o toma de decisiones. El debate contemporáneo sobre neuroderechos surgió principalmente frente a neurotecnologías capaces de registrar, interpretar o modificar la actividad cerebral. Ienca y Andorno (2017) propusieron repensar los derechos humanos ante los avances de la neurociencia y la neurotecnología, especialmente en relación con la libertad cognitiva, la privacidad mental, la integridad mental y la continuidad psicológica. Yuste et al. (2017) advirtieron que la convergencia entre neurotecnologías e inteligencia artificial puede afectar la privacidad, la identidad, la agencia y la igualdad. Farahany (2023), por su parte, ha defendido la libertad cognitiva como una condición fundamental de la autonomía humana frente a tecnologías capaces de monitorear, inferir o influir sobre la actividad mental.

Limitar la protección jurídica de la mente a los datos cerebrales obtenidos mediante interfaces cerebro-computadora, sensores neuronales o dispositivos biométricos especializados deja fuera un conjunto de riesgos cada vez más visibles en contextos educativos. Las universidades incorporan plataformas de gestión del aprendizaje, herramientas de escritura generativa, sistemas de retroalimentación automatizada, software de detección de similitud o uso de IA, modelos predictivos de abandono, asistentes conversacionales y recursos de seguimiento estudiantil. Estas tecnologías no generan necesariamente neurodatos en sentido estricto. Sin embargo, pueden construir interpretaciones funcionalmente próximas a la vida cognitiva, académica o conductual de una persona cuando infieren riesgo, motivación, honestidad académica, atención, desempeño probable, vulnerabilidad o necesidad de intervención institucional.

El problema no consiste en afirmar que toda inteligencia artificial educativa amenaza los neuroderechos. Una tesis de ese tipo sería conceptualmente excesiva y normativamente poco útil. El argumento de este artículo es más delimitado: determinadas inferencias producidas por sistemas de IA, aunque derivadas de datos no neuronales, pueden adquirir relevancia jurídica reforzada cuando afectan la privacidad mental, la autonomía cognitiva, la autodeterminación informativa, la igualdad de oportunidades, el debido proceso o la posibilidad de revisar decisiones institucionales. La fuente del dato no agota la pregunta jurídica. También importa qué interpretación se produce, con qué finalidad, bajo qué grado de opacidad y con qué efectos sobre la trayectoria académica, profesional y personal de quienes participan en la educación superior.

Esta discusión resulta especialmente relevante para América Latina. La región

combina aceleración tecnológica, adopción de plataformas externas, desigualdades sociales persistentes, brechas de conectividad, diversidad cultural y capacidades institucionales limitadas para auditar sistemas algorítmicos. En este escenario, las universidades latinoamericanas pueden incorporar herramientas de IA sin contar siempre con marcos robustos de evaluación de impacto, supervisión humana, explicabilidad, impugnación o reparación. El riesgo no es únicamente técnico. También es jurídico, pedagógico y político, porque la educación superior distribuye oportunidades, certifica saberes, produce reputación académica y participa en procesos de movilidad social.

Los estándares internacionales de gobernanza de IA ofrecen una base relevante para abordar este problema. El Reglamento (UE) 2024/1689 adopta un enfoque basado en riesgos e incluye determinados usos de IA en educación y formación profesional dentro de la categoría de alto riesgo cuando inciden en el acceso, la admisión, la evaluación o la trayectoria educativa. El mismo instrumento restringe ciertos sistemas de reconocimiento de emociones en contextos educativos y laborales. De manera complementaria, el Convenio Marco del Consejo de Europa sobre Inteligencia Artificial, Derechos Humanos, Democracia y Estado de Derecho sitúa el ciclo de vida de los sistemas de IA dentro de obligaciones vinculadas con derechos humanos, rendición de cuentas y recursos efectivos. La Recomendación de la UNESCO sobre la Ética de la Inteligencia Artificial y los Principios de IA de la OCDE refuerzan criterios de dignidad, transparencia, equidad, supervisión humana, trazabilidad, seguridad y responsabilidad.

En América Latina, Chile constituye un antecedente especialmente relevante por la reforma constitucional introducida mediante la Ley 21.383, que ordena resguardar la actividad cerebral y la información proveniente de ella frente al desarrollo científico y tecnológico. El caso Emotiv Inc., conocido por la Corte Suprema chilena en 2023, trasladó esta preocupación al terreno judicial y mostró la centralidad de la privacidad mental frente a dispositivos comerciales de neurotecnología. Ecuador, aunque no ha constitucionalizado expresamente los neuroderechos, cuenta con una base normativa significativa en materia de protección de datos personales, autodeterminación informativa, principios de educación superior y regulación emergente de IA desde la Superintendencia de Protección de Datos Personales. Ambos casos permiten formular una agenda regional que no se limite a proteger neurodatos cuando existan, sino que también reconozca los riesgos derivados de inferencias cognitivas indirectas en sectores sensibles como la educación superior.

La pregunta de investigación que orienta este artículo es la siguiente: ¿cómo debe estructurarse un marco jurídico-conceptual de gobernanza de inteligencia artificial para proteger la privacidad mental, la autonomía cognitiva y la justicia algorítmica frente a inferencias generadas en la educación superior latinoamericana? El objetivo general es proponer un marco de gobernanza de IA para universidades latinoamericanas, con énfasis en deberes institucionales frente a inferencias que inciden en decisiones educativas relevantes.

La contribución del artículo es triple. En primer lugar, delimita la categoría de inferencias cognitivas como una zona intermedia entre datos personales ordinarios, datos sensibles y neurodatos estrictos. En segundo lugar, sitúa la educación superior como sector sensible de regulación algorítmica, no solo por el volumen de datos tratados, sino

por el impacto de las decisiones universitarias sobre la autonomía intelectual, la igualdad de oportunidades y las trayectorias académicas. En tercer lugar, propone una matriz de deberes institucionales aplicable a universidades que implementan IA en procesos académicos, administrativos o disciplinarios. Esta propuesta busca evitar dos extremos: la expansión indiscriminada de los neuroderechos hacia cualquier uso educativo de tecnología y la protección insuficiente frente a formas indirectas de perfilamiento mental, académico o conductual.

El artículo se organiza en nueve secciones. Tras esta introducción, la segunda sección presenta la metodología doctrinal-conceptual y comparativa. La tercera desarrolla los fundamentos conceptuales de los neuroderechos, la privacidad mental y las inferencias cognitivas. La cuarta analiza estándares internacionales de gobernanza de IA. La quinta examina la educación superior como sector sensible de regulación algorítmica. La sexta aborda los casos de Chile y Ecuador. La séptima formula la propuesta normativa de deberes institucionales. La octava discute los alcances, tensiones y límites de la propuesta. La novena sección presenta las conclusiones.

METODOLOGIA

El estudio adopta un enfoque cualitativo, doctrinal-conceptual y comparativo, orientado a construir un marco jurídico-educativo para analizar los riesgos derivados del uso de sistemas de inteligencia artificial en la educación superior latinoamericana. La elección de este enfoque responde a la naturaleza del problema investigado, que no busca medir empíricamente la frecuencia de uso de tecnologías específicas ni evaluar plataformas concretas, sino examinar categorías normativas, conceptuales y educativas vinculadas con la privacidad mental, la autonomía cognitiva, la protección de datos personales, la justicia algorítmica y la gobernanza institucional de la IA.

Desde la perspectiva doctrinal, la investigación analiza normas, estándares internacionales, principios jurídicos y desarrollos jurisprudenciales relacionados con neuroderechos, protección de datos personales, derechos humanos, inteligencia artificial y educación superior. Este análisis permite identificar los bienes jurídicos comprometidos por el uso de sistemas automatizados en contextos universitarios, especialmente cuando dichos sistemas producen perfiles, predicciones o valoraciones sobre rendimiento, atención, conducta, riesgo académico, integridad académica o vulnerabilidad estudiantil. Desde la perspectiva conceptual, el estudio delimita la noción de inferencias cognitivas como una categoría intermedia de riesgo jurídico y educativo, situada entre los datos personales ordinarios, los datos sensibles y los neurodatos en sentido estricto. Esta categoría se utiliza para analizar aquellas predicciones, clasificaciones o valoraciones algorítmicas que, sin provenir directamente de la actividad cerebral, pueden incidir sobre dimensiones vinculadas con la vida mental, la autonomía intelectual, la trayectoria académica y las oportunidades educativas de estudiantes y docentes.

Desde la perspectiva comparativa, el artículo examina el diálogo entre los estándares internacionales de gobernanza de IA y dos desarrollos latinoamericanos relevantes: por un lado, la constitucionalización chilena de la protección de la actividad cerebral y de la información proveniente de ella; por otro, la regulación ecuatoriana emergente de la IA desde el derecho a la protección de datos personales y los principios propios de la educación superior. La comparación Chile-Ecuador no pretende establecer equivalencias normativas entre ambos países, sino identificar rutas complementarias para la protección jurídica de la mente en contextos educativos mediados por inteligencia

artificial.

El corpus de análisis se organizó en cuatro grupos documentales. El primer grupo comprende literatura académica sobre neuroderechos, privacidad mental, libertad cognitiva, neurodatos, integridad mental e identidad personal, con especial atención a los aportes de Ienca y Andorno (2017), Yuste et al. (2017), Farahany (2023), Ligthart et al. (2023), Oliveira e Pacífico (2025) y Cornejo-Plaza et al. (2024). El segundo grupo reúne literatura sobre gobernanza de IA, opacidad algorítmica, vigilancia, justicia, evaluación de impacto, derechos humanos y rendición de cuentas, en diálogo con Pasquale (2015), Binns (2018), Floridi y Cowls (2019), Zuboff (2019), Crawford (2021), Mantelero (2022) y Wachter y Mittelstadt (2019). El tercer grupo incorpora literatura educativa sobre IA, big data, analítica de aprendizaje, ética institucional y automatización educativa, particularmente Holmes et al. (2019), Slade y Prinsloo (2013), Williamson (2017), Selwyn (2019) y Zawacki-Richter et al. (2019). El cuarto grupo corresponde a instrumentos normativos internacionales, regionales y nacionales relevantes para la protección de derechos en sistemas automatizados.

La selección del corpus respondió a criterios de pertinencia temática, relevancia doctrinal y actualidad normativa. Se priorizaron fuentes académicas y jurídicas que abordaran, de manera directa o indirecta, la relación entre IA, derechos fundamentales, neuroderechos, protección de datos, educación superior, vigilancia algorítmica, evaluación automatizada y toma de decisiones institucionales. Asimismo, se incluyeron instrumentos normativos y de soft law que ofrecen criterios aplicables a sistemas de IA en sectores sensibles, tales como educación, trabajo, salud, justicia y administración pública. Se excluyeron textos centrados exclusivamente en aspectos técnicos de programación, rendimiento computacional o diseño de plataformas cuando no ofrecían una discusión jurídica, ética o educativa relevante para el objeto del estudio.

El análisis normativo internacional se centró en el Reglamento (UE) 2024/1689 sobre inteligencia artificial, el Convenio Marco del Consejo de Europa sobre Inteligencia Artificial, Derechos Humanos, Democracia y Estado de Derecho, la Recomendación sobre la Ética de la Inteligencia Artificial de la UNESCO y los Principios de IA de la OCDE. Estos instrumentos fueron seleccionados porque expresan criterios convergentes de gobernanza basada en riesgos, derechos humanos, transparencia, supervisión humana, trazabilidad, no discriminación, evaluación de impacto, seguridad y rendición de cuentas. Aunque su fuerza jurídica no es homogénea, permiten identificar un núcleo normativo útil para examinar sistemas de IA aplicados a sectores sensibles, entre ellos la educación superior.

En el caso del Reglamento (UE) 2024/1689, el análisis se concentró especialmente en el artículo 5, apartado 1, letra f), relativo a la prohibición de comercializar, poner en servicio o utilizar sistemas de IA destinados a inferir emociones de una persona en el ámbito laboral y en instituciones educativas, salvo cuando su uso responda a razones médicas o de seguridad. Asimismo, se examinó el artículo 6, apartados 1 a 3, sobre las reglas de clasificación de los sistemas de IA de alto riesgo, así como el Anexo III, punto 3, letras a)–d), relativo a los sistemas de IA de alto riesgo en educación y formación profesional. Este último resulta especialmente relevante para el objeto del estudio, en cuanto incluye sistemas utilizados para la admisión o asignación de personas a instituciones o programas educativos, la evaluación de resultados de aprendizaje, la determinación del nivel educativo accesible y el monitoreo o detección de conductas prohibidas durante evaluaciones.

En el caso ecuatoriano, se revisaron de manera específica los artículos 1, 4, 5, 6, 7, 9 y 10 de la Resolución N.º SPDP-SPD-2026-0009-R, que expide la Norma General para la Garantía del Derecho de Protección de Datos Personales en el Uso de Sistemas de

Inteligencia Artificial. Estos artículos fueron seleccionados por su relación con el objeto y ámbito de aplicación de la norma, la garantía de los derechos de los titulares, el deber de información clara y transparente, la gestión de riesgos, la evaluación de impacto, las medidas de seguridad, el registro de actividades de tratamiento, las auditorías y las medidas correctivas aplicables a sistemas de IA que tratan datos personales. Dado el carácter emergente de esta regulación administrativa, el análisis se realizó sobre la versión normativa consultada al momento de elaboración del estudio.

La selección de Chile y Ecuador responde a una lógica de contraste regional. Chile ofrece un caso de reconocimiento explícito de la protección de la actividad cerebral y de la información proveniente de ella, además de un precedente judicial relevante sobre neurodatos comerciales. Ecuador, en cambio, ofrece un caso de protección constitucional y legal de datos personales, regulación administrativa reciente sobre IA y principios sectoriales de educación superior vinculados con igualdad de oportunidades, autonomía responsable, calidad, pertinencia y autodeterminación para la producción del pensamiento y conocimiento. Esta comparación permite observar dos rutas jurídicas complementarias: una vía neuroconstitucional, centrada en la actividad cerebral y los neurodatos, y una vía de gobernanza de datos e IA aplicada a contextos educativos.

El procedimiento analítico se desarrolló en cuatro momentos. En primer lugar, se identificaron los bienes jurídicos y educativos comprometidos por el uso de IA en educación superior: privacidad mental, autonomía cognitiva, autodeterminación informativa, igualdad de oportunidades, debido proceso, integridad académica, justicia educativa y autonomía intelectual. En segundo lugar, se diferenciaron las principales categorías de información y riesgo involucradas: datos personales ordinarios, datos sensibles, neurodatos e inferencias cognitivas. En tercer lugar, se examinaron los estándares internacionales y regionales de IA para determinar qué principios podían ser trasladados al ámbito universitario. En cuarto lugar, dichos principios fueron traducidos en deberes institucionales aplicables a universidades que emplean sistemas automatizados o asistidos por IA en procesos académicos, administrativos, predictivos, evaluativos o disciplinarios.

El análisis se orientó por una lectura transversal de tres dimensiones: la fuente del dato, el tipo de inferencia producida y el efecto institucional de dicha inferencia. Esta estrategia permitió evitar una reducción del problema a la naturaleza técnica del dato. En lugar de preguntar únicamente si una información constituye o no neurodato, el estudio examina qué tipo de interpretación algorítmica se produce, con qué finalidad, bajo qué condiciones de transparencia, con qué grado de supervisión humana y con qué consecuencias para la trayectoria académica, la reputación, la autonomía intelectual o las oportunidades educativas de las personas afectadas.

La propuesta resultante tiene carácter normativo-conceptual. No pretende medir empíricamente la frecuencia de uso de estos sistemas ni evaluar plataformas específicas, sino ofrecer criterios jurídicos y educativos para examinarlos, regularlos y someterlos a garantías institucionales. Esta delimitación es importante porque el artículo no sostiene que las inferencias cognitivas sean idénticas a los neurodatos, ni que toda IA educativa constituya una amenaza para los neuroderechos. Su propósito es formular una categoría intermedia que permita evitar tanto la reducción de la protección de la mente a la neurotecnología invasiva como la expansión indiscriminada de los neuroderechos hacia cualquier dato educativo.

La investigación presenta tres límites principales. El primero es de alcance regional: no analiza de forma exhaustiva todas las legislaciones latinoamericanas sobre protección de datos, inteligencia artificial, neuroderechos o educación superior. La comparación Chile-Ecuador funciona como dispositivo analítico para abrir una agenda

regional, no como mapa normativo completo de América Latina. El segundo límite es de alcance empírico: el artículo no estudia plataformas universitarias concretas ni recoge percepciones de estudiantes, docentes o autoridades. El tercero es de alcance conceptual: la noción de inferencias cognitivas se propone como categoría de riesgo reforzado, no como sustituto de los neurodatos ni como ampliación automática de los neuroderechos. Su utilidad radica en permitir la evaluación de sistemas algorítmicos que, sin acceder directamente al cerebro, pueden incidir sobre la vida mental, la autonomía intelectual y las trayectorias educativas de estudiantes y docentes.

RESULTADOS

Neuroderechos, privacidad mental e inferencias cognitivas: delimitación conceptual

Los neuroderechos pueden entenderse como un conjunto emergente de garantías orientadas a proteger dimensiones de la persona particularmente expuestas por el avance de la neurociencia, la neurotecnología y la inteligencia artificial. Su núcleo se relaciona con la libertad cognitiva, la privacidad mental, la integridad mental, la identidad personal, la agencia y la igualdad. Ienca y Andorno (2017) formularon una de las propuestas más influyentes al sostener que los derechos humanos existentes requieren actualización frente a tecnologías capaces de acceder, registrar, interpretar o modificar procesos mentales. Yuste et al. (2017) reforzaron esta preocupación al vincular la convergencia entre neurotecnologías e IA con riesgos para la privacidad, la identidad, la agencia y la equidad. La privacidad mental no se reduce a la protección de datos cerebrales como información sensible. Implica preservar un espacio de control sobre pensamientos, procesos cognitivos, estados mentales, decisiones internas y posibilidades de autodeterminación. Farahany (2023) propone la libertad cognitiva como una dimensión indispensable para pensar y decidir sin vigilancia o interferencia indebida. Esta perspectiva desplaza la atención desde el dato aislado hacia la libertad de la persona para sostener una vida mental no capturada permanentemente por sistemas de monitoreo, clasificación, predicción o influencia.

El debate exige precisión conceptual. Ligthart et al. (2023) advierten que la categoría de neuroderechos puede perder fuerza si se utiliza para designar cualquier problema digital vinculado con privacidad o autonomía. Esta advertencia resulta especialmente relevante para el campo educativo. No toda herramienta de IA que personaliza contenidos, recomienda recursos o entrega retroalimentación compromete la privacidad mental. Tampoco toda inferencia algorítmica debe considerarse neurodato. Una regulación conceptualmente sólida debe distinguir fuentes, objetos, finalidades, grados de intrusión y efectos institucionales.

Los neurodatos, en sentido estricto, se refieren a información derivada de la estructura o actividad del sistema nervioso, especialmente del cerebro, obtenida mediante neurotecnologías, sensores neuronales, interfaces cerebro-computadora u otros dispositivos capaces de registrar señales cerebrales. Su relevancia jurídica proviene de que pueden revelar información extremadamente íntima y, en ciertos escenarios, permitir inferencias sobre emociones, atención, intención, salud, preferencias o patrones cognitivos. El caso chileno Emotiv Insight muestra de forma clara esta problemática, pues

involucró un dispositivo comercial capaz de registrar actividad eléctrica cerebral y generó discusión judicial sobre privacidad mental, consentimiento, neurodatos y protección constitucional.

Las inferencias cognitivas, en cambio, no provienen necesariamente de señales cerebrales. En este artículo se definen como predicciones, clasificaciones, valoraciones o interpretaciones generadas por sistemas de IA sobre estados, capacidades, comportamientos o procesos asociados con la vida mental, académica o conductual de una persona, a partir de datos no neuronales. En educación superior, pueden derivarse de patrones de escritura, tiempos de conexión, participación en foros, historial académico, interacción con recursos digitales, respuestas en evaluaciones, alertas de integridad académica, uso de asistentes generativos o datos de permanencia estudiantil.

La diferencia entre neurodatos e inferencias cognitivas debe mantenerse. Equipararlas automáticamente debilitaría la especificidad de los neuroderechos y podría generar una expansión conceptual difícil de sostener. Sin embargo, negar relevancia jurídica a las inferencias indirectas también resulta insuficiente. Una clasificación algorítmica de “bajo compromiso”, “riesgo de abandono”, “escritura sospechosa”, “baja atención”, “vulnerabilidad académica” o “probabilidad de incumplimiento” puede afectar reputación, oportunidades, acceso a apoyos, trato institucional o procesos disciplinarios. En estos casos, el riesgo no se ubica en la fuente cerebral del dato, sino en el efecto jurídico, pedagógico y relacional de la interpretación producida.

La literatura sobre inferencias algorítmicas permite reforzar esta tesis. Wachter y Mittelstadt (2019) sostienen que las inferencias producidas por sistemas de big data e IA pueden ser altamente invasivas y decisivas, incluso cuando no hayan sido entregadas directamente por la persona. Esta idea es particularmente relevante para el derecho educativo: un estudiante puede aceptar una plataforma sin comprender que su comportamiento digital será usado para producir perfiles de rendimiento, riesgo, honestidad académica, atención o necesidad de intervención. El consentimiento formal, en estos casos, no equivale necesariamente a comprensión, control significativo ni posibilidad real de oposición.

La categoría de inferencias cognitivas opera como una zona intermedia entre datos personales ordinarios, datos sensibles y neurodatos. No toda inferencia de aprendizaje requiere protección reforzada, pero sí aquellas que cumplen cuatro condiciones: se refieren a aspectos vinculados con atención, emoción, conducta, rendimiento, vulnerabilidad, capacidad, motivación o toma de decisiones; se generan en contextos de asimetría institucional; pueden incidir en decisiones relevantes sobre admisión, evaluación, permanencia, integridad académica, tutoría, becas o apoyos; y no cuentan con información suficiente, supervisión humana efectiva, posibilidad de corrección o mecanismos de impugnación.

Esta delimitación permite evitar dos riesgos teóricos. El primero consiste en reducir la protección de la mente únicamente a neurotecnologías invasivas, dejando fuera formas indirectas de perfilamiento cognitivo que ya operan en sistemas educativos digitales. El segundo consiste en denominar neuroderecho a cualquier problema de datos, lo que podría vaciar de precisión la categoría. Entre ambos riesgos, las inferencias cognitivas permiten formular un criterio de protección reforzada: cuando una

interpretación algorítmica se aproxima a dimensiones de la vida mental y produce consecuencias institucionales relevantes, debe someterse a garantías especiales de finalidad legítima, proporcionalidad, transparencia, supervisión humana, trazabilidad, revisión y reparación.

Estándares internacionales de gobernanza de IA: riesgo, derechos humanos y supervisión humana

La gobernanza jurídica de la inteligencia artificial ha dejado de ser una cuestión exclusivamente ética o técnica. Los instrumentos internacionales recientes muestran una convergencia hacia enfoques basados en riesgos, derechos humanos, transparencia, trazabilidad, supervisión humana y mecanismos de rendición de cuentas. Aunque estos estándares no siempre utilizan el lenguaje de los neuroderechos, ofrecen bases normativas para abordar inferencias cognitivas en educación superior cuando estas afectan privacidad mental, autonomía cognitiva, igualdad de oportunidades o justicia educativa.

El Reglamento (UE) 2024/1689 constituye uno de los referentes más influyentes. Su lógica principal consiste en graduar obligaciones según el nivel de riesgo que un sistema de IA puede generar para la salud, la seguridad o los derechos fundamentales. En educación y formación profesional, el Anexo III incluye sistemas utilizados para determinar acceso, admisión o asignación a instituciones educativas; evaluar resultados de aprendizaje; valorar el nivel educativo apropiado; o monitorear conductas prohibidas durante exámenes. Esta clasificación reconoce que la IA educativa puede afectar oportunidades y trayectorias vitales, no solo procesos administrativos.

El mismo reglamento restringe prácticas particularmente problemáticas, entre ellas ciertos sistemas de reconocimiento de emociones en contextos laborales y educativos, salvo excepciones delimitadas. Este punto es relevante para la privacidad mental: la inferencia automatizada de emociones, atención o estados internos en espacios de asimetría institucional puede convertirse en vigilancia desproporcionada. La prohibición europea no debe trasladarse mecánicamente a América Latina, pero sí ofrece un criterio de prudencia regulatoria: cuanto más íntima o psicológica sea la inferencia y mayor sea su impacto institucional, más fuertes deben ser las garantías.

El Convenio Marco del Consejo de Europa sobre Inteligencia Artificial, Derechos Humanos, Democracia y Estado de Derecho refuerza esta orientación al ubicar la IA dentro del ciclo de vida completo de diseño, desarrollo, implementación y uso. Su aporte consiste en desplazar la discusión desde la autorregulación voluntaria hacia obligaciones públicas e institucionales de prevención, evaluación, mitigación de riesgos, transparencia, rendición de cuentas y acceso a recursos. Para las universidades, esto implica que la responsabilidad no termina al contratar una plataforma externa: la institución que despliega el sistema conserva deberes frente a su comunidad académica.

La Recomendación sobre la Ética de la Inteligencia Artificial de la UNESCO resulta especialmente relevante para el campo educativo. Al enfatizar dignidad humana, derechos humanos, diversidad, inclusión, equidad, protección de datos, transparencia, explicabilidad y supervisión humana, la recomendación ofrece un marco normativo de soft law útil para países latinoamericanos. En educación superior, estos principios exigen

que la IA no se presente únicamente como innovación, eficiencia o modernización institucional, sino como una tecnología que debe evaluarse según su impacto en la autonomía intelectual, la justicia educativa y la participación humana significativa.

Los Principios de IA de la OCDE complementan este núcleo al vincular la IA confiable con derechos humanos, Estado de derecho, privacidad, autonomía, igualdad, equidad, robustez, seguridad, trazabilidad y responsabilidad. Estos criterios permiten articular gobernanza técnica y garantía jurídica. La transparencia no se limita a revelar la existencia de un sistema, sino que debe facilitar comprensión sobre finalidad, funcionamiento general, datos empleados, resultados producidos y vías de cuestionamiento. La supervisión humana no consiste en presencia simbólica, sino en capacidad real de intervenir, corregir, contextualizar y asumir responsabilidad.

De estos estándares se derivan cinco exigencias aplicables a sistemas de IA en sectores sensibles. La primera es la evaluación previa y periódica de riesgos e impactos. La segunda es la información clara y comprensible para las personas afectadas. La tercera es la supervisión humana efectiva, especialmente en decisiones de alto impacto. La cuarta es la trazabilidad y explicabilidad suficiente para auditar e impugnar resultados. La quinta es la existencia de mecanismos de revisión, reparación y rendición de cuentas. Trasladadas a la educación superior, estas exigencias permiten distinguir innovación legítima de perfilamiento opaco.

La utilidad de estos estándares para el debate sobre neuroderechos se encuentra en su capacidad para proteger bienes asociados con la vida mental sin depender exclusivamente de la existencia de neurodatos. Si un sistema universitario infiere riesgo académico, atención, motivación, vulnerabilidad, honestidad o capacidad a partir de datos digitales, debe ser analizado desde la proporcionalidad, la finalidad legítima, la minimización de datos e inferencias, el debido proceso y la posibilidad de revisión humana. De este modo, la gobernanza de IA funciona como puente entre protección de datos, derechos humanos, justicia educativa y privacidad mental.

La educación superior como sector sensible de regulación algorítmica

La educación superior debe considerarse un sector sensible para la gobernanza de IA porque en ella se producen decisiones con efectos directos sobre derechos, oportunidades y trayectorias de vida. Las universidades no solo transmiten contenidos; también admiten, evalúan, certifican, sancionan, orientan, asignan apoyos y construyen reputación académica. Cuando sistemas automatizados intervienen en estos procesos, sus resultados pueden afectar el acceso al conocimiento, la movilidad social, la identidad profesional, la autonomía intelectual y las posibilidades de inserción laboral.

La literatura sobre IA en educación superior muestra beneficios potenciales. Holmes et al. (2019) señalan que los sistemas inteligentes pueden apoyar procesos de retroalimentación, tutoría, personalización y evaluación. No obstante, Zawacki-Richter et al. (2019) advierten que buena parte de la investigación se ha concentrado en aplicaciones técnicas, mientras que las implicaciones pedagógicas, éticas y sociales han recibido menor atención. Esta brecha es jurídicamente relevante: una herramienta puede ser eficaz desde el punto de vista predictivo y, al mismo tiempo, problemática si opera sin

explicación, reproduce sesgos o genera consecuencias institucionales injustas.

La analítica de aprendizaje transforma la forma en que las instituciones observan y gobiernan el aprendizaje. Williamson (2017) sostiene que el Big Data educativo convierte prácticas formativas en objetos de medición, predicción y administración. Bajo esta lógica, el estudiante puede aparecer ante la institución como un conjunto de registros, probabilidades y alertas. Slade y Prinsloo (2013) advierten que la analítica de aprendizaje plantea dilemas éticos sobre consentimiento, transparencia, responsabilidad institucional, privacidad y posibilidad de daño, especialmente porque las universidades no solo observan a los estudiantes, sino que actúan sobre ellos a partir de los datos producidos.

La sensibilidad del sector universitario puede explicarse mediante tres dimensiones. La primera es decisional: la IA puede influir en admisión, evaluación, permanencia, becas, integridad académica, titulación o acceso a apoyos. La segunda es informacional: las universidades procesan datos personales, académicos, conductuales y, en algunos casos, datos o inferencias sobre vulnerabilidad, salud, discapacidad, emociones o condiciones socioeconómicas. La tercera es formativa: la educación superior debe promover pensamiento crítico, autonomía intelectual, producción de conocimiento y deliberación. Cuando una tecnología incide simultáneamente en decisiones, información personal y formación de la autonomía, las garantías deben ser reforzadas.

La integridad académica constituye un campo particularmente delicado. Los sistemas de detección de similitud, *proctoring*, análisis de escritura o identificación de uso de IA generativa pueden producir reportes con consecuencias disciplinarias. Si operan de manera opaca, sin comunicar márgenes de error o sin revisión humana suficiente, pueden afectar la presunción de buena fe, la reputación académica, el debido proceso y el derecho a impugnar. En estos casos, el problema no es únicamente la privacidad; también está en juego la justicia procedimental dentro de la universidad.

La predicción de riesgo académico también exige cautela. Los modelos predictivos pueden facilitar apoyo oportuno, pero también pueden producir etiquetamiento temprano, discriminación indirecta o profecías autocumplidas. Clasificar a un estudiante como “en riesgo” sin considerar conectividad, responsabilidades familiares, territorio, lengua, discapacidad, situación socioeconómica o trayectoria previa puede transformar desigualdades contextuales en atributos individuales. En América Latina, donde estas desigualdades atraviesan la experiencia universitaria, la gobernanza de IA debe incorporar criterios de justicia educativa y no solo de precisión estadística.

La supervisión humana es indispensable, pero debe ser pedagógica y sustantiva. Selwyn (2019) cuestiona las narrativas que presentan la automatización como sustituto del juicio docente. La educación implica interpretación contextual, cuidado, deliberación y responsabilidad humana. Un docente, autoridad o equipo institucional que revisa una alerta algorítmica no debe limitarse a aceptar el resultado, sino examinar su base, margen de error, sesgo posible, pertinencia pedagógica y consecuencias para la persona afectada.

La educación superior es sensible porque combina datos, decisiones, poder institucional y formación de la autonomía. Los sistemas de IA no solo median procesos pedagógicos o administrativos; pueden participar en la construcción de perfiles sobre conducta, desempeño, riesgo o capacidad. Las universidades que adoptan IA requieren mecanismos institucionales de evaluación de impacto, información clara, minimización

de inferencias, auditoría, supervisión humana y revisión de decisiones, especialmente cuando las inferencias generadas pueden afectar privacidad mental, autonomía cognitiva, trayectoria académica o justicia educativa.

América Latina: Chile, Ecuador y la necesidad de una gobernanza situada

América Latina enfrenta la expansión de la IA desde condiciones estructurales específicas: desigualdad social, brechas digitales, dependencia tecnológica, diversidad cultural y lingüística, y capacidades desiguales de auditoría pública e institucional. Estas condiciones hacen insuficiente una simple importación de modelos regulatorios externos. La región necesita una gobernanza situada que traduzca estándares internacionales a contextos educativos concretos, especialmente cuando universidades públicas y privadas adoptan plataformas diseñadas por proveedores globales con lógicas comerciales, contractuales y técnicas poco transparentes.

Chile ocupa un lugar central en el debate regional por la Ley 21.383, que modificó la Constitución para establecer que el desarrollo científico y tecnológico debe realizarse con respeto a la vida y a la integridad física y psíquica, y que la ley debe resguardar especialmente la actividad cerebral y la información proveniente de ella. Este reconocimiento es relevante porque introduce la actividad cerebral como objeto explícito de protección constitucional y posiciona a Chile como referente internacional en neuroderechos.

La decisión de la Corte Suprema chilena en el caso *Emotiv Insight* profundizó esta discusión. El caso involucró un dispositivo comercial capaz de registrar actividad eléctrica cerebral y permitió discutir judicialmente privacidad mental, neurodatos y protección constitucional. Cornejo-Plaza et al. (2024) sostienen que este precedente muestra el tránsito de los neuroderechos desde la bioética hacia el derecho constitucional, la protección de datos y la litigación estratégica. Su importancia para este artículo radica en que muestra cómo una tecnología comercial puede convertir actividad cerebral en información tratable, transferible y jurídicamente sensible.

La experiencia chilena también revela un límite: el foco en neurodatos directos no agota los riesgos contemporáneos para la vida mental. Muchas tecnologías educativas no registran actividad cerebral, pero sí infieren patrones de atención, conducta, rendimiento, riesgo o vulnerabilidad. La ruta chilena ofrece un antecedente fuerte para neurodatos y, al mismo tiempo, abre la discusión sobre formas indirectas de perfilamiento cognitivo que pueden producir efectos semejantes sobre privacidad mental y autonomía cognitiva sin operar mediante neurotecnologías invasivas.

Ecuador ofrece una ruta distinta. La Constitución reconoce el derecho a la protección de datos personales, incluyendo acceso y decisión sobre información de este carácter, así como su protección frente a recolección, archivo, procesamiento, distribución o difusión sin autorización del titular o mandato legal. La Ley Orgánica de Protección de Datos Personales desarrolla esta garantía mediante principios de finalidad, proporcionalidad, transparencia, minimización, seguridad, confidencialidad, responsabilidad proactiva y derechos de los titulares. Este marco no usa la terminología de neuroderechos, pero ofrece una base para controlar tratamientos de datos que pueden

alimentar sistemas de IA y producir inferencias sobre estudiantes, docentes o comunidades académicas.

El desarrollo ecuatoriano más específico proviene de la Resolución N.º SPDP-SPD-2026-0009-R, expedida por la Superintendencia de Protección de Datos Personales, que establece una norma general para garantizar el derecho a la protección de datos personales en el uso de sistemas de inteligencia artificial. Su relevancia radica en que incorpora la IA al campo regulatorio ecuatoriano desde un enfoque de protección de datos, riesgos y obligaciones para responsables y encargados. Este enfoque puede ser especialmente útil en educación superior si se combina con evaluación de impacto, protección desde el diseño, auditoría y mecanismos de supervisión institucional.

La Ley Orgánica de Educación Superior agrega otra dimensión. Al establecer principios como autonomía responsable, igualdad de oportunidades, calidad, pertinencia, integralidad y autodeterminación para la producción del pensamiento y conocimiento, ofrece una base sectorial para vincular IA, autonomía cognitiva y finalidad universitaria. La protección de la mente en educación superior no se limita a evitar filtraciones de datos; también exige preservar condiciones para pensar, aprender, investigar y producir conocimiento sin quedar reducido a perfiles algorítmicos opacos.

La comparación Chile-Ecuador permite identificar dos rutas complementarias. Chile aporta una vía neuroconstitucional, centrada en actividad cerebral y neurodatos. Ecuador aporta una vía de gobernanza de datos e IA, articulada con educación superior y autodeterminación del pensamiento. Ambas rutas son incompletas si se consideran de forma aislada: la primera puede ser demasiado estrecha frente a inferencias no neuronales; la segunda puede ser insuficiente si no reconoce los riesgos específicos para privacidad mental y autonomía cognitiva. Su articulación permite construir una agenda regional más robusta para proteger la mente en contextos educativos mediados por IA.

Tabla 1. Rutas normativas para la protección de la mente en Chile y Ecuador

País	Base normativa principal	Aporte al debate	Límite o desafío
Chile	Ley 21.383 y jurisprudencia sobre Emotiv Insight	Reconoce la necesidad de resguardar la actividad cerebral y la información proveniente de ella; desplaza los neuroderechos hacia el derecho constitucional y la protección judicial.	Su foco principal está en neurodatos directos; requiere ampliarse con cautela hacia riesgos algorítmicos indirectos.
Ecuador	Constitución, Ley Orgánica de Protección de Datos Personales, Resolución N.º SPDP-SPD-2026-0009-R y Ley Orgánica de Educación Superior	Permite articular protección de datos, regulación emergente de IA y autodeterminación para la producción del pensamiento y conocimiento.	Debe traducir principios generales en mecanismos universitarios concretos de evaluación de impacto, auditoría y revisión.

Nota. Elaboración propia a partir del análisis doctrinal y normativo.

Propuesta normativa: deberes institucionales para universidades que implementan IA

A partir del análisis desarrollado, este artículo propone un marco de deberes institucionales para universidades latinoamericanas que incorporan sistemas automatizados o asistidos por inteligencia artificial en procesos académicos, administrativos, predictivos o disciplinarios. La propuesta no busca prohibir la IA educativa ni negar sus beneficios potenciales para la retroalimentación, el acompañamiento estudiantil, la personalización del aprendizaje o la gestión institucional (Holmes et al., 2019; Zawacki-Richter et al., 2019). Su finalidad es establecer condiciones jurídicas y pedagógicas mínimas para que la innovación tecnológica sea compatible con los derechos fundamentales, la protección de datos, la justicia educativa, la autonomía cognitiva, la privacidad mental y el debido proceso.

La premisa central es que las universidades no son usuarias pasivas de tecnologías desarrolladas por terceros. Cuando incorporan sistemas de IA en admisión, evaluación, permanencia, integridad académica, tutoría, escritura, retroalimentación, apoyo estudiantil o predicción de riesgo, asumen deberes frente a estudiantes, docentes y demás miembros de la comunidad académica. Esta responsabilidad se intensifica cuando los sistemas producen inferencias sobre atención, conducta, rendimiento, vulnerabilidad, motivación, capacidad, escritura, honestidad académica o toma de decisiones. En tales casos, la cuestión no se reduce al tratamiento de datos personales, sino que alcanza dimensiones vinculadas con la vida mental, la autonomía intelectual y las trayectorias educativas.

El primer deber es la finalidad legítima y proporcionalidad. La institución debe justificar qué problema educativo busca resolver, por qué la IA resulta adecuada para ese fin y si existen alternativas menos intrusivas. No basta invocar eficiencia, innovación, modernización o personalización del aprendizaje. Si el sistema puede afectar evaluación, admisión, permanencia, becas, integridad académica o acceso a apoyos, la finalidad debe ser concreta, verificable y compatible con derechos fundamentales. Este deber se relaciona con el enfoque basado en riesgos del Reglamento (UE) 2024/1689, que reconoce determinados usos de IA en educación y formación profesional como sistemas de alto riesgo cuando inciden en oportunidades o trayectorias educativas.

El segundo deber es la información clara y comprensible. Estudiantes y docentes deben conocer qué sistema se utiliza, con qué finalidad, qué datos procesa, qué inferencias genera, quién accede a los resultados, durante cuánto tiempo se conserva la información y qué consecuencias institucionales pueden producirse. Las políticas generales de privacidad o los términos extensos de uso no satisfacen este deber si no explican las inferencias de salida y sus posibles efectos académicos, administrativos o disciplinarios. La transparencia, en este contexto, no consiste únicamente en revelar que existe IA, sino en permitir que la persona comprenda cómo el sistema puede afectar su trayectoria educativa, su reputación académica o su relación con la institución (UNESCO, 2022; OECD, 2024).

El tercer deber es el consentimiento informado reforzado. En contextos educativos obligatorios o altamente asimétricos, el consentimiento formal puede ser insuficiente, porque el estudiante o docente no siempre dispone de condiciones reales para rechazar el uso de una plataforma sin sufrir desventajas. Cuando un sistema produce inferencias sobre vida académica, conducta, rendimiento, vulnerabilidad, atención, escritura o riesgo,

la persona debe recibir información específica, comprensible y contextualizada. Cuando sea posible, deben ofrecerse opciones razonables, mecanismos de oposición o alternativas pedagógicas. Este consentimiento no puede justificar prácticas desproporcionadas ni reemplazar la evaluación institucional de impacto. Su función es fortalecer la autodeterminación informativa y la autonomía cognitiva, no legitimar formas opacas de perfilamiento.

El cuarto deber es la evaluación de impacto algorítmico y de derechos fundamentales. Antes del despliegue de sistemas de IA en procesos educativos relevantes, la universidad debe analizar la finalidad del sistema, los datos utilizados, las inferencias producidas, los grupos afectados, los posibles sesgos, las consecuencias institucionales, las medidas de mitigación, la supervisión humana, la trazabilidad y los mecanismos de revisión. Mantelero (2022) sostiene que las evaluaciones de impacto permiten superar una comprensión meramente procedimental de la protección de datos, al incorporar dimensiones éticas, sociales y jurídicas. En educación superior, esta evaluación debe considerar también efectos pedagógicos: etiquetamiento temprano, autocensura, ansiedad evaluativa, debilitamiento de la autonomía intelectual o afectación de la relación docente-estudiante. En usos de alto impacto, la evaluación debe actualizarse periódicamente y después de incidentes, reclamaciones o cambios sustantivos en el sistema.

El quinto deber es la minimización de datos e inferencias. La protección no debe limitarse a recolectar menos información; también debe limitar la producción de perfiles innecesarios sobre la vida académica, conductual o cognitiva de las personas. No todo lo que puede inferirse debe inferirse. Una universidad puede necesitar indicadores de desempeño para acompañar aprendizajes, pero no necesariamente debe inferir emociones, atención permanente, personalidad, vulnerabilidad o probabilidad de incumplimiento sin una justificación pedagógica estricta. Este deber adapta el principio clásico de minimización al campo de la privacidad mental: reducir no solo los datos de entrada, sino también las interpretaciones algorítmicas que pueden convertirse en vigilancia cognitiva o etiquetamiento institucional.

El sexto deber es la supervisión humana pedagógica. La intervención humana debe ser real, informada y contextual. Quien revisa una alerta, clasificación o recomendación generada por IA debe comprender los límites del sistema, su margen de error, sus posibles sesgos y sus consecuencias para la persona afectada. La IA puede apoyar el juicio docente o institucional, pero no sustituirlo en decisiones que afecten evaluación, integridad académica, permanencia, becas, sanciones o acceso a oportunidades. Selwyn (2019) advierte que la educación no puede reducirse a automatización, predicción o eficiencia, porque implica interpretación, cuidado, deliberación y responsabilidad humana. Desde esta perspectiva, la supervisión humana pedagógica protege no solo el debido proceso, sino también la finalidad formativa de la educación superior.

El séptimo deber es la trazabilidad, explicabilidad y auditoría. Las instituciones deben documentar qué sistemas usan, quién los provee, qué datos emplean, qué resultados generan, cómo se interpretan esos resultados y de qué manera influyen en decisiones institucionales. La explicabilidad debe ser suficiente para que una persona afectada comprenda razonablemente por qué se produjo una clasificación, alerta o recomendación.

Pasquale (2015) ha mostrado que los sistemas algorítmicos opacos pueden funcionar como cajas negras que dificultan la auditoría, la impugnación y la rendición de cuentas. En el contexto universitario, esta opacidad puede afectar de manera directa la reputación académica, la integridad, la permanencia o el acceso a apoyos. La auditoría debe ser técnica, ética, jurídica y pedagógica, especialmente cuando intervienen proveedores externos cuyas lógicas de diseño, entrenamiento o clasificación no son plenamente conocidas por la institución.

El octavo deber es la revisión, impugnación y reparación. Toda persona afectada por una decisión automatizada o asistida por IA debe poder conocer la existencia de la inferencia, solicitar explicación, aportar contexto, corregir información errónea y obtener revisión humana significativa. Este deber es especialmente relevante cuando la inferencia activa sospechas de fraude académico, clasificaciones de bajo rendimiento, alertas de riesgo, exclusión de beneficios, restricciones de acceso o decisiones disciplinarias. Wachter y Mittelstadt (2019) advierten que las inferencias algorítmicas pueden ser invasivas y decisivas incluso cuando no han sido proporcionadas directamente por la persona. Si una inferencia produce daño, discriminación, exclusión o afectación reputacional, la universidad debe corregir la decisión, eliminar etiquetas indebidas, reparar consecuencias y modificar el sistema para evitar repetición.

Tabla 2. Matriz de deberes institucionales para proteger la privacidad mental y la autonomía cognitiva en educación superior mediada por IA

Deber institucional	Riesgo que previene	Aplicación universitaria	Bien jurídico protegido
Finalidad legítima y proporcionalidad	Uso expansivo o injustificado de IA	Justificar sistemas aplicados a evaluación, tutoría, permanencia, becas o integridad académica.	Autonomía, igualdad y debido proceso
Información clara	Opacidad sobre datos, inferencias y consecuencias	Explicar finalidad, datos procesados, perfiles generados, destinatarios, conservación y vías de revisión.	Transparencia y autodeterminación informativa
Consentimiento informado reforzado	Perfilamiento invisible o aceptación meramente formal	Informar inferencias sobre conducta, rendimiento, atención, vulnerabilidad o escritura; ofrecer alternativas cuando sea posible.	Privacidad mental y autonomía cognitiva
Evaluación de impacto	Daños no anticipados, sesgos o discriminación	Analizar riesgos antes del despliegue y periódicamente en usos de alto impacto.	Derechos fundamentales y justicia educativa
Minimización de datos e inferencias	Vigilancia cognitiva o psicológica innecesaria	Evitar perfiles sobre emociones, atención, personalidad o vulnerabilidad sin necesidad pedagógica estricta.	Privacidad, proporcionalidad y dignidad

Deber institucional	Riesgo que previene	Aplicación universitaria	Bien jurídico protegido
Supervisión humana pedagógica	Delegación acrítica de decisiones educativas	Exigir revisión contextual por docentes o autoridades con capacidad de corregir el resultado algorítmico.	Agencia humana y finalidad educativa
Trazabilidad, explicabilidad y auditoría	Cajas negras, imposibilidad de control o dependencia del proveedor	Documentar sistemas, datos, resultados, criterios de interpretación, acceso y decisiones derivadas.	Rendición de cuentas y control institucional
Revisión, impugnación y reparación	Indefensión frente a alertas, sanciones o etiquetas erróneas	Permitir explicación, corrección, revisión humana significativa, reparación y eliminación de etiquetas indebidas.	Tutela efectiva, reputación y debido proceso

Nota. La matriz operacionaliza estándares internacionales de IA y principios de protección de datos en el contexto universitario latinoamericano.

DISCUSIÓN

El artículo propone una extensión prudente del debate sobre neuroderechos hacia sistemas de inteligencia artificial que no producen neurodatos en sentido estricto, pero que pueden generar interpretaciones sobre dimensiones cognitivas, conductuales o académicas con efectos institucionales relevantes. Esta propuesta dialoga con la advertencia de Lighthart et al. (2023) sobre la necesidad de delimitar cuidadosamente los neuroderechos, a fin de evitar tanto su expansión indiscriminada como su reducción exclusiva a tecnologías neuronales invasivas. La cuestión central no es convertir toda inferencia educativa en un problema de neuroderechos, sino reconocer que ciertas inferencias algorítmicas pueden afectar bienes jurídicos próximos a la privacidad mental, la autonomía cognitiva y la justicia educativa cuando inciden en decisiones universitarias de alto impacto.

La primera contribución teórica consiste en desplazar parcialmente la atención desde la fuente del dato hacia el efecto de la inferencia. En el derecho de protección de datos, la clasificación suele iniciar con la naturaleza de la información tratada: dato personal, dato sensible, dato biométrico, dato genético o neurodato. Esa distinción sigue siendo necesaria, pero resulta insuficiente en entornos algorítmicos donde una inferencia generada a partir de datos aparentemente ordinarios puede producir consecuencias más invasivas que el dato original. Wachter y Mittelstadt (2019) han mostrado que las inferencias producidas por sistemas de big data e IA pueden ser decisivas, opacas y difíciles de controlar, incluso cuando no han sido proporcionadas directamente por la persona. Desde esta perspectiva, la gobernanza de IA debe considerar tanto los datos de entrada como los perfiles de salida y sus efectos sobre la vida académica.

La segunda contribución se ubica en la educación superior. Buena parte del debate sobre IA educativa se ha concentrado en innovación pedagógica, eficiencia administrativa, personalización del aprendizaje o prevención del fraude académico. Sin negar estos usos, el artículo sostiene que la universidad debe analizarse también como

espacio de poder jurídico y pedagógico. Las decisiones universitarias distribuyen oportunidades, producen reputación, condicionan trayectorias académicas y profesionales, y pueden activar consecuencias disciplinarias. Por ello, cuando un sistema de IA participa en procesos de admisión, evaluación, permanencia, integridad académica, becas o acompañamiento estudiantil, no basta con evaluar su precisión técnica. También deben examinarse sus efectos sobre igualdad, autonomía intelectual, debido proceso, confianza institucional y justicia educativa.

Esta discusión se vincula con las advertencias de Slade y Prinsloo (2013) sobre los dilemas éticos de la analítica de aprendizaje, especialmente en relación con consentimiento, transparencia, responsabilidad institucional, privacidad y posibilidad de daño. También dialoga con Zawacki-Richter et al. (2019), quienes muestran que la investigación sobre IA en educación superior ha privilegiado con frecuencia el desarrollo técnico por encima de sus implicaciones pedagógicas, éticas y sociales. La categoría de inferencias cognitivas permite conectar ambos debates: no se trata solo de preguntar si una herramienta funciona, sino qué tipo de estudiante construye, qué decisiones habilita, qué riesgos invisibiliza y qué posibilidades de revisión ofrece.

La tercera contribución es regional. América Latina no puede limitarse a importar modelos regulatorios europeos, aunque estos resulten útiles como referencia normativa. La región requiere una lectura situada que considere brechas digitales, desigualdad socioeconómica, dependencia tecnológica, diversidad lingüística y cultural, y capacidades institucionales limitadas para auditar sistemas algorítmicos. Un modelo predictivo entrenado, validado o comercializado desde contextos distintos puede etiquetar como riesgo individual aquello que expresa exclusión territorial, conectividad precaria, responsabilidades de cuidado, desigualdad lingüística, discapacidad, pobreza o trayectorias educativas históricamente desiguales. La justicia algorítmica universitaria, en este marco, debe ser también justicia contextual.

La comparación entre Chile y Ecuador muestra que la protección de la mente puede construirse por vías diferentes. Chile ofrece una formulación explícita de resguardo de la actividad cerebral y de la información proveniente de ella, lo que fortalece la respuesta frente a neurotecnologías directas y neurodatos comerciales. Ecuador, en cambio, permite explorar una ruta basada en protección de datos personales, regulación administrativa emergente sobre IA y principios de educación superior vinculados con igualdad de oportunidades, autonomía responsable y autodeterminación para la producción del pensamiento y conocimiento. Esta segunda vía puede ser particularmente fértil para abordar inferencias cognitivas indirectas, siempre que no se limite al cumplimiento documental de políticas de privacidad.

La propuesta también presenta límites. El concepto de inferencias cognitivas requiere mayor desarrollo empírico, normativo y metodológico. Futuras investigaciones podrían estudiar casos concretos de analítica de aprendizaje, proctoring, sistemas de detección de IA generativa, asistentes de escritura o modelos predictivos de permanencia en universidades latinoamericanas. También sería necesario analizar cómo estudiantes, docentes, autoridades y proveedores tecnológicos comprenden los riesgos asociados con privacidad mental, autonomía cognitiva y perfilamiento académico. La matriz de deberes

institucionales, además, debe adaptarse a las capacidades, tamaños, recursos y marcos normativos de cada institución.

El enfoque propuesto no implica que todos los usos de IA requieran el mismo nivel de control. Un sistema de apoyo gramatical, una herramienta de organización de contenidos o un asistente de retroalimentación formativa no tienen el mismo impacto que un sistema utilizado para detectar fraude, predecir abandono, restringir apoyos, clasificar vulnerabilidad o activar procedimientos disciplinarios. La intensidad de la regulación debe ajustarse al nivel de riesgo, al tipo de inferencia producida, al grado de opacidad del sistema, a la posibilidad de revisión humana y a las consecuencias institucionales de la decisión.

La categoría de inferencias cognitivas permite anticipar problemas antes de que se consoliden prácticas universitarias difíciles de revertir. La historia de la protección de datos muestra que el derecho suele responder tarde a tecnologías ya normalizadas. En educación superior, esa demora puede traducirse en estudiantes etiquetados, sancionados, excluidos o tratados de manera diferenciada por sistemas que no comprenden, no pueden auditar y no logran impugnar. Una gobernanza preventiva, participativa y pedagógicamente informada permitiría reducir estos riesgos sin bloquear usos legítimos de la IA.

La discusión sobre privacidad mental en educación superior no debe interpretarse como resistencia a la innovación. Una gobernanza robusta puede aumentar la confianza institucional, mejorar la calidad de las decisiones, proteger a estudiantes y docentes, y orientar el desarrollo tecnológico hacia fines educativos legítimos. La IA puede apoyar tutorías, retroalimentación, accesibilidad, acompañamiento y gestión académica, pero su incorporación será compatible con una universidad democrática solo si permanece subordinada a la dignidad humana, la autonomía cognitiva, la justicia educativa, la supervisión humana significativa y el derecho a cuestionar decisiones que afectan la vida académica.

CONCLUSIONES

Este artículo analizó la relación entre neuroderechos, inteligencia artificial y educación superior en América Latina, con énfasis en la protección de la privacidad mental, la autonomía cognitiva y la justicia algorítmica frente a inferencias producidas por sistemas automatizados en contextos formativos. El argumento central sostuvo que ciertos sistemas de IA, aun sin acceder directamente a la actividad cerebral, pueden producir clasificaciones, predicciones o valoraciones sobre conducta, rendimiento, vulnerabilidad, atención, escritura, honestidad académica o toma de decisiones, con efectos relevantes sobre estudiantes y docentes.

La principal contribución del trabajo consiste en proponer las inferencias cognitivas como una categoría de riesgo reforzado. Estas inferencias no son neurodatos en sentido estricto y no deben ser tratadas automáticamente como tales. Su relevancia surge cuando se generan en contextos de asimetría institucional y pueden incidir en admisión, evaluación, permanencia, integridad académica, becas, tutorías, apoyos o procedimientos disciplinarios. En tales casos, requieren garantías específicas de finalidad legítima, proporcionalidad, transparencia, minimización, supervisión humana, trazabilidad, revisión e impugnación.

El análisis permitió situar la educación superior como un sector sensible de regulación algorítmica. La universidad no solo procesa datos; también forma sujetos, certifica saberes, distribuye oportunidades, produce reputación académica y adopta decisiones con impacto personal, profesional y social. Por ello, los sistemas de IA utilizados en este ámbito deben evaluarse desde criterios de derechos fundamentales, justicia educativa y finalidad pedagógica, y no únicamente desde parámetros de eficiencia, precisión predictiva o innovación tecnológica.

La lectura latinoamericana mostró la necesidad de una gobernanza situada. Chile aporta un antecedente relevante mediante la protección constitucional de la actividad cerebral y de la información proveniente de ella, así como a través del debate judicial sobre neurodatos comerciales. Ecuador muestra el potencial de articular protección de datos personales, regulación emergente de IA y principios de educación superior orientados a la autodeterminación para la producción del pensamiento y conocimiento. Ambos casos permiten sostener que la protección de la mente en la región debe atender tanto a neurotecnologías directas como a inferencias algorítmicas indirectas en sectores sensibles.

La propuesta normativa formulada en el artículo plantea ocho deberes institucionales para universidades que implementan IA: finalidad legítima y proporcionalidad, información clara, consentimiento informado reforzado, evaluación de impacto, minimización de datos e inferencias, supervisión humana pedagógica, trazabilidad, explicabilidad y auditoría, y revisión, impugnación y reparación. Estos deberes no pretenden obstaculizar la innovación educativa, sino asegurar que su despliegue sea compatible con dignidad humana, privacidad mental, autonomía cognitiva, debido proceso e igualdad de oportunidades.

Proteger la mente en la educación superior mediada por IA exige evitar dos riesgos paralelos: sobredimensionar los neuroderechos hasta perder precisión conceptual y reducir la gobernanza algorítmica a un cumplimiento formal de protección de datos. Entre ambos extremos, la categoría de inferencias cognitivas ofrece una vía intermedia para anticipar daños, orientar políticas universitarias y preservar la capacidad de las personas para pensar, aprender, decidir y producir conocimiento sin quedar reducidas a perfiles algorítmicos opacos, descontextualizados o inimpugnables.

REFERENCIAS

AGUERO, L. de O.; PACÍFICO, M. Neuroeducação e os Benefícios do Bilinguismo. **Observatório De La Economía Latinoamericana**, [S. l.], v. 23, n. 11, p. e12318, 2025. DOI: 10.55905/oelv23n11-130. Disponível em: <https://ojs.observatoriolatinoamericano.com/ojs/index.php/olel/article/view/12318>. Acesso em: 6 jul. 2026.

BINNS, R. Fairness in machine learning: lessons from political philosophy. **Proceedings of Machine Learning Research**, v. 81, p. 149-159, 2018. Disponível em: <http://proceedings.mlr.press/v81/binns18a.html>. Acesso em: 07 jul. 2026.

CHILE. **Ley n.º 21.383, de 25 de octubre de 2021**. Modifica la Carta Fundamental para establecer el desarrollo científico y tecnológico al servicio de las personas. Biblioteca del Congreso Nacional de Chile, Santiago, 2021. Disponível em: <https://www.bcn.cl/leychile/navegar?idNorma=1166983>. Acesso em: 07 jul. 2026.

CONSTITUCIÓN DE LA REPÚBLICA DEL ECUADOR. **Registro Oficial n.º 449**, Quito, 20 oct. 2008.

CORNEJO-PLAZA, M. I.; CIPPITANI, R.; PASQUINO, V. Chilean Supreme Court ruling on the protection of brain activity: neurorights, personal data protection, and neurodata. **Frontiers in Psychology**, Lausanne, v. 15, art. 1330439, 2024. <https://doi.org/10.3389/fpsyg.2024.1330439>. Acesso em: 07 jul. 2026.

COUNCIL OF EUROPE. **Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law**. Council of Europe Treaty Series, n. 225. Strasbourg: Council of Europe, 2024. Disponível em: <https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence>. Acesso em: 07 jul. 2026.

CRAWFORD, K. **Atlas of AI: Power, Politics, and the Planetary Costs Of Artificial intelligence**. New Haven: Yale University Press, 2021.

ECUADOR. **Ley Orgánica de Educación Superior**. Registro Oficial Suplemento n.º 298, Quito, 12 oct. 2010.

ECUADOR. **Ley Orgánica de Protección de Datos Personales**. Registro Oficial Suplemento n.º 459, Quito, 26 mayo 2021.

ECUADOR. Superintendencia de Protección de Datos Personales. **Resolución N.º SPDP-SPD-2026-0009-R, de 12 de febrero de 2026**. Norma General para la Garantía del Derecho de Protección de Datos Personales en el Uso de Sistemas de Inteligencia Artificial. Quito: Superintendencia de Protección de Datos Personales, 2026.

FARAHANY, N. A. **The Battle for your Brain: Defending the Right to Think Freely in the Age of Neurotechnology**. New York: St. Martin's Press, 2023.

FLORIDI, L.; COWLS, J. A unified framework of five principles for AI in society. **Harvard Data Science Review**, Cambridge, v. 1, n. 1, 2019. <https://doi.org/10.1162/99608f92.8cd550d1>. Acesso em: 07 jul. 2026.

HOLMES, W.; BIALIK, M.; FADEL, C. **Artificial intelligence in education: promises and implications for teaching and learning**. Boston: Center for Curriculum Redesign, 2019.

IENCA, M.; ANDORNO, R. Towards new human rights in the age of neuroscience and neurotechnology. **Life Sciences, Society and Policy**, London, v. 13, art. 5, 2017. <https://doi.org/10.1186/s40504-017-0050-1>. Acesso em: 07 jul. 2026.

LIGTHART, S.; IENCA, M.; MEYNEN, G.; MOLNAR-GABOR, F.; ANDORNO, R.; BUBLITZ, C.; CATLEY, P.; CLAYDON, L.; DOUGLAS, T.; FARAHANY, N.; FINS, J. J.; GOERING, S.; HASELAGER, P.; JOTTERAND, F.; LAVAZZA, A.; MCCAY, A.; WAJNERMAN PAZ, A.; RAINEY, S.; RYBERG, J.; KELLMEYER, P. Minding rights: mapping ethical and legal foundations of “neurorights”. **Cambridge Quarterly of Healthcare Ethics**, Cambridge, v. 32, n. 4, p. 461-481, 2023. <https://doi.org/10.1017/S0963180123000245>. Acesso em: 07 jul. 2026.

MANTELERO, A. **Beyond data: Human Rights, Ethical and Social Impact Assessment in AI**. The Hague: T.M.C. Asser Press, 2022. <https://doi.org/10.1007/978-94-6265-531-7>. Acesso em: 07 jul. 2026.

NATIONAL COMMISSION FOR THE PROTECTION OF HUMAN SUBJECTS OF BIOMEDICAL AND BEHAVIORAL RESEARCH. **The Belmont Report: Ethical Principles and Guidelines for the Protection of Human Subjects of Research**. Washington, DC: U.S. Department of Health, Education, and Welfare, 1979. Disponível em: <https://www.hhs.gov/ohrp/regulations-and-policy/belmont-report/index.html>. Acesso em: 07 jul. 2026.

OECD. **OECD AI Principles**. Paris: OECD.AI Policy Observatory, 2024. Disponível em: <https://oecd.ai/en/ai-principles>. Acesso em: 07 jul. 2026.

PARLAMENTO EUROPEO; CONSEJO DE LA UNIÓN EUROPEA. **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. Diario Oficial de la Unión Europea, 2024.

PASQUALE, F. **The black box society: the secret algorithms that control money and information**. Cambridge: Harvard University Press, 2015.

SELWYN, N. **Should Robots Replace Teachers? AI and the Future of Education**. Cambridge: Polity Press, 2019.

SLADE, S.; PRINSLOO, P. Learning analytics: ethical issues and dilemmas. **American Behavioral Scientist**, Thousand Oaks, v. 57, n. 10, p. 1510-1529, 2013. <https://doi.org/10.1177/0002764213479366>. Acesso em: 07 jul. 2026.

UNESCO. **Recomendación sobre la ética de la inteligencia artificial**. Paris: UNESCO, 2022. Disponível em: https://unesdoc.unesco.org/ark:/48223/pf0000380455_spa. Acesso em: 07 jul. 2026.

UNITED STATES GOVERNMENT PRINTING OFFICE. **Trials of war criminals before the Nuernberg Military Tribunals under Control Council Law No. 10: vol. 2. The Medical Case**. Washington, DC: United States Government Printing Office, 1949. Disponível em: https://www.loc.gov/rr/frd/Military_Law/pdf/NT_war-criminals_Vol-II.pdf. Acesso em: 07 jul. 2026.

WACHTER, S.; MITTELSTADT, B. A right to reasonable inferences: re-thinking data protection law in the age of Big Data and AI. **Columbia Business Law Review**, New York, v. 2019, n. 2, p. 494-620, 2019. <https://doi.org/10.7916/cblr.v2019i2.3424>. Acesso em: 07 jul. 2026.

WILLIAMSON, B. **Big Data in Education: The Digital Future of Learning, Policy and Practice**. London: SAGE Publications, 2017.

WORLD MEDICAL ASSOCIATION. World Medical Association Declaration of Helsinki: ethical principles for medical research involving human subjects. **JAMA**,

Chicago, v. 310, n. 20, p. 2191-2194, 2013. <https://doi.org/10.1001/jama.2013.281053>. Acesso em: 07 jul. 2026.

YUSTE, R.; GOERING, S.; AGÜERA Y ARCAS, B.; BI, G.; CARMENA, J. M.; CARTER, A.; FINS, J. J.; FRIESEN, P.; GALLANT, J.; HUGGINS, J. E.; ILLES, J.; KELLMEYER, P.; KLEIN, E.; MARBLESTONE, A.; MITCHELL, C.; PARENS, E.; PHAM, M.; RUBEL, A.; SADATO, N.; WOLPAW, J. Four ethical priorities for neurotechnologies and AI. **Nature**, London, v. 551, n. 7679, p. 159-163, 2017. <https://doi.org/10.1038/551159a>. Acesso em: 07 jul. 2026.

ZAWACKI-RICHTER, O.; MARÍN, V. I.; BOND, M.; GOUVERNEUR, F. Systematic review of research on artificial intelligence applications in higher education: where are the educators? **International Journal of Educational Technology in Higher Education**, Cham, v. 16, art. 39, 2019. <https://doi.org/10.1186/s41239-019-0171-0>. Acesso em: 07 jul. 2026.

ZUBOFF, S. **The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power**. New York: Public Affairs, 2019.