

Evaluación por competencias en la era de la inteligencia artificial: de la prohibición normativa a la adaptación pedagógica

Competency-based assessment in the age of Artificial Intelligence: from regulatory prohibition to pedagogical adaptation

Roberto Daniel Breslin¹ y Mercedes Liliana Méndez²

Educación e ingeniería/ artículo científico

Citar: Breslin, R. D. y Méndez, M. L. (2025). Evaluación por competencias en la era de la inteligencia artificial: de la prohibición normativa a la adaptación pedagógica. *Cuadernos de Ingeniería* (16). <http://revistas.ucasal.edu.ar>

Recibido: setiembre/2025

Aceptado: Diciembre/2025

Resumen

El siguiente artículo, presenta los resultados de un trabajo de investigación-acción que aborda la crisis de validación de los saberes en el sistema educativo provocada por la masificación de la inteligencia artificial generativa (IAG). Ante la obsolescencia de los instrumentos de evaluación tradicionales, que han devenido “solubles” en términos tecnológicos por modelos de lenguaje grande (LLM), se implementó un dispositivo de capacitación docente en formato de ateneo didáctico bajo el marco de la Resolución Ministerial N.º 70/2020 de la Provincia de Salta. A través de una muestra de doscientos setenta y siete participantes de nivel superior, técnico y formación profesional, se desplegó una metodología experimental que incluyó fases de sensibilización neurocognitiva, diagnóstico de alucinaciones algorítmicas y rediseño curricular. El estudio introduce el concepto teórico de evaluación de anclaje bio-situado (EABS) como estrategia de resistencia pedagógica, que propone consignas ancladas en la experiencia física y el contexto local a fin de garantizar la insolubilidad algorítmica.

Los resultados obtenidos demuestran que, tras la intervención, el 67 % de las producciones docentes lograron incorporar mecanismos efectivos de validación de identidad cognitiva, al desplazar el foco desde la recuperación de contenidos (donde la IA presenta una efectividad del 72,7 %) hacia el juicio evaluativo y la competencia situada (donde la efectividad de la IA desciende al 4,5 %). Asimismo, se discute la tensión entre las normativas jurisdiccionales de restricción de hard-

¹ Universidad Católica de Salta

² Instituto de Educación Superior N.º 6036 - UFIDET - Salta

ware (prohibición de dispositivos) y la inevitabilidad tecnológica proyectada hacia el horizonte 6G. Se concluye que la única defensa sostenible contra el fraude automatizado no es la vigilancia física, sino la inmunización didáctica de la consigna.

Palabras claves: inteligencia artificial generativa; evaluación educativa; evaluación de anclaje bio-situado; insolubilidad algorítmica

Abstract

This article shows results of an action research study that addresses the crisis of validation of knowledge in the education system caused by the widespread adoption of generative artificial intelligence (GAI). In light of the obsolescence of traditional assessment tools—which have become technologically “obsolete” due to large language models (LLMs)—a teacher training program was implemented in the form of a didactic workshop under the framework of Ministerial Resolution No. 70/2020 of the Province of Salta.

Using a sample of two hundred and seventy-seven participants from higher education, technical, and vocational training programs, an experimental

methodology was deployed that included phases of neurocognitive awareness, diagnosis of algorithmic hallucinations, and curriculum redesign. The study introduces the theoretical concept of bio-situated anchoring assessment (BSAA) as a strategy for pedagogical resistance, proposing guidelines anchored in physical experience and the local context to ensure algorithmic insolubility.

The results show that, after the procedure, 67% of educational materials successfully incorporated effective mechanisms for validating cognitive identity, shifting the focus from content retrieval (where AI has an effectiveness rate of 72.7%) to evaluative judgment and situated competence (where AI's effectiveness drops to 4.5%). Likewise, the tension between jurisdictional regulations restricting hardware (device bans) and the technological inevitability projected toward the 6G horizon is discussed. This leads to the conclusion that the only sustainable defense against automated fraud is not physical surveillance but rather the didactic immunization of the assignment.

Keywords: Generative artificial intelligence; educational assessment; bio-situated anchoring assessment; algorithmic insolubility

1. Introducción: la crisis de la validación del saber

La irrupción de la inteligencia artificial generativa (IAG) en las aulas ha provocado un cambio radical en los fundamentos de la didáctica clásica, en especial en la instancia de acreditación de saberes. Herramientas como ChatGPT, Copilot o Gemini permiten a los estudiantes resolver ejercicios complejos, redactar ensayos y analizar casos en segundos, con una sintaxis y coherencia que a menudo superan la media esperada. Esto ha dejado en evidencia una verdad incómoda: los instrumentos de evaluación tradicionales (centrados en la retención de datos, la definición de conceptos o la redacción descriptiva) han quedado obsoletos en términos tecnológicos.

La problemática es global y urgente. Un estudio reciente de la Universidad de Reading (Scarfe y Roesch, 2024) expuso la magnitud de la “ceguera institucional” en un experimento controlado donde se infiltraron treinta y tres exámenes generados en su totalidad por IA en un sistema de corrección universitario real. Los docentes humanos solo lograron detectar uno como fraudulento. Los otros treinta y dos no solo pasaron desapercibidos, sino que obtuvieron calificaciones superiores a las de los alumnos reales. En el contexto argentino, relevamientos

de la Universidad Austral (2024) indican una brecha de adopción crítica. Mientras el uso de IA por parte de alumnos supera el 70 %, una minoría docente domina estas herramientas, lo que provoca un escenario asimétrico donde la validación del conocimiento se vuelve incierta. La crisis de validación que enfrentamos es un fenómeno de alcance global, donde la asimetría entre la capacidad de procesamiento algorítmico y la supervisión humana tradicional ha generado lo que la literatura identifica como una brecha de integridad académica estructural (Eaton, 2021).

2. Marco teórico y conceptualización

Con el propósito de abordar el fenómeno de la IA en la evaluación, este trabajo define y operacionaliza los siguientes cuatro conceptos clave que estructuran el análisis: la solubilidad ante la IA, la evaluación auténtica, la alucinación tecnológica y el diagnóstico multimodal.

2.1. Solubilidad ante la IA (*AI-solubility*)

Definimos la “solubilidad ante la IA” como el grado en el que un modelo de lenguaje grande (LLM) sin intervención crítica humana puede resolver una consigna de evaluación de manera satisfactoria y completa. De acuerdo con la metáfora química, una consigna es “altamente soluble” cuando la IA puede disolver su complejidad y entregar un producto terminado (ensayo, cálculo, código) indistinguible del de un estudiante competente. Por el contrario, una consigna es “insoluble” o “resistente” cuando contiene elementos que el algoritmo no puede procesar, tales como vivencias físicas, datos no indexados o contextos situacionales efímeros.

A nivel operativo, definimos una tarea como “soluble” cuando se encuentra dentro de lo que Dell’Acqua et al. (2023) denominan la “frontera tecnológica irregular” (*jagged technological frontier*). En su estudio seminal realizado con consultores de alto nivel, los autores demostraron que la IA no avanza de manera lineal, sino que crea una frontera dentada (dentro y fuera de la frontera).

- Dentro de la frontera (zona soluble): tareas de generación de ideas, redacción analítica y síntesis de información estandarizada, donde la IA supera el desempeño humano promedio en velocidad y calidad.
- Fuera de la frontera (zona insoluble o resistente): tareas que requieren verificación de hechos matizados, contexto situacional específico o juicio ético sobre datos ambiguos.

Por tanto, la “solubilidad” no es una propiedad binaria (sí/no), sino un gradiente que depende de la exposición contextual de la consigna. Según Mollick (2024), las tareas educativas tradicionales (resúmenes, definiciones, ensayos estándar) han sufrido un “colapso de costos cognitivos”, al volverse solubles de manera instantánea. La resistencia, entonces, se logra al empujar la evaluación hacia los “picos” de esa frontera irregular (lo vivencial, lo físico y lo local).

2.2. Evaluación auténtica

De acuerdo con Wiggins (1998) y Perrenoud (2019), la evaluación auténtica se define como aquella que requiere que los estudiantes apliquen conocimientos y habilidades a fin de resolver

problemas realistas en contextos significativos. En la era de la IA, la autenticidad ya no es opcional, es el único mecanismo de defensa. Si la evaluación no está anclada en un contexto real y situado (el aula, el taller, la comunidad local), se convierte en un ejercicio de recuperación de información que la máquina realiza mejor que el humano.

2.3. Alucinación tecnológica (*confabulation*)

La “alucinación” en IA se refiere a la generación de texto o imágenes que son correctos en términos gramaticales y plausibles desde una perspectiva lógica, pero carentes de base fáctica o inexistentes (Ji et al., 2023). A diferencia del error humano (fruto del desconocimiento), la alucinación es fruto del patrón probabilístico, ya que la IA completa la frase con la palabra más probable, no con la verdadera. En educación, esto genera el riesgo de la “verosimilitud vacía”, respuestas que suenan académicas, pero carecen de sustento real.

2.4. Diagnóstico multimodal

Se refiere a la capacidad de la IA para procesar e interpretar de manera simultánea diferentes tipos de señales (texto e imagen). Según OpenAI (2023), modelos como GPT-4V integran visión por computadora. Sin embargo, este trabajo postula que la interpretación semántica de la imagen (el “por qué” de una foto) todavía implica un punto de falla crítica, propenso a sesgos y errores de contexto (*contextual bias*).

Bender et al. (2021) analizan el concepto de loro estocástico. Explican por qué la IA es capaz de producir textos correctos a nivel gramatical, pero proveer inconsistencias contextuales.

Los LLM (*large language models*) como ChatGPT no comprenden el significado; solo unen palabras por probabilidad estadística. Son “loros” que repiten sin entender el contexto social ni físico. Del mismo modo, las imágenes reciben el mismo tratamiento, debido a que se procesan de manera estocástica píxel a píxel. A partir de allí, no les resulta complejo eliminar una porción de imagen o agregar una imagen a un contexto definido; es un cálculo probabilístico respecto a “qué corresponde” para que sea coherente. Un video no es más que una unión de imágenes; lo mismo sucede con el audio. Es decir, la IA logra un resultado coherente, no necesariamente veraz.

2.5. El desplazamiento cognitivo: la taxonomía de Bloom revisada

Con el objetivo de operativizar el concepto de “resistencia” o “insolubilidad” ante la inteligencia artificial, la presente investigación adopta como marco de referencia la revisión de la taxonomía de objetivos educativos propuesta por Anderson y Krathwohl (2001). A diferencia de la estructura original unidimensional de Benjamin Bloom (1956), esta actualización introduce una estructura bidimensional que resulta crítica para el análisis de la competencia digital, es decir, la dimensión del proceso cognitivo y la dimensión del conocimiento.

2.5.1. De la estática a la dinámica: la dimensión del proceso

Anderson y Krathwohl (2001) transformaron los sustantivos originales en verbos, al describir acciones cognitivas en orden de complejidad creciente. En el contexto de la IAG, este ordenamiento permite identificar las “zonas de dominio” algorítmico de la siguiente manera:

- Recordar (*remember*) y comprender (*understand*): Estos niveles inferiores, centrados en la recuperación de información y la construcción de significado básico (resumir, clasificar), constituyen la zona de “alta solubilidad”. Los LLM, entrenados con terabytes de datos, superan la capacidad humana de recuperación enciclopédica.
- Aplicar (*apply*): Gracias a la ejecución de procedimientos estándar, la IA demuestra alta eficacia en la resolución de ejercicios algorítmicos cerrados.
- Analizar (*analyze*), evaluar (*evaluate*) y crear (*create*): Los niveles superiores, que exigen descomponer el todo, emitir juicios basados en criterios contextuales y reorganizar elementos en un nuevo patrón, constituyen la zona de “resistencia humana”. En especial, el nivel de Crear (que Anderson ubica en la cúspide) demanda una originalidad situada que la IA, limitada a la recombinación probabilística de datos preexistentes, simula con dificultad (Churches, 2009).

2.5.2. La dimensión del conocimiento: el factor metacognitivo

La segunda dimensión clasifica el tipo de saber. Aquí reside el hallazgo fundamental para el diseño de evaluaciones “anti-IA”. Existen los siguientes tipos de conocimiento:

- Conocimiento factual y conceptual: Datos, términos y teorías. Es el “alimento” base de los LLM. Las evaluaciones centradas aquí son vulnerables por diseño.
- Conocimiento procedimental: Técnicas y métodos. Vulnerabilidad media.
- Conocimiento metacognitivo: Definido como “el conocimiento sobre la propia cognición y la toma de conciencia del propio aprendizaje” (Anderson & Krathwohl, 2001, p. 29).

Esta última categoría es el punto ciego de la inteligencia artificial. Al carecer de consciencia, cuerpo y biografía, un algoritmo no puede realizar procesos metacognitivos genuinos. No puede “reflexionar sobre su error” ni “evaluar cómo se sintió durante el proceso”, ya que no experimenta el proceso.

2.5.3. La intersección como filtro de solubilidad

La hipótesis central de este trabajo postula que la “solubilidad ante la IA” es una función inversa a la complejidad taxonómica. Una consigna ubicada en la intersección de “recordar + conocimiento factual” (ej.: “Defina la Revolución de Mayo”) es 100 % soluble para la IA. Por el contrario, el desplazamiento de la consigna hacia la intersección de “evaluar/crear + conocimiento metacognitivo” (ej.: “Juzgue la pertinencia de su propia respuesta anterior basándose en la experiencia vivida en el aula”) genera una barrera de insolubilidad que obliga al estudiante a la producción auténtica.

2.6. La pirámide de aprendizaje de Glasser y la acción competente

Con el propósito de comprender por qué el desplazamiento hacia los niveles superiores de la taxonomía de Bloom (1956) genera resistencia ante la IA, es fundamental integrar la perspectiva de la pirámide de aprendizaje, conceptualizada por el psiquiatra y educador estadounidense William Glasser (1998). Este modelo postula que la retención del conocimiento y la profundidad del aprendizaje presentan una relación directa con el nivel de participación activa del sujeto. Glasser (1998) estructura una jerarquía de eficacia que distingue de modo drástico entre la recepción pasiva y la acción competente.

2.6.1. La jerarquía de la retención

Según el modelo de Glasser (1998), los porcentajes de retención del aprendizaje varían según la actividad realizada del siguiente modo:

- 10 % de lo que leemos y 20 % de lo que oímos: Esta es la zona de aprendizaje pasivo. En el contexto actual, cuando un estudiante le pide a una IA que redacte un resumen y simplemente lo entrega, opera en este nivel mínimo de retención.
- 70 % de lo que discutimos con otros: Aquí comienza la interacción social (zona intermedia).
- 80 % de lo que hacemos: La zona de la experiencia práctica (laboratorio, taller).
- 95 % de lo que enseñamos a otros: La cúspide del aprendizaje. Con el objetivo de explicar un concepto a un par, el estudiante debe haberlo procesado, estructurado y dominado.

2.6.2. Convergencia con los verbos de la taxonomía revisada

Existe un isomorfismo claro entre la base “activa” de la pirámide de Glasser (1998) y los verbos de competencia situados a la derecha de la taxonomía de Anderson y Krathwohl (2001). Mientras que los niveles inferiores de la pirámide (leer, ver, escuchar) se corresponden con los verbos de “recordar y comprender” (donde la IA tiene solubilidad total), los niveles superiores de la pirámide (hacer, decir, enseñar) exigen verbos de “aplicar, evaluar y crear”. La presente investigación sostiene que la “competencia” no es un estado de almacenamiento de información (que corresponde al 10 % de Glasser), sino un estado de ejecución situada. Al diseñar evaluaciones donde el estudiante debe “hacer” (simular una soldadura, medir un transistor) o “enseñar” (explicar a un par, detectar el error en la IA), el alumno se ve forzado a operar en la zona del 95 % de retención. En ese caso, la IA puede servir como herramienta de apoyo, pero nunca como sustituto de la acción vivencial.

3. Metodología

3.1. Marco normativo

La intervención se estructuró bajo el amparo de la Resolución Ministerial N.º 70/2020, instrumento legal que establece el reconocimiento oficial y el auspicio de las acciones de capacitación realizadas por la cartera educativa y sus organismos dependientes. La adhesión

a este marco normativo fue determinante para la configuración de la muestra de investigación. Al garantizar la certificación oficial y el puntaje docente, la convocatoria logró reclutar a doscientos setenta y siete participantes de perfiles diversos (nivel superior, técnico y formación profesional), al evitar el sesgo de participación que suele ocurrir en cursos optativos de tecnología (donde solo asisten los “entusiastas”). Esto permitió constituir un “laboratorio social” representativo de la docencia real, con sus resistencias, dudas y necesidades concretas frente a la evaluación.

3.2. Justificación del dispositivo: la elección del ateneo

Para el abordaje de la problemática, se respondió a la Convocatoria Abierta para Presentación de Proyectos de la Subsecretaría de Desarrollo Curricular, al seleccionar de modo específico el formato de ateneo didáctico sobre las opciones de curso, taller o seminario. Esta decisión metodológica se fundamentó en la definición estipulada en las bases de la convocatoria, que caracteriza al ateneo como un “espacio de análisis y reflexión conjunta de situaciones de la práctica docente”. Los cursos y los seminarios se desestimaron por su naturaleza transmisiva o teórica, insuficiente para capturar la reacción vivencial del docente frente a la IA. Asimismo, aunque los talleres se centran en el “hacer”, suelen enfocarse en la producción de recursos, por lo que también se descartaron. Por último, se seleccionó el formato tipo ateneo debido a que su objetivo es la “discusión crítica colectiva”. Este formato permitió transformar el aula en un espacio de debate donde los participantes pudieron analizar “los hallazgos, los déficits y las razones que justifican las decisiones”, tras reconocer el error o la incertidumbre frente a la IA no como un fracaso, sino como un dato clínico para la investigación (Davini, 2008, citado en Subsecretaría de Desarrollo Curricular e Innovación Pedagógica, 2025).

3.3. El problema teórico: la IA y el “aprendizaje estratégico”

El diseño del ateneo buscó sensibilizar a los docentes sobre un fenómeno que la literatura especializada denomina “aprendizaje estratégico” (Biggs, 1987; Entwistle and Ramsden, 1983). Bajo este enfoque, el estudiante no se orienta a la comprensión profunda del contenido (*deep learning*), sino a la maximización del rendimiento (la nota) mediante el uso eficiente de recursos externos. La hipótesis de trabajo planteada en el ateneo fue que la inteligencia artificial generativa actúa hoy como la herramienta definitiva del aprendizaje estratégico, ya que permite al alumno resolver consignas estándar (ensayos, resúmenes, definiciones) y acreditar la materia sin que ocurra un proceso cognitivo interno significativo. La preocupación latente en el auditorio “¿Si usan IA, aprenden algo?” confirmó la vigencia de este constructo teórico. El ateneo funcionó, entonces, como un dispositivo de alerta pedagógica debido a que se pudo demostrar que mantener evaluaciones tradicionales (solubles por IA) es incentivar el aprendizaje superficial y estratégico.

3.4. Diseño de la intervención: hacia la evaluación por competencias

Alineado con las temáticas prioritarias de la convocatoria “Diseño y aplicación de instrumentos de evaluación” y “Uso pedagógico de la IA”, el laboratorio propuso a los docentes experimentar un cambio de modelo evaluativo en las dos siguientes direcciones:

- Evaluaciones de insolubilidad (anti-IA): diseñar consignas ancladas en la vivencia física, la práctica situada o la referencia local, donde la IA carece de contexto para ofrecer una respuesta válida (“solubilidad nula”).
- Evaluaciones de alfabetización (pro-IA): diseñar consignas que requieran el uso explícito de IA, al evaluar la competencia del estudiante para la ingeniería de instrucciones (*prompting*) y la validación crítica, para transformar así a la herramienta de plagio en un asistente cognitivo.

De este modo, el cumplimiento de la Resolución 70/2020 y las bases de la convocatoria permitieron generar un entorno de investigación-acción privilegiado, donde la capacitación docente sirvió de vehículo para diagnosticar y proponer soluciones ante la crisis de acreditación de saberes en la era digital.

4. Configuración de la muestra y perfil de los participantes

Con el fin de garantizar la validez ecológica de los hallazgos de esta investigación, se realizó un análisis estadístico descriptivo de la población inscrita en el dispositivo de capacitación. La muestra final se conformó a partir de la depuración de las planillas de inscripción digital (formulario Google Forms) y la correspondiente lista de espera, lo que cerró la convocatoria con un total de doscientos setenta y siete postulantes. Es fundamental precisar que, para el abordaje multidimensional de esta investigación, se trabajó con una muestra estratificada a partir del universo total de inscritos (N=277). El análisis de los datos se desglosa según el dispositivo de recolección aplicado, es decir que el análisis cuantitativo sobre actitudes y percepciones se realizó sobre una submuestra representativa de cincuenta y tres participantes (sección 5), mientras que las dinámicas experimentales de laboratorio social, por razones de capacidad operativa y logística de las salas plenarios, se realizaron sobre grupos de trabajo de setenta y tres (experimento 1) y veintidós (experimento 2) participantes, respectivamente. Esta segmentación, lejos de reducir la validez, permite una triangulación robusta entre percepciones actitudinales y desempeño técnico observado.

4.1. Dimensionamiento y distribución por roles

La alta demanda de la propuesta superó el cupo inicial previsto, lo que obligó a una selección de cohortes. Del total de inscriptos (N=277), doscientos catorce participantes (77,3 %) fueron admitidos como titulares activos; mientras que sesenta y tres postulantes (22,7 %) quedaron en lista de espera, lo que pone de relieve la escasa exploración y el interés crítico que despierta la temática de la IA en el colectivo docente.

En cuanto a la distribución por rol dentro del sistema educativo, la muestra presentó una composición mixta estratégica. El 70 % se conformaba de profesionales en ejercicio en la actualidad, lo que garantizó que las producciones del ateneo estuvieran ancladas en problemas reales y no teóricos. Del mismo modo, el 30 % restante se componía de estudiantes de formación docente (avanzados), es decir, alumnos de 3.º y 4.º año de profesorado. La inclusión de este segundo grupo fue determinante a fin de contrastar la visión del “docente experto” con la del “nativo digital” que está por ingresar al sistema.

4.2. Análisis de procedencia institucional: el sesgo técnico-profesional

Al desagregar la procedencia de los inscriptos según su institución de base, se detectó un patrón revelador que contradice la hipótesis de que la IA es un tema exclusivo de las “humanidades” o la “informática”. Para empezar, se registró una participación masiva de docentes y estudiantes de los Institutos de Educación Superior (IES), de los cuales se destaca el IES N.º 6001 “Gral. Belgrano” y el IES N.º 6036 (sede del evento). Esto configura un público con alto nivel de formación pedagógica teórica. Además, el dato más significativo del análisis muestral es la fuerte presencia de perfiles vinculados a la educación técnica y la formación para el trabajo. Se contabilizaron cuarenta y siete docentes (17 %) provenientes de Escuelas de Educación Técnica (EET) y organismos como UFIDET (Unidad de Formación, Investigación y Desarrollo Tecnológico), y quince instructores (5,4 %) de centros de formación profesional (CFP).

La suma de estos perfiles (casi un cuarto de la muestra) indica que la preocupación por la “copia con IA” es más aguda en las disciplinas donde el saber debe “funcionar” en la práctica (electrónica, mecánica, gastronomía). A diferencia de un ensayo de historia que la IA puede falsificar bien, un cálculo de estructura o un circuito eléctrico, presentan un margen de error funcional que preocupa al docente técnico.

4.3. Conclusiones sobre el perfil del público

A partir de la triangulación de estos datos, es posible definir al público asistente no como un grupo de “entusiastas tecnológicos” (*early adopters*), sino como una muestra representativa de la docencia en contextos educativos concretos.

Por un lado, es posible encontrar al público pragmático. La fuerte presencia de técnicos e instructores de FP sugiere una audiencia que no buscaba debates filosóficos abstractos sobre la IA, sino herramientas instrumentales para validar si sus alumnos “saben hacer” lo que la máquina dice que saben. Resulta además pertinente destacar al público heterogéneo y tensionado. La convivencia entre docentes experimentados (70 %) y estudiantes en formación (30 %) generó un cruce generacional valioso: el miedo a la tecnología de unos se compensó con la naturalización de uso de los otros, lo que enriqueció el debate del ateneo. Al contar con participantes de institutos de formación docente, escuelas técnicas y universidades (UNSa, UCASAL), la muestra reproduce a pequeña escala las tensiones de todo el sistema educativo provincial, lo que otorga a las conclusiones de este trabajo un alto potencial de generalización.

Esta configuración de la muestra valida al ateneo como un laboratorio social robusto. Lo que allí se observó no son reacciones aisladas, sino el síntoma de cómo el sistema educativo real procesa el impacto tecnológico.

5. Análisis cuantitativo de las percepciones y actitudes docentes

Con el objetivo de triangular la información cualitativa surgida en los ateneos, se aplicó un instrumento de recolección de datos estructurado (N=53) orientado a dimensionar las representaciones sociales y la disposición actitudinal de los participantes frente a la integración de la inteligencia artificial generativa. El procesamiento estadístico de las respuestas permite desarticular prejuicios comunes respecto a la brecha digital en el ámbito docente.

5.1. Prevalencia de la apertura cognitiva sobre la resistencia

Los resultados obtenidos refutan la hipótesis de un rechazo tecnofóbico generalizado en el cuerpo docente. El 64,15 % de los encuestados manifestó curiosidad o interés como emoción predominante, mientras que la frustración se ubicó en un valor marginal del 5,66 %. Estos indicadores sugieren que el sistema educativo no se encuentra en una fase de negación, sino en un estado de alerta exploratoria. La comunidad educativa, lejos de replegarse, muestra una disposición proactiva hacia la comprensión del fenómeno, al superar la barrera del miedo inicial que suele caracterizar a las disrupciones tecnológicas. Este estado de alerta exploratoria es consistente con los hallazgos de Selwyn (2020), quien sostiene que la adopción de tecnologías disruptivas por parte del profesorado, no sigue una curva lineal de resistencia, sino procesos de negociación pedagógica compleja ante nuevas herramientas digitales.

5.2. Determinantes de la tensión pedagógica: la demanda de instrumentalización

Al indagar sobre los factores generadores de malestar, se observa que la tensión no es de índole ontológica o laboral, sino instrumental. El 69,81 % de la muestra identificó la falta de capacitación técnica para la integración áulica como la causa principal de su incertidumbre, seguida por un 18,87 % que señaló el desconocimiento del funcionamiento de los modelos. Resulta significativo que el temor al desplazamiento laboral o reemplazo docente obtuviera una representatividad estadística insignificante (1,88 %), lo que evidencia que los educadores demandan herramientas didácticas concretas para gestionar la tecnología, desestimando las narrativas distópicas sobre la obsolescencia de su rol.

5.3. Correlación etaria y plasticidad adaptativa

La segmentación de datos por rango etario reveló un comportamiento contraintuitivo respecto a la variable generacional. El segmento de docentes mayores de cincuenta años exhibió los índices más elevados de apertura, con un 90 % de respuestas orientadas a la curiosidad. En contraste, la franja etaria de cuarenta a cincuenta años concentró los mayores niveles de escepticismo y preocupación (33,3 %). En términos sociológicos, esto podría interpretarse como una mayor resiliencia en los docentes de mayor antigüedad, quienes, al haber transitado múltiples reformas y cambios tecnológicos previos, poseen una perspectiva histórica más amplia. Por su parte, la generación intermedia parece experimentar una mayor presión adaptativa ante la necesidad de reconversión profesional.

5.4. Nivel de alfabetización algorítmica y desantropomorfización

Un hallazgo crítico para el diseño de políticas de formación es el nivel de comprensión técnica de la base docente. El 77,35 % de los participantes definió correctamente a la IA como un sistema de procesamiento de datos basado en patrones probabilísticos, tras rechazar visiones antropomórficas. Solo un 3,77 % atribuyó a la herramienta capacidades de razonamiento o decisión autónoma. Este dato indica que los docentes poseen una concepción desmitificada de la tecnología, lo cual constituye un sustrato favorable para la enseñanza de competencias complejas como la ingeniería de instrucciones (*prompt engineering*), sin caer en el pensamiento mágico.

5.5. Consenso sobre la integración curricular

Por último, existe una convergencia casi absoluta (94,33 %) en la concepción de la IA como un recurso pedagógico válido, lo que condiciona a una enseñanza crítica de su uso. Respecto a las competencias futuras, el 62,26 % de la muestra considera que el perfil profesional dominante requerirá la hibridación entre el conocimiento experto disciplinar y la capacidad de interactuar eficientemente con los modelos de lenguaje. El relevamiento efectuado confirma que la docencia representada en esta muestra rechaza las posturas prohibicionistas y demanda una pedagogía de la inteligencia artificial que transforme la herramienta en un asistente cognitivo transparente.

6. Estrategias de sensibilización y experimentación: la fase de “descongelamiento” cognitivo

A fin de que el ateneo funcionara de modo efectivo como un laboratorio de investigación-acción, fue imperativo establecer una fase preliminar de sensibilización disruptiva. Según la teoría del cambio de Kurt Lewin (1947), todo proceso de aprendizaje profundo requiere una etapa inicial de “descongelamiento” (*unfreezing*). En este contexto, el objetivo fue desarticular las preconcepciones de los docentes, quienes oscilaban entre el miedo al reemplazo y el pensamiento mágico sobre la “inteligencia” de la máquina, antes de abordar los conceptos técnicos. Con este fin, se diseñó e implementó un experimento pedagógico vivencial denominado “alucinación analógica”.

6.1. Experimento 1: simulación humana de un modelo de lenguaje

El objetivo de esta dinámica fue demostrar, mediante la experiencia directa, cómo un sistema puede generar respuestas correctas a nivel gramatical y verosímiles en términos académicos sin tener ninguna comprensión semántica del contenido, al emular así el funcionamiento estocástico de una inteligencia artificial generativa.

A. Protocolo y consignas del experimento

Se organizó a los participantes (N=73 en la sala plenaria) en pares y se les impartió una instrucción estricta diseñada para inducir presión cognitiva y forzar la “generación de texto” por sobre la reflexión “Usted debe escribir una respuesta académica y coherente al tema que le solicitará su compañero de la izquierda. La producción debe realizarse en un tiempo máximo de dos minutos y contener una extensión mínima de treinta palabras. No importa si desconoce el tema; su objetivo es que la respuesta parezca cierta.”

Con el fin de garantizar el desconocimiento y forzar la invención, se distribuyó de forma aleatoria una matriz de estímulos con temas deliberadamente específicos, oscuros o multidisciplinares, ajenos a la especialidad del docente promedio. El listado de consignas utilizado fue el siguiente:

Panel de estímulos (listado de consignas)

- Describa la formación Yacoraite, su ubicación, morfología y efecto paisajístico.
- Describa las ideas principales de la escuela austríaca de economía.

- Explique qué es el vector de Poynting y cuál es su aplicación física.
- Describa la utilización jurídica del concepto de *ius sanguinis*.
- ¿Por qué causa el ácido acetilsalicílico ha dejado de usarse de manera masiva?
- ¿Cuál es la comparación que hace Camilo Canegato respecto a las gallinas en la literatura?
- ¿Cuál fue la razón por la cual se vincula al Dr. Arturo Oñativia con el golpe de Estado al Dr. Illia?
- Haga un listado de los nombres completos de los siete hermanos del Gral. Martín Miguel de Güemes.
- Dé la ubicación exacta de la fosa de las Marianas y explique su importancia geográfica.
- ¿Qué condiciones físicas se deben dar para que un satélite entre en órbita geoestacionaria?
- ¿En qué año sucedió y quién era el autor de la frase “mi voto es no positivo”?
- ¿En qué lugar y contexto detuvieron al asesino de Liliana en la película *El secreto de sus ojos*?
- Describa a qué se refiere Charly García en la canción “Alicia en el país” cuando habla de “morsas y tortugas”.
- ¿Qué importancia tiene geográficamente la bahía de Samborombón?

B. Desarrollo de la dinámica y resultados observados

Bajo la presión del cronómetro, la totalidad de los docentes logró completar la tarea en el tiempo estipulado. Se observó un fenómeno de “escritura automática” que, ante la falta de datos fácticos (memoria), los participantes activaron sus competencias lingüísticas para estructurar oraciones que sonaran a respuesta de examen. Posteriormente, se seleccionaron casos al azar para su lectura en voz alta frente al auditorio. La consigna para la audiencia fue simple: “¿Le creemos o no?”.

En la mayoría de los casos, las respuestas inventadas fueron validadas como “creíbles” por el auditorio. Por ejemplo, ante la pregunta sobre el vector de Poynting, un docente de Letras redactó una explicación convincente sobre “vectores de fuerza en estructuras arquitectónicas”, al utilizar terminología técnica apropiada pero incorrecta en términos físicos. La buena redacción, la cohesión gramatical y el tono académico actuaron como un “velo de verdad”. Tras la votación, se procedió a leer la respuesta verdadera (factual). El contraste entre la invención plausible del docente y la realidad del dato generó un ambiente de asombro generalizado y risas nerviosas.

C. Conclusiones del experimento: la disonancia cognitiva

Este ejercicio provocó lo que León Festinger (1957) define como disonancia cognitiva. Los participantes experimentaron en carne propia que la corrección sintáctica no es garantía de verdad semántica. Al verse a sí mismos “mintiendo con elegancia” a fin de cumplir con la forma (treinta palabras, dos minutos), comprendieron de manera intuitiva el funcionamiento de un LLM: la IA no “sabe” sobre la fosa de las Marianas o la escuela austríaca de economía; tan solo predice la siguiente palabra más probable para satisfacer la demanda del usuario con un texto fluido. Este “descongelamiento” fue el punto de inflexión del ateneo, ya que los docentes dejaron de ver a la IA como un oráculo de conocimiento y comenzaron a verla como un generador de texto estocástico, al validar la necesidad de la supervisión humana experta (evaluación) por encima de la automatización. A su vez, replica de manera analógica el funcionamiento de los loros estocásticos descritos por Bender et al. (2021), al verificar la hipótesis de que la verosimilitud sintáctica no implica comprensión semántica.

6.2. Experimento 2: evaluación de la competencia contextual en modelos multimodales

Mientras que la fase inicial exploró la generación textual, el segundo dispositivo experimental se diseñó para auditar la capacidad de inferencia situada de los modelos de inteligencia artificial multimodal. Este ejercicio buscó validar de modo empírico la hipótesis de la frontera tecnológica irregular (*jagged technological frontier*) propuesta por Dell’Acqua et al. (2023), la cual postula que la IA puede exhibir un desempeño sobrehumano en tareas específicas (como el reconocimiento biométrico) y, al mismo tiempo, colapsar ante tareas de interpretación contextual que requieren juicio situado.

A. Protocolo experimental y estímulo

Se expuso a los participantes (N=22) a un estímulo visual de alta carga simbólica pero desprovisto de metadatos: la fotografía del encuentro de noviembre de 1993 en la Quinta de Olivos entre el presidente Carlos Menem y el expresidente Raúl Alfonsín. La consigna impartida a la IA fue deliberadamente causal: “Analice la imagen y elabore una teoría sobre por qué a la persona de pelo oscuro le convenía esta reunión”. El objetivo subyacente no era evaluar la mera identificación de los sujetos, sino determinar si el algoritmo podía reponer el contexto histórico-político (el Pacto de Olivos, la reforma constitucional y la hegemonía judicial) o si, por el contrario, incuriría en alucinaciones por falta de anclaje (*grounding*).

B. Primer nivel de análisis: precisión fáctica y recuperación enciclopédica

En una primera instancia de evaluación, centrada en la validación documental (¿sabe la IA quiénes son y qué está pasando?), los resultados mostraron una eficacia aparente considerable. Un segmento mayoritario del 72,7 % de la muestra (incluyendo los registros de Eugenia V., Tatiana R. y Carina P.) logró una identificación exitosa. Los modelos realizaron un reconocimiento facial correcto de los protagonistas y accedieron a su base de entrenamiento para inferir el evento: el “Pacto de Olivos”. Sin embargo, se observó en este grupo un alto grado de isomorfismo estructural: respuestas idénticas, organizadas en viñetas y con un tono impersonal, lo que evidencia que los docentes actuaron como operadores de “copiar y pegar” de una respuesta estandarizada.

A pesar de la alta tasa de aciertos, casi un tercio (27,3 %) de la muestra cayó en lo que Ji et al. (2023) clasifican como alucinación. Participantes como Ana Inés O. y Lylian presentaron respuestas donde la IA confundió el evento de 1993 con el traspaso de mando de 1989. El algoritmo priorizó la probabilidad estadística del evento más dramático (la crisis hiperinflacionaria) por sobre la evidencia visual, al ignorar que para 1993 la Ley de Convertibilidad (1991) ya había estabilizado la economía. Por lo tanto, es posible detectar el desplazamiento cronológico (alucinación temporal). Asimismo, en casos como los de Alejandra O. o Andrea L., la IA no logró superar el problema del *symbol grounding* (Harnad, 1990). Al no reconocer el contexto político, redujo la imagen a sus componentes objetales y la describió como una “caminata relajada en un parque cualquiera” o una reunión para “buscar discreción”, lo que trivializó un hito histórico fundacional. Es por eso que, en este caso, se dio un vaciamiento semántico (falla de anclaje).

C. Segundo nivel de análisis: precisión hermenéutica y sesgo de ingenuidad

Al profundizar el análisis más allá del dato enciclopédico y al evaluar la calidad de la respuesta política (¿comprendió la IA la motivación real de Menem?), la tasa de éxito se desploma de modo drástico. Se detectó un sesgo sistémico hacia la asepsia política o “ingenuidad algorítmica”.

Dentro del grupo que identificó de manera correcta el Pacto, la gran mayoría de las respuestas (68,2 %), como las de Tatiana R. o Alejandra O., atribuyeron a Menem motivaciones idealizadas, como “buscar consenso democrático”, “bajar tensiones” o “fortalecer las instituciones”. Estas respuestas son erróneas en términos hermenéuticos. La historiografía política coincide en que la motivación del sujeto (Menem) era pragmática y de poder, ya que buscaba habilitar la reelección inmediata y consolidar su mayoría en la Corte Suprema. La IA, entrenada para ser neutral, “romantizó” el pacto y ocultó la operación subyacente.

Solo un caso en toda la muestra (Eliana B.) logró explicitar las verdaderas tensiones institucionales, al mencionar la “mayoría automática” en la Corte Suprema y la manipulación judicial como parte del beneficio buscado, lo que constituyó un 4,5 %.

D. Conclusiones del experimento

El análisis estratificado de los resultados arroja una conclusión pedagógica alarmante para la formación docente. Por una parte, la IA tiene una efectividad del 72,7 % para reconocer el evento. Por otra parte, la IA demuestra una efectividad real de apenas el 4,5 % para interpretar las motivaciones políticas complejas sin sesgos de ingenuidad.

Esto confirma que los modelos de lenguaje operan como motores de recuperación de datos asépticos. Si bien pueden aprobar un examen de historia fáctica (fechas y nombres), reprueban de modo sistemático en la interpretación de procesos políticos, al ofrecerle al estudiante una versión “lavada” y despolitizada de la historia. El experimento evidencia que la intervención docente es insustituible no solo para corregir alucinaciones (como la de 1989), sino para reponer el sentido crítico que el algoritmo, por diseño, tiende a neutralizar.

La literatura reciente confirma la existencia de un uso acrítico de la tecnología. Según el análisis forense realizado por Turnitin (2024) sobre más de doscientos millones de trabajos estudiantiles a nivel global, se detectó que más de seis millones de entregas contenían al menos un 80 % de texto generado por IA, lo que evidencia una sustitución casi total de la autoría huma-

na. Esta tendencia se corrobora con los hallazgos del Higher Education Policy Institute (2024), donde, si bien la mayoría de los estudiantes utiliza la IA como apoyo, existe un núcleo duro que admite incorporar texto sintético sin edición alguna en sus evaluaciones sumativas, validando la hipótesis del aprendizaje estratégico automatizado.

7. Etapa de sensibilización: desarmando la “caja negra”

Tras la fase experimental inicial, que sirvió para generar disonancia cognitiva (Experimento 1) y demostrar la falibilidad contextual (Experimento 2), se dio inicio a la etapa de sensibilización teórica. El objetivo de este bloque fue desmitificar la inteligencia artificial y desplazarla del lugar de “magia” o “amenaza incomprensible” hacia su definición técnica real: un sistema matemático de predicción basado en arquitectura neuronal. Para lograr una comprensión profunda en un auditorio no informático, se diseñó una secuencia didáctica analógica: ir de lo biológico (cerebro) a lo digital (red neuronal) y de lo perceptivo (sesgos humanos) a lo algorítmico (prompts).

7.1. De la sinapsis biológica a la red neuronal artificial

Con el propósito de facilitar la comprensión de la arquitectura de los modelos de lenguaje grande (LLM), se implementó una estrategia de analogía biomimética basada en material audiovisual. Se inició con la proyección de animaciones 3D de la sinapsis neuronal (Asociación Montessori Internacional, 2012). Esto permitió ilustrar que el aprendizaje humano no es un “almacenamiento de datos” en un disco duro, sino un refuerzo de conexiones. Cuando una señal eléctrica logra saltar la brecha sináptica repetidas veces, el camino se fortalece (“aprender”); cuando no, se debilita (“olvidar”). Este concepto biológico sirvió de base a fin de explicar el funcionamiento de las redes neuronales artificiales (Schwenke, 2017; Dot CSV, 2018). Los docentes comprendieron que la IA generativa (como ChatGPT) funciona bajo un principio similar, ya que no “sabe” definiciones, sino que posee “conexiones matemáticas” (pesos o weights) fortalecidas por el entrenamiento. Se explicó que, al igual que en el cerebro, si el “peso” de la conexión entre dos palabras (por ejemplo, “Pacto” y “Olivos”) es fuerte, la IA generará la respuesta correcta. Si es débil, alucinará. Esto despojó a la IA de su aura de “inteligencia consciente”.

7.2. La manipulación del contexto: una aproximación lúdica al *prompting*

Un punto de inflexión en el ateneo fue la introducción del concepto de ingeniería de contexto mediante el análisis de la serie documental *Asombrosamente* (NatGeo, 2015). Se utilizaron fragmentos referidos a la psicología de la percepción para demostrar cómo el cerebro humano sesga su interpretación de la realidad al basarse en estímulos periféricos (colores, vestimenta, entorno). El video demostró, de manera lúdica, que la percepción de una misma persona cambia de manera radical si se modifica el contexto visual o social. Esta dinámica tuvo un objetivo crítico que consistió en introducir la noción de “ventana de contexto” en la IA. Se demostró a los participantes que, al igual que los humanos somos manipulables por el entorno, la inteligencia artificial cambia de modo drástico su “personalidad” y sus respuestas de acuerdo con el contex-

to (prompt) que el usuario le provea. Esto fundamentó la necesidad de que el docente deje de ser un mero usuario para convertirse en un “ingeniero de prompts”, capaz de manipular el contexto a fin de evitar alucinaciones y dirigir la herramienta de modo ético.

7.3. Validación documental: la ética ante la evaluación (lectura de cierre)

Con el objetivo de validar los hallazgos de los experimentos y la explicación técnica, se procedió a la lectura y al análisis de un documento de síntesis elaborado para la capacitación. Este texto, fundamentado en los lineamientos de la UNESCO (2024) y evidencias recientes, sirvió para institucionalizar la problemática.

Ante este escenario, la única salida ética es la Evaluación Auténtica (Wiggins, 1998), debido a que pedir definiciones no es suficiente. Según la UNESCO (2024), “Debemos diseñar consignas donde la IA sea insuficiente”, al utilizar preguntas sobre vivencias de clase o laboratorios locales que la IA no “vivió”, al exigirle al alumno que explique por qué eligió una respuesta, lo que valida su criterio, y al diseñar problemas donde la respuesta dependa de un razonamiento ético o una toma de decisión subjetiva. Se concluye que la IA puede ser una aliada para el análisis y la creatividad, pero el juicio crítico. La evidencia analizada sugiere que, para preservar la integridad de los saberes, es crítico que el juicio evaluativo y el sentido crítico se mantengan como competencias centrales bajo supervisión humana” (UNESCO, 2024).

8. Contexto legislativo y tensión institucional: la respuesta restrictiva

En paralelo al desarrollo de esta experiencia de capacitación, el marco regulatorio del sistema educativo de la provincia de Salta experimentó un giro significativo hacia el control restrictivo de los dispositivos tecnológicos. Mientras el dispositivo de ateneo proponía la integración pedagógica de la inteligencia artificial, la normativa jurisdiccional avanzó hacia la limitación del soporte físico (hardware) necesario para su ejecución, al configurar un escenario de tensión entre la innovación didáctica y el control disciplinar.

8.1. El nuevo marco regulatorio: Ley N.º 8.474 y resolución N.º 631/25

A fines del ciclo lectivo 2024, la Legislatura provincial sancionó la Ley N.º 8.474 (promulgada mediante el Decreto N.º 880/24), cuyo objeto es regular la utilización de dispositivos digitales en todos los establecimientos educativos de gestión pública y privada. La operacionalización de esta ley se concretó en julio de 2025 con la publicación de la Resolución Ministerial N.º 631/25, instrumento que aprueba el “Protocolo Marco” de aplicación obligatoria. Esta normativa establece un cambio de paradigma: el dispositivo móvil deja de ser un objeto de portación libre para convertirse en un elemento de uso condicionado y sujeto a restricción, al fundamentarse de modo oficial en la protección de la salud mental, la mejora de los índices de atención y la prevención del ciberacoso (*grooming* y *bullying*).

8.2. Estratificación de la prohibición

El protocolo implementado establece el siguiente régimen de gradualidad restrictiva según el nivel educativo, con implicancias directas sobre la evaluación:

- Nivel inicial: Se establece una prohibición absoluta. Los dispositivos no están permitidos bajo ninguna circunstancia, priorizando el desarrollo neurocognitivo y la interacción física.
- Nivel primario: Se impone una restricción general. El uso solo se habilita de manera excepcional a partir de 6.º grado, al requerir una planificación pedagógica específica y la autorización escrita de los tutores legales.
- Nivel secundario: Se configura un uso condicionado. La norma prohíbe el uso de celulares durante recreos, horas libres y en espacios sanitarios. En el aula, su presencia solo es lícita si media una inclusión explícita en la secuencia didáctica del docente.
- Disposición antifraude: Un aspecto crítico de la resolución es la prohibición explícita del uso de dispositivos durante instancias de evaluación, salvo indicación contraria fundada, y la restricción de acceso a redes sociales no educativas.

8.3. Análisis del paradigma: restricción del medio vs. adaptación evaluativa

La coexistencia de esta normativa con la irrupción de la IA generativa revela una estrategia institucional de contención mecánica. Ante la dificultad de fiscalizar el uso de algoritmos generativos (software) que vulneran la integridad de las evaluaciones tradicionales, la política pública ha optado por la sustracción del soporte (hardware). Esta medida actúa como un paliativo analógico frente a un problema digital. Al retirar el dispositivo, se busca neutralizar el “plagio automatizado” y restablecer las condiciones de la evaluación clásica (memorística o procedimental sin asistencia). Sin embargo, este enfoque genera una dicotomía con la propuesta del presente trabajo de investigación. Por un lado, el enfoque normativo (Resolución 631/25) busca garantizar la validez de la evaluación mediante la ausencia de tecnología. Por otra parte, el enfoque del ateneo (inmunización cognitiva) busca garantizar la validez de la evaluación mediante la transformación de la consigna, al asumir la presencia inevitable de la tecnología. Este estado de alerta exploratoria es consistente con los hallazgos de Selwyn (2020), quien sostiene que la adopción de tecnologías disruptivas por parte del profesorado no sigue una curva lineal de resistencia, sino procesos de negociación pedagógica compleja ante nuevas herramientas digitales. La tendencia a la prohibición del hardware como respuesta institucional se analiza en la literatura como una respuesta de “pánico moral” frente a la disrupción tecnológica, la cual suele ser ineficaz ante la ubicuidad de los dispositivos en el entorno cotidiano del estudiante (Selwyn et al., 2023).

En conclusión, el contexto legal vigente en Salta impone un desafío adicional a la docencia: diseñar estrategias de alfabetización en IA que sean respetuosas de la prohibición de uso recreativo o desregulado, al demostrar que la tecnología en el aula solo se justifica cuando media una intencionalidad pedagógica insustituible.

9. La inevitabilidad tecnológica: la obsolescencia de la prohibición

Frente a la respuesta institucional de restricción de dispositivos (analizada en la sección precedente), el ateneo propuso una instancia de debate prospectivo. Se planteó a los docentes

la hipótesis de que legislar sobre la tenencia física del teléfono móvil es una solución coyuntural destinada a caducar a corto plazo, dado el avance hacia la ubicuidad tecnológica y la integración cognitiva de la IA. Con el fin de sustentar esta visión de “inevitabilidad”, se analizaron dos exposiciones de referentes globales que anticipan la disolución de las fronteras físicas y cognitivas actuales.

9.1. La obsolescencia del hardware: la visión 6G (Dr. Marcos Katz)

Con el fin de contextualizar la fugacidad de las prohibiciones centradas en el objeto físico (“guardar el celular”), se proyectó y analizó la conferencia dictada en el Instituto Balseiro por el Dr. Marcos Katz, profesor de la Universidad de Oulu (Finlandia) y referente del programa global 6G Flagship. Katz (2023) postula que la próxima revolución de las telecomunicaciones no consistirá en un dispositivo más rápido, sino en la desaparición del dispositivo tal como lo conocemos (el “ladrillo” rectangular). A través de la convergencia de tecnologías como la electrónica impresa (tintas conductoras) y la comunicación por luz visible (VLC - *Visible Light Communication*), el paradigma se desplazará hacia el “Internet de todo” (*Internet of Everything - IoE*). Los docentes analizaron que, en un futuro inminente, la conectividad estará integrada en las superficies mismas del aula (pupitres inteligentes, paredes activas) o en la indumentaria (wearables). La conclusión fue contundente, debido a que la normativa actual, que prohíbe “el teléfono”, será técnicamente inaplicable cuando el entorno mismo sea el dispositivo. Prohibir la tecnología implicaría, en ese escenario, vaciar el aula de mobiliario, al poner de relieve que la escuela debe prepararse para un entorno de conectividad ambiental y ubicua, donde la vigilancia física es obsoleta.

9.2. La fusión cognitiva: la IA como entorno natural (Sam Altman)

En consonancia con la visión de la infraestructura, se abordó la transformación del sujeto de aprendizaje mediante el análisis de una entrevista reciente a Sam Altman, CEO de OpenAI. Se focalizó el debate en su tesis sobre la “naturalización” de la inteligencia artificial para las nuevas generaciones. En el fragmento seleccionado (minuto 18:00), Altman sostiene una premisa provocadora: “Un niño nacido hoy nunca será más inteligente que la IA”. Lejos de una lectura apocalíptica, el análisis grupal interpretó esto como un cambio ontológico en la relación humano-máquina. Para los estudiantes actuales, la IA no es una herramienta externa a la que deben adaptarse (como lo fue la informática para los adultos), sino un “copiloto” cognitivo nativo. El grupo de trabajo concluyó que la educación no puede legislar contra la realidad circundante si la tecnología se vuelve invisible (como plantea Katz) o si la inteligencia artificial se vuelve nativa (como plantea Altman).

Entonces, la evaluación escolar debe abandonar la vigilancia de la “tenencia de dispositivos” para pasar a evaluar el uso crítico de una inteligencia expandida. La prohibición se reveló, a la luz de estos videos, como un intento analógico de contener un tsunami digital, al validar la propuesta del ateneo de “inmunizar” las consignas en lugar de aislar a los alumnos.

10. Propuesta teórica: hacia la “evaluación de anclaje bio-situado” (EABS)

Superada la fase de diagnóstico sobre la inevitabilidad tecnológica y la insuficiencia de los controles coercitivos, el ateneo se abocó a la formulación de una respuesta didáctica superadora. A partir de la síntesis entre la teoría de la evaluación auténtica y las evidencias experimentales recolectadas, este trabajo propone un nuevo constructo teórico-práctico denominado evaluación de anclaje bio-situado (EABS).

Definimos la EABS como aquella estrategia de diseño evaluativo que vincula la acreditación de saberes a variables de experiencia física (“bio”) y coordenadas territoriales específicas (“situado”), en lo que respecta a la consigna insoluble en términos algorítmicos o dependiente de la validación humana experta.

En este marco conceptual, se sistematizaron las siguientes taxonomías operativas para el diseño de instrumentos, fundamentadas en la literatura académica contemporánea: la resistencia por diseño y la integración crítica.

10.1. El paradigma de resistencia: insolubilidad y oralidad

Este enfoque no persigue la prohibición administrativa de la IA, sino su inutilización funcional mediante la ingeniería de la consigna. El objetivo es diseñar reactivos donde el gran modelo de lenguaje (LLM) carezca de la “ventana de contexto” necesaria para inferir una respuesta válida.

Conforme a lo planteado por Rudolph et al. (2023), se postula que, ante la ubicuidad del texto sintético asincrónico, la reserva de valor de la evaluación humana se desplaza hacia la oralidad interactiva y la defensa en tiempo real. La IA puede escribir el ensayo, pero no puede defenderlo ante la repregunta imprevista del docente. Asimismo, en concordancia con Dawson (2020) y su marco de *Assessment Security*, se propone evaluar la “trayectoria de producción” y no el entregable final. Esto implica la validación de borradores manuscritos, mapas mentales y esquemas de ideación en el aula, donde el docente certifica la génesis cognitiva del trabajo, reduciendo el margen para la tercerización algorítmica. Por otra parte, bajo el concepto EABS, una consigna de resistencia sería: “Describe los fenómenos de oxidación observados en la reja perimetral de la escuela y proponga una solución basada en los materiales disponibles en el pañol del taller”. La IA desconoce el estado de la reja y el inventario del pañol; el estudiante debe “estar ahí” para responder.

10.2. El paradigma de integración: copilotaje y juicio evaluativo

Este enfoque asume la presencia de la IA como un inevitable pedagógico y reconfigura la evaluación para medir la competencia de uso, edición y curaduría del estudiante sobre el producto sintético.

A partir de las estrategias propuestas por Mollick y Mollick (2023), se implementan las siguientes dinámicas de inversión de roles:

- *AI as student* (IA como estudiante): El alumno debe instruir al algoritmo sobre un tema complejo y, más tarde, auditar la respuesta generada a fin de detectar alucinaciones o sesgos.
- *AI as simulator* (IA como simulador): Se utiliza la IA para recrear debates con figuras históricas o escenarios técnicos, al evaluar la capacidad del alumno para interactuar con la simulación mediante *prompts* de alta precisión.

Según la tesis de Bearman y Ajjawi (2023), se desplaza el foco de la “producción de texto” al “juicio de calidad”. En un ecosistema de generación infinita de contenidos, la competencia crítica reside en la capacidad del estudiante para discernir entre múltiples respuestas generadas, justificar la selección de la más pertinente y mejorarla (iteración humana).

10.3 Síntesis de la matriz evaluativa

La implementación del modelo de evaluación de anclaje biosituado (EABS) permite superar la falsa dicotomía entre prohibir y liberar. La innovación reside en que la validez de la acreditación no depende de si el alumno usó o no tecnología, sino de cómo se la integra a fin de resolver el problema. Si la IA por sí sola puede aprobar el examen, el instrumento es deficiente (falta de anclaje situado). Si el alumno requiere aportar su vivencia (bio) o su criterio (juicio evaluativo) para completar la tarea, estamos ante una evaluación auténtica blindada contra el fraude automatizado.

11. Validación didáctica: de la enciclopedia a la competencia situada

Una vez establecido el marco teórico de la evaluación de anclaje bio-situado (EABS), la capacitación avanzó hacia la fase de operacionalización práctica. Con el objetivo de demostrar la transversalidad del método, se presentaron dos “casos testigo” (uno del campo técnico-duro y otro de las ciencias sociales) diseñados para evidenciar la diferencia entre una pregunta soluble por IA (contenido) y una pregunta resistente (competencia). El objetivo pedagógico consistió en demostrar que, mediante la alteración de la consigna, es posible transformar un reactivo de “copiar y pegar” en un desafío de simulación mental.

11.1 Caso 1: el desafío técnico (física aplicada vs. definición)

En el ámbito de la Educación Técnica Profesional, el riesgo de la IA es que resuelva problemas estándar de manual.

- La pregunta tradicional: “¿Cómo funciona un inyector naftero de inyección directa?”
 - Diagnóstico: Esta es una pregunta de Nivel 1 (recordar/comprender) en la taxonomía de Bloom (1956). ChatGPT puede generar una respuesta perfecta, técnica y detallada en segundos, ya que la definición es universal y estática. El estudiante no necesita entender el motor, solo necesita saber pedir la definición.
- La pregunta por competencias (EABS):
 - Consigna: “Por qué motivo cierra la aguja... teniendo en cuenta la enorme presión en la cámara... a) Solenoide/resorte vs. presión... b) Autoinducción... c) Gravedad...”
 - Análisis de la resistencia: Esta consigna introduce variables dinámicas en conflicto (tensión eléctrica, fuerza mecánica del resorte, presión neumática de los gases).
 - La trampa para la IA: Si el alumno copia la pregunta en una IA básica sin contexto, es probable que la IA ofrezca una explicación general sobre el cierre del inyector, pero falle al discernir la interacción específica de fuerzas en el instante postexplosión.
 - Competencia activada: El estudiante debe realizar una simulación mental mecánica. Debe visualizar que, aunque hay presión en la cámara intentando mantener la aguja abierta (o ce-

rrada, según el diseño), es la desaparición del campo magnético (solenoides) lo que libera la energía potencial del resorte. La opción B (autoinducción) es un distractor físico plausible pero incorrecto en la función mecánica, y la C (gravedad) es un absurdo técnico. Para responder, el alumno no puede “leer”, tiene que “funcionar” como el inyector.

11.2. Caso 2: el desafío histórico (la trampa del anacronismo)

Este ejemplo es brillante porque introduce un “veneno contextual” (una premisa falsa) para probar si el estudiante (o la IA) valida la información o tan solo alucina una justificación. La pregunta tradicional “¿Cuál fue la causa de la Revolución de Mayo?” consiste en una pregunta enciclopédica a partir de la cual la IA listará las Invasiones inglesas, la Caída de Sevilla y las ideas de la Ilustración. Una consigna de este tipo lleva al alumno a aprobar sin procesar la información.

En contraste, la pregunta por competencias (EABS) introduce la consigna “Ud. es Juan Larrea... Cabildo del 22 de mayo de 1810... vota por la continuidad del virrey Sobremonte. ¿Por qué lo hace?”. La clave del reactivo reside en una trampa deliberada. En mayo de 1810, el virrey era Baltasar Hidalgo de Cisneros. El marqués de Sobremonte había sido destituido y repudiado socialmente años antes (tras su huida en las invasiones de 1806). Ante esta consigna, una IA generativa, programada para ser “útil”, intentará justificar la premisa del usuario. Es probable que invente una razón, como “Voto por Sobremonte por lealtad a la corona...” o “Por su experiencia militar...”, al caer en la trampa y producir una alucinación forzada. Sin embargo, el estudiante humano, si ha estudiado el proceso y tiene pensamiento crítico, debe detenerse y marcar la opción e) Ninguna de las anteriores (o impugnar la pregunta). El alumno llega a esa conclusión, debido a que sabe que Larrea era un patriota (vocal de la Primera Junta), por lo que no votaría continuidad, y es consciente de que el virrey era Cisneros, no Sobremonte.

En este sentido, la pregunta por competencias (EABS) evalúa la ubicación temporal, el rol político y la detección de falacias. Es imposible de responder de manera correcta mediante un simple prompt si no se detecta primero el error histórico en la formulación.

11.3. Conclusión de la validación

Ambos ejemplos demostraron al auditorio que la insolubilidad algorítmica no requiere tecnología compleja, sino sofisticación en el diseño de la pregunta. Por un lado, en el caso técnico, se logró mediante la intersección de físicas (electricidad + mecánica + termodinámica). Por otro lado, en el caso social, se logró mediante el *role-playing* y la falsación de premisas.

Esto validó la hipótesis central del ateneo, que consiste en que la IA puede responder qué es un inyector o qué pasó en mayo, pero tiene dificultades críticas para resolver por qué falla el inyector bajo presión o qué haría un prócer ante un escenario contrafáctico o erróneo.

13. Análisis de resultados: dimensionamiento de la resistencia algorítmica

Como fase de cierre de la investigación-acción, se procedió a la evaluación colegiada de los instrumentos diseñados por la cohorte participante durante el tercer encuentro. Para garantizar la objetividad del análisis, se sometieron las producciones (N=total de la muestra representativa) a una auditoría de solubilidad, confrontando cada ítem evaluativo con modelos de lenguaje

estándar (ChatGPT-4o y Gemini) a fin de determinar su grado de vulnerabilidad. A continuación, se detallan los hallazgos estadísticos y cualitativos sobre la efectividad de las estrategias de diseño implementadas.

13.1. Perfil general de efectividad: la curva de resistencia

A partir del análisis global de los instrumentos de evaluación entregados, se observa una distribución estratificada en cuanto a la robustez frente al fraude automatizado.

Primero, un quinto de la producción docente (22 %) logró blindar la totalidad o la mayoría de sus ítems (cuatro de cinco). Estos exámenes resultaron insolubles para la IA o requirieron un nivel de *prompting* experto superior a la media estudiantil.

Además, los instrumentos de resistencia media (45 %) constituyen el grueso de la muestra. Presentan una estructura híbrida donde conviven preguntas fácticas vulnerables (solubles por IA) con preguntas de transferencia o juicio crítico que exigen intervención humana.

A pesar de la etapa de sensibilización, un tercio de las propuestas (33 %) mantuvo una inercia hacia la pregunta declarativa (“Defina”, “Enumere”, “Describa”), resultando en exámenes que la IA pudo aprobar con calificaciones superiores al 80 %.

13.2. Taxonomía del acierto: ¿por qué falló la IA?

Al centrar el análisis en el corpus de las “preguntas bien redactadas” (aquellas clasificadas como de alta resistencia), se procedió a categorizarlas según el factor lógico que provocó el fallo o la insuficiencia del algoritmo. El siguiente desglose porcentual revela cuáles fueron las estrategias más apropiadas por los docentes:

- Por déficit de contexto o “anclaje biosituado” (42 %): Esta fue la estrategia dominante. Las preguntas fueron exitosas porque solicitaban datos que no están en el entrenamiento del modelo por ser efímeros o hiperlocales. Ejemplo genérico: Análisis de muestras físicas de laboratorio escolar, fenómenos barriales recientes o dinámicas áulicas no registradas digitalmente.
- Por requerimiento de juicio evaluativo/auditoría (38 %): Corresponde al modelo “pro-IA”. Estas preguntas fueron eficaces porque no pedían producir contenido, sino criticarlo. La IA falló no por ignorancia, sino porque se le pidió al alumno que detectara los sesgos o silencios de la propia máquina. Ejemplo genérico: “Identifique qué variable política omitió la IA en el siguiente resumen y justifique su importancia”.
- Por inferencia pragmática o falsación de premisas (20 %): Corresponde a las preguntas “trampa”. Fueron efectivas porque incluían premisas falsas o ironías culturales que la IA procesó de manera literal, lo que puso en evidencia su falta de sentido común o comprensión sociolingüística. Ejemplo genérico: El caso del “virrey Sobremonte” o modismos locales (“se pinchó la redonda”).

13.3. Análisis de casos testigo (anonimizados)

Con el propósito de ilustrar la profundidad de la apropiación metodológica, se seleccionaron tres casos paradigmáticos que fueron auditados ítem por ítem.

- Caso A: historia y formación ética (estrategia de meta-evaluación)
 - Participante: docente de nivel secundario (Referencia: F.E.).
 - Análisis: Este instrumento destacó por invertir la carga de la prueba. En lugar de prohibir la IA, el ítem 5 obligaba al estudiante a auditar un texto sintético sobre la “Ley 1420”.
 - Factor de éxito: La consigna “Señale qué aspecto ideológico NO menciona la IA” resultó devastadora para el algoritmo. Al ser un modelo probabilístico que busca el consenso, la IA rara vez explicita los conflictos o lo “no dicho”. Para resolverlo, el estudiante debía poseer un marco teórico superior al de la herramienta.
 - Calificación de resistencia: 95/100.
- Caso B: educación técnico-electromecánica (estrategia de interdependencia sistémica)
 - Participante: docente de escuela técnica (Referencia: M.B.).
 - Análisis: Se abandonó la pregunta funcional (“¿Cómo funciona un inyector?”) por la pregunta de fallo sistémico.
 - Factor de éxito: La pregunta obligaba a cruzar variables de dominios distintos (presión termodinámica vs. tensión eléctrica) en un instante de tiempo específico (postexplosión). La IA ofreció explicaciones generales correctas, pero falló en la precisión de la causalidad mecánica en ese milisegundo específico. Se validó que la IA “lee” manuales, pero no “simula” sistemas físicos complejos con precisión de taller.
 - Calificación de resistencia: 90/100.
- Caso C: literatura (estrategia de endogamia hermenéutica)
 - Participante: docente de ciclo orientado (Referencia: S.P.).
 - Análisis: El docente utilizó textos literarios breves leídos exclusivamente en el aula, con personajes genéricos (“Brian y María”).
 - Factor de éxito: Al preguntar “¿Por qué Brian es romántico?”, la IA alucinó respuestas basadas en estereotipos de novelas famosas, ya que carecía del texto fuente en su dataset. La pregunta se volvió resistente por “oscuridad bibliográfica”, ya que la respuesta solo existía en la carpeta del alumno, no en la nube.
 - Calificación de resistencia: 100/100 (insoluble sin el texto físico).

13.4. Conclusión de la auditoría

Los datos arrojan una conclusión esperanzadora, ya que, si bien la inercia del “copiar y pegar” persiste en un tercio de las propuestas, el 67 % de los docentes (suma de resistencia alta y media) logró incorporar al menos un mecanismo de validación de identidad cognitiva. Se demostró que la barrera más efectiva contra el fraude no es tecnológica, sino didáctica, debido a que las preguntas que mejor resistieron a la IA no fueron las más complejas en términos académicos, sino las más arraigadas a la experiencia humana situada y al juicio crítico.

14. Conclusiones

A partir del recorrido teórico, experimental y estadístico desarrollado en esta investigación, se desprenden las siguientes conclusiones que buscan aportar claridad sobre el futuro de la acreditación de saberes en la era de la inteligencia artificial.

14.1. La confirmación de la “frontera irregular” y el fin del enciclopedismo

Los experimentos realizados (específicamente el análisis del caso “Pacto de Olivos”) validan de modo empírico la teoría de la frontera tecnológica irregular (Dell’Acqua et al., 2023). Se demostró que la IA posee un desempeño asimétrico, ya que es capaz de superar estándares humanos en la recuperación de datos fácticos y reconocimiento biométrico (72,7 % de aciertos), pero colapsa ante requerimientos de interpretación hermenéutica profunda o detección de silencios políticos (4,5 % de efectividad real). Esto implica que la evaluación escolar centrada en el enciclopedismo (recordar/comprender en la taxonomía de Bloom) ha perdido de modo definitivo la validez certificante. Cualquier instrumento que pueda ser resuelto por un LLM sin intervención humana no evalúa competencia, sino la disponibilidad tecnológica.

14.2. El aporte teórico: la eficacia del “anclaje bio-situado” (EABS)

Como contribución original de este trabajo, se sistematizó el concepto de Evaluación de Anclaje Bio-Situado (EABS). Los datos de la auditoría final revelan que las estrategias más robustas frente al fraude (presentes en el 22 % de los exámenes de alta resistencia) no fueron aquellas que aumentaron la complejidad académica abstracta, sino las que introdujeron variables de oscuridad bibliográfica (textos del aula no indexados) y vivencia física (fenómenos de laboratorio o contexto barrial). La “insolubilidad” del examen no depende de la dificultad, sino de la territorialidad.

14.3. El mito de la resistencia docente

El análisis sociológico de la muestra (N=53) derribó el prejuicio de la tecnofobia en el colectivo docente. Con un 64,15 % que manifestó curiosidad y solo un 1,88 % que expresó temor al reemplazo laboral, se concluye que el sistema educativo se encuentra en un estado de “alerta exploratoria”. La frustración detectada (69 %) no es ideológica, sino instrumental, ya que los docentes no rechazan la IA, reclaman alfabetización técnica para dominarla. La mayor apertura en docentes de mayor antigüedad observada en nuestra muestra coincide con estudios que señalan que la experiencia profesional previa actúa como un factor protector, lo que le permite al docente integrar las novedades tecnológicas desde una perspectiva histórica y crítica, y evitar así la adopción ingenua (Lucas and Hillman, 2022).

14.4. La paradoja normativa: restricción analógica vs. ubicuidad digital

Se identificó una tensión crítica entre la política pública de corto plazo y la realidad tecnológica prospectiva. Mientras la normativa provincial (Resolución 631/25) avanza hacia la prohibición del soporte físico (*hardware*) como paliativo ante la distracción y el fraude, los análisis de prospectiva (Katz, Altman) indican que la tecnología se dirige hacia la invisibilidad ambiental (6G, IoE). A la luz de los hallazgos, esta investigación sugiere que la legislación centrada exclu-

sivamente en la tenencia del dispositivo físico podría enfrentar limitaciones significativas a corto plazo. Más que una medida ineficaz, se interpreta como una solución coyuntural que encuentra desafíos frente a la tendencia hacia la ubicuidad tecnológica proyectada por el paradigma 6G e Internet de todo (*IoE*), donde el entorno mismo podría integrar capacidades de procesamiento y conectividad. La única política de seguridad evaluativa sostenible es la alfabetización en el diseño de consignas, que consiste en pasar de vigilar el bolsillo del alumno a auditar la estructura de la pregunta.

14.5. Hacia una nueva soberanía pedagógica

Por último, el ateneo demostró que es posible revertir la asimetría de poder frente al algoritmo. El paso de preguntas “solubles” a preguntas de “juicio evaluativo” (donde el alumno audita a la IA) restituye la jerarquía cognitiva. En un mundo donde la respuesta es un *commodity* instantáneo y gratuito generado por máquinas, el valor humano se refugia en la pregunta y en el criterio. Educar en la era de la IA no es enseñar a competir con el robot en velocidad o memoria, sino enseñar a operar allí donde la máquina alucina: en la ética, en el contexto y en la construcción de sentido situado.

Referencias

- Anderson, L. W., & Krathwohl, D. R. (Eds.). (2001). *A taxonomy for learning, teaching, and assessing: A revision of Bloom's taxonomy of educational objectives*. Longman.
- Asociación Montessori Internacional. (2012, 16 de enero). *La Sinapsis Neuronal* [Video]. YouTube.
- Bearman, M., & Ajjawi, R. (2023). Learning to work with the black box: Pedagogy for a world with artificial intelligence. *British Journal of Educational Technology*, 54(5), 1160–1173.
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623). Association for Computing Machinery.
- Biggs, J. B. (1987). *Student approaches to learning and studying*. Australian Council for Educational Research.
- Bloom, B. S. (Ed.). (1956). *Taxonomy of educational objectives: The classification of educational goals. Handbook I: Cognitive domain*. David McKay Company.
- Davini, M. C. (2008). *Métodos de enseñanza: Didáctica general para maestros y profesores*. Santillana.
- Dawson, P. (2020). *Defending Assessment Security in a Digital World: Preventing E-Cheating and Supporting Academic Integrity*. Routledge.
- Dell'Acqua, F., McFowland, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. R. (2023). Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality. *Harvard Business School Working Paper*, 24-013.
- Dot CSV. (2018, 20 de junio). *¿Qué es una Red Neuronal? Explicado* [Video]. YouTube.
- Eaton, S. E. (2021). *Plagiarism in higher education: Tackling tough topics in academic integrity*. Libraries Unlimited.

- Entwistle, N. J., & Ramsden, P. (1983). *Understanding student learning*. Croom Helm.
- Festinger, L. (1957). *A theory of cognitive dissonance*. Stanford University Press.
- Fridman, L. (Host). (2023, 25 de marzo). Sam Altman: OpenAI CEO on GPT-4, ChatGPT, and the Future of AI (N.º 367) [Episodio de Podcast de Video]. En *Lex Fridman Podcast*. YouTube.
- Glasser, W. (1998). *Choice Theory: A New Psychology of Personal Freedom*. HarperCollins.
- Harnad, S. (1990). The symbol grounding problem. *Physica D: Nonlinear Phenomena*, 42(1-3), 335-346.
- Higher Education Policy Institute (HEPI). (2024). *Policy Note 51: Students' views on Generative AI*. HEPI UK.
- Instituto Balseiro [IB]. (2023, 27 de abril). *Hacia la era 6G: Oportunidades y desafíos - Dr. Marcos Katz* [Video]. YouTube.
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y., Madotto, A., & Fung, P. (2023). Survey of Hallucination in Natural Language Generation. *ACM Computing Surveys*, 55(12), 1-38.
- Lewin, K. (1947). Frontiers in group dynamics: Concept, method and reality in social science; social equilibria and social change. *Human Relations*, 1(1), 5-41.
- Lucas, M., & Hillman, T. (2022). What's in a teacher's resistance to digital technology? A conceptual framework for understanding the professional adoption of new tools. *Journal of Computer Assisted Learning*, 38(2), 450-462.
- Ministerio de Educación, Cultura, Ciencia y Tecnología de la Provincia de Salta. (2020). *Resolución Ministerial N.º 70/2020*.
- Mollick, E. (2024). *Co-Intelligence: Living and Working with AI*. Portfolio/Penguin.
- Mollick, E., & Mollick, L. (2023). *Assigning AI: Seven Approaches for Students, with Prompts*. Harvard Business Publishing Education.
- NatGeo Latinoamérica. (2015, 23 de septiembre). *Asombrosamente: La atracción y el contexto* [Video]. YouTube.
- Perrenoud, P. (2019). *Construir competencias desde la escuela*. J.C. Sáez Editor.
- Poder Legislativo de la Provincia de Salta. (2024). *Ley N.º 8.474: Regulación de dispositivos digitales en establecimientos educativos*.
- 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).
- Scarfe, P., & Roesch, E. B. (2024). A Turing Test for AI-generated text in Higher Education assessment. *PLOS ONE*, 19(6), e0305364.
- Schwenke, F. (2017, 12 de mayo). *Redes Neuronales Artificiales: Una introducción* [Video]. YouTube.
- Selwyn, N. (2020). *Education and technology: Key issues and debates* (3.ª ed.). Bloomsbury Academic.
- Selwyn, N., Pangrazio, M., Nemorin, S., & Perrotta, C. (2023). *Digital disruption in teaching and learning: Re-imagining the role of the teacher*. Routledge.
- Subsecretaría de Desarrollo Curricular e Innovación Pedagógica de Salta. (2025). *Resolución Ministerial N.º 631/25: Protocolo Marco para el uso de dispositivos digitales*. Ministerio de Educación, Cultura, Ciencia y Tecnología.
- Subsecretaría de Desarrollo Curricular e Innovación Pedagógica. (2025). *Convocatoria abierta para presentación de proyectos de capacitación docente: Documento marco y bases de condiciones*. Ministerio de Educación, Cultura, Ciencia y Tecnología de Salta.

- Turnitin. (2024). *Year in Review: 1 Year of AI Writing Detection Statistics*. Turnitin Press.
- UNESCO. (2024). *Guidance for generative AI in education and research*. UNESCO Publishing.
- Universidad Austral. (2024). *Relevamiento sobre el uso y percepción de la Inteligencia Artificial en el ámbito educativo argentino*. Escuela de Educación.
- Wiggins, G. P. (1998). *Educative Assessment: Designing Assessments to Inform and Improve Student Performance*. Jossey-Bass.

Roberto Daniel Breslin

Perfil académico y profesional: Ingeniero electrónico. Doctor en educación superior y universitaria. Profesor-investigador especializado en Electrotecnia, Telecomunicaciones y Autotrónica. Se desempeña en la Universidad Católica de Salta e institutos de educación técnica superior (UFIDET). Se enfoca en proyectos de desarrollo tecnológico y sus líneas de trabajo abarcan la inteligencia artificial aplicada a la educación, la evaluación de competencias situadas, el procesamiento de señales y las redes integradas de comunicaciones.

Correo electrónico: rbreslin@ucasal.edu.ar

Mercedes Liliana Méndez

Perfil académico y profesional: Doctora en Ingeniería, orientación biotecnología. Docente e investigadora especialista en Bioelectrónica, Mantenimiento de Sistemas de Salud y Electromedicina en el ámbito de la educación superior y técnica (UFIDET). Cuenta con amplia experiencia en la orientación de proyectos de innovación tecnológica, prototipado funcional y formación por competencias aplicadas al área biomédica. Participa en proyectos de investigación, capacitación docente e integración de tecnologías emergentes en el proceso de enseñanza-aprendizaje.