

Modelo de Redes Neuronales Inteligentes Alfa-Pi-Mi: optimización del aprendizaje en ingeniería aplicando inteligencia artificial

Alfa-Pi-Mi Intelligent Neural Network Model: optimizing learning in engineering through artificial intelligence

Diógenes Álvarez-Solórzano¹

Álvarez-Solórzano, D. Modelo de redes neuronales inteligentes Alfa-Pi-Mi: optimización del aprendizaje en ingeniería aplicando inteligencia artificial. *Tecnología en Marcha*. Vol. 39 N° especial sobre Inteligencia Artificial. Febrero, 2026. Pág. 338-351.

 <https://doi.org/10.18845/tm.v39i5.8506>

Palabras clave

Redes neuronales; aprendizaje personalizado; inteligencia artificial en educación; Alfa-Pi-Mi; optimización educativa; Industria 5.0; innovación centrada en el ser humano.

Resumen

Este artículo presenta el diseño e implementación del modelo de redes neuronales Alfa-Pi-Mi, desarrollado para personalizar y optimizar el aprendizaje de estudiantes de Ingeniería Industrial mediante la aplicación de inteligencia artificial (IA). El modelo identifica patrones de aprendizaje y genera estrategias educativas adaptadas a necesidades individuales. Los hallazgos deben interpretarse como evidencia inicial de factibilidad (estudio piloto), no como resultados generalizables. Se detalla la arquitectura del sistema (variables de entrada, funciones de activación y métodos de optimización) y se reportan resultados de una prueba piloto con $n = 30$ estudiantes durante cuatro semanas: el desempeño académico mejoró respecto a la línea base en Optimización y Modelos Matemáticos (+60%), Estadística y Métodos Cuantitativos (+55%), Gestión de Producción (+50%), Gestión de Proyectos (+45%) e Ingeniería de Calidad (+50%). El 62% de los participantes alcanzó mejoras moderadas o altas, y la aceleración del progreso se observó a partir de la semana 3. Adicionalmente, como validación metodológica del componente predictivo, se evaluó una red neuronal mínima viable —una red multicapa (multilayer perceptron, MLP)—, en la que el error absoluto medio (mean absolute error, MAE) se redujo 23% y la raíz del error cuadrático medio (root mean squared error, RMSE) 19% respecto a la línea base (configuración inicial sin ajuste), tras normalización/estandarización y ajuste de hiperparámetros.

Keywords

Neural networks; personalized learning; artificial intelligence in education; Alfa-Pi-Mi; educational optimization; industry 5.0; human-centered innovation.

Abstract

This paper presents the design and implementation of the Alfa-Pi-Mi neural network model, developed to personalize and optimize learning for Industrial Engineering students through artificial intelligence (AI). The model identifies learning patterns and generates tailored instructional strategies. The findings should be interpreted as initial feasibility evidence from a pilot study, not as generalizable effects. We detail the system architecture (input variables, activation functions, and optimization methods) and report pilot results with $n = 30$ students over four weeks: academic performance improved relative to baseline in Optimization and Mathematical Models (+60%) and Statistics and Quantitative Methods (+55%), with additional gains in Production Management (+50%), Project Management (+45%), and Quality Engineering (+50%). Overall, 62% of participants achieved moderate-to-high improvement, and learning trajectories steepened from week 3. In addition, as a methodological validation of the predictive component, a minimal multilayer perceptron (MLP) task reduced mean absolute error (MAE) by 23% and root mean squared error (RMSE) by 19% versus baseline (initial, untuned configuration), after normalization/standardization and hyperparameter tuning.

Introducción

La personalización del aprendizaje es un reto persistente en programas de Ingeniería Industrial, donde la heterogeneidad de perfiles (cognitivos, motivacionales, contextuales) dificulta enfoques instruccionales homogéneos [7]. La literatura reciente sobre analíticas de aprendizaje con IA y aprendizaje profundo en educación abierta destaca beneficios en predicción temprana, personalización y retroalimentación formativa [4], [6]. En paralelo, el marco Industria 5.0 resalta una innovación centrada en el ser humano, resiliente y sostenible [7], convocando a integrar IA sin desplazar la agencia del docente/estudiante.

Este trabajo introduce un modelo basado en redes neuronales multicapa (MLP), denominado *Alfa-Pi-Mi* (APM), propuesto por el autor para identificar patrones de aprendizaje y recomendar estrategias personalizadas en contextos de Ingeniería. El fundamento técnico se apega a buenas prácticas consolidadas —normalización/estandarización para estabilidad numérica [1], selección de funciones de activación (*Rectified Linear Unit*, *ReLU*), Tangente hiperbólica (*tanh*) y Función logística (*sigmoide*) según su comportamiento y riesgos [2], y optimización con Adam y ajuste de hiperparámetros (η , β_1 , β_2 , épocas, batch) [3]—, mientras que la motivación y el beneficio pedagógico se alinean con la evidencia reciente en analíticas de aprendizaje, personalización y retroalimentación formativa [4], [6]. El objetivo es diseñar, implementar y evaluar un prototipo que mejore resultados académicos y sustente decisiones didácticas con trazabilidad ética. El estudio corresponde a una evaluación preliminar mediante una prueba piloto de muestra pequeña ($n = 30$) con una duración de cuatro semanas; por tanto, los hallazgos se interpretan como evidencia inicial de factibilidad y no como resultados generalizables.

Materiales y métodos

A. Arquitectura y variables

El modelo Alfa-Pi-Mi es una red multicapa (MLP) [5] con:

- Capa de entrada (6 variables): exámenes, estilo de aprendizaje, historial de rendimiento, participación, estrés e intervenciones previas.
- Capas ocultas: dos capas (8–8 neuronas) para capturar relaciones no lineales complejas.
- Capa de salida: cabezales para recomendación de estrategias (p. ej., simulaciones, andamiaje, intervención).

Para que el modelo aprenda de forma estable, primero “ponemos en la misma escala” todas las variables de entrada (normalización/estandarización). Esto significa, por ejemplo, transformar notas de 0–100 y niveles de estrés de 0–10 a rangos comparables, evitando que una variable “pese” más solo por tener números más grandes. En algunos casos también aplicamos este ajuste dentro de la red (sobre ciertas activaciones). Con ello se evitan inestabilidades numéricas (como valores no válidos o *Not a Number*, *NaN*), algo que puede ocurrir cuando los datos tienen distribuciones muy distintas o rangos muy dispares [1].

En figura 1, se puede observar la arquitectura del modelo.

B. Funciones de activación (selección y justificación)

- ReLU en capas ocultas por eficiencia y menor saturación en la región positiva; Tanh en bloques donde se desea activación centrada alrededor de 0; Sigmoide en salidas binarias (p. ej., probabilidad de intervención). Esta selección y sus riesgos/mitigaciones siguen evidencia comparativa de activaciones [2].
- Definiciones y derivadas (véase Sección II-D, Ecs. A1–A6) [2].

C. Entrenamiento, optimización y validación

El entrenamiento se realizó con un conjunto inicial de 30 estudiantes (curso piloto), división train/test 70/30 y 100 ciclos (variantes de 200–300 épocas) con tasa de aprendizaje 0.001, batch 16–32 y optimizador Adam; la pérdida utilizada fue MSE. La elección de Adam y pautas de ajuste de hiperparámetros (η , β_1 , β_2 , épocas, batch) siguen recomendaciones de encuestas recientes sobre algoritmos de optimización en redes modernas [3]. Se reportan MAE/RMSE en train/test y gráfico de residuales para inspección de sobreajuste.

D. Fórmulas y expresiones clave (activaciones)

Las expresiones que siguen corresponden a las definiciones canónicas de las funciones de activación más utilizadas en aprendizaje profundo. Para una discusión aplicada y comparativa de su idoneidad, ventajas y limitaciones en redes profundas, véase [2].

1. Rectified Linear Unit (ReLU)

$$\text{Definición: } \text{ReLU}(z) = \max(0, z) \quad (\text{A1})$$

$$\text{Derivada: } \frac{d}{dz} \text{ReLU}(z) = \begin{cases} 0, & z < 0 \\ 1, & z > 0 \end{cases} \text{ (indef. en } z = 0) \quad (\text{A2})$$

2. Tangente hiperbólica (Tanh)

$$\text{Definición: } \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}} \quad (\text{A3})$$

$$\text{Derivada: } \frac{d}{dz} \tanh(z) = 1 - \tanh^2(z) \quad (\text{A4})$$

3. Sigmoide (logística)

$$\text{Definición: } \sigma(z) = \frac{1}{1 + e^{-z}} \quad (\text{A5})$$

$$\text{Derivada: } \frac{d}{dz} \sigma(z) = \sigma(z)(1 - \sigma(z)) \quad (\text{A6})$$

4. Buenas prácticas:

- Normalización previa de las entradas (y, si aplica, de activaciones): pone las variables en escalas comparables, evita saturación de *tanh/sigmoide*, estabiliza gradientes y reduce riesgos de valores NaN durante el entrenamiento. Mejora la convergencia del modelo. Ver [1].
- Inicialización acorde a la activación: en capas con ReLU usamos inicialización He/Kaiming para conservar la varianza a través de las capas, evitar “neuronas muertas” y mantener gradientes informativos. Ver [2].

Resultado: entrenamiento más estable y rápido, con menos problemas numéricos y mejor desempeño global [1], [2].

inicialización acorde (p. ej., He para ReLU), normalización previa (entradas/activaciones) para reducir saturación y dispersión numérica [1], [2].

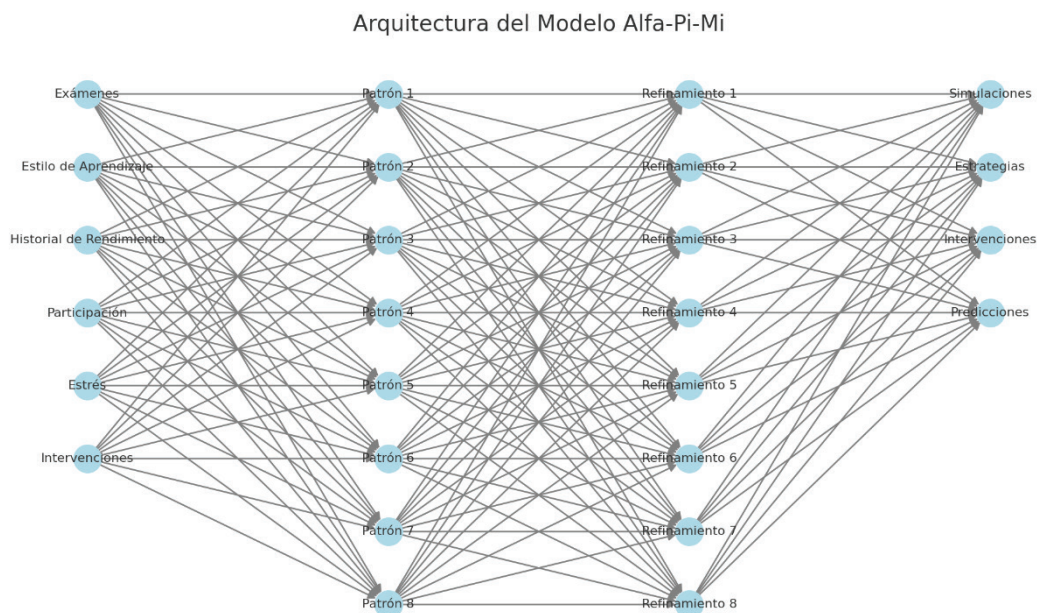


Figura 1. Arquitectura del Modelo Alfa-Pi-Mi.

La Figura 1 presenta la arquitectura del modelo Alfa-Pi-Mi (MLP) como un flujo de procesamiento desde seis variables de entrada —exámenes, estilo de aprendizaje, historial de rendimiento, participación, estrés e intervenciones previas— hacia dos capas ocultas (8-8 neuronas) y una capa de salida que combina predicción del desempeño con recomendaciones pedagógicas (simulaciones, andamiaje e intervención). Se explicita el preprocesamiento mediante normalización/estandarización para estabilidad numérica y la selección de funciones de activación (ReLU/Tanh/Sigmoide) conforme al rol de cada submódulo, lo cual reduce saturación, evita NaN y acelera la convergencia. El valor de la figura es clarificar cómo los insumos estudiantiles se transforman en señales accionables para la docencia, integrando buenas prácticas de inicialización y optimización que sustentan la robustez del sistema.

E. Contexto educativo y diseño del piloto

E1. Piloto empírico (n=30; 4 semanas; S0/S4)

El estudio se implementó como piloto empírico en un contexto educativo de Ingeniería Industrial, con una cohorte de $n = 30$ estudiantes y una duración de cuatro semanas. Para evaluar el efecto del enfoque en competencias disciplinares, las mediciones se mapearon a las siguientes áreas del plan de estudios: Optimización y Modelos Matemáticos, Estadística y Métodos Cuantitativos, Gestión de Producción y Operaciones, Gestión de Proyectos y Análisis de Riesgos, e Ingeniería de Calidad y Mejora Continua. Se aplicaron dos cortes por estudiante, correspondientes a línea base (Semana 0, S0) y post-intervención (Semana 4, S4), mediante instrumentos equivalentes y con escala 0–100.

Como parte del preprocesamiento, se utilizó normalización/estandarización de variables de entrada (y cuando aplicó, tratamiento equivalente para asegurar estabilidad numérica), con el fin de evitar inestabilidades (p. ej., NaN) y mejorar la convergencia del entrenamiento [1]. Los resultados reportados en el Cuadro 2 y en las Figuras 2–3 corresponden exclusivamente a este piloto empírico, y deben interpretarse como evidencia preliminar de factibilidad (no generalizable) dada la escala y duración del estudio.

E2. Validación metodológica (simulación Monte Carlo y consultas agregadas; no empírica)

Adicionalmente, con el fin de evaluar robustez, así como sensibilidad y estabilidad interna del enfoque (validación metodológica), se ejecutaron réplicas independientes mediante simulación Monte Carlo, introduciendo perturbaciones controladas sobre las variables de entrada para caracterizar la variabilidad del desempeño del modelo y la consistencia del procedimiento. En este componente se incorporó un módulo de consultas interuniversitarias orientado exclusivamente al entrenamiento y validación a nivel de distribución, sin acceso ni procesamiento de información personal individual.

El módulo interuniversitario operó con resultados académicos comparables en forma agregada/anónima (p. ej., resúmenes estadísticos), preservando la privacidad y evitando identificación individual. En la validación metodológica, dicha señal agregada se utilizó como regularización externa y como mecanismo de alineación de distribuciones (domain alignment) entre contextos institucionales. Los productos derivados de esta sección se reportan como evidencia metodológica (no empírica) y no deben interpretarse como estimación de impacto real del piloto sobre el desempeño académico de estudiantes.

En el escenario simulado, el modelo se implementó como una red neuronal multicapa (MLP) entrenada con Adam (pérdida MSE; $\eta = 0.001$; batch 16–32; 100–300 épocas), con selección de funciones de activación (ReLU/tanh/sigmoide) y ajuste de hiperparámetros conforme a la evidencia reportada [2], [3]. Cada réplica ejecutó un ciclo de entrenamiento/validación (p. ej., k-fold con $k = 5$) y, tras integrar la señal agregada del módulo interuniversitario, se recalibraron parámetros para reducir sobreajuste y mejorar la consistencia del desempeño fuera del dominio local.

En síntesis: Cuadro 2 y Figuras 2 y 3 provienen del piloto empírico; la simulación Monte Carlo y el módulo interuniversitario se emplean únicamente como validación metodológica.

E3. Ética y privacidad

Las consultas interuniversitarias se manejaron a un nivel agregado, no identificable, y con fines exclusivamente metodológicos de simulación; no se accedió ni procesó información personal individual. El procedimiento se diseñó para cumplir principios de privacidad, minimización de datos y uso responsable en investigación educativa. Por otro lado, el uso de datos simulados/anonimizados, registro de prompts y decisiones, no delegación de la decisión pedagógica a la IA; control docente en todo momento, alineado con prácticas de transparencia y agencia promovidas por Industria 5.0 [7].

F. Fuentes de datos y naturaleza de la evidencia

Para asegurar transparencia metodológica, se diferencian explícitamente las fuentes de evidencia del estudio: piloto empírico ($n = 30$; 4 semanas), simulación Monte Carlo (validación metodológica) y caso didáctico (material ilustrativo). El cuadro 1 sintetiza su naturaleza, propósito y forma de reporte.

Cuadro 1. Fuentes de Datos y naturaleza de la evidencia (piloto, simulación y caso didáctico).

Tipo de evidencia / componente	Fuente y naturaleza de los datos	Propósito metodológico	Salida reportada en el manuscrito
Piloto empírico (n = 30; 4 semanas)	Datos observados del piloto educativo (instrumentos de evaluación y registros académicos). Escala 0–100. Cortes temporales S0 (línea base) y S4 (cierre); seguimiento intersemanal S1–S3.	Estimar factibilidad y señales iniciales de mejora bajo implementación real en aula; evaluar cambio de desempeño por área y trayectoria temporal (semanas).	Resultados principales: mejoras por área (p. ej., Optimización, Estadística, etc.), distribución de mejora (p. ej., % moderado-alto), patrón temporal (aceleración desde semana 3).
Simulación / réplicas Monte Carlo (evidencia metodológica)	Datos sintéticos generados para análisis de sensibilidad/robustez (réplicas) a partir de supuestos/parametrización definida. No corresponden a mediciones reales de estudiantes individuales.	Evaluar robustez, varianza y estabilidad del método; analizar sensibilidad a hiperparámetros/optimizador y estabilidad numérica; apoyar validación interna del enfoque.	Resultados metodológicos (se reportan por separado): métricas de estabilidad/variabilidad, sensibilidad a parámetros, evidencia de consistencia interna. No se presentan como “impacto real” en desempeño académico.
Caso didáctico a posteriori (ilustrativo/pedagógico)	Material de diseño instruccional y aplicación guiada (caso) desarrollado para apoyar replicabilidad. Puede usar datos simulados/ejemplificados o descripciones procedimentales.	Asegurar transferibilidad y guía de implementación (cómo llevar el modelo al aula); aportar un “how-to” pedagógico sin contaminar la evidencia empírica.	Se traslada a Anexo A (material complementario). En el cuerpo del artículo se incluye una síntesis breve y referencia: “ver Anexo A”.

Nota 1 (línea base): S0 corresponde a la medición inicial previa a la intervención; S4 a la medición final al concluir la semana 4.

Nota 2 (separación de evidencia): Los resultados de simulación se reportan únicamente como validación metodológica y no deben interpretarse como efectos empíricos.

El caso didáctico aplicado se presenta como material complementario en el Anexo A, con fines ilustrativos y de transferibilidad, y no constituye evidencia empírica adicional del piloto.

Resultados

A. Resultados del piloto: desempeño por áreas (evidencia empírica)

A1. Desempeño por áreas (S0 vs S4, escala 0-100)

El Cuadro 2 resume los resultados del piloto empírico (promedios de cohorte). Las dimensiones evaluadas se mapearon a áreas de conocimiento para reflejar la transferencia de habilidades hacia dominios cuantitativos y de gestión. Se observan mejoras en Optimización y Modelos Matemáticos, Estadística y Métodos Cuantitativos, Gestión de Producción y Operaciones, Gestión de Proyectos y Análisis de Riesgos, e Ingeniería de Calidad y Mejora Continua.

El Puntaje inicial (S0, escala 0–100) corresponde al promedio de cohorte en Semana 0. La Mejora relativa (% de S0) representa el incremento porcentual calculado respecto al valor inicial; para facilitar la lectura en la escala 0-100, el cambio absoluto puede expresarse como:

$$\Delta(\text{puntos}) = S0 \times \frac{\text{Mejora relativa}}{100}$$

Cuadro 2. Resultados Observados por Área.

Área de Conocimiento	Puntaje Inicial (S0, escala 0-100)	Mejora relativa (% de S0)
Optimización y Modelos Matemáticos	70	60
Estadística y Métodos Cuantitativos	50	55
Gestión de Producción y Operaciones	65	50
Gestión de Proyectos y Análisis de Riesgos	60	45
Ingeniería de Calidad y Mejora Continua	40	50

Nota: La “mejora relativa” se calcula respecto a S0. El cambio absoluto puede expresarse como

$$\Delta(\text{puntos}) = S0 \times \frac{\text{mejora relativa}}{100} .$$

B. Resultados del piloto: progresión temporal y distribución de mejora (evidencia empírica).

B.1 Progresión temporal (S1-S4 vs S0)

La Figura 2 muestra la progresión semanal de la mejora relativa por área del plan de estudios (S1-S4 respecto a la línea base S0), evidenciando un incremento sostenido con aceleración a partir de la semana 3. En Optimización y Modelos Matemáticos y en Estadística y Métodos Cuantitativos se observan las pendientes más pronunciadas, coherentes con los resultados agregados reportados (+60% y +55%, respectivamente), mientras que Producción, Proyectos y Calidad presentan mejoras consistentes, pero de menor magnitud relativa. Esta visual permite identificar el periodo de “calentamiento” del modelo (≈ 2 semanas) antes de que las recomendaciones personalizadas desplieguen su mayor efecto.

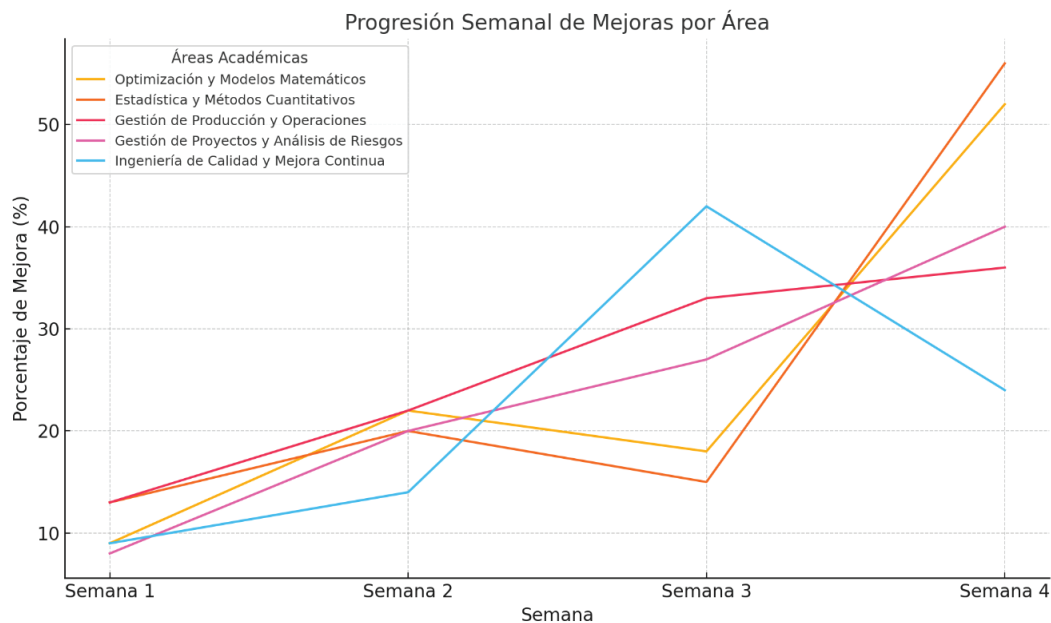


Figura 2. Progresión semanal de mejoras por área. Muestra incrementos graduales en el desempeño, especialmente en las semanas 3 y 4.

B2. Distribución de estudiantes por nivel de mejora (S4 vs S0)

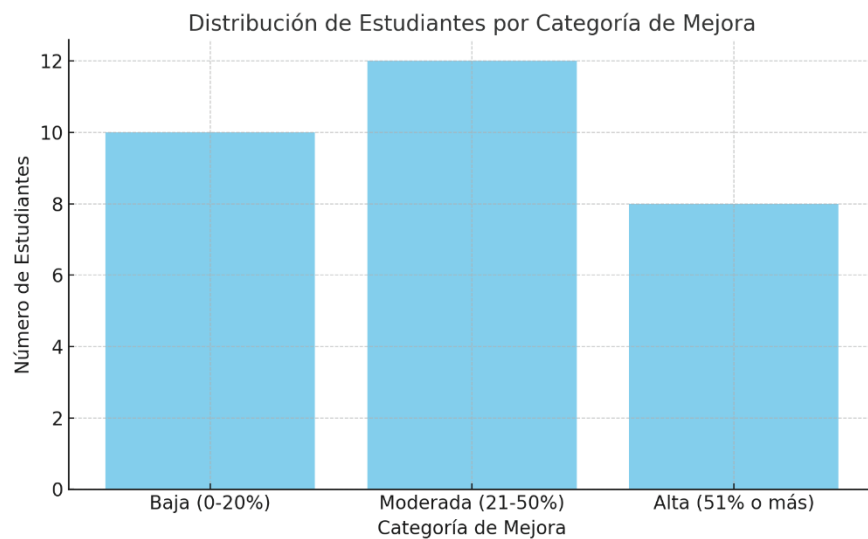


Figura 3. Distribución de estudiantes por nivel de mejora.

La Figura 3 sintetiza la distribución de estudiantes por nivel de mejora al cierre del piloto (S4 vs S0), clasificando en baja (0–20%), moderada (21–50%) y alta ($\geq 51\%$), y destacando que el 62% se ubica en los rangos moderado-alto. Esta representación orienta decisiones focalizadas: en los casos de baja mejora se observaron, de manera descriptiva, patrones concurrentes como mayor variabilidad en indicadores de estrés y menor participación registrada. Estas

observaciones se reportan como señales exploratorias (no causales) para orientar refuerzos específicos (p. ej., acompañamiento, hábitos de estudio, micro-intervenciones) y maximizar la efectividad en subgrupos.

C. Resultados del piloto: trazabilidad didáctica de la personalización

Estos hallazgos son coherentes con literatura reciente sobre personalización, predicción de desempeño y beneficios pedagógicos de IA/AA profunda en contextos educativos [4], [6]. Los resultados evidencian que el modelo Alfa-Pi-Mi personaliza el aprendizaje con trazabilidad didáctica, al generar recomendaciones explícitas (p. ej., simulaciones, andamiaje e intervenciones puntuales) asociadas a patrones detectados en las variables de entrada. En términos de impacto, el piloto mostró mejoras cuantitativas por área (ver Cuadro 2) y una dinámica temporal consistente con el aumento de desempeño observado desde la semana 3 (Figuras 2 y 3).

En particular, la aceleración de la curva de aprendizaje a partir de la semana 3 se interpreta como una observación descriptiva del piloto, compatible con un periodo inicial de estabilización del sistema (2 semanas) previo a que las recomendaciones personalizadas alcancen mayor consistencia operativa; este comportamiento requiere replicación para confirmar su estabilidad en otras cohortes y contextos. En conjunto, la proporción de estudiantes en rangos moderado/alto de mejora sugiere que las recomendaciones generadas por la red aportan soporte analítico para la toma de decisiones pedagógicas (ajuste de actividades, diferenciación de apoyos y secuenciación de contenidos) [4], [6].

Ejemplos de recomendaciones, plantillas de trazabilidad y guía operativa de implementación se incluyen en el Anexo A.

D. Resultados del componente predictivo (RN mínima viable: MLP y métricas)

D1. Métricas predictivas (MAE, RMSE) vs línea base

En la tarea de predicción mínima viable con una red neuronal multicapa (multilayer perceptron, MLP), el error absoluto medio (MAE) se redujo en 23% y la raíz del error cuadrático medio (RMSE) en 19% respecto a la línea base, definida como la configuración inicial del modelo sin ajuste de hiperparámetros (parámetros por defecto y sin calibración). Tras aplicar normalización/estandarización y optimización de hiperparámetros, estas mejoras respaldan la validez interna del enfoque y su utilidad para anticipar desempeños y priorizar intervenciones.

D2. Verificación de estabilidad (residuales y entrenamiento-validación)

Las métricas anteriores se complementaron con inspecciones de residuales para descartar patrones sistemáticos de error. Adicionalmente, la evolución entrenamiento-validación no mostró divergencias relevantes, lo cual se alinea con buenas prácticas de estabilidad numérica y selección de funciones de activación y optimizador [1]–[3].

E. Limitaciones observadas (alcance e interpretación de resultados)

Si bien los hallazgos son consistentes con la literatura sobre analíticas de aprendizaje con IA [4], [6], el tamaño muestral ($n = 30$) recomienda cautela en la generalización. La infraestructura computacional requerida y la sensibilidad a hiperparámetros/optimizador constituyen restricciones prácticas, atenuadas mediante validación y lectura sistemática de residuales. Asimismo, el desempeño del modelo depende de la calidad de los datos de entrada y de la normalización/estandarización para evitar inestabilidades numéricas (p. ej., NaN) [1]. En consecuencia, se sugiere interpretar las mejoras como evidencia de factibilidad y eficacia inicial, sujeta a replicación con muestras más amplias y estratificación por perfiles.

En particular, el alcance interpretativo está delimitado por:

- Muestra: piloto con $n = 30$ estudiantes.
- Duración: intervención de 4 semanas, lo que limita inferencias sobre estabilidad del efecto en el mediano plazo.
- Contexto: implementación en un contexto institucional y curricular específico, por lo que la transferibilidad a otros programas requiere replicación.

F. Consideraciones éticas y de privacidad (resultado operacional del piloto)

El piloto se ejecutó con datos simulados/anonimizados y decisiones bajo control docente en todo momento. El uso de IA se limitó a apoyar la detección de patrones y la generación de recomendaciones, sin delegar la conceptualización, el juicio pedagógico ni la evaluación del aprendizaje. En las simulaciones con réplicas y consultas agregadas interuniversitarias (LatAm), se aplicaron criterios de minimización de datos, no identificación y uso responsable; este marco operativo se alinea con los principios de Industria 5.0 orientados a soluciones human-centric, resilientes y sostenibles [7]. Estos resguardos no solo son un requisito de integridad, sino que constituyen un resultado operativo del estudio: demuestran que la integración de IA puede implementarse con garantías éticas y de privacidad en contextos académicos reales.

Conclusiones y recomendaciones

El modelo Alfa-Pi-Mi mostró potencial para personalizar el aprendizaje y apoyar mejoras en el desempeño académico en dominios cuantitativos y de gestión, aportando trazabilidad didáctica mediante recomendaciones explícitas y soporte analítico para la toma de decisiones pedagógicas [4], [6]. En conjunto, los resultados constituyen evidencia inicial de factibilidad del enfoque y de su utilidad como herramienta de apoyo docente en contextos formativos de Ingeniería Industrial.

Dado el carácter piloto del estudio, los hallazgos deben interpretarse con cautela y no como efectos generalizables. Como líneas de continuidad se prioriza: (i) replicación con muestras más amplias y mayor horizonte temporal, (ii) estratificación por perfiles para evaluar robustez y equidad del desempeño, (iii) fortalecimiento de interpretabilidad y trazabilidad para docentes, y (iv) consolidación de gobernanza ética y resguardos de privacidad en el ciclo analítico. El caso didáctico aplicado y la hoja de ruta complementaria se presentan como material de apoyo en el Anexo A.

Referencia al caso didáctico (material complementario)

Con el fin de mantener el foco en el aporte metodológico y los resultados del estudio piloto, el caso didáctico aplicado y sus líneas de desarrollo se presentan como material complementario en el Anexo A.

Agradecimientos

Se agradece al CURSO TALLER PIDEA: Tecnologías Emergentes. Recursos creados con IA (CEDA) y al facilitador Guillermo Bautista por la guía metodológica que inspiró la construcción del recurso didáctico y su implementación en el aula.”

Comité organizador de la II Convención Internacional de Innovación en Ingeniería y V Congreso Internacional de Ingeniería: Industria 5.0 (Universidad Continental, Arequipa, Perú), por el espacio para la difusión y discusión académica de esta ponencia.

Referencias

- [1] L. Huang et al., "Normalization Techniques in Training DNNs: Methodology, Analysis, and Application," IEEE Trans. Pattern Anal. Mach. Intell., vol. 45, no. 8, pp. 10173–10196, Aug. 2023, doi: 10.1109/TPAMI.2023.3250241.
- [2] S. R. Dubey, S. K. Singh, and B. B. Chaudhuri, "Activation Functions in Deep Learning: A Comprehensive Survey and Benchmark," Neurocomputing, 2022, doi: 10.1016/j.neucom.2022.06.111.
- [3] A. Abdulkadirov et al., "Survey of Optimization Algorithms in Modern Neural Networks," Mathematics, vol. 11, no. 4, p. 831, 2023, doi: 10.3390/math11040831.
- [4] O. Estrada-Molina, J. Mena, and A. López-Padrón, "The Use of Deep Learning in Open Learning: A Systematic Review (2019–2023)," IRRODL, vol. 25, no. 3, Aug. 2024.
- [5] I. D. Mienye and P. Sun, "A Comprehensive Review of Deep Learning: Architectures, Applications, and Challenges," Information, vol. 15, no. 12, p. 755, 2024, doi: 10.3390/info15120755.
- [6] N. Cannistrà et al., "Machine Learning and Generative AI in Learning Analytics: A Review," Applied Sciences, vol. 15, no. 9, p. 3769, 2025, doi: 10.3390/app15093769.
- [7] European Commission, "Industry 5.0: Towards a sustainable, human-centric and resilient European industry," 2021. (Official report). [Publications Office of the EU](https://publications.ec.europa.eu/publication-detail/?id=646022213199747781&lang=en)
- [8] Universidad Continental, "II Convención Internacional... Industria 5.0," Arequipa, Perú, 2023. (Evento; válido como referencia contextual.)

Declaración de uso de herramientas de Inteligencia Artificial (IA)

En este trabajo se utilizaron herramientas de IA como apoyo a la redacción y edición (claridad/consistencia), traducción puntual (ES↔EN) y generación de figuras (diagramas y gráficos). Todas las salidas fueron revisadas críticamente por el autor para verificar exactitud y coherencia. Para análisis y simulación se emplearon hojas de cálculo y/o scripts propios; los asistentes de IA se usaron solo como apoyo instrumental (sensibilidad, escenarios), con validación humana de resultados. No se delegó en la IA la conceptualización, el diseño metodológico, la interpretación ni las conclusiones. Las figuras generadas con asistencia de IA se limitaron a la ilustración de conceptos ya descritos y fueron verificadas para corresponder a los datos y teorías presentadas. Esta declaración se incluye en cumplimiento de las políticas editoriales de uso de IA de la revista (integridad, transparencia y responsabilidad).

Anexos

Anexo A. Caso didáctico aplicado (a posteriori) y líneas de desarrollo

A1. Propósito del anexo

Este anexo presenta un caso didáctico aplicado y una guía operativa de implementación para facilitar la transferibilidad del modelo Alfa-Pi-Mi (APM) a contextos docentes similares. El material se incluye con fines ilustrativos y de replicabilidad; no constituye evidencia adicional del piloto empírico reportado en el cuerpo del artículo.

A2. Contexto. En el marco del Curso-Taller PIDEA:

"Tecnologías Emergentes. Recursos creados con IA", coordinado por el Centro de Desarrollo Académico (CEDA) del TEC y facilitado por Guillermo Bautista, se diseñó e implementó en aula un recurso didáctico con integración de IA. El taller tiene como propósito formar al profesorado en el uso pedagógico de tecnologías emergentes, priorizando (i) el diseño instruccional del recurso, (ii) la alineación con resultados de aprendizaje y (iii) la justificación explícita del valor agregado de la IA en términos de accesibilidad, retroalimentación y personalización. Como

evidencia, se elaboró un artefacto aplicable en clase y una memoria técnica que documenta objetivos, herramientas empleadas, consideraciones éticas y criterios de evaluación del aprendizaje. [4], [6].

Título del recurso. *De la fórmula al criterio: simulación financiera con IA para decisiones de inversión.*

Propósito. Transitar del cálculo mecánico (VAN/TIR/VAE) al criterio profesional, integrando sensibilidad, escenarios coherentes y una RN simple para predecir demanda/costos bajo incertidumbre [4], [6].

Duración. 4 semanas (integrado al curso).

Asignatura/Público. Ingeniería → PI-4505 Análisis Económico (Licenciatura, 7.º semestre). 25–30 estudiantes por grupo.

Objetivos de aprendizaje. (1) Aplicar VAN/TIR/VAE con supuestos explícitos; (2) ejecutar sensibilidad y escenarios coherentes; (3) integrar y comparar (método clásico, IA y RN) para justificar la decisión; (4) documentar uso ético y trazable de la IA (datos, sesgos, límites, decisiones no delegables) [4], [6].

Aplicaciones de IA y tecnología

- Asistente de IA para: (i) tablas “¿qué pasa si...?”; (ii) escenarios base/optimista/pesimista con verificación de coherencia; (iii) chequeo de consistencia del flujo (signos, unidades, inflación, inversión vs. gastos operativos) [4], [6].
- Excel para flujos de caja y percentiles p50/p85/p95.
- RN simple (MLP): 1 capa oculta (8–16 neuronas), ReLU en ocultas, salida lineal; Adam; pérdida MSE; 200–300 épocas; *batch* 16–32; *split* 70/30 (train/test). Reporte: MAE/RMSE (train/test) + gráfico de residuales + 4–6 líneas de interpretación (sobreajuste, tamaño muestral, variables omitidas). Se mantiene normalización/estandarización para estabilidad numérica y evitar NaN [1]; la selección de activaciones (ReLU/Tanh/Sigmoide) se fundamenta en [2]; el uso de Adam y criterio de ajuste de hiperparámetros se apoya en [3].

Prompts de referencia

- **Sensibilidad:** “Genera una tabla de sensibilidad del VAN variando: demanda ($\pm 10\%$, $\pm 20\%$), costo unitario ($\pm 5\%$, $\pm 10\%$) y tasa (8%, 10%, 12%, 14%). Ordena resultados e indica la variable más influyente.”
- **Escenarios coherentes:** “Propón 3 escenarios (base/optimista/pesimista) para demanda, costo y precio, con justificación breve; devuelve tabla de supuestos y verifica coherencia (capacidad, costos asociados, restricciones).”
- **Chequeo de coherencia:** “Revisa este flujo de caja e identifica inconsistencias (signos, unidades, inflación duplicada, inversión vs. gastos). Entrega hallazgos y correcciones.”

Secuencia didáctica (4 semanas)

- **S1:** Mini-caso “botón rojo/verde”; construcción del flujo; percentiles p50/p85/p95; IA para sensibilidad; bitácora ética ($\frac{1}{2}$ pág) [4], [6].
- **S2:** IA para escenarios + chequeo de coherencia; contraste en VAN/TIR; ruta “Solo Excel” disponible [4], [6].

- **S3:** RN para predecir demanda/costo (normalización previa [1], activaciones [2], Adam/hiperparámetros [3]); integración al VAN; tabla de triangulación (Clásico vs IA vs RN) e interpretación [4], [6].
- **S4:** Informe individual (≤ 2 págs) + anexos (Excel, capturas IA/RN, ficha ética, triangulación) y pitch (2 min) [4], [6].

Criterios de accesibilidad y ética. Ruta “Solo Excel” con evidencias equivalentes; trabajo en parejas/tercias (solo ejecución técnica); datos anonimizados; ficha de integridad (origen de datos, sesgos, límites, decisiones no delegables); registro de prompts/configuración con fecha/versión; declaración de autoría individual [4], [6].

Evaluación del aprendizaje (rúbrica, 100 pts)

- Modelo financiero trazable (25)
- Sensibilidad y escenarios (25)
- Criterio profesional + triangulación (30)
- Ética/transparencia (10)
- Comunicación (10)

Evidencias: archivo Excel, capturas del asistente de IA, resultados de RN, ficha ética (1 pág), tabla de triangulación, informe (2 págs), pitch (2 min) [4], [6].

Anexos obligatorios (entrega individual)

A) Prompts usados (con fecha/versión y limitaciones).

B) Configuración RN (hiperparámetros, corridas, MAE/RMSE, residuales) [1]–[3].

C) Ficha de integridad/ética (1 pág) [4], [6].

D) Tabla de triangulación (Clásico vs IA vs RN) [4], [6].

E) Ruta “Solo Excel” (si aplica): hoja con fórmulas visibles y justificación equivalente.

Resultado esperado. Un recurso didáctico individual, ético y trazable que demuestre con claridad por qué y cómo la IA conduce al aprendizaje: acelerando exploración de supuestos/escenarios, fomentando argumentación y defensa de decisiones bajo incertidumbre, y dejando trazabilidad técnica y ética [4], [6].