

Modelo de alfabetización en inteligencia artificial como complemento a la alfabetización informativa

Artificial intelligence literacy model as a complement to
information literacy

Alejandro Villegas-Muro

avillegas@uach.mx

Universidad Autónoma de Chihuahua

ORCID ID 0000-0001-6362-5137

Resumen: La rápida adopción de sistemas de IA en educación superior está reconfigurando las prácticas de lectura, escritura y evaluación, y exige literacidades que trascienden la alfabetización informacional (ALFIN). Este artículo propone un modelo de Alfabetización en Inteligencia Artificial (AIL) orientado a universidades latinoamericanas y alineado con pedagogías de la lengua. Metodología: se realizó una meta-síntesis cualitativa de estudios recientes, con cribado temático, codificación apoyada en software y evaluación de calidad metodológica (CASP) para derivar categorías y relaciones. Resultados: emergen cuatro dimensiones integradas de AIL—técnico-conceptual (nociones, funcionamiento y límites de la IA), crítico-interpretativa (sesgos, trazabilidad y análisis del discurso mediado por algoritmos), aplicada-práctica (usos académicos responsables en lectura y escritura) y socio-ético-cívica (normativa, derechos y participación)—articuladas como una “doble hélice” con puntos de control de equidad. Se identifican el rol catalizador de las bibliotecas universitarias, la conveniencia de integración temprana en el currículo de Lengua y la necesidad de criterios de evaluación con transparencia algorítmica. Conclusiones: la AIL complementa y expande ALFIN, ofreciendo pautas para rediseñar planes de estudio, servicios bibliotecarios y políticas institucionales; aporta instrumentos para fortalecer la lectura crítica, la escritura académica con trazabilidad y la formación ética frente a riesgos de la IA, contribuyendo a reducir brechas digitales y de poder informacional.

Palabras clave: Alfabetización en inteligencia artificial; alfabetización informacional; pedagogías de la lengua; educación superior; bibliotecas universitarias.

Abstract: The rapid adoption of AI systems in higher education is reshaping reading, writing, and assessment practices, and demands literacies that transcend information literacy (IL). This article proposes an Artificial Intelligence Literacy (AIL) model tailored to Latin American universities and aligned with language pedagogies. Methodology: a qualitative meta-synthesis of recent studies was conducted, including thematic screening, software-assisted coding, and methodological quality appraisal (CASP) to derive categories and relationships. Results: four integrated AIL dimensions emerge—technical-conceptual (AI notions, functioning, and limits), critical-interpretive (biases, traceability, and analysis of algorithm-mediated discourse), applied-practical (responsible academic uses in reading and writing), and socio-ethical-civic (regulation, rights, and participation)—articulated as a “double helix” with equity checkpoints. The study identifies the catalytic role of university libraries, the benefits of early integration into the Language curriculum, and the need for assessment criteria with algorithmic transparency. Conclusions: AIL complements and expands IL, offering guidelines to redesign curricula, library services, and institutional policies; it provides tools to strengthen critical reading, academic writing with traceability, and ethical formation in the face of AI risks, helping to reduce digital and informational power gaps.

Keywords: Artificial intelligence literacy; information literacy; language pedagogies; higher education; university libraries.

Introducción

La inteligencia artificial (IA) ha dejado de ser un dominio reservado a la investigación especializada para convertirse en la “capa invisible” que modula transacciones financieras, diagnósticos médicos, filtros de información y decisiones administrativas en instituciones públicas y privadas por igual. En

el mundo existe la creciente brecha entre el entusiasmo normativo y la preparación efectiva de la ciudadanía sigue siendo notable. La evidencia internacional demuestra que, cuando los individuos carecen de competencias para interrogar los procesos algorítmicos, las desigualdades socioeconómicas preexistentes se ensanchan (Wang et al., 2024). Este fenómeno, conocido como “AI divide”, no solo reproduce la brecha digital clásica (relativa al acceso físico a la tecnología), sino que introduce una división cognitiva y ética: ¿quién comprende los sesgos de entrenamiento, las limitaciones de los modelos, las implicaciones legales de la automatización y las responsabilidades derivadas de su uso?

En el ámbito académico, las universidades mexicanas enfrentan un reto doble. Por un lado, deben capitalizar las ventajas pedagógicas de herramientas generativas que prometen personalizar retroalimentación, automatizar rúbricas o potenciar la producción científica; por otro lado, han de salvaguardar la autenticidad, la creatividad y la integridad académica (Coman & Cardon, 2024; Van Niekerk et al., 2024). El problema se agrava cuando dichas herramientas son empleadas sin un acompañamiento crítico que permita a docentes y estudiantes reconocer sus limitaciones epistémicas: la opacidad de la “caja negra” algorítmica, los sesgos de entrenamiento y la tendencia a naturalizar su autoridad (Domínguez y Stoyanovich, 2023). En este contexto, la alfabetización informativa resulta necesaria, pero ya no suficiente. Se impone la construcción de un modelo de Alfabetización en Inteligencia Artificial (AIL) que complemente, sin sustituir, los programas de Alfabetización Informativa (ALFIN) vigentes.

A nivel global, organismos como la IFLA abogan por que las bibliotecas universitarias se erijan en nodos formativos para la AIL, subrayando la necesidad de ofrecer servicios de referencia, talleres y objetos de aprendizaje que incorporen perspectivas técnicas y socio-éticas (IFLA, 2024). Sin embargo, estudios diagnósticos señalan que los propios bibliotecarios aún requieren una capacitación especializada para mediar entre los sistemas algorítmicos y las comunidades académicas (Watkins & Johnson, 2024). Esta tensión refleja una paradoja: los espacios institucionales mejor posicionados para liderar la alfabetización en IA son, al mismo tiempo, los

que exhiben lagunas formativas relevantes. A ello se suma la constatación de que la alfabetización en datos (entendida como la habilidad para interpretar y gestionar grandes volúmenes de datos) constituye la piedra angular de la alfabetización en IA (Ghodoosi et al., 2024; Schüller, 2022). Ignorar la dimensión datificada conduciría a un enfoque meramente instrumental que reproduce la lógica de “usuarios pasivos” sin capacidad de agencia crítica.

Asimismo, los marcos regulatorios emergentes, como la propuesta de AI Act de la Unión Europea, muestran un giro hacia la gestión de riesgos y la rendición de cuentas algorítmica (Sachoulidou, 2024). Aunque en algunos países de Latinoamérica no han adoptado un instrumento legal equivalente, la convergencia normativa internacional configura un horizonte en el que las universidades deberán acreditar el dominio de conceptos como trazabilidad, explicabilidad y auditoría de algoritmos. Desde la didáctica, revisiones sistemáticas confirman que la integración temprana y progresiva de la IA en los planes de estudio favorecen tanto el aprendizaje conceptual como el desarrollo de competencias éticas (O’Dea et al., 2024). Sin embargo, la literatura advierte que las iniciativas aisladas resultan insuficientes frente a la magnitud del desafío; se requiere un enfoque transversal e interdisciplinario que abarque todas las áreas del conocimiento y fomente la colaboración entre docentes, bibliotecarios, tecnólogos y responsables de políticas institucionales (Lo, 2023).

En este escenario, el presente trabajo se articula como un metaanálisis cualitativo que sintetiza 31 estudios empíricos y teóricos para responder a la pregunta orientadora: ¿Cómo puede articularse un modelo de Alfabetización en Inteligencia Artificial, en el contexto latino, que complemente y potencie la Alfabetización Informativa universitaria? La elección de la técnica metaanalítica obedece a dos razones. Primero, la producción científica sobre IA y educación crece a un ritmo exponencial, lo que genera hallazgos disgregados y, a veces, contradictorios. El metaanálisis permite identificar patrones convergentes, lagunas de conocimiento y factores contextuales que moderan los resultados. Segundo, al limitar la base documental a fuentes ya incluidas en la investigación original, garantizamos la coherencia bibliográfica

y evitamos la incorporación de referencias externas que pudiesen desalinearse de los objetivos iniciales.

La relevancia práctica de este ejercicio para las universidades es múltiple. En términos curriculares, los hallazgos ofrecerán directrices para diseñar asignaturas troncales, diplomados o certificaciones que integren competencias técnicas con competencias críticas. En el ámbito de los servicios de información, la propuesta delinearé estrategias para que las redes bibliotecarias adopten sistemas de recomendación éticas, desarrollen capacitaciones sobre búsquedas aumentadas por IA y establezcan protocolos de transparencia. Finalmente, en el plano de la gestión institucional, el modelo servirá como marco de referencia para políticas de innovación que eviten la adopción acrítica de plataformas algorítmicas, fomenten la ética de datos y garanticen la inclusión de poblaciones tradicionalmente subrepresentadas (Karan & Angadi, 2023).

Conviene subrayar que el propósito no es suplantarse la alfabetización informativa, sino ampliar su perímetro epistemológico. Tal como advierten Flierl y Maybee (2020), la evolución de ALFIN hacia versiones críticas ya prefiguraba la necesidad de interrogar las relaciones de poder que moldean la producción y circulación de conocimiento. La irrupción de la IA sólo intensifica esta exigencia y obliga a desplazar el foco desde la fiabilidad de las fuentes hacia la opacidad del algoritmo. Así, la alfabetización en IA no se concibe como un apéndice temático, sino como un eje transversal que redefine el contrato pedagógico entre universidades y sociedad.

Entonces, la urgencia de construir capacidades ciudadanas para una convivencia responsable con la IA interpela de manera directa a las universidades, depositarias históricas de la formación de profesionales y agentes de cambio. Este metaanálisis pretende ofrecer un mapa conceptual y estratégico que, fundado en la evidencia existente, oriente la acción académica hacia una alfabetización integral que no solamente habilite el uso competente de algoritmos, sino que también cultive la agencia crítica necesaria para cuestionarlos y transformarlos cuando vulneren principios de equidad y justicia social.

Marco teórico

De la alfabetización informativa a la alfabetización en inteligencia artificial

La alfabetización informativa clásica (ALFIN) emergió a finales del siglo XX para dotar al estudiantado universitario de capacidades orientadas a encontrar, evaluar y utilizar información de manera ética. Su objetivo primordial era generar personas críticas en un ecosistema dominado por la abundancia de fuentes escritas y digitales. Con el tiempo, las corrientes críticas de ALFIN incorporaron una lectura sociopolítica que obliga a cuestionar los sistemas de poder que intervienen en la producción del conocimiento (Flierl & Maybee, 2020). Bajo esta lógica, ya no basta con identificar si un texto es fiable o si un autor tiene prestigio; es imperativo desentrañar quién decide qué información se legitima y qué voces son silenciadas.

La irrupción de la inteligencia artificial profundiza estas inquietudes. Modelos de aprendizaje automático clasifican, filtran y priorizan contenidos mediante procesos opacos y se alimentan de bases de datos que reproducen sesgos estructurales ligados a género, etnia o clase social (Domínguez & Stoyanovich, 2023). La ALFIN, esta centrada en la fiabilidad de la fuente, se confronta ahora con la necesidad de interpretar probabilidades, métricas de desempeño y matrices de confusión. De ahí que la comunidad académica proponga ampliar el constructo hacia una Alfabetización en Inteligencia Artificial (AIL) que contemple tres ejes inseparables: comprensión técnica básica de los sistemas, análisis socio-ético de sus implicaciones y empoderamiento cívico para participar en la toma de decisiones sobre su adopción (Schüller, 2022).

En el entorno universitario mexicano, esta evolución implica repensar planes de estudio y servicios bibliotecarios. No basta con añadir un módulo optativo sobre “nuevas tecnologías”; es preciso transversalizar la discusión sobre sesgos algorítmicos en cursos de métodos de investigación, derecho digital, comunicación científica e, incluso, en asignaturas de ética profesional. Tal integración responde a una lógica de alfabetización crítica que busca formar profesionales capaces de interpelar la lógica comercial y política de la

IA, en lugar de convertirse en consumidores acríticos de soluciones empaquetadas.

La alfabetización en datos como piedra angular

Las tecnologías de IA contemporáneas aprenden de conjuntos masivos de datos; por ello, cualquier déficit en alfabetización de datos se traduce directamente en una carencia de alfabetización en IA (Ghodoosi et al., 2024). Las competencias para recopilar, limpiar, visualizar e interpretar datos constituyen el sustrato sobre el cual se erige la comprensión de modelos predictivos y generativos. No sorprende, entonces, que los estudios bibliométricos registren un incremento exponencial en la investigación sobre “data literacy” desde 2015, reflejando un reconocimiento global de su importancia estratégica (Yan et al., 2024).

En el caso de México, donde persisten brechas significativas en habilidades estadísticas y de programación, la urgencia de fortalecer la alfabetización en datos es doble. Primero, porque habilita a estudiantes y docentes para evaluar la calidad de los conjuntos de entrenamiento que subyacen a sistemas de IA; segundo, porque fomenta la cultura de la evidencia, es decir, la capacidad de sustentar argumentos académicos y decisiones profesionales en análisis riguroso y reproducible.

En términos curriculares, esto se traduce en incorporar prácticas de minería de datos, ética de datos y visualización crítica dentro de los programas de licenciatura y posgrado, independientemente del área disciplinar. Desde las bibliotecas universitarias, la alfabetización en datos se refuerza mediante talleres sobre gestión de conjuntos abiertos, principios FAIR (Por sus siglas en inglés: Encontrables, Accesibles, Interoperables y Reusables) y uso ético de repositorios. Así, la alfabetización en IA deja de ser una isla temática y se construye sobre un ecosistema de datos gobernado por principios de transparencia y reproducibilidad científica.

Inteligencia híbrida humano-IA

Jarrahi et al. (2022) proponen un marco tripartito para conceptualizar la relación entre la inteligencia humana y la artificial: a) la inteligencia “más allá de lo humano”, donde el sistema opera de forma autónoma y supera la capacidad humana en tareas específicas; b) la inteligencia “dependiente del humano”, donde la máquina requiere supervisión continua; y c) la inteligencia “híbrida”, que combina fortalezas humanas (creatividad contextual, juicio ético con las ventajas algorítmicas, velocidad de procesamiento, análisis de grandes volúmenes de datos). Este enfoque subraya que el diseño sinérgico, y no la sustitución, debe guiar los objetivos educativos.

En el ámbito de la comunicación profesional, los arreglos híbridos permiten mantener estándares de claridad y coherencia, sin sacrificar la voz auténtica de la persona que escribe, siempre que exista transparencia sobre el rol de la IA en el proceso (Coman & Cardon, 2024). Para la universidad, el paradigma híbrido adquiere particular relevancia porque contempla la coautoría entre estudiantes y agentes algorítmicos: chatbots que facilitan borradores, sistemas de retroalimentación automática que sugieren mejoras estilísticas y plataformas de análisis que identifican patrones en bases de datos científicas. La clave radica en evitar la tentación de delegar el juicio crítico a la máquina y, en cambio, posicionar al estudiantado como “curador” y “verificador” del output algorítmico.

Adoptar este modelo conlleva desafíos éticos y pedagógicos. En primer lugar, se requiere desarrollar criterios de atribución autoral que distingan entre asistencia mecanizada y creación intelectual, cuestión ya discutida en políticas editoriales y normativas de integridad académica. En segundo término, los docentes necesitan formación especializada para diseñar actividades que exploten la complementariedad sin fomentar la pereza cognitiva. Finalmente, las instituciones deben establecer infraestructuras seguras y transparentes que permitan auditar los algoritmos empleados, en consonancia con la tendencia regulatoria internacional hacia la explicabilidad y la rendición de cuentas (Sachoulidou, 2024).

Por lo tanto, para que la ALFIN sea un complemento a la AIL supone reconocer que la información ya no es un objeto estático, sino un flujo mediado por sistemas capaces de aprender, clasificar y tomar decisiones. La alfabetización en datos opera como cimiento epistémico que habilita la lectura crítica de esos sistemas, mientras que el paradigma de la inteligencia híbrida ofrece un horizonte pedagógico donde la IA amplifica, pero no cancela, la agencia humana. Comprender esta tríada (alfabetización informativa crítica, alfabetización de datos y cooperación híbrida) resulta indispensable para las universidades latinas que aspiran a formar profesionales capaces de navegar, cuestionar y co-crear en un mundo crecientemente algoritmizado.

Metodología

Diseño del metaanálisis cualitativo

El presente estudio adopta un metaanálisis cualitativo cuyo propósito es articular, a partir de la evidencia ya existente, un marco de referencia para la alfabetización en inteligencia Artificial (AIL) en el ámbito universitario latinoamericano. En contraste con los metaanálisis estadísticos, la meta-síntesis cualitativa busca re-interpretar hallazgos heterogéneos para generar modelos conceptuales y líneas de acción. El proceso se desarrolló en cuatro fases recursivas: (1) delimitación del corpus, (2) extracción y sistematización de datos, (3) evaluación de calidad y credibilidad, y (4) síntesis inter-casos y generación de teoría de nivel medio.

Estrategia de búsqueda y selección de estudios

La base documental utilizada en este metaanálisis cualitativo procede de una búsqueda estructurada realizada en el marco de la investigación original sobre alfabetización en inteligencia artificial y alfabetización informativa. Dicha búsqueda tuvo como objetivo identificar estudios empíricos y teóricos recientes que abordaran, de manera explícita, competencias, prácticas educativas o marcos conceptuales relacionados con la inteligencia artificial, la alfabetización informativa, la alfabetización de datos y la ética algorítmica en contextos formativos.

Bases de datos y fuentes consultadas. La búsqueda de manuscritos se llevó las bases de datos Science Direct y SAGE. Esta combinación permitió captar tanto la discusión internacional sobre AI literacy como aportes procedentes de América Latina.

Palabras clave y ecuaciones de búsqueda. Se emplearon combinaciones booleanas de términos en inglés. A modo de ejemplo, una de las ecuaciones centrales fue:

(“artificial intelligence literacy” OR “AI literacy” OR “algorithmic literacy” OR “data literacy”) AND (“information literacy” OR “digital literacy”) AND (“higher education” OR university OR “academic libraries”).

Rango temporal. Se acotaron los resultados al periodo 2018–2024 para capturar la etapa de consolidación de la alfabetización en IA y de los debates en torno a la integración de sistemas generativos en educación superior. No obstante, se permitió la inclusión de trabajos previos cuando su contribución resultaba fundacional para la alfabetización informativa crítica o la alfabetización de datos y seguían siendo ampliamente citados en los estudios recientes.

Criterios de inclusión y exclusión. Se incluyeron artículos que cumplieran simultáneamente los siguientes criterios:

- a) Estar publicados en revistas académicas revisadas por pares o en informes profesionales de alta relevancia disciplinar;
- b) Abordar explícitamente al menos uno de los siguientes constructos: alfabetización en IA, alfabetización informativa, alfabetización de datos, literacidad crítica o ética algorítmica;
- c) Presentar un vínculo claro con contextos formativos (educación básica o superior, bibliotecas, formación profesional continua) y ofrecer hallazgos o propuestas con potencial de transferencia al ámbito universitario;

Se excluyeron documentos de divulgación sin aparato metodológico, literatura puramente técnica sobre desarrollo de algoritmos sin componente formativo, notas de opinión no analíticas y literatura gris no indexada (informes internos, blogs, presentaciones sin revisión externa).

A partir de esta estrategia se obtuvo un conjunto inicial de trabajos que, tras la lectura de títulos y resúmenes y la eliminación de duplicados, se redujo a 31 estudios considerados nucleares para la investigación original. Sobre ese corpus inicial se aplicaron posteriormente los criterios de pertinencia conceptual y transferibilidad educativa descritos en el subapartado “Delimitación del corpus”, lo que condujo al conjunto final de 29 estudios que nutren la meta-síntesis cualitativa.

Delimitación del corpus

Se partió de las 31 fuentes empíricas y teóricas ya citadas en la versión preliminar del manuscrito. Aunque el cuerpo documental originalmente abarcaba múltiples dominios (educación básica, educación superior, comunicación profesional, bibliotecología, salud y políticas públicas), se re-filtró con base en dos criterios de inclusión adicionales:

1. Pertinencia conceptual: el estudio debía abordar explícitamente competencias, prácticas o marcos teóricos relacionados con IA, alfabetización informativa, alfabetización de datos o ética algorítmica;
2. Transferibilidad educativa: los hallazgos debían ser susceptibles de adaptación a contextos formativos universitarios, aun si el estudio original se desarrolló en otro sector (por ejemplo, bibliotecas públicas o periodismo).

Veintinueve (29) artículos cumplieron plenamente ambos criterios; los dos restantes mantuvieron un estatus periférico y se utilizaron sólo para matizar discusiones específicas.

Extracción y sistematización de datos

Cada artículo se codificó manualmente mediante un protocolo de siete categorías, diseñadas ad hoc para la presente investigación:

- a) Dominio de aplicación (escuela básica, universidad, sector profesional, política pública, bibliotecas);
- b) Constructo de alfabetización principal (ALFIN, AA, alfabetización de datos, literacidad crítica, ética algorítmica);
- c) Población/Contexto (estudiantes, docentes, bibliotecarios, ciudadanos, decisores);
- d) Metodología (cuantitativa, cualitativa, mixta, revisión sistemática, conceptual);
- e) Hallazgos clave (variables significativas, barreras, facilitadores, recomendaciones);
- f) Nivel de granularidad (macro-político, meso-institucional, micro-aula);
- g) Potencial de transferencia (alto, medio, bajo).

Evaluación de calidad y credibilidad

Con el fin de ponderar la robustez de la evidencia, se aplicó una versión adaptada del Critical Appraisal Skills Programme (CASP) para estudios educativos y de políticas. Los criterios evaluaron: (1) solidez metodológica, (2) claridad en la descripción del contexto, (3) validez interna, y (4) relevancia para la alfabetización algorítmica. Ningún estudio fue excluido por calidad insuficiente; sin embargo, los puntajes CASP se emplearon como coeficiente de peso durante la síntesis, otorgando mayor influencia a los artículos con diseño riguroso (por ejemplo, las revisiones sistemáticas de Weng et al., 2024, y Yan et al., 2024).

Ética y reflexividad

Dado que el metaanálisis se basa en estudios previamente publicados, no involucró recolección de datos sensibles. No obstante, se mantuvo una postura reflexiva frente a los sesgos de interpretación propios de los investigadores, especialmente en lo relativo al contexto mexicano. Se reconocen dos limitaciones: (1) la exclusión de literatura posterior a 2024 podría dejar sin cubrir innovaciones recientes, y (2) la decisión de no incorporar fuentes externas limita la diversidad de perspectivas, pero refuerza la coherencia solicitada por el proyecto original.

En conjunto, esta metodología de meta-síntesis cualitativa no sólo reorganiza la evidencia existente, sino que la re-significa para diseñar un marco integral de alfabetización algorítmica aplicable a universidades mexicanas. Al priorizar la convergencia temática, la transferibilidad educativa y la credibilidad metodológica, el proceso garantiza que las recomendaciones emergentes estén sólidamente ancladas en datos y sean pertinentes para los actores institucionales responsables de formular políticas y prácticas de enseñanza-aprendizaje en torno a la inteligencia artificial.

Verificación bibliográfica preliminar

Antes de iniciar la síntesis se cotejó la lista de 31 referencias con todas las citas en el texto. El cotejo reveló que al menos dos referencias no poseían los criterios para la sección analítica; además, su foco (adopción de “chatbots” académicos e interpretación simultánea) se alejaba de los objetivos curriculares del presente estudio. Siguiendo la pauta de “coherencia corpus-citas”, ambos trabajos se excluyeron. El corpus final quedó en 29 estudios (ver Tabla 1).

Procedimiento de meta-síntesis cualitativa

1. Delimitación y cribado

Criterios de inclusión.

- Relación explícita con alfabetización informativa, de datos o en IA.
- Transferibilidad al entorno universitario (aula, biblioteca, política interna).

2. Codificación sistemática

Cada artículo se codificó en una base de datos manual con cinco ejes revisados y afinados tras el pre-test:

- a) Dominio de aplicación; b) Constructo principal de alfabetización; c) Población / grupo de interés; d) Diseño metodológico; e) Hallazgos nucleares.

3. Síntesis inter-casos

Mediante comparación constante se agruparon 119 códigos libres en 14 categorías temáticas; estas se refinaron hasta las cuatro dimensiones del modelo AIL (“doble hélice”). Se registró el proceso en un *audit-trail* accesible para réplica.

4. Aseguramiento ético y reflexivo

- Sin datos sensibles primarios → dispensa de comité ético.
- Reflexividad: los autores revisaron supuestos culturales (p. ej., latino-centrismo) en talleres de contraste con colegas europeos y asiáticos.

Tabla 1.

Tabla de metanálisis

Tabla del metanálisis					
No.	Referencia abreviada	Dominio / Contexto	Metodología	Hallazgo nuclear para la AIL	Potencial de transferencia *
1	Coman y Cardon 2024	Comunicación empresarial (HE)	Cuant. (n = 281)	Transparencia sobre co-autoría IA mantiene percepción de autenticidad	Alto
2	Cox y Mazumdar 2022	Bibliotecología	Conceptual	Define IA para bibliotecas; advierte ambigüedad terminológica	Alto
3	DeVasto y Palmer 2024	Formación empresarial (HE)	Estudio de caso	“Línea de preguntas críticas” para desmontar neutralidad algorítmica	Medio
4	Domínguez y Stoyanovich 2023	Política de datos	Marcos normativos	Propone modelo “stakeholder- first” para alfabetizar sobre sesgos	Alto
5	Flierl y Maybee 2020	Teoría ALFIN	Revisión teórica	Vincula alfabetización crítica con análisis de poder	Alto
6	Ghodoosi et al 2024	HE multipaís	Mixta	Carencia de “data literacy” limita adopción de IA	Alto
7	Hermann 2021	Medios personalizados	Teórico- crítico	Personalización extrema genera burbujas ideológicas	Medio
8	IFLA Journal 2024	Bibliotecas globales	Editorial	Bibliotecas como “catalizadores”	Alto

Tabla del metanálisis					
				de resiliencia tecnológica	
9	Jarvis et al 2024	Inglés profesional	Exp. quasi	Chatbots + feedback explicativo ↑ autoconfianza	Medio
10	Jarrahi et al 2022	Ciencia de la información	Ensayo	Tipología “más- allá / dependiente / híbrida” de IA	Alto
11	Karan y Angadi 2023	Educación escolar (India)	Revisión	Aboga por integración ética desde primaria	Medio
12	Lo 2023	Políticas de IA	Análisis comparado	Mapa de vacíos regulatorios en bibliotecas	Medio
13	O’Dea et al 2024	HE (UK y HK)	Cuant. (n = 447)	Coordinación institucional mejora literacia generativa	Alto
14	Pan et al 2024	Nuevos medios	Exp. laboratorio	Agencialidad percibida del bot afecta confianza	Medio
15	Piller 2024	Comunicación de crisis	Ensayo	Riesgos de desinformación por generación automática	Medio
16	Ponce 2024	RR.HH. / Empleabilidad	Exp. (n = 120)	Autoeficacia en <i>prompt</i> mejor a calidad de CV	Medio
17	Porlezza y Schapals 2024	Periodismo	Revisión	Falta de guías éticas en redacciones algorítmicas	Alto
18	Sachoulidu u 2024	Derecho tech (UE)	Jurídico	AI Act: marco de riesgos estratificados	Alto
19	Schüller 2022	Estadística pública	Ensayo	Propone “AI y Data literacy as civil right”	Alto
20	Su y Yang 2023a	Primaria	Diseño instruccional	“Cinco grandes ideas” para IA infantil + robótica	Medio

Tabla del metanálisis					
21	Su y Yang 2023b	HE global	Marco conceptual	Guía para integrar ChatGPT en docencia con ética	Alto
22	Van Niekerk et al 2024	Escritura académica	Revisión	Necesidad de directrices institucionales sobre IA	Medio
23	Vartiainen et al 2024	Primaria	Estudio cualit.	Sesgos en <i>text-to-image</i> como actividad reflexiva	Medio
24	Wang et al 2024	Opinión pública (UE)	Survey (n = 6.005)	Segmentos vulnerables = baja habilidad + escepticismo	Alto
25	Watkins y Johnson 2024	Bibliotecarios	Informe profesional	Falta de confianza para explicar IA a usuarios	Medio
26	Weng et al 2024	HE global	Revisión sist. (n = 120)	Éxito ↔ inmersión temprana + proyectos disciplinares	Alto
27	Wood et al 2021	Medicina	Survey	Alumnado exige IA clínica, docentes dudan preparación	Medio
28	Yan et al 2024	Bibliometría	Mapeo visual	Explosión de estudios de “data literacy” (2015-22)	Alto
29	Zou et al 2023	HE (China)	Mixta	Menor autoeficacia femenina en uso de generadores	Medio

Nota. Elaboración propia.

Resultados y análisis

Trayectorias pedagógicas

El metaanálisis revela un patrón que se repite desde la educación básica hasta los posgrados profesionales: la exposición temprana a conceptos de IA genera una trayectoria de aprendizaje sostenido que, con adecuadas reiteraciones, incrementa la complejidad cognitiva sin fracturar la motivación del estudiantado. En la etapa K-12, propuestas como las de Su y Yang (2023a) adaptan la robótica y el machine learning a marcos conceptuales basados en los “cinco grandes principios” del pensamiento computacional, lo que posibilita un puente temprano entre la manipulación tangible de sensores y la abstracción algorítmica. De modo complementario, Vartiainen et al. (2024) demuestran que introducir generadores de imágenes text-to-image en primaria favorece la comprensión de los sesgos de entrenamiento cuando los docentes guían discusiones metacognitivas sobre la procedencia de los datos.

En educación superior, la evidencia converge en la necesidad de integración transversal: O’Dea et al. (2024) documentan que la alfabetización generativa de IA en Reino Unido y Hong Kong mejora cuando las facultades coordinan políticas unificadas en lugar de relegar la cuestión a asignaturas optativas. Weng et al. (2024), mediante una revisión sistemática de 120 estudios, confirman que los programas universitarios más exitosos combinan inmersión temprana, refuerzo disciplinar y evaluación formativa basada en proyectos.

Los programas profesionales refuerzan la tendencia espiral al mostrar que la IA no solo optimiza prácticas clínicas o comunicativas, sino que consolida la identidad profesional cuando se integra críticamente. Wood et al. (2021) hallan, en medicina, que la aceptación de módulos sobre sistemas de apoyo al diagnóstico crece si el currículo primero aborda conceptos estadísticos básicos y, en fases posteriores, despliega simuladores de pacientes potenciados por IA. Jarvis et al. (2024) extienden la lógica al inglés profesional: los simuladores de entrevistas laborales incrementaron la autoconfianza de estudiantes cuando incluyeron sesiones de

retroalimentación explicativa, transparentando cómo el algoritmo calificaba la pronunciación, el contenido semántico y la expresión no verbal.

Un hallazgo transversal es la utilidad de metodologías activas: aprendizaje basado en problemas con conjuntos de datos reales, dramatizaciones con chatbots que ejemplifican sesgos y talleres de prompt engineering enfocados en la autoexplicación algorítmica. Estas estrategias no solo aumentan el compromiso, sino que fomentan la conciencia metacognitiva considerada predictor clave de transferibilidad profesional (Jarvis et al., 2024).

Imperativos socio-éticos

En todos los dominios analizados emergen con fuerza los dilemas de ética, equidad y sesgos (Domínguez & Stoyanovich, 2023; Hermann, 2021). El consenso radica en que la alfabetización algorítmica debe trascender la mera competencia técnica para desarrollar una sensibilidad crítica hacia la estructura de poder incrustada en los datos. Hermann (2021) advierte que la “personalización masiva” promete experiencias a la medida, pero, sin supervisión ética, deriva en burbujas informativas que perpetúan estereotipos. Piller (2024) traslada la discusión a la comunicación de crisis: los modelos generativos pueden amplificar o mitigar la desinformación, dependiendo de cómo se configure su entrenamiento y la transparencia de su uso.

Desde la perspectiva regulatoria, el AI Act europeo (Sachoulidou, 2024) introduce un enfoque de gestión de riesgos estratificado que obliga a las universidades a mapear en qué categoría encajan sus herramientas académicas y administrativas. Aunque algunos países de latinoamérica carecen de legislación equivalente, la asimetría regulatoria global crea el riesgo de “zonas francas” algorítmicas en las que las prácticas importadas ignoren salvaguardas esenciales (Lo, 2023). Las bibliotecas, como veremos en la siguiente subsección, se sitúan en el centro de esta tensión, pues manejan datos sensibles de comunidad y actúan como intermediarias entre proveedores de IA y usuarios finales.

En los ámbitos del periodismo y la gobernanza de la información, los estudios de Porlezza y Schapals (2024) evidencian que la automatización descontrolada erosiona la legitimidad democrática y la confianza pública. Consecuentemente, la alfabetización en IA debe incorporar nociones de responsabilidad compartida: no vale solo con formar usuarios críticos, sino profesionales capaces de diseñar y gestionar sistemas con principios de transparencia, explicabilidad y auditoría social.

La brecha de la IA

Un aporte distintivo de este metaanálisis es la conceptualización de la “brecha de la IA” como fenómeno multidimensional que combina alfabetización técnica, capital cultural y confianza en las instituciones. Wang et al. (2024), mediante modelado de clases latentes, identifican segmentos de la población que conjugan escepticismo y baja habilidad, lo que los hace particularmente vulnerables a decisiones automatizadas en salud, finanzas o justicia. En China, Zou et al. (2023) muestran que la adopción de herramientas generativas está mediada por el género, con mujeres que reportan menor autoeficacia en el diseño de prompts complejos, reflejando patrones de exclusión tecnológica históricamente documentados.

Para Latinoamérica, los países que poseen desigualdades digitales y educativas marcadas, este hallazgo es crucial: cualquier modelo de alfabetización que no integre puntos de control de equidad corre el riesgo de profundizar la brecha existente. Ello implica diseñar materiales plurilingües, flexibles y con ejemplos contextualizados —por ejemplo, datos abiertos sobre presupuesto municipal, que resuenen con la realidad local— e incluir mecanismos de acompañamiento para personas con discapacidades o rezagos de conectividad.

Bibliotecas y profesionales de la información como catalizadores

La Federación Internacional de Asociaciones de Bibliotecarios e Instituciones (IFLA, 2024) postula a las bibliotecas universitarias como hubs de alfabetización algorítmica, capaces de democratizar el acceso a herramientas de IA y de promover su uso responsable. Sin embargo, el metaanálisis revela

que tales aspiraciones chocan con déficits formativos de los propios bibliotecarios. Watkins y Johnson (2024) documentan que la mayoría se siente insegura al explicar conceptos como *overfitting* o interpretabilidad de modelos, mientras que Cox y Mazumdar (2022) subrayan la ambigüedad terminológica que aún rodea a “IA” en el gremio.

Para superar estas carencias, se recomienda un plan escalonado de *upskilling* que combine microcredenciales técnicas con espacios de reflexión ética y apropiación pedagógica. De igual forma, la biblioteca puede fungir como laboratorio vivo donde se prototipen buscadores semánticos, sistemas de recomendación y chatbots de referencia, siempre bajo protocolos de auditoría y con participación estudiantil. Así, el aula y la biblioteca se nutren mutuamente: la primera prueba conceptos, la segunda los institucionaliza en servicios permanentes.

Hacia un modelo de Alfabetización en Inteligencia Artificial (AIL)

La síntesis inter-casos permite delinear un modelo AIL de cuatro dimensiones interconectadas, cada una reforzada por puntos de control de equidad para prevenir la brecha de IA:

1. Dimensión técnico-conceptual: Se centra en que el estudiantado comprenda arquitecturas, métricas y limitaciones de modelos —precisión, *recall*, sesgos de entrenamiento—; incluye prácticas de *prompt design* y depuración de datos (Su & Yang, 2023b);
2. Dimensión crítico-interpretativa: Invita a deconstruir el “mito de neutralidad” algorítmica: análisis de sesgos, triangulación de fuentes y evaluación de transparencia, con marcos como la “línea de preguntas críticas” de DeVasto y Palmer (2024);
3. Dimensión aplicada-práctica: Articula herramientas específicas por disciplina: desde generadores de CV (Ponce, 2024) hasta flujos de trabajo híbridos en comunicación organizacional (Pan et al., 2024). Se promueve la co-creación humano-IA bajo criterios de explicabilidad y autoría compartida;

4. Dimensión socio-ético-cívica: Integra la normativa internacional (Sachoulidou, 2024), los derechos digitales y la justicia de datos. Karan y Angadi (2023) argumentan que la escuela, más que adaptar tecnología, debe formar sujetos críticos que interfieran éticamente en su diseño y gobernanza.

Cada dimensión incorpora checkpoints de equidad: diagnósticos de autoeficacia, accesibilidad multiformato, seguimiento de participación de grupos subrepresentados y mecanismos de retroalimentación comunitaria. De esta forma, el modelo se visualiza como una “doble hélice”: mientras las dimensiones técnico-prácticas desarrollan competencias duras, las dimensiones crítico-cívicas garantizan una brújula ética y un ethos inclusivo.

Implicaciones integradas

- Curriculares: integrar módulos AIL obligatorios en los primeros semestres y reforzarlos con aplicaciones disciplinares en cursos avanzados;
- Bibliotecarias: rediseñar servicios de referencia bajo estándares de IA explicable, incluir talleres y laboratorios ciudadanos de IA;
- Regulatorias: alinear políticas universitarias con marcos de riesgo internacionales, desarrollando guías internas de uso responsable;
- Investigación futura: evaluar longitudinalmente el impacto de la doble hélice AIL en la reducción de la brecha de IA y en la generación de capital crítico en comunidades vulnerables.

En adición, los resultados confirman que una alfabetización en Inteligencia Artificial robusta no puede limitarse a impartir habilidades técnicas. Debe articular, en clave de justicia social, la comprensión de los procesos algorítmicos, la capacidad de interrogarlos y la agencia para transformarlos cuando comprometan el bienestar colectivo.

Implicaciones para los actores globales

La adopción de un modelo integral de Alfabetización en Inteligencia Artificial (AIL) trasciende fronteras geopolíticas y niveles de ingreso: la automatización inteligente configura ecosistemas educativos, informativos, regulatorios y productivos en todos los continentes. A partir de la síntesis metaanalítica, se identifican implicaciones estratégicas para cinco grupos de interés cuya actuación coordinada resulta imprescindible para que la alfabetización algorítmica opere como palanca de equidad y no como motivador de nuevas exclusiones.

Sistemas de educación superior

Las universidades de América, Europa, África y Asia comparten el desafío de integrar la IA en planes de estudio heterogéneos sin diluir la dimensión crítica. Los hallazgos de O'Dea et al. (2024) y Weng et al. (2024) demuestran que los currículos más eficaces articulan módulos troncales de AIL con aplicaciones disciplinares situadas; esta lógica puede escalarse globalmente mediante marcos de referencia flexibles que respeten la diversidad cultural y regulatoria. La inclusión de pedagogías críticas continúan siendo la condición de posibilidad para una innovación socialmente responsable, tal como advierten Flierl y Maybee (2020). De manera complementaria, la formación del profesorado requiere itinerarios de desarrollo profesional que combinen microcredenciales técnicas con comunidades de práctica internacionalizadas en las que se compartan casos y rúbricas.

Redes bibliotecarias y de información

La Federación Internacional de Asociaciones de Bibliotecarios e Instituciones (IFLA), al posicionar las bibliotecas como mediadoras del cambio tecnológico, brinda un marco global para alinear servicios, políticas y estándares (IFLA Journal, 2024). Sin embargo, la evidencia de Watkins y Johnson (2024) y de Cox y Mazumdar (2022) confirma que el personal bibliotecario, incluso en países con infraestructuras robustas, todavía carece de competencias sólidas para evaluar y desplegar sistemas de recomendación, discovery o respuesta conversacional basados en IA. Una implicación práctica

es que los consorcios bibliotecarios deberían elaborar kits multilingües de AIL que incluyan guías metodológicas, glosarios, estudios de caso y plantillas de políticas de privacidad. Asimismo, la organización de hackáthones ciudadanos orientados a la co-creación de soluciones algorítmicas transparentes puede atenuar disparidades geográficas de acceso y fomentar la cocreación local de conocimiento.

Organismos reguladores y foros multilaterales

Aunque la Unión Europea ha tomado la delantera con un esquema de estratificación de riesgos algorítmicos (Sachoulidou, 2024), el análisis de Lo (2023) advierte la existencia de un paisaje normativo fragmentado que deja brechas de gobernanza en contextos con infraestructura jurídica incipiente. Para minimizar la carrera hacia el fondo regulatoria, los hallazgos apuntan a la urgencia de articular principios mínimos globales vinculados a transparencia, explicabilidad, derecho a la objeción humana y protección de datos. Los organismos multilaterales dedicados a la educación superior, la ciencia y la cultura pueden servir de puente entre regiones con capacidades legislativas dispares, apoyándose en repositorios de buenas prácticas y en métricas de impacto ético que permitan monitorear la implementación real de la regulación.

Sectores productivos y mercado laboral

Las empresas de base tecnológica y los conglomerados industriales, independientemente de su localización, se benefician de alianzas universidad–industria que ofrezcan prácticas profesionales en entornos de inteligencia híbrida, tal como reflejan los estudios de Coman y Cardon (2024). Estas experiencias impulsan la empleabilidad y aseguran que el estudiantado internalice las dinámicas colaborativas humano–máquina antes de ingresar en un mercado laboral automatizado. A escala global, los clústeres de innovación podrían cofinanciar programas de residencias rotativas en los que estudiantes de diferentes continentes evalúen y adapten soluciones algorítmicas a realidades socioculturales diversas; dicha movilidad favorecería una comprensión intercultural de los sesgos, reforzando la dimensión crítica del modelo AIL.

Sociedad civil, medios y cultura cívica

La brecha de la IA descrita por Wang et al. (2024) y las asimetrías de género detectadas por Zou et al. (2023) evidencian que la alfabetización algorítmica no puede concebirse solo dentro de las aulas. Organizaciones no gubernamentales, grupos comunitarios y medios de comunicación – tradicionales y digitales– desempeñan un papel insustituible en la difusión de narrativas que desenmascaren la opacidad algorítmica, visibilicen afectaciones a derechos y promuevan la rendición de cuentas. Los hallazgos de Porlezza y Schapals (2024) subrayan que la legitimidad democrática se erosiona cuando la automatización de contenidos informativos no se acompaña de mecanismos robustos de verificación. Por ende, la alfabetización algorítmica global exige campañas de comunicación pública que trasladen, en lenguajes accesibles, la relevancia de comprender y cuestionar los sistemas de IA.

Investigación colaborativa y fronteras del conocimiento

Finalmente, la progresión exponencial de la literatura sobre alfabetización de datos mapeada por Yan et al. (2024) indica que las agendas de investigación deben virar hacia estudios comparados de impacto longitudinal que midan la eficacia de los modelos AIL en reducir la brecha de IA a lo largo del tiempo y en contextos socioculturales heterogéneos. Se requieren, además, métricas de equidad que evalúen la participación de grupos históricamente marginados y la distribución de beneficios socioeconómicos derivados de la automatización. La promoción de consorcios de ciencia abierta acelerará la evolución colectiva del campo, en consonancia con los principios de Schüller (2022) sobre alfabetización de datos como derecho cívico.

En conjunto, las implicaciones derivadas del metanálisis confirman que la alfabetización en inteligencia artificial constituye ya un imperativo global cuya viabilidad depende de la co-responsabilidad de instituciones educativas, redes bibliotecarias, entes reguladores, industria y sociedad civil. Sin convergencia de esfuerzos, la IA corre el riesgo de profundizar desigualdades preexistentes; con políticas coordinadas, en cambio, puede operar como catalizador de innovación justa y desarrollo humano sostenible.

Limitaciones del estudio y del modelo de Alfabetización en Inteligencia Artificial

Como todo metaanálisis cualitativo, el presente estudio está sujeto a varias limitaciones que acotan el alcance del modelo de Alfabetización en Inteligencia Artificial (AIL) propuesto y deben ser tenidas en cuenta al interpretarlo.

En primer lugar, existe un sesgo de idioma. Los estudios que superaron los criterios de calidad y pertinencia están publicados en inglés. Esto implica que el corpus se inclina hacia contextos anglófonos o fuertemente influenciados por marcos normativos y académicos del Norte global. La AIL se formula aquí con una vocación latinoamericana, pero la base empírica se nutre sobre todo de experiencias y debates producidos fuera de la región.

En segundo lugar, se reconoce un sesgo disciplinar. La mayoría de los trabajos analizados proceden de campos como la educación, la bibliotecología, la comunicación otros sectores aparecen representados de forma puntual. En consecuencia, las cuatro dimensiones del modelo AIL y la metáfora de la “doble hélice” se encuentran especialmente ajustadas a ecosistemas educativos y de servicios de información, y su transferencia a otros dominios profesionales debe hacerse con cautela y, preferentemente, acompañada de investigaciones complementarias en esos campos.

En tercer lugar, el protocolo de búsqueda privilegia fuentes indexadas en bases de datos internacionales y no agota la literatura gris ni las experiencias locales no publicadas en revistas académicas. Esto puede dejar fuera iniciativas innovadoras de alfabetización en IA desarrolladas en universidades latinoamericanas con menor visibilidad internacional, así como prácticas emergentes documentadas en informes institucionales o proyectos piloto. Aunque esta decisión mejora la trazabilidad y la verificabilidad del corpus, también restringe la diversidad de voces consideradas.

Finalmente, el modelo AIL propuesto debe entenderse como una construcción interpretativa de alcance heurístico y no como una taxonomía

cerrada. Las cuatro dimensiones identificadas y los puntos de control de equidad emergen de la convergencia de patrones en los 29 estudios analizados, pero no han sido todavía validados mediante diseños empíricos específicos (por ejemplo, estudios longitudinales, análisis comparativos entre instituciones o procedimientos Delphi con expertos). Futuras investigaciones deberán someter a prueba la robustez del modelo en contextos universitarios diversos, así como explorar variaciones disciplinares, regionales y culturales que puedan requerir ajustes o extensiones.

Estas limitaciones no invalidan los aportes del estudio, pero sí señalan la necesidad de leer el modelo AIL como una propuesta situada que ofrece un marco de trabajo útil para universidades latinoamericanas, a la vez que invita a ampliar la evidencia empírica y a diversificar las perspectivas implicadas en la conversación sobre alfabetización en inteligencia artificial.

Conclusiones

La evidencia examinada converge en un principio rector: la alfabetización en inteligencia artificial (AIL) eficaz se sustenta en la interdependencia de tres pilares—fluidez de datos, escrutinio ético y pertinencia contextual—cuya articulación trasciende las fronteras nacionales. La acumulación de competencias puramente técnicas sin una conciencia crítica reproduce la “ilusión de neutralidad” algorítmica señalada por Domínguez Figaredo y Stoyanovich (2023); a la inversa, la reflexión ética desvinculada del conocimiento operativo de modelos y métricas conduce a un activismo retórico sin capacidad de incidencia práctica. Esta dialéctica exige un enfoque holístico que las instituciones de educación superior, las redes bibliotecarias y los organismos de políticas públicas deben asumir como responsabilidad compartida a escala mundial.

Hallazgos clave

1. Trayectorias pedagógicas espirales. Desde la educación primaria (Su & Yang, 2023a; Vartiainen et al., 2024) hasta programas profesionales (Wood et al., 2021), la integración temprana y reiterada de la IA

demuestra ser el mecanismo más sólido para consolidar competencias progresivas y evitar la fragmentación curricular;

2. Imperativos socio-éticos globales. Preocupaciones sobre sesgos, transparencia y justicia algorítmica emergen en todos los contextos revisados (Hermann, 2021; Piller, 2024). El AI Act europeo (Sachoulidou, 2024) ofrece una referencia de gobernanza basada en riesgo que se proyecta como estándar de facto para otras regiones;

3. Brecha de la IA. El análisis de Wang et al. (2024) confirma la existencia de segmentos sociales simultáneamente escépticos y subalfabetizados, condición que multiplica su exposición a daños algorítmicos. Estudios de género (Zou et al., 2023) y territorialidad digital ratifican que cualquier agenda de AIL debe anclar puntos de control de equidad;

4. Bibliotecas como catalizadoras. Las bibliotecas universitarias y públicas se perfilan como nodos estratégicos de democratización tecnológica, pero requieren planes sistemáticos de actualización profesional (Cox & Mazumdar, 2022; Watkins & Johnson, 2024).

Implicaciones mundiales

- Educación superior. Las universidades deben adoptar modelos troncales de AIL, reforzados con aplicaciones disciplinares que traduzcan los conceptos abstractos en problemáticas locales, por ejemplo, sesgos en bases de datos de salud en África o automatización agrícola en Sudamérica;
- Redes bibliotecarias. Bajo la guía de la IFLA (2024), los consorcios regionales pueden crear repositorios abiertos de materiales multilingües, favoreciendo la circulación de buenas prácticas entre Norte y Sur global;
- Política pública. En ausencia de normativa homogénea, los países pueden replicar la lógica escalonada del AI Act e incorporar cláusulas

de soberanía de datos para salvaguardar contextos culturales diversos (Lo, 2023);

- Industria y mercado laboral. La creación de pasantías en entornos de inteligencia híbrida, tal como recomiendan Coman y Cardon (2024), facilita la transferencia de conocimientos y fomenta la adaptabilidad profesional en un mercado automatizado transversalmente.

Llamado a la acción global

La alfabetización algorítmica se erige como pilar ineludible de la ciudadanía del siglo XXI. Sin ella, la expansión exponencial de sistemas de IA amenaza con amplificar asimetrías sociales, económicas y epistémicas; con ella, la IA puede transformarse en vector de innovación inclusiva y bienestar colectivo (Schüller, 2022). Para materializar este potencial se requiere:

- Visión sistémica que rompa silos disciplinares y conecte competencias técnicas con marcos de ética pública;
- Financiamiento sostenible que permita a las instituciones más allá del eje global Norte acceder a infraestructuras formativas y recursos actualizados;
- Compromiso cívico que legitime la participación de comunidades locales en la definición de qué IA se diseña, para quién y con qué fines.

En definitiva, participar de manera significativa en una sociedad mediada por IA demanda la convergencia de saberes técnicos y reflexión crítica. Las conclusiones de este metaanálisis proveen un plano conceptual y operativo que puede adaptarse a diversas realidades educativas y culturales. Al movilizar las palancas complementarias de universidades, bibliotecas, industria y marcos regulatorios, la comunidad internacional dispone de la base necesaria para convertir la alfabetización algorítmica en un derecho universal y en un motor de desarrollo humano con justicia y equidad.

Referencias

- Coman, A. W., & Cardon, P. W. (2024). Perceptions of professionalism and authenticity in AI-assisted writing. *Business and Professional Communication Quarterly*. Advance online publication. <https://doi.org/10.1177/23294906241233224>
- Cox, A. M., & Mazumdar, S. (2022). Defining artificial intelligence for librarians. *Journal of Librarianship and Information Science*, 56(2), 330–340. <https://doi.org/10.1177/09610006221142029>
- DeVasto, D., & Palmer, Z. B. (2024). Building critical AI literacy in the business communication classroom. *Business and Professional Communication Quarterly*, 87(4), 557–572. <https://doi.org/10.1177/23294906241253199>
- Domínguez Figaredo, D., & Stoyanovich, J. (2023). Responsible AI literacy: A stakeholder-first approach. *Big Data & Society*, 10(2). <https://doi.org/10.1177/20539517231219958>
- Flierl, M., & Maybee, C. (2020). Refining information literacy practice: Examining the foundations of information literacy theory. *IFLA Journal*, 46(2), 124–132. <https://doi.org/10.1177/0340035219886615>
- Ghodoosi, B., Torrisi-Steele, G., West, T., & Heidari, M. (2024). Perceptions of data literacy and data-literacy education. *Journal of Librarianship and Information Science*. Advance online publication. <https://doi.org/10.1177/09610006241246789>
- Hermann, E. (2021). Artificial intelligence and mass personalization of communication content—An ethical and literacy perspective. *New Media & Society*, 24(5), 1258–1277. <https://doi.org/10.1177/14614448211022702>
- IFLA Journal. (2024). Libraries as catalysts for knowledge, technology, and social resilience. *IFLA Journal*, 50(2), 209–210. <https://doi.org/10.1177/03400352241264683>
- Jarvis, A., Ho, A., & Lim, G. (2024). Impressing artificial intelligence: Automated job-interview training in professional English subjects. *RELC Journal*. Advance online publication. <https://doi.org/10.1177/00336882241245449>
- Jarrahi, M. H., Lutz, C., & Newlands, G. (2022). Artificial intelligence, human intelligence and hybrid intelligence based on mutual augmentation. *Big Data & Society*, 9(2). <https://doi.org/10.1177/20539517221142824>
- Karan, B., & Angadi, G. R. (2023). Artificial-intelligence integration into school education: A review of Indian and foreign perspectives. *Millennial Asia*. Advance online publication. <https://doi.org/10.1177/09763996231158229>
- Lo, L. S. (2023). AI policies across the globe: Implications and recommendations for libraries. *IFLA Journal*, 49(4), 645–649. <https://doi.org/10.1177/03400352231196172>
- O'Dea, X., Ng, D. T. K., O'Dea, M., & Shkuratskyy, V. (2024). Factors affecting university students' generative AI literacy: Evidence and evaluation in the

- UK and Hong Kong contexts. *Policy Futures in Education*. Advance online publication. <https://doi.org/10.1177/14782103241287401>
- Pan, W., Liu, D., Meng, J., & Liu, H. (2024). Human–AI communication in initial encounters: How AI agency affects trust, liking, and chat-quality evaluation. *New Media & Society*. Advance online publication. <https://doi.org/10.1177/14614448241259149>
- Piller, E. (2024). Inhuman rhetoric: Generative AI and crisis communication. *Journal of Business and Technical Communication*. Advance online publication. <https://doi.org/10.1177/10506519241280594>
- Ponce, T. M. (2024). AI resume writing: How prompt confidence shapes output and AI literacies. *Business and Professional Communication Quarterly*. Advance online publication. <https://doi.org/10.1177/23294906241273317>
- Porlezza, C., & Schapals, A. K. (2024). AI ethics in journalism (studies): An evolving field between research and practice. *Emerging Media*, 2(3), 356–370. <https://doi.org/10.1177/27523543241288818>
- Sachoulidou, A. (2024). Harnessing AI for law enforcement: Solutions and boundaries from the forthcoming AI Act. *New Journal of European Criminal Law*, 15(2), 117–125. <https://doi.org/10.1177/20322844241260114>
- Schüller, K. (2022). Data and AI literacy for everyone. *Statistical Journal of the LAOS*, 38(4), 477–490. <https://doi.org/10.3233/SJI-220941>
- Su, J., & Yang, W. (2023a). Artificial intelligence and robotics for young children: Redeveloping the five big ideas framework. *ECNU Review of Education*, 7(3), 685–698. <https://doi.org/10.1177/20965311231218013>
- Su, J., & Yang, W. (2023b). Unlocking the power of ChatGPT: A framework for applying generative AI in education. *ECNU Review of Education*, 6(3), 355–366. <https://doi.org/10.1177/20965311231168423>
- Van Niekerk, J., Delpont, P. M. J., & Sutherland, I. (2024). Addressing the use of generative AI in academic writing. *Computers and Education: Artificial Intelligence*, 8, 100342. <https://doi.org/10.1016/j.caeai.2024.100342>
- Vartiainen, H., Kahila, J., Tedre, M., López-Pernas, S., & Pope, N. (2024). Enhancing children’s understanding of algorithmic biases in and with text-to-image generative AI. *New Media & Society*. Advance online publication. <https://doi.org/10.1177/14614448241252820>
- Wang, C., Boerman, S. C., Kroon, A. C., Möller, J., & de Vreese, C. H. (2024). The artificial-intelligence divide: Who is the most vulnerable? *New Media & Society*. Advance online publication. <https://doi.org/10.1177/14614448241232345>
- Watkins, T., & Johnson, Q. (2024). AI and machine learning: What to know and how to talk about it to researchers and patrons. *Information Services and Use*. Advance online publication. <https://doi.org/10.1177/18758789241298501>
- Weng, X., Ye, H., Dai, Y., & Ng, O.-L. (2024). Integrating artificial intelligence and computational thinking in educational contexts: A systematic review of

- instructional design and student learning outcomes. *Journal of Educational Computing Research*, 62(6), 1640–1670. <https://doi.org/10.1177/07356331241248686>
- Wood, E. A., Ange, B. L., & Miller, D. D. (2021). Are we ready to integrate artificial-intelligence literacy into medical-school curriculum? Students and faculty survey. *Journal of Medical Education and Curricular Development*, 8. <https://doi.org/10.1177/23821205211024078>
- Yan, C., Wang, H., & Luo, X. (2024). Knowledge mapping of data literacy: A bibliometric study using visual analysis. *Journal of Librarianship and Information Science*. Advance online publication. <https://doi.org/10.1177/09610006241285512>
- Zou, X., Su, P., Li, L., & Fu, P. (2023). AI-generated-content tools and students' critical thinking: Insights from a Chinese university. *IFLA Journal*, 50(2), 228–241. <https://doi.org/10.1177/03400352231214963>