

La Integración de la Inteligencia Artificial en la Academia: Superando la Dependencia de Herramientas de Detección en el Nuevo Contexto Normativo Venezolano

Integrating Artificial Intelligence into Academia: Overcoming
Dependence on Detection Tools in the New Venezuelan
Regulatory Context

Recibido: 21 - 01 - 2026

Aceptado: 16 - 03 - 2026

Miguel Subero
DOCENTE
Universidad de Los Andes (ULA)
Mérida, Venezuela.
msubero2014@gmail.com
<https://orcid.org/0009-0005-2380-0494>

Resumen

El ensayo analiza el desafío de la gobernanza académica en la educación superior venezolana frente a la irrupción de la Inteligencia Artificial Generativa (IAG). A partir de un análisis crítico documental, se contrasta la evidencia técnica sobre la baja precisión y el sesgo de los detectores de IA (OpenAI, 2023; Liang et al., 2023) con el emergente marco regulatorio nacional, destacando la reciente aprobación en primera discusión del Proyecto de Ley de Inteligencia Artificial (Asamblea Nacional, 2024) y el Código de Ética para el Desarrollo y Aplicación Responsable de la IA (Mincyt, 2024). Se argumenta que los modelos prohibitivos contradicen los principios de equidad e inclusión del Estado venezolano. En consecuencia, se propone transitar hacia un modelo human-in-the-loop, en sintonía con propuestas internacionales (UNESCO, 2024), donde la evaluación oral y un Programa Nacional de Alfabetización Algorítmica Docente ocupen un lugar central, transformando la IA de amenaza a andamiaje cognitivo.

Palabras clave: Inteligencia artificial generativa; Integridad académica; Evaluación auténtica; Gobernanza universitaria; Ética tecnológica; Formación docente.

Abstract

This essay analyzes the challenge of academic governance in Venezuelan higher education facing the emergence of Generative Artificial Intelligence (GAI). Based on a critical documentary analysis, technical evidence on AI detectors' low accuracy and bias (OpenAI, 2023; Liang et al., 2023) is contrasted with the emerging national regulatory framework, highlighting the recent first-round approval of the Artificial Intelligence Bill (National Assembly, 2024) and the Code of Ethics for the Responsible Development and Application of AI (Mincyt, 2024). It argues that prohibitive models contradict the Venezuelan State's principles of equity and inclusion. Consequently, it proposes transitioning towards a human-in-the-loop model, aligned with international frameworks (UNESCO, 2024), where oral assessment and a National Program for Teacher Algorithmic Literacy take center stage, reframing AI from a threat into cognitive scaffolding.

Keywords: Generative artificial intelligence; Academic integrity; Authentic assessment; University governance; Technological ethics; Teacher training.

Introducción

En la tradición universitaria hispana, cuando existe duda normativa en el uso del lenguaje, la Real Academia Española funge como árbitro último; la irrupción de la Inteligencia Artificial Generativa (IAG) descoloca ese esquema, pues no existe una "RAE de la IA". Ningún organismo supranacional ofrece hoy criterios cerrados para delimitar con certeza dónde termina la autoría humana y dónde comienza la asistencia algorítmica, cuestión que impacta directamente la manera en que enseñamos y evaluamos en el nivel superior.

En la región latinoamericana, el panorama regulatorio ha comenzado a tomar forma. Mientras Chile consolidó una Política Nacional de Inteligencia Artificial (2021) con ejes de habilitadores, adopción y gobernanza, la República Bolivariana de Venezuela ha dado recientemente pasos históricos y determinantes para abandonar el vacío normativo. En noviembre de 2024, la Asamblea Nacional aprobó en primera discusión el Proyecto de Ley de Inteligencia Artificial, con el objetivo de promover y fortalecer las capacidades de los sectores productivos del país (Asamblea Nacional, 2024). Paralelamente, el Ministerio del Poder Popular para Ciencia y Tecnología (Mincyt) publicó el Código de Ética para el Desarrollo y Aplicación Responsable de la Inteligencia Artificial, estableciendo nueve principios rectores ineludibles, entre los que destacan la IA Humanista, la Equidad, Igualdad y No Discriminación, y la Ciencia Abierta (Mincyt, 2024).

En este panorama de efervescencia legislativa estatal, las instituciones de educación superior enfrentan un dilema. Algunas casas de estudio, como la Universidad Católica Andrés Bello (UCAB) (2023), han emitido lineamientos internos tempranos para incorporar la IA bajo criterios de ética; sin embargo, una gran parte de la academia ha optado por la vía reactiva: la prohibición y la vigilancia a través de detectores automatizados de plagio sintético.

El presente ensayo tiene como propósito cuestionar la eficacia técnica y ética de estos modelos prohibitivos en la educación superior, argumentando que no solo son falibles, sino que entran en contradicción directa con los nuevos principios éticos del Estado venezolano. Para ello, se analizará en primer lugar la ineficacia de los detectores de IA; en segundo lugar, se expondrán las implicaciones éticas y de desigualdad que estos generan; en tercer lugar, se propondrá la transición hacia un modelo de evaluación auténtica bajo el enfoque human-in-the-loop; y, finalmente, se justificará la necesidad urgente de un programa nacional de formación docente liderado por las universidades tradicionales.

Desarrollo

El Espejismo Tecnológico: Limitaciones y Sesgos de los Detectores de IA

La confianza institucional en detectores automáticos de "texto de IA" no resiste el escrutinio técnico. El problema de fondo se asemeja a una carrera armamentista donde las herramientas de detección basan su funcionamiento en métricas estadísticas predictivas, principalmente dos conceptos emergentes en la lingüística computacional: la perplejidad

(perplexity) y la variabilidad o ráfaga (burstiness). La perplejidad mide qué tan predecible es la elección de palabras en una oración; una perplejidad baja indica que el texto es predecible y, por ende, estadísticamente probable de haber sido escrito por una máquina. Por su parte, la variabilidad evalúa los cambios en la longitud y complejidad de las oraciones a lo largo del texto; los humanos tienden a mezclar oraciones largas y cortas de forma caótica, mientras que los modelos generativos producen estructuras más monótonas y estandarizadas.

No obstante, basar la integridad académica y el futuro de un estudiante en estas métricas es un grave error de gobernanza. OpenAI, la empresa creadora de ChatGPT, discontinuó su propio AI Text Classifier por su "baja tasa de precisión" y advierte categóricamente que no debe usarse como herramienta decisoria principal (OpenAI, 2023). Peor aún, técnicas de ofuscación y parafraseo mediante segunda IA degradan de forma notable el rendimiento de los detectores, minando su utilidad forense. Weber-Wulff et al. (2023), en una evaluación comparativa amplia, concluyen que estas herramientas no son precisas ni fiables, advirtiendo que los detectores arrojan falsos positivos sistemáticos frente a la escritura académica formal, la cual, por su propia naturaleza estructurada, tiende a poseer baja perplejidad y variabilidad. Confiar ciegamente en estas herramientas viola el principio de "Transparencia y Rendición de Cuentas" exigido por el nuevo Código de Ética venezolano (Mincyt, 2024), ya que las instituciones estarían delegando decisiones punitivas a algoritmos de caja negra que ni siquiera sus creadores logran auditar con certeza.

Implicaciones Éticas: La Brecha Digital y el Choque con la "IA Humanista"

La ineficacia técnica es grave, pero el impacto ético y social de los detectores en la región latinoamericana es aún mayor. La desconfianza hacia los detectores no es tecnofobia, sino prudencia institucional, pues existe evidencia empírica de sesgos algorítmicos severos. Liang et al. (2023) demuestran que diversos detectores penalizan desproporcionadamente a escritores no nativos del inglés o a estudiantes con menor riqueza de vocabulario, aumentando los falsos positivos en redacciones con sintaxis más rígida.

La aplicación acrítica de estos detectores amenaza con amplificar desigualdades estructurales, generando lo que en sociología se denomina el "Efecto Mateo"¹ (Merton, 1968; American Psychological Association, 2018). Este fenómeno describe procesos de ventaja acumulativa: en el aula, quienes disponen de recursos económicos para pagar versiones Premium de IA o herramientas de "humanización" textual (algoritmos diseñados para inyectar burstiness artificial evadiendo radares) escapan de la detección con facilidad. Por el contrario, los estudiantes con recursos limitados, que utilizan versiones gratuitas o que escriben con estructuras sintácticas más básicas producto de deficiencias formativas previas, quedan expuestos a la sospecha de fraude.

¹ En ciencias sociales, el "Efecto Mateo" describe procesos de ventaja acumulativa por los cuales quienes ya poseen más recursos o reconocimiento tienden a recibir proporcionalmente más con el tiempo, mientras que quienes parten con menos acumulan desventajas. El término fue acuñado por Robert K. Merton (1968) para explicar la asignación desigual de crédito en la ciencia, y se ha extendido a ámbitos como la educación, donde puede manifestarse en brechas que se amplían si no median políticas de equidad y diseño pedagógico. (Véanse Merton, 1968; American Psychological Association, 2018).

Esta realidad choca frontalmente con el principio de "Equidad, Igualdad y No Discriminación" del Código de Ética del Mincyt (2024), el cual exige que los sistemas de IA no segmenten, no categoricen y protejan a los grupos vulnerables. Estudios recientes en América Latina (OIT, 2024; BID, 2025) advierten que la divisoria digital actúa como un cuello de botella, obstaculizando especialmente a colectivos vulnerables. Sancionar a un estudiante basándose en un detector sesgado es perpetuar una exclusión tecnológica inaceptable en el marco de una "IA Humanista".

Hacia una Evaluación Auténtica: El Modelo Human-in-the-Loop

Convertir la evaluación en una carrera de obstáculos contra la tecnología alimenta una "pedagogía de la sospecha". Frente a esto, instituciones de prestigio internacional han comenzado a retroceder. La Universidad de Vanderbilt (2023), la University of Waterloo (2025) y la Curtin University (2025) anunciaron la deshabilitación definitiva de la detección de IA en plataformas como Turnitin, alegando falta de transparencia, sesgos y un costo moral inaceptable para su comunidad.

Es preferible reorientar la academia hacia lo que históricamente ha valorado: indagar, argumentar, contrastar fuentes y sostener oralmente la validez de lo dicho. En esa óptica, la IA es un instrumento intelectual cuya potencia depende de la pericia con que se opera; el mérito reside menos en la "fuerza bruta" de escribir sin apoyo y más en el criterio para plantear buenos prompts (instrucciones), interpretar resultados, auditar alucinaciones del modelo y triangular con literatura confiable. Éticamente, su intervención es equiparable al acompañamiento tradicional de un tutor metodológico (Torres Vargas, 2023). El fraude no reside en apoyarse en la herramienta para potenciar el pensamiento complejo, sino en la opacidad de no declarar su uso o en la negligencia de no validar sus resultados.

Operacionalizar este enfoque, alineado con las directrices de la UNESCO (2024), exige métodos concretos basados en un ecosistema human-in-the-loop (donde la máquina sugiere, pero el humano audita, valida y decide). Primero, la evaluación de proceso y trazabilidad: exigir la genealogía del trabajo (borradores, prompts utilizados y una reflexión metacognitiva sobre cómo la IA modificó la idea original). Segundo, la defensa oral híbrida para que el estudiante justifique elecciones y responda contraejemplos, asegurando que el conocimiento ha sido internalizado. Tercero, el análisis comparativo crítico: solicitar una versión generada con IA y otra contrastada y corregida por el estudiante, enseñándole a ser el editor crítico de la máquina.

Alfabetización Algorítmica y Formación Docente: Un imperativo Nacional

El principio de "Excelencia" estipulado en el Código de Ética venezolano (Mincyt, 2024) exige impulsar el desarrollo del talento humano al más alto nivel. No se puede exigir a los estudiantes un uso ético y avanzado de la inteligencia artificial si el cuerpo docente carece de las competencias instrumentales y pedagógicas para guiar este proceso. La prohibición de la IA en las aulas suele ser, en muchas ocasiones, un reflejo del desconocimiento técnico del propio claustro de profesores frente a herramientas disruptivas.

Por consiguiente, el verdadero reto para la educación superior venezolana no radica en adquirir licencias de software anti-plagio, sino en la urgente creación de un Programa Nacional de Alfabetización Algorítmica y Formación Docente. Este programa debe estar orientado a enseñar a los profesores a rediseñar sus instrumentos de evaluación, a comprender la lógica subyacente de los Modelos de Lenguaje Grande (LLM) y a integrar la IA como un co-piloto pedagógico.

Las universidades autónomas y tradicionales, como la Universidad de los Andes (ULA), por su histórico peso en la investigación y generación de conocimiento, están llamadas a ser los epicentros de esta formación. Aprovechando el principio de "Ciencia Abierta" (Mincyt, 2024), estas casas de estudio pueden liderar laboratorios de innovación educativa donde se socialicen mejores prácticas, se construyan repositorios de prompts académicos de libre acceso y se capacite a formadores de formadores. Solo a través de la capacitación masiva del docente se logrará que la integración de la IA en Venezuela trascienda el debate del plagio y se convierta en un motor genuino de desarrollo científico y productivo, tal como lo persigue el naciente marco legal del Estado.

Conclusiones

En primer lugar, la evidencia técnica y documental demuestra de manera concluyente que las herramientas de detección automatizada de Inteligencia Artificial carecen de la fiabilidad, transparencia y solidez necesarias para actuar como mecanismos sancionatorios en la educación superior. La dependencia de métricas estadísticas predictivas resulta ineficaz frente a la rápida evolución de los modelos generativos, convirtiendo la fiscalización algorítmica en una estrategia obsoleta, propensa al error y contraria al principio de rendición de cuentas.

En segundo lugar, el uso de detectores prohibitivos entraña un profundo riesgo ético que agudiza la brecha digital en el contexto universitario. Al propiciar el "Efecto Mateo", estas herramientas castigan de manera desproporcionada a los estudiantes con menores competencias lingüísticas o recursos limitados, entrando en directa contradicción con los principios de Equidad, Igualdad y No Discriminación establecidos en el nuevo Código de Ética para el Desarrollo y Aplicación Responsable de la IA de la República Bolivariana de Venezuela.

Finalmente, el avance legislativo evidenciado en el Proyecto de Ley de Inteligencia Artificial (2024) demuestra que el Estado venezolano apuesta por el aprovechamiento productivo de esta tecnología. Las universidades no pueden quedarse rezagadas en una cultura de prohibición. La ruta institucional sostenible exige abandonar la "pedagogía de la sospecha" para abrazar la evaluación auténtica (human-in-the-loop) y, fundamentalmente, invertir recursos y esfuerzos en un Programa Nacional de Formación Docente. Solo capacitando a nuestros profesores en los espacios de excelencia, como la Universidad de los Andes, lograremos que la IA deje de ser una amenaza y se convierta en un andamiaje cognitivo indispensable para el futuro del país.

Referencias

- American Psychological Association. (2018). APA dictionary of psychology (entrada Matthew effect'). <https://dictionary.apa.org/matthew-effect>
- Asamblea Nacional de la República Bolivariana de Venezuela. (2024, 19 de noviembre). AN aprueba en primera discusión Proyecto de Ley de Inteligencia Artificial. <https://www.asambleanacional.gob.ve/noticias/an-aprueba-en-primera-discusion-proyecto-de-ley-de-inteligencia-artificial>
- Curtin University. (2025, 4 de septiembre). Update on Turnitin AI-Detection Tool (GenAI detection disabled from 1 Jan 2026). <https://www.curtin.edu.au/news/oasis-news/update-on-turnitin-ai-detection-tool/>
- Inter-American Development Bank. (2025). Generative AI in education: A framework for leveraging digital tools in Latin American classrooms (IDB-TN-3199). <https://doi.org/10.18235/0013853>
- Liang, W., Yuksekogonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779. <https://doi.org/10.1016/j.patter.2023.100779>
- Merton, R. K. (1968). The Matthew effect in science. *Science*, 159(3810), 56–63. <https://garfield.library.upenn.edu/merton/matthew1.pdf>
- Ministerio de Ciencia, Tecnología, Conocimiento e Innovación (Chile). (2021). Política Nacional de Inteligencia Artificial. <https://www.minciencia.gob.cl/areas/inteligencia-artificial/politica-nacional-de-inteligencia-artificial/>
- Ministerio del Poder Popular para Ciencia y Tecnología [Mincyt]. (2024). Código de ética para el desarrollo y aplicación responsable de la inteligencia artificial de la República Bolivariana de Venezuela. https://mincyt.gob.ve/wp-content/uploads/2026/02/Codigo_de_Etica_de_Inteligencia_Artificial_de_la_Republica_Bolivar
- OpenAI. (2023, 31 de enero). New AI classifier for indicating AI-written text. <https://openai.com/index/new-ai-classifier-for-indicating-ai-written-text/>
- Organización Internacional del Trabajo y Banco Mundial. (2024). Buffer or bottleneck? Employment exposure to generative AI and the digital divide in Latin America (ILO Working Paper 121). https://www.ilo.org/sites/default/files/2024-07/WP121_web.pdf
- Torres Vargas, J. D. (2023). La Inteligencia Artificial (IA) en la Educación Superior: Retos y oportunidades. *Dialéctica*, 21, 376–388.
- UNESCO. (2024). AI competency framework for teachers. <https://doi.org/10.54675/ZJTE2084>
- Universidad Católica Andrés Bello. (2023, 20 de junio). La UCAB anunció lineamientos para incorporar la inteligencia artificial en todos sus procesos. <https://elucabista.com/2023/06/20/la-ucab-anuncio-lineamientos-para-incorporar-la-inteligencia-artificial-en-todos-sus-procesos/>
- University of Waterloo. (2025). Discontinuing use of AI detection functionality in Turnitin. <https://uwaterloo.ca/associate-vice-president-academic/discontinuing-use-ai-detection-functionality-turnitin>

Vanderbilt University. (2023,16 de agosto). Guidance on AI detection and why we're disabling Turnitin's AI detector.
<https://www.vanderbilt.edu/brightspace/2023/08/16/guidance-on-ai-detection-and-why-were-disabling-turnitins-ai-detector/>

Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltýnek, T., Guerrero-Dib, J., Popoola, O., Šigut, P., & Waddington, L. (2023). Testing of detection tools for AI-generated text. *International Journal for Educational Integrity*, 19, 26.
<https://doi.org/10.1007/s40979-023-00146-z>

Para citar este ensayo:

Subero, M. (2026). La Integración de la Inteligencia Artificial en la Academia: Superando la Dependencia de Herramientas de Detección en el Nuevo Contexto Normativo Venezolano. *Revista Aprendizaje Digital*. Vol. 8, Número 1 enero-junio, 136 - 143.

