

# INTELIGENCIA ARTIFICIAL GENERATIVA, PENSAMIENTO CRÍTICO E INTEGRIDAD ACADÉMICA EN LA EVALUACIÓN UNIVERSITARIA

## GENERATIVE ARTIFICIAL INTELLIGENCE, CRITICAL THINKING AND ACADEMIC INTEGRITY IN UNIVERSITY ASSESSMENT

Dr. C. Leopoldo Vinicio Venegas Loo<sup>1\*</sup>

<sup>1</sup> Docente de la Carrera de Pedagogía de los Idiomas Nacionales y Extranjeros, Facultad de Ciencias Sociales, Humanísticas y de la Educación; docente tutor de la Maestría en Educación, Instituto de Posgrado, Universidad Estatal del Sur de Manabí, Jipijapa, Ecuador. ORCID: <https://orcid.org/0000-0002-3100-6320>. Correo: [leopoldo.venegas@unesum.edu.ec](mailto:leopoldo.venegas@unesum.edu.ec)

Dra. C. Paola Yadira Moreira Aguayo<sup>2</sup>

<sup>2</sup> Docente de la Carrera de Pedagogía de los Idiomas Nacionales y Extranjeros, Facultad de Ciencias Sociales, Humanísticas y de la Educación; docente tutora de la Maestría en Educación, Instituto de Posgrado, Universidad Estatal del Sur de Manabí, Jipijapa, Ecuador. ORCID: <https://orcid.org/0000-0002-6380-2794>. Correo: [paola.moreira@unesum.edu.ec](mailto:paola.moreira@unesum.edu.ec)

Econ. Liliana Maribel Figueroa Soledispa, Mg<sup>3</sup>

<sup>3</sup> Responsable de Titulación y docente de la Maestría en Educación, Instituto de Posgrado, Universidad Estatal del Sur de Manabí, Jipijapa, Ecuador. ORCID: <https://orcid.org/0000-0001-9785-2581>. Correo: [liliana.figueroa@unesum.edu.ec](mailto:liliana.figueroa@unesum.edu.ec)

Jefferson Ricardo Álvarez Indacochea<sup>4</sup>

<sup>4</sup> Estudiante de la Carrera de Derecho, Facultad de Ciencias Sociales, Humanísticas y de la Educación, Universidad Estatal del Sur de Manabí, Jipijapa, Ecuador. ORCID: <https://orcid.org/0009-0006-3456-5589>. Correo: [jefferson8025@unesum.edu.ec](mailto:jefferson8025@unesum.edu.ec)

\* Autor para correspondencia: [leopoldo.venegas@unesum.edu.ec](mailto:leopoldo.venegas@unesum.edu.ec)

### Resumen

La expansión de la inteligencia artificial generativa ha modificado las condiciones bajo las cuales las universidades interpretan la autoría, el aprendizaje y la evidencia evaluativa. El objetivo de este artículo es

analizar críticamente sus implicaciones para el pensamiento crítico, la integridad académica y la evaluación en educación superior, y proponer criterios para su integración responsable. Se realizó una revisión documental crítica e integrativa de literatura científica y documentos institucionales publicados principalmente entre 2021 y 2026. Los resultados muestran que los efectos educativos dependen del diseño de las tareas, la alfabetización de docentes y estudiantes, la transparencia de uso y la capacidad de verificar el proceso intelectual. Las estrategias centradas exclusivamente en prohibición y detección resultan insuficientes. Como contribución conceptual se propone el Marco EVALUAR-IA, compuesto por siete dimensiones: evidencia del proceso, verificación, autoría demostrable, lectura crítica, uso ético, aprendizaje reflexivo y responsabilidad humana. Se concluye que la evaluación universitaria debe desplazarse del producto final hacia la comprensión, la justificación, la verificación y la trazabilidad del razonamiento.

**Palabras clave:** inteligencia artificial generativa; educación superior; evaluación; pensamiento crítico; integridad académica

### Abstract

*The expansion of generative artificial intelligence has changed the conditions under which universities interpret authorship, learning and assessment evidence. This article critically analyses its implications for critical thinking, academic integrity and assessment in higher education and proposes criteria for responsible integration. A critical and integrative documentary review of scientific literature and institutional documents published mainly between 2021 and 2026 was conducted. Findings show that educational effects depend on task design, teacher and student literacy, transparency of use and the ability to verify the intellectual process. Strategies focused exclusively on prohibition and detection are insufficient. As a conceptual contribution, the EVALUAR-AI Framework is proposed, consisting of seven dimensions: evidence of process, verification, demonstrable authorship, critical reading, ethical use, reflective learning and human responsibility. University assessment should therefore move beyond the final product toward comprehension, justification, verification and traceability of reasoning.*

**Keywords:** generative artificial intelligence; higher education; assessment; critical thinking; academic integrity

**Fecha de recibido:** 24/06/2026

**Fecha de aceptado:** 26/08/2026

**Fecha de publicado:** 07/09/2026

### Introducción

La educación superior atraviesa una transformación que no puede interpretarse únicamente como la incorporación de una nueva herramienta digital. Los sistemas de inteligencia artificial (IA) generativa producen textos, imágenes, códigos, explicaciones, resúmenes y respuestas argumentativas mediante interfaces de lenguaje natural. Su disponibilidad modifica una relación histórica de la universidad: la

asociación entre la producción observable del estudiante y la inferencia sobre aquello que realmente sabe, comprende y puede realizar.

Durante décadas, un ensayo, un informe o una tarea domiciliaria podían considerarse evidencias aproximadas de lectura, escritura, búsqueda y razonamiento. La inteligencia artificial generativa debilita esa equivalencia porque un producto lingüísticamente correcto ya no demuestra, por sí solo, que el razonamiento expresado haya sido construido por la persona evaluada (Rodríguez et al., 2020), (Soledispa et al., 2025). El problema, por tanto, no se limita al uso de una plataforma; obliga a revisar el concepto de evidencia de aprendizaje.

Kasneci et al. (2023) señalan que los grandes modelos de lenguaje pueden favorecer personalización, generación de materiales y retroalimentación, pero también producir errores factuales y sesgos. Tlili et al. (2023) identifican beneficios potenciales junto con problemas de honestidad académica, privacidad y confiabilidad. La UNESCO (Miao & Holmes, 2023) propone una integración centrada en el ser humano, la agencia, la equidad y la responsabilidad. La misma tecnología puede, entonces, apoyar el aprendizaje o sustituir el proceso cognitivo que una actividad pretende desarrollar.

La literatura reciente ha desplazado el debate desde la prohibición hacia problemas más complejos. Moorhouse et al. (2023) evidencian que las universidades han empezado a desarrollar orientaciones específicas para evaluación y uso de IA. Xia et al. (2024) muestran que la transformación debe involucrar estudiantes, docentes e instituciones. Bittle y El-Gayar (2025) concluyen que la integridad académica requiere combinar políticas, alfabetización y rediseño pedagógico, en lugar de depender únicamente de mecanismos de detección.

El pensamiento crítico ocupa un lugar central. Los sistemas generativos pueden estimular comparación, refutación y análisis, pero también favorecer descarga cognitiva y aceptación acrítica. Helal et al. (2025) describen este carácter dual, mientras Cong-Lem y Thuy-Dung (2026) muestran que el pensamiento crítico en ambientes con IA incluye detectar errores, evaluar credibilidad y verificar fuentes. La cuestión central deja así de ser tecnológica y se convierte en epistemológica, pedagógica y ética (Rodríguez Rodríguez & Solórzano Álava, 2024).

El objetivo de este artículo es analizar críticamente las implicaciones de la inteligencia artificial generativa para el pensamiento crítico, la integridad académica y la evaluación en educación superior, y formular principios para un modelo evaluativo que integre responsablemente estas tecnologías sin sustituir la agencia intelectual humana. Se plantean tres preguntas: cómo modifica la IA las condiciones de evaluación, qué riesgos y oportunidades genera y qué principios deberían orientar el rediseño evaluativo. Como aporte conceptual se propone el Marco EVALUAR-IA.

## **Materiales y métodos**

La investigación se desarrolló mediante una revisión documental crítica e integrativa, seleccionada porque el propósito no fue calcular un efecto agregado de una intervención, sino articular literatura procedente de tecnología educativa, evaluación, alfabetización en inteligencia artificial, integridad académica, pensamiento

crítico, ética y gobernanza. Se priorizaron trabajos publicados entre 2021 y 2026 para mantener actualidad frente a un fenómeno de rápida evolución.

Se incluyeron artículos científicos sobre inteligencia artificial generativa en educación superior, evaluación, autoría, integridad, pensamiento crítico, alfabetización en IA y políticas institucionales, así como documentos de organismos internacionales. Se priorizaron fuentes con identificación bibliográfica completa, DOI o enlace institucional verificable (Cornelio et al., 2024), (Álava et al., 2023). Se excluyeron materiales promocionales, documentos sin trazabilidad académica y publicaciones sin relación directa con el objeto de estudio.

El análisis se organizó en seis categorías: evidencia de aprendizaje; autoría e integridad; pensamiento crítico y delegación cognitiva; alfabetización; rediseño de la evaluación; y gobernanza institucional. A partir de la síntesis se elaboró el Marco EVALUAR-IA. La propuesta debe entenderse como un marco conceptual susceptible de posterior juicio de expertos, pilotaje y validación empírica, no como una escala psicométrica ya validada.

## Resultados y discusión

### Evidencia de aprendizaje, agencia y autoría intelectual

Uno de los efectos más relevantes de la inteligencia artificial generativa es el debilitamiento de la equivalencia entre producto entregado y competencia demostrada. Un ensayo puede continuar siendo pedagógicamente útil, pero ha perdido parte de su capacidad para actuar aisladamente como evidencia inequívoca de lectura, argumentación y comprensión. Rudolph et al. (2023), Overono y Ditta (2023) y Nikolic et al. (2023) coinciden en que la capacidad de los modelos generativos obliga a reconsiderar tareas diseñadas bajo supuestos anteriores a su disponibilidad.

La conducta deshonesta continúa siendo responsabilidad de quien la realiza; sin embargo, desde el aseguramiento del aprendizaje, una evaluación vulnerable a sustitución completa también requiere revisión. Si una tarea puede ser resuelta satisfactoriamente mediante una instrucción breve, el docente debe analizar si el resultado de aprendizaje continúa siendo observable. La respuesta no consiste necesariamente en eliminar el ensayo o el informe, sino en incorporar evidencias de formulación inicial, selección de fuentes, decisiones, versiones, verificación y defensa.

La inteligencia artificial obliga además a distinguir producción de autoría. Una persona puede entregar un producto sin haber construido sus componentes intelectuales principales; de manera inversa, puede utilizar herramientas tecnológicas y conservar una autoría sustantiva sobre las decisiones y conclusiones. Perkins (2023) advierte que el uso de un modelo no determina automáticamente plagio: la valoración depende de las reglas, la transparencia y la función desempeñada por la herramienta.

Puede distinguirse una gradación entre asistencia instrumental, apoyo cognitivo, co-construcción, sustitución parcial y sustitución integral. La primera comprende corrección o tareas mecánicas autorizadas; la segunda ofrece explicaciones o retroalimentación; la tercera supone interacción iterativa seguida de selección y verificación humana; la cuarta incorpora contenido sustantivo con escaso procesamiento; y la quinta delega prácticamente toda la tarea. La legitimidad de cada nivel depende del propósito evaluativo. La pregunta

relevante deja de ser únicamente si se utilizó IA y pasa a ser qué función desempeñó y qué responsabilidad intelectual conservó el estudiante.

### Integridad académica: de la detección a la integridad por diseño

Los modelos generativos pueden producir secuencias textuales originales que no representan el trabajo intelectual de quien las presenta. Por ello, la similitud textual deja de ser suficiente como indicador único de integridad. Bittle y El-Gayar (2025) identifican riesgos de deshonestidad junto con la necesidad de alfabetización y políticas institucionales; Cotton et al. (2024) señalan que la facilidad de generación automatizada desafía los enfoques tradicionales de prevención del plagio.

Una política basada exclusivamente en detectores resulta frágil por varias razones: ningún clasificador debe considerarse prueba infalible de autoría; la tecnología evoluciona rápidamente; la detección describe características del producto, pero no reconstruye el aprendizaje; y una cultura centrada exclusivamente en vigilancia puede desplazar recursos que deberían dirigirse al diseño de tareas más auténticas. Esto no implica eliminar mecanismos de control, sino integrarlos dentro de una estrategia más amplia.

La integridad académica puede evolucionar hacia un enfoque de integridad por diseño: reglas claras, evidencias durante el proceso, transparencia sobre herramientas, defensa de decisiones, prevención y formación. Este enfoque mantiene la responsabilidad individual, pero mejora la capacidad institucional de demostrar aprendizaje y reduce la dependencia de predicciones algorítmicas sobre el origen de un texto.

**Tabla 1:** Del paradigma de detección al paradigma de integridad por diseño.

| Paradigma tradicional             | Paradigma emergente                          |
|-----------------------------------|----------------------------------------------|
| Producto final                    | Producto, proceso y defensa                  |
| Detectar uso de IA                | Comprender función y nivel de intervención   |
| Detector como evidencia principal | Evidencias múltiples y diálogo académico     |
| Prohibición general               | Reglas según resultado de aprendizaje        |
| Autoría implícita                 | Declaración de asistencia tecnológica        |
| Sanción posterior                 | Prevención, alfabetización y responsabilidad |

### Pensamiento crítico y descarga cognitiva

La inteligencia artificial puede fortalecer o debilitar el pensamiento crítico según la forma de interacción. Puede estimularlo cuando el estudiante identifica errores, contrasta fuentes, formula contraargumentos, detecta sesgos o justifica por qué rechaza una recomendación. En estos casos, la IA funciona como objeto de pensamiento. El efecto cambia cuando reemplaza lectura, escritura o resolución y sus respuestas son aceptadas sin contraste; entonces funciona como sustituto del pensamiento.

Helal et al. (2025) describen un impacto dual mediado por autorregulación, confianza y crítica metacognitiva. Cong-Lem y Thuy-Dung (2026) muestran que la conceptualización contemporánea del pensamiento crítico incorpora capacidades específicas de detección de errores de IA, evaluación de credibilidad y verificación de fuentes. Pensar críticamente en la universidad implica ahora examinar contenidos humanos y algorítmicos.

Esta competencia exige una actitud epistemológica de verificación: comprobar si la fuente existe, si sostiene la afirmación, si el sistema omite información relevante, reproduce sesgos o presenta seguridad sobre asuntos inciertos. Enseñar a utilizar IA sin enseñar a desconfiar metodológicamente de sus respuestas constituye una integración incompleta.

La descarga cognitiva no es necesariamente negativa. Calculadoras, buscadores y software estadístico permiten delegar operaciones mecánicas para concentrar recursos en actividades complejas. El problema surge cuando se externaliza precisamente la competencia que debía desarrollarse. Si el resultado de aprendizaje consiste en argumentar, formular hipótesis o resolver un problema, delegar íntegramente esa actividad impide ejercitarla. La pregunta pedagógica adecuada es qué componentes pueden delegarse sin destruir el resultado de aprendizaje.

Esta distinción conduce a una pedagogía de la fricción cognitiva productiva. La IA puede reducir tareas improductivas, pero la universidad debe conservar aquellas dificultades que obligan a recuperar, comparar, argumentar, interpretar y decidir. La eficiencia tecnológica no equivale automáticamente a eficacia pedagógica.

### **Alfabetización en inteligencia artificial**

La alfabetización en IA no puede reducirse a redactar prompts. Ng et al. (2021) la vinculan con conocer, utilizar, evaluar y abordar cuestiones éticas. Un estudiante puede obtener resultados sofisticados y desconocer cómo funcionan los modelos, confiar excesivamente en ellos, introducir información confidencial o aceptar referencias inexistentes. La habilidad operativa no garantiza alfabetización crítica.

Una formación universitaria integral debería incluir una dimensión técnica básica, una dimensión informacional de verificación, una dimensión epistemológica que diferencie fluidez de verdad, una dimensión ética relacionada con privacidad y sesgos, una dimensión pedagógica que distinga apoyo de sustitución, una dimensión disciplinar ajustada a cada profesión y una dimensión metacognitiva para reconocer dependencia.

Los marcos de competencias de la Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura (UNESCO) para docentes y estudiantes (Miao & Cukurova, 2024; Miao, Shiohira, & Lao, 2024) refuerzan esta visión al incorporar enfoque humano, ética, fundamentos, aplicaciones y desarrollo progresivo de capacidades. La alfabetización en IA debe entenderse como competencia transversal y no como destreza tecnológica aislada.

La alfabetización docente es igualmente necesaria. Un profesor no necesita convertirse en ingeniero de inteligencia artificial, pero sí reconocer capacidades y límites, diseñar actividades resistentes a sustitución cognitiva, establecer reglas, interpretar usos aceptables, proteger datos y evitar acusaciones apoyadas únicamente en detectores (Yero et al., 2018). Sin formación docente, la política institucional se vuelve inconsistente.

### **Rediseño de la evaluación y reglas diferenciadas de uso**

La inteligencia artificial no creó las debilidades de la evaluación memorística o descontextualizada; las hizo más visibles. Farrokhnia et al. (2024) muestran que los beneficios de ChatGPT coexisten con riesgos de



dependencia e información incorrecta. Rediseñar la evaluación implica aumentar la proximidad entre la evidencia solicitada y el proceso cognitivo que se pretende medir.

Una primera estrategia consiste en evaluar procesos mediante borradores, mapas argumentales, registros de búsqueda, decisiones metodológicas y versiones. Una segunda es incorporar defensas orales para explorar comprensión y justificar fuentes. Una tercera es personalizar tareas mediante casos locales, datos producidos por el estudiante o decisiones disciplinarias específicas. Una cuarta consiste en evaluar verificación, pidiendo auditar respuestas de IA. Finalmente, la reflexión metacognitiva permite explicar qué cambió, qué se aprendió y qué papel desempeñó la herramienta.

No todas las actividades requieren el mismo régimen de uso. Puede establecerse un escenario sin IA cuando se necesita demostrar una habilidad individual; un escenario de asistencia limitada para funciones específicas; un escenario de uso permitido con declaración transparente; y un escenario donde la IA es obligatoria como objeto de aprendizaje y el estudiante debe contrastarla y corregirla. Estas categorías transforman reglas ambiguas en expectativas explícitas.

La evaluación auténtica cobra especial relevancia. En Derecho, un estudiante puede verificar jurisprudencia sugerida por IA; en Educación, evaluar una planificación; en Salud, identificar riesgos en recomendaciones sin usar datos confidenciales; en Administración, contrastar análisis con indicadores; y en Idiomas, comparar traducciones y justificar decisiones. La competencia profesional se desplaza de competir con la máquina a dirigirla, limitarla, cuestionarla y asumir responsabilidad por el resultado.

La retroalimentación constituye otro uso prometedor. Xia et al. (2024) identifican posibilidades de retroalimentación inmediata y diversa, pero la respuesta automática no debe confundirse con evaluación docente. Una estrategia útil es pedir al estudiante que primero intente resolver, registre su razonamiento, consulte después el sistema, compare, corrija y explique qué aprendió. La IA actúa entonces como segundo interlocutor y no como primer sustituto del esfuerzo.

## Transparencia, privacidad y equidad

La utilización responsable requiere una cultura de declaración. Las instituciones pueden solicitar que el estudiante indique herramienta, finalidad, partes del proceso en que intervino, grado de modificación y mecanismos de verificación. La finalidad no es estigmatizar el uso legítimo, sino convertir una práctica invisible en una práctica evaluable. Chan (2023) sostiene que las políticas universitarias deben integrar dimensiones pedagógicas, de gobernanza y operativas.

La equidad también debe formar parte de la política. Los estudiantes no acceden necesariamente a las mismas versiones, funcionalidades o recursos de pago. Una universidad que obliga a utilizar sistemas comerciales sin garantizar alternativas puede generar desigualdad. Deben considerarse disponibilidad, accesibilidad, diversidad lingüística, sesgos y dependencia de proveedores privados.

La privacidad es especialmente crítica cuando se trabaja con datos clínicos, expedientes jurídicos, información de menores, bases de investigación o documentos reservados. La disponibilidad técnica para introducir información en una IA no equivale a autorización ética o jurídica. Este principio debe enseñarse de forma explícita.

Los errores o alucinaciones constituyen otro riesgo. Dwivedi et al. (2023) y Lo (2023) destacan problemas de credibilidad e información incorrecta. La respuesta académica debe incluir protocolos de verificación: localizar la publicación, comprobar autores, año y DOI, consultar la fuente original y determinar si realmente sostiene la afirmación. La fabricación de referencias no puede ser normalizada como un simple error tecnológico.

## Marco EVALUAR-IA

A partir de la síntesis se propone el Marco EVALUAR-IA como herramienta conceptual para revisar actividades en contextos de disponibilidad de inteligencia artificial generativa. Sus siete dimensiones son Evidencia del proceso, Verificación, Autoría demostrable, Lectura crítica, Uso ético y transparente, Aprendizaje reflexivo y Responsabilidad humana. No pretende imponer la misma intensidad en todas las tareas, sino ofrecer preguntas de control según el resultado de aprendizaje.

- **Evidencia del proceso**

La evaluación debe observar parte del trayecto seguido mediante borradores, fuentes, decisiones, versiones o defensa. Pregunta de control: ¿se observa cómo se construyó el conocimiento o solo el resultado final?

- **Verificación**

Toda información generada debe poder contrastarse con fuentes y datos independientes. Pregunta de control: ¿el estudiante demuestra capacidad para verificar aquello que la herramienta produce?

- **Autoría demostrable**

La autoría implica control intelectual: explicar decisiones, reconstruir argumentos y reconocer el aporte de la tecnología. Pregunta de control: ¿puede el estudiante defender aquello que presenta como propio?

- **Lectura crítica**

El estudiante debe identificar sesgos, generalizaciones, contradicciones y posibles alucinaciones. Pregunta de control: ¿utiliza la IA como autoridad o como información que también debe ser juzgada?

- **Uso ético y transparente**

Las reglas deben ser explícitas y el uso declarado. Pregunta de control: ¿puede reconstruirse con claridad qué papel desempeñó la inteligencia artificial?

- **Aprendizaje reflexivo**

El estudiante debe explicar qué sabía, qué cambió y qué aprendió. Pregunta de control: ¿la interacción produjo aprendizaje consciente o solamente un producto?

- **Responsabilidad humana**

La responsabilidad final no puede trasladarse al algoritmo. Pregunta de control: ¿está claramente identificada la persona responsable de las decisiones y del resultado?



**Tabla 2:** Matriz operativa del Marco EVALUAR-IA.

| Dimensión       | Riesgo                  | Evidencia            | Resultado           |
|-----------------|-------------------------|----------------------|---------------------|
| Evidencia       | Sustitución invisible   | Borradores/versiones | Trazabilidad        |
| Verificación    | Información falsa       | Fuentes comprobadas  | Rigor               |
| Autoría         | Delegación integral     | Defensa              | Agencia intelectual |
| Lectura crítica | Sesgo de automatización | Análisis de errores  | Pensamiento crítico |
| Uso ético       | Opacidad                | Declaración de IA    | Transparencia       |
| Aprendizaje     | Dependencia             | Reflexión            | Autorregulación     |
| Responsabilidad | Transferencia           | Decisión argumentada | Ética profesional   |

### Implicaciones institucionales y agenda futura

La transformación debe operar en varios niveles. Las universidades necesitan políticas que definan usos permitidos y restringidos, transparencia, protección de datos, integridad y procedimientos ante presuntas infracciones. Los perfiles de egreso deben revisar competencias transversales de verificación, ética, privacidad y responsabilidad. La formación estudiantil debería comenzar desde etapas iniciales y la capacitación docente ser sostenida, no limitarse a talleres puntuales sobre herramientas.

Los trabajos de titulación constituyen un espacio especialmente sensible. Una política equilibrada puede incluir declaración de uso de IA, certificación de verificación de referencias y datos, defensa fortalecida y evidencia del proceso mediante tutorías y entregas parciales. El propósito no es construir un sistema policial, sino proteger la legitimidad del grado académico y la autoría del estudiante.

La investigación universitaria necesita criterios equivalentes. La IA puede apoyar organización, revisión lingüística o generación de palabras clave, pero genera riesgos cuando inventa referencias, interpreta resultados sin supervisión o recibe bases confidenciales. La responsabilidad científica nunca puede atribuirse al sistema: corresponde al investigador que firma y presenta el trabajo.

La agenda futura requiere estudios longitudinales sobre escritura, lectura profunda, razonamiento y dependencia cognitiva; comparaciones disciplinares; instrumentos que diferencien pensamiento crítico apoyado por IA, pensamiento crítico sobre IA y pensamiento sustituido por IA; y estudios sobre alfabetización docente. También es necesario ampliar evidencia en América Latina, considerando desigualdad digital, diversidad institucional y marcos regulatorios.

El Marco EVALUAR-IA debería ser sometido a juicio de expertos, estudios Delphi, pilotaje y análisis de confiabilidad. Su valor potencial reside en ofrecer una estructura común adaptable a diferentes disciplinas, pero su utilidad práctica debe ser contrastada empíricamente antes de convertirlo en instrumento institucional estandarizado.

### Conclusiones

La inteligencia artificial generativa modifica la relación entre producto académico, autoría y evidencia de aprendizaje. Determinadas formas tradicionales de evaluación pierden capacidad para demostrar por sí solas lo que un estudiante comprende, porque una parte sustantiva del producto puede ser generada externamente.

La respuesta universitaria no debería ser exclusivamente prohibitiva ni indiscriminadamente permisiva. La IA puede facilitar explicaciones, retroalimentación y autoevaluación, pero también favorecer sustitución cognitiva, deshonestidad y aceptación acrítica. El efecto depende del diseño de la actividad y del nivel de agencia humana conservado.

En integridad académica, la detección automatizada es insuficiente como estrategia central. La universidad necesita avanzar hacia integridad por diseño: reglas explícitas, transparencia, evidencia de proceso, defensa intelectual y formación. La responsabilidad individual se mantiene, pero se fortalece la capacidad institucional para demostrar aprendizaje.

El pensamiento crítico adquiere una dimensión específica frente a contenidos algorítmicos. Verificar referencias, identificar errores, reconocer sesgos y distinguir fluidez de verdad deben formar parte de la alfabetización universitaria. Esta alfabetización incluye además privacidad, ética, conocimiento disciplinar y capacidad metacognitiva para reconocer dependencia.

El Marco EVALUAR-IA propone siete dimensiones para orientar el rediseño: evidencia del proceso, verificación, autoría demostrable, lectura crítica, uso ético, aprendizaje reflexivo y responsabilidad humana. Su propósito es reemplazar la pregunta limitada de si se utilizó IA por preguntas sobre qué aprendió el estudiante, qué verificó, qué decisiones tomó, qué produjo intelectualmente y si puede defender el resultado.

El futuro de la evaluación universitaria no debe centrarse en competir con máquinas para producir textos más rápidamente. Debe orientarse hacia capacidades humanas de comprensión, juicio, verificación y responsabilidad. La pregunta decisiva es si el estudiante puede comprender, justificar y asumir responsabilidad intelectual por las decisiones que adopta en un entorno académico y profesional acompañado por inteligencia artificial.

## Referencias

- Álava, W. L. S., Rodríguez, A. R., Cornelio, O. M., & Fonseca, B. B. (2023). El papel de la inteligencia artificial en la transformación digital de las empresas. Tono, Revista Técnica de la Empresa de Telecomunicaciones de Cuba SA, 19(1), 23-42. <https://www.revistatono.etcscu.cu/tono/article/download/377/680>
- Bittle, K., & El-Gayar, O. (2025). Generative AI and academic integrity in higher education: A systematic review and research agenda. Information, 16(4), 296. <https://doi.org/10.3390/info16040296>
- Chan, C. K. Y. (2023). A comprehensive AI policy education framework for university teaching and learning. International Journal of Educational Technology in Higher Education, 20, 38. <https://doi.org/10.1186/s41239-023-00408-3>
- Cong-Lem, N., & Thuy-Dung, N. T. (2026). Towards a framework for understanding and assessing critical thinking in generative AI-enhanced learning environments: A scoping review. European Journal of Psychology of Education, 41(2), Article 41. <https://doi.org/10.1007/s10212-026-01089-y>

- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, Cornelio, O. M., Rodríguez, A. R., Álava, W. L. S., Mora, P. G. A., Mera, L. M. S., & Bravo, B. J. P. (2024). *La Inteligencia Artificial: desafíos para la educación*. Editorial Internacional Alema. <https://editorialalema.org/libros/index.php/alema/article/download/34/33>
- A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Farrokhnia, M., Banihashemi, S. K., Noroozi, O., & Wals, A. (2024). A SWOT analysis of ChatGPT: Implications for educational practice and research. *Innovations in Education and Teaching International*, 61(3), 460–474. <https://doi.org/10.1080/14703297.2023.2195846>
- Helal, M. Y. I., Elgendy, I. A., Albashrawi, M. A., Dwivedi, Y. K., Al-Ahmadi, M. S., & Jeon, I. (2025). The impact of generative AI on critical thinking skills: A systematic review, conceptual framework and future research directions. *Information Discovery and Delivery*. <https://doi.org/10.1108/IDD-05-2025-0125>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Lo, C. K. (2023). What is the impact of ChatGPT on education? A rapid review of the literature. *Education Sciences*, 13(4), 410. <https://doi.org/10.3390/educsci13040410>
- Miao, F., & Cukurova, M. (2024). AI competency framework for teachers. UNESCO. <https://doi.org/10.54675/ZJTE2084>
- Miao, F., & Holmes, W. (2023). Guidance for generative AI in education and research. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000386693>
- Miao, F., Shiohira, K., & Lao, N. (2024). AI competency framework for students. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000391105>
- Moorhouse, B. L., Yeo, M. A., & Wan, Y. W. (2023). Generative AI tools and assessment: Guidelines of the world’s top-ranking universities. *Computers and Education Open*, 5, 100151. <https://doi.org/10.1016/j.caeo.2023.100151>

- Ng, D. T. K., Leung, J. K. L., Chu, K. W. S., & Qiao, M. S. (2021). AI literacy: Definition, teaching, evaluation and ethical issues. *Proceedings of the Association for Information Science and Technology*, 58(1), 504–509. <https://doi.org/10.1002/pr2.487>
- Nikolic, S., Daniel, S., Haque, R., Belkina, M., Hassan, G. M., Grundy, S., Lyden, S., Neal, P., & Sandison, C. (2023). ChatGPT versus engineering education assessment: A multidisciplinary and multi-institutional benchmarking and analysis of this generative artificial intelligence tool to investigate assessment integrity. *European Journal of Engineering Education*, 48(4), 559–614. <https://doi.org/10.1080/03043797.2023.2213169>
- Overono, A. L., & Ditta, A. S. (2023). The rise of artificial intelligence: A clarion call for higher education to redefine learning and reimagine assessment. *College Teaching*, 73(2), 123–126. <https://doi.org/10.1080/87567555.2023.2233653>
- Perkins, M. (2023). Academic integrity considerations of AI large language models in the post-pandemic era: ChatGPT and beyond. *Journal of University Teaching & Learning Practice*, 20(2), Article 7. <https://doi.org/10.53761/1.20.02.07>
- Rudolph, J., Tan, S., & Tan, S. (2023). ChatGPT: Bullshit spewer or the end of traditional assessments in higher education? *Journal of Applied Learning & Teaching*, 6(1), 342–363. <https://doi.org/10.37074/jalt.2023.6.1.9>
- Rodríguez, A., Lucas, H. B. D., Álava, W. L. S., & Ávila, X. L. A. (2020). La minería de datos y algunas de sus aplicaciones contextuales. *Serie Científica de la Universidad de las Ciencias Informáticas*, 13(11), 17-25. <https://dialnet.unirioja.es/download/articulo/8590365.pdf>
- Rodríguez Rodríguez, A., & Solórzano Álava, W. L. (2024). Tecnologías inteligentes y sus aplicaciones en la educación. *Serie Científica de la Universidad de las Ciencias Informáticas*, 17(1), 138-153. <http://scielo.sld.cu/pdf/sc/v17n1/2306-2495-sc-17-01-138.pdf>
- Soledispa, L. M. F., Rodríguez, A. R., Figueroa, R. O. A., & Álava, W. L. S. (2025). Modelo estratégico de integración de funciones sustantivas universitarias desde los Problemas Sociales de la Ciencia y la Tecnología. *Serie Científica de la Universidad de las Ciencias Informáticas*, 18(2), 505-528. <https://publicaciones.uci.cu/index.php/serie/article/download/1895/1520>
- Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., & Agyemang, B. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart Learning Environments*, 10, 15. <https://doi.org/10.1186/s40561-023-00237-x>
- Xia, Q., Weng, X., Ouyang, F., Lin, T. J., & Chiu, T. K. F. (2024). A scoping review on how generative artificial intelligence transforms assessment in higher education. *International Journal of Educational Technology in Higher Education*, 21, 40. <https://doi.org/10.1186/s41239-024-00468-z>
- Yero, L. C., Reinaldo, A. C., Rodríguez, A. R., García, J. L. G., & Merino, J. M. (2018). Falacias que atentan contra el desarrollo del pensamiento lógico matemático. *Didasc@lia: Didáctica y Educación*, 9(4), 227-238.