

Diabetes Prediction Using the Ghost Deep Learning Model

Darwin Patiño-Pérez, Ph.D.¹; Luis Armijos-Valarezo, MSc¹; Luis Chóez-Acosta, MSc¹;

Sara Falconí-SanLucas, MSc²; **Celia Munive-Mora, BS³**

¹Universidad de Guayaquil, Facultad de Ciencias Matemáticas y Física,
Grupo de Investigación de Inteligencia Artificial, Ecuador, darwin.patinop@ug.edu.ec,
luis.armijosv@ug.edu.ec, luis.choeza@ug.edu.ec,

²Universidad de Guayaquil, Facultad de Ciencias Médicas, Ecuador, sara.falconis@ug.edu.ec,

³St Luke's University Hospital Network, PA, United States, celia.munive@sluhn.org

Abstract— *The article presents an innovative approach for diabetes prediction by applying a phantom deep learning model, designed to optimize accuracy and efficiency in early diagnosis. This model uses a robust clinical data set, integrating advanced artificial intelligence techniques with a lightweight and efficient architecture, allowing computational costs to be significantly reduced without sacrificing predictive performance. Key features of the model include its ability to handle imbalanced data sets, common in medical environments, and its adaptability to various clinical contexts, making it especially versatile. Experimental results demonstrate that the phantom deep learning model greatly outperforms conventional methods such as artificial neural networks (ANNs) in terms of accuracy and efficiency. Specifically, the phantom model achieved an accuracy of 79.24% without overfitting in identifying patients at risk of diabetes, compared to the 88.57% obtained by the conventional ANN model with overfitting. In addition, the loss of the phantom model has a greater generalization capacity and lower error in predictions. These results highlight the superiority of the phantom model in terms of accuracy and stability. Its scalability and low resource requirements make it a viable option for implementation in public and private health systems; The phantom deep learning model represents a significant advance in the application of artificial intelligence in the healthcare field.*

Keywords— *Diabetes prediction, Ghost deep learning model, artificial intelligence, key features, early detection.*

Predicción de Diabetes usando el Modelo de Aprendizaje Profundo Fantasma

Darwin Patiño-Pérez, Ph.D.¹; Luis Armijos-Valarezo, MSc¹; Luis Chóez-Acosta, MSc¹;

Sara Falconí-SanLucas, MSc²; Celia Munive-Mora, BS³

¹Universidad de Guayaquil, Facultad de Ciencias Matemáticas y Física,
Grupo de Investigación de Inteligencia Artificial, Ecuador, darwin.patinop@ug.edu.ec,
luis.armijosv@ug.edu.ec, luis.choeza@ug.edu.ec,

²Universidad de Guayaquil, Facultad de Ciencias Médicas, Ecuador, sara.falconis@ug.edu.ec,

³St Luke's University Hospital Network, PA, United States, celia.munive@sluhn.org

Resumen— El artículo presenta un enfoque innovador para la predicción de diabetes mediante la aplicación de un modelo de aprendizaje profundo fantasma, diseñado para optimizar la precisión y eficiencia en el diagnóstico temprano. Este modelo utiliza un conjunto de datos clínicos robustos, integrando técnicas avanzadas de inteligencia artificial con una arquitectura ligera y eficiente, lo que permite reducir significativamente los costes computacionales sin sacrificar el rendimiento predictivo. Entre las características clave del modelo destacan su capacidad para manejar conjuntos de datos desequilibrados, común en entornos médicos, y su adaptabilidad a diversos contextos clínicos, lo que lo hace especialmente versátil. Los resultados experimentales demuestran que el modelo de aprendizaje profundo fantasma supera ampliamente a los métodos convencionales, como las redes neuronales artificiales (RNA), en términos de precisión y eficiencia. En concreto, el modelo fantasma alcanzó un accuracy del 79.24% sin sobreajuste en la identificación de pacientes con riesgo de diabetes, en comparación con el 88.57% obtenido por el modelo convencional de RNA con sobreajuste. Además, el loss (pérdida) del modelo fantasma tiene una mayor capacidad de generalización y menor error en las predicciones. Estos resultados resaltan la superioridad del modelo fantasma en términos de precisión y estabilidad. Su escalabilidad y bajo requerimiento de recursos lo convierten en una opción viable para su implementación en sistemas de salud públicos y privados; el modelo de aprendizaje profundo fantasma representa un avance significativo en la aplicación de la inteligencia artificial en el ámbito de la salud.

Palabras clave— predicción de diabetes, modelo de aprendizaje profundo fantasma, inteligencia artificial, características clave, detección temprana.

I. INTRODUCCIÓN

La diabetes es una enfermedad crónica que se caracteriza por niveles elevados de glucosa en la sangre, resultado de una producción insuficiente de insulina o de la incapacidad del cuerpo para utilizar eficazmente esta hormona[1]. La insulina, producida por el páncreas, es esencial para regular el azúcar en la sangre y permitir que las células la utilicen como fuente de energía. Cuando este proceso falla, la glucosa se acumula en el torrente sanguíneo, lo que puede derivar en complicaciones graves a largo plazo, como enfermedades cardiovasculares, daño renal, problemas de visión y neuropatías[2].

La diabetes es una de las principales causas de mortalidad a nivel mundial y su prevalencia ha aumentado significativamente en las últimas décadas, debido en parte a

factores como el sedentarismo, la obesidad y los hábitos alimenticios poco saludables. La diabetes mellitus está considerada como un grupo de trastornos metabólicos caracterizados por niveles elevados de glucosa en la sangre, debido a defectos en la producción o acción de la insulina[3].

Entre los diferentes tipos que existen la más común es la denominada diabetes tipo II según la Fig. 1, cuya forma está asociada a la resistencia a la insulina y a una producción insuficiente de esta hormona, siendo frecuente en adultos con factores de riesgo como obesidad, sedentarismo y antecedentes familiares. La diabetes mellitus es una enfermedad crónica que afecta a millones de personas en todo el mundo, y su detección temprana es crucial para prevenir complicaciones graves[4].

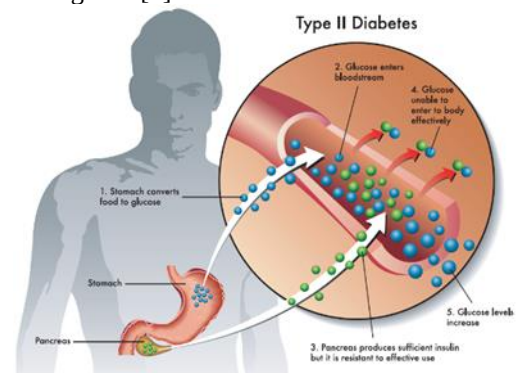


Fig. 1 Diabetes Tipo II.

En los últimos años, los modelos de aprendizaje profundo han demostrado ser herramientas efectivas para la predicción de enfermedades. Sin embargo, estos modelos suelen ser computacionalmente costosos y propensos al sobreajuste, lo que limita su aplicabilidad en entornos clínicos[5]. Los modelos de aprendizaje profundo fantasma surgen como una solución innovadora, enfocándose en la eficiencia y generalización mediante la reducción de redundancias y la optimización de recursos computacionales[6].

Este estudio propone un modelo de red neuronal artificial ver la Fig. 2, que utilice neuronas fantasmas, así como técnicas de regularización L1 o L2, dropout y estratificación de datos para mejorar la calidad del modelo en la predicción de diabetes. El modelo de aprendizaje profundo fantasma

representa una innovadora aproximación en el campo de las redes neuronales artificiales, caracterizada por su capacidad única de emplear "neuronas fantasmas" que no existen físicamente en la red pero que participan activamente en el proceso de aprendizaje[7].

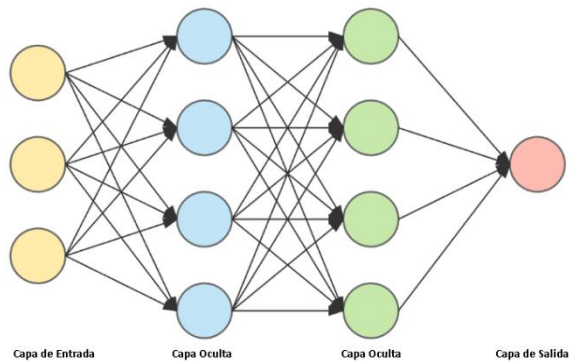


Fig. 2 Red Neuronal Artificial.

Esta arquitectura dinámica permite que la red se adapte y evolucione según las necesidades específicas del problema, activando temporalmente estas neuronas virtuales cuando son necesarias para mejorar el rendimiento del modelo. A diferencia de las redes neuronales artificiales tradicionales, donde la arquitectura permanece estática durante el entrenamiento, el modelo fantasma puede modificar su estructura de manera fluida y eficiente[8].

Una de las ventajas más significativas del aprendizaje profundo fantasma, es su capacidad para optimizar los recursos computacionales ver Fig. 3, mientras mantiene e incluso mejora la eficacia del aprendizaje. Mediante la implementación de conexiones dinámicas y la activación selectiva de neuronas fantasma, el modelo logra reducir significativamente el riesgo de sobreajuste, un problema común en las redes neuronales tradicionales. Además, esta aproximación permite una mejor generalización del conocimiento adquirido, lo que se traduce en un rendimiento superior cuando el modelo se enfrenta a datos nuevos y desconocidos[9].



Fig. 3 Ghost Deep Learning Model.

Las aplicaciones prácticas del modelo fantasma son diversas y prometedoras, abarcando campos como la visión

por computador, el procesamiento de lenguaje natural y el análisis de series temporales. En cada una de estas áreas, la capacidad del modelo para ajustar dinámicamente su complejidad representa una ventaja significativa sobre los enfoques tradicionales. Esta adaptabilidad no solo mejora la precisión de las predicciones, sino que también optimiza el proceso de entrenamiento, reduciendo el tiempo necesario para alcanzar resultados satisfactorios y permitiendo una utilización más eficiente de los recursos computacionales disponibles[10].

II. MATERIALES Y MÉTODOS

A. Dataset

El estudio se basó en la recopilación de un conjunto de datos destinado al diagnóstico de pacientes con diabetes mellitus tipo-2, obtenidos de varios centros de salud privados en Guayaquil, Ecuador. Este dataset, compuesto por 2768 registros de pacientes, incluye variables clave como el número de embarazos (pregnant_times), niveles de glucosa (glucose), presión arterial (blood_pressure), grosor del pliegue cutáneo (tst), niveles de insulina (insulin), índice de masa corporal (bmi), función de pedigrí de diabetes (dpf), edad (age) y una etiqueta que indica si el paciente es diabético (is_diabetic)

Estas variables fueron seleccionadas por su relevancia clínica y su capacidad para predecir el riesgo de diabetes, lo que permite un análisis exhaustivo de los factores asociados a esta enfermedad. La Tabla. I proporciona una descripción detallada de las unidades y los intervalos de las características de riesgo presentes en el conjunto de datos.

TABLA I
CARACTERÍSTICAS Y ETIQUETA

Variables	Descripción	Unidad
Embarazos	numero de veces de embarazo	-
Sexo	sexo M(1),F(0)	-
Glucosa	concentración de la glucosa	ml/dl
PresionSanguinea	presion arterial diastólica	mm.Hg
PliegueCutaneo	grosor pliegue cutaneo triceps	mm
Insulina	insulina sérica a 2-horas	μU/ml
IndiceDeMasaCorporal	índice de masa corporal	kg/m^2
PedigriDiabetesFuncion	función pedigrí diabetes	-
Edad	edad de la persona en años	-
Etiqueta		
DiabetesResultado	Sí(1),No(0)	-

Esta información es crucial para comprender la distribución y el rango de valores de cada variable, lo que facilita la interpretación de los resultados y la aplicación de modelos predictivos. Además, el análisis de estos datos puede contribuir a identificar patrones y tendencias en la población estudiada[11], lo que podría ser útil para mejorar las estrategias de prevención y tratamiento de la diabetes mellitus tipo-2 en la región. La disponibilidad de este dataset representa una valiosa herramienta para la investigación médica y el desarrollo de soluciones basadas en datos en el ámbito de la salud.

B. Técnicas de Aprendizaje Profundo

Para la implementación de las técnicas de aprendizaje profundo, tanto con redes neuronales artificiales (RNA) convencionales como con modelos basados en Ghost Deep Learning, se utilizará Python como lenguaje de programación principal[12]. Este entorno se ejecutará sobre una máquina virtual de Google llamada Colab, la cual está preconfigurada con todas las librerías necesarias para el desarrollo de modelos de machine learning y deep learning, como TensorFlow, Keras, PyTorch y Scikit-learn[13]. Estas herramientas permiten la creación y entrenamiento de RNA convencionales, así como la implementación de arquitecturas más avanzadas, como los modelos Ghost Deep Learning, que se caracterizan por su eficiencia computacional y su capacidad para reducir el costo de recursos sin comprometer el rendimiento predictivo.

Además, para garantizar un flujo de trabajo óptimo, es fundamental contar con una conexión a internet estable y de amplio ancho de banda, ya que Colab opera en la nube y requiere una interacción constante con sus servidores. Esto es especialmente importante cuando se trabaja con modelos de deep learning, como los basados en Ghost Deep Learning, que pueden involucrar grandes volúmenes de datos y procesos de entrenamiento intensivos en recursos[14]. La combinación de estas tecnologías y herramientas proporciona un entorno robusto y flexible para explorar y comparar el desempeño de las RNA convencionales con enfoques más innovadores, como los modelos Ghost Deep Learning, en tareas de clasificación o regresión dentro del ámbito del aprendizaje supervisado[15].

C. Estratificación

La estratificación es una técnica fundamental en estadística y análisis de datos que permite organizar una población o conjunto de datos en subgrupos homogéneos, conocidos como estratos, basándose en características o variables específicas. Estos estratos se construyen de manera que los elementos dentro de cada grupo compartan similitudes, mientras que existan diferencias significativas entre los distintos estratos. El principal objetivo de esta técnica es mejorar la precisión de los análisis, asegurar que los subgrupos relevantes estén adecuadamente representados y minimizar el sesgo en los resultados[16]. Por ejemplo, en estudios de mercado, la estratificación puede utilizarse para dividir a los consumidores según su edad, ingresos o preferencias, lo que permite un análisis más detallado y específico de cada segmento.

En el ámbito del machine learning, la estratificación juega un papel crucial durante la preparación de los datos, especialmente al dividir el dataset en conjuntos de entrenamiento y prueba. Esta división estratificada garantiza que ambos conjuntos mantengan una distribución proporcional de las clases o categorías de la variable objetivo, lo que es esencial para evitar desequilibrios que puedan afectar el rendimiento del modelo[17]. Este enfoque es particularmente importante en problemas de clasificación con clases desbalanceadas, donde una distribución desigual podría llevar a un modelo sesgado hacia la clase mayoritaria. Por ejemplo,

en un modelo predictivo para diagnosticar una enfermedad rara, la estratificación asegura que los casos positivos y negativos estén representados de manera equitativa en los conjuntos de entrenamiento y prueba, mejorando así la capacidad del modelo para generalizar y realizar predicciones precisas.

D. Marco Teórico

Modelo de Aprendizaje Profundo Convencional

Las redes neuronales artificiales (RNA) convencionales son modelos computacionales inspirados en el funcionamiento del cerebro humano, diseñados para aprender patrones y realizar tareas complejas a partir de datos. Están compuestas por capas interconectadas de nodos, llamadas neuronas artificiales, que se organizan en una estructura jerárquica. Estas capas incluyen una capa de entrada, una o varias capas ocultas y una capa de salida. La capa de entrada recibe los datos iniciales, como características de un dataset, mientras que las capas ocultas procesan esta información mediante operaciones matemáticas. Finalmente, la capa de salida produce el resultado del modelo, como una predicción o clasificación según la Fig 4. Cada neurona recibe señales de entrada, las combina con pesos ajustables, aplica una función de activación y transmite la salida a las neuronas de la siguiente capa[18].

El funcionamiento de una RNA convencional se basa en dos procesos principales: el feedforward (propagación hacia adelante) y el backpropagation (retropropagación). Durante el feedforward, los datos se mueven desde la capa de entrada a la capa de salida, generando una predicción. Luego, en la fase de backpropagation, el modelo compara su predicción con el valor real y ajusta los pesos de las conexiones entre neuronas para minimizar el error, utilizando algoritmos de optimización como el descenso de gradiente. Este proceso iterativo permite que la red "aprenda" a mejorar su precisión con el tiempo[19].

Las RNA convencionales son ampliamente utilizadas en tareas como reconocimiento de imágenes, procesamiento de lenguaje natural y predicción de series temporales, gracias a su capacidad para modelar relaciones no lineales y adaptarse a una variedad de problemas[20].

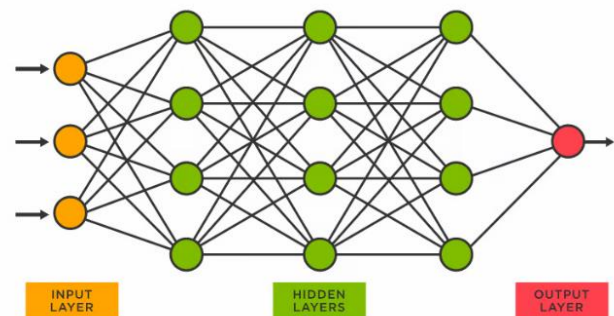


Fig. 4 Red Neuronal Artificial.

Modelo de Aprendizaje Profundo Fantasma

Las redes neuronales basadas en el modelo de aprendizaje profundo fantasma *Ghost Deep Learning Model* son una evolución de las redes neuronales convencionales, diseñadas para mejorar la eficiencia computacional y reducir el costo de recursos sin sacrificar el rendimiento predictivo[21]. Estas redes utilizan un enfoque innovador que consiste en generar "neuronas fantasmas" (ghost neurons) a partir de un conjunto reducido de neuronas originales, en lugar de crear todas las neuronas de manera independiente. Esto se logra mediante transformaciones lineales o no lineales aplicadas a las neuronas existentes, lo que permite expandir la capacidad de representación de la red sin aumentar significativamente el número de parámetros[22].

Este proceso es especialmente útil en tareas que requieren modelos profundos, pero donde los recursos computacionales son limitados. El funcionamiento técnico de estos modelos comienza con la creación de un conjunto base de neuronas en cada capa, que luego se transforman para generar las neuronas fantasmas. Estas transformaciones pueden incluir operaciones como desplazamientos, rotaciones o combinaciones lineales, que permiten diversificar la información capturada por la red[23]. Por ejemplo, en una capa convolucional, un filtro base puede ser transformado para producir múltiples filtros fantasma, ampliando así la capacidad de extracción de características sin incrementar el costo computacional. Este enfoque no solo reduce el número de parámetros para entrenamiento, sino que también disminuye el tiempo de entrenamiento y el consumo de memoria, lo que hace que estos modelos sean ideales para aplicaciones en dispositivos con recursos limitados, como móviles o sistemas embebidos[24].

Finalmente, durante el entrenamiento, las redes basadas en *Ghost Deep Learning Model* utilizan algoritmos de optimización similares a los de las redes neuronales tradicionales, como el descenso de gradiente, pero con una estructura más eficiente[25]. La retropropagación se aplica tanto a las neuronas originales como a las transformaciones que generan las neuronas fantasmas, lo que permite ajustar los pesos de manera óptima. Esto resulta en modelos que mantienen un alto rendimiento en tareas como clasificación de imágenes, detección de objetos o procesamiento de lenguaje natural, pero con una huella computacional significativamente menor. Esta combinación de eficiencia y precisión ha posicionado a los modelos *Ghost Deep Learning* como una alternativa prometedora en el campo del aprendizaje profundo[26].

Diferencias de los Modelos

Una de las principales diferencias entre un modelo de aprendizaje profundo fantasma y un modelo de aprendizaje profundo tradicional radica en su arquitectura y eficiencia computacional. Los modelos tradicionales suelen basarse en redes neuronales densas y complejas, con un alto número de parámetros y capas, lo que requiere un gran poder de procesamiento y tiempo de entrenamiento. En cambio, los modelos fantasmas utilizan una arquitectura optimizada que

reduce significativamente la cantidad de parámetros y operaciones necesarias, sin comprometer el rendimiento. Esto se logra mediante técnicas como la reutilización de características y la creación de "mapas fantasmas", que imitan los mapas de características de las capas tradicionales, pero con un coste computacional mucho menor. Como resultado, los modelos fantasmas son más ligeros, rápidos y adecuados para entornos con recursos limitados, como dispositivos móviles o sistemas embebidos[27].

Otra diferencia clave es la capacidad de los modelos fantasma para manejar conjuntos de datos desequilibrados y adaptarse a diferentes contextos. Mientras que los modelos tradicionales pueden requerir un gran volumen de datos balanceados para lograr un buen rendimiento, los modelos fantasmas están diseñados para extraer información relevante incluso con datos escasos o desequilibrados, lo que los hace ideales para aplicaciones en el ámbito médico, donde los datos suelen ser heterogéneos y de difícil obtención[28]. Además, los modelos fantasmas suelen mostrar una mayor generalización y menor sobreajuste (overfitting) en comparación con los modelos tradicionales, lo que se traduce en una mayor precisión y estabilidad en las predicciones. Estas ventajas hacen que los modelos fantasmas sean una alternativa atractiva para tareas complejas, como la predicción de enfermedades, donde la eficiencia y la precisión son críticas[29].

Métricas de Evaluación

Para los modelos basados en RNA, las métricas de evaluación son fundamentales para medir el rendimiento y la capacidad de generalización del modelo. Una de las métricas más utilizadas es la Curva ROC (Receiver Operating Characteristic), que evalúa la capacidad del modelo para distinguir entre clases, especialmente en problemas binarios.

Esta curva representa la tasa de verdaderos positivos (sensibilidad) frente a la tasa de falsos positivos (1 - especificidad) para diferentes umbrales de clasificación. El área bajo la curva ROC (AUC) es un valor clave que resume el rendimiento del modelo[30]: un AUC cercano a 1 indica un modelo con alta capacidad de discriminación, mientras que un AUC cercano a 0.5 sugiere un rendimiento similar al azar según la Fig. 5.

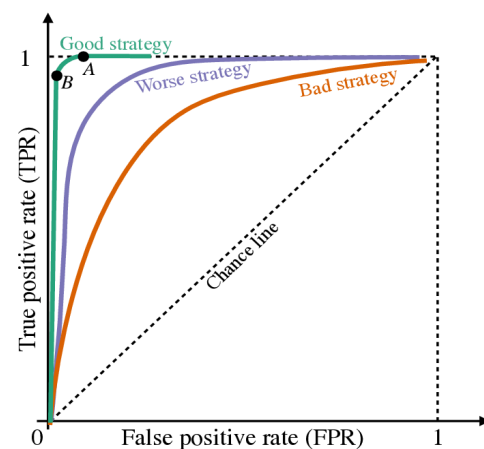


Fig. 5 Curva ROC.

La Curva ROC es especialmente útil cuando las clases están desbalanceadas, ya que proporciona una visión más completa del rendimiento del modelo que la precisión simple. Otras métricas importantes son las curvas históricas de *accuracy* (precisión) y *loss* (pérdida), que se generan durante el entrenamiento del modelo. La curva de *accuracy* muestra la evolución de la precisión del modelo en los conjuntos de entrenamiento y validación a lo largo de las épocas, permitiendo identificar si el modelo está aprendiendo correctamente o si existe sobreajuste (*overfitting*)[31]. Por otro lado, la curva de *loss* representa la reducción de la función de pérdida (como la entropía cruzada) durante el entrenamiento, lo que indica cómo el modelo está minimizando el error en sus predicciones. Una curva de *loss* que disminuye de manera estable en el conjunto de entrenamiento, pero se estanca o aumenta en el conjunto de validación es un indicador claro de sobreajuste. Estas curvas son herramientas esenciales para ajustar hiperparámetros, como la tasa de aprendizaje o el número de épocas, y para garantizar que el modelo generalice bien a datos no vistos[32].

III. RESULTADOS

En este trabajo se aplicó un método de investigación experimental con un enfoque cuantitativo para su desarrollo y mediante aprendizaje profundo se crearon 2 modelos de redes neuronales artificiales donde se aplicaron los siguientes pasos:

- 1) Tratamiento de la Data
- 2) Creación del Modelo
- 3) Entrenamiento de los Modelos
- 4) Pruebas de Funcionabilidad
- 5) Evaluación de los Modelos

Paso 1)

Se tomó el dataset que contiene 2768 registros de información de pacientes repartidos en clases, hay 1816 pacientes con diabetes(1) y 952 pacientes sin diabetes(0), se procedió a usar un modelo con las nueve columnas de entrada sin ningún tipo de tratamiento y en otros modelos se tomaron las 6 columnas de entrada más significativas o relevantes las cuales fueron estandarizadas y escaladas, en todas las muestras se tomó el 80% para entrenamiento(train) y el 20% para prueba(test), además se tomó una muestra estratificada y otra no estratificada para realizar el comparativo.

Pasos 2) al 5)

En esta fase se exponen las etapas que van desde la creación del modelo, entrenamiento, prueba, evaluación y aplicación de estos, por lo que destaca de todos los modelos la red neuronal artificial.

Modelo de Aprendizaje Profundo Convencional: Para este modelo, se usó el dataset original que tiene las 9 columnas de características de entrada de donde se tomó el 80% para train y

20% para test, sin aplicarse ningún proceso de estandarización y/o normalización y sin haberse aplicado la estratificación de los datos según se aprecia en la Fig.6. El *accuracy* obtenido fue del 87.57% con una pérdida(*loss*) del 0.3089

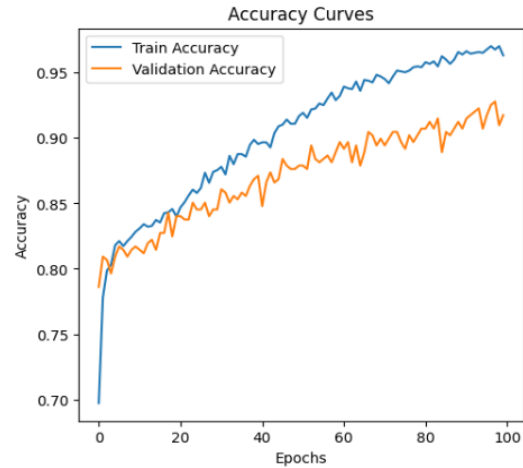


Fig. 6 Curvas Accuracy

La separación de las dos curvas de *accuracy* a partir de la época 20, donde una tiende hacia arriba y la otra hacia abajo, sugiere un fenómeno común en el entrenamiento de modelos de aprendizaje profundo: sobreajuste (*overfitting*). La curva que aumenta(curva azul) representa el *accuracy* en el conjunto de entrenamiento, lo que indica que el modelo sigue mejorando su capacidad para predecir correctamente los datos con los que está siendo entrenado. Sin embargo, la curva que disminuye(curva naranja), que corresponde al *accuracy* en el conjunto de validación, muestra que el modelo comienza a perder generalización.

Esto significa que, aunque el modelo se vuelve más preciso con los datos de entrenamiento, su capacidad para predecir correctamente datos nuevos o no vistos se reduce, lo que es un signo claro de sobreajuste. Este comportamiento es una señal de que el modelo está aprendiendo patrones específicos del conjunto de entrenamiento que no son generalizables a otros datos.

Según la Fig. 7 las curvas de pérdida (*loss*) muestran el comportamiento del modelo de aprendizaje profundo convencional a lo largo de 100 épocas, diferenciando entre la pérdida de entrenamiento (Total Loss) y la pérdida de validación (Validation Loss). Ambas curvas presentan una tendencia descendente, lo que indica que el modelo está aprendiendo de manera efectiva y reduciendo el error en sus predicciones a medida que avanza el entrenamiento.

En las primeras épocas, se observa una disminución rápida en ambas curvas, lo que refleja que el modelo está capturando los patrones principales de los datos y ajustando sus parámetros para minimizar la función de pérdida. Este comportamiento inicial es esperado y deseable en un proceso de entrenamiento bien configurado.

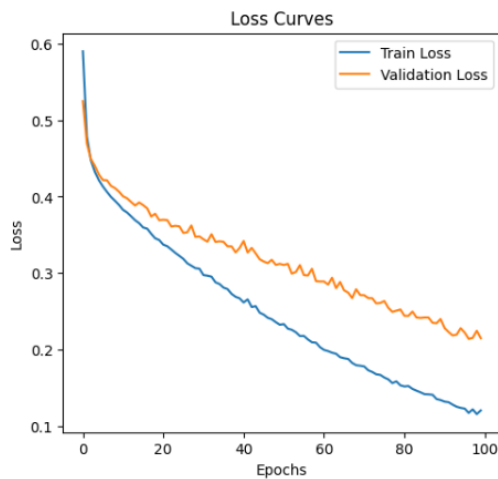


Fig. 7 Curvas de Loss

Sin embargo, a partir de la época 20, se nota una ligera divergencia entre las dos curvas: mientras que la pérdida de entrenamiento continúa descendiendo de manera más pronunciada, la pérdida de validación disminuye a un ritmo más lento e incluso se estabiliza. Esta divergencia puede ser un indicio de sobreajuste (overfitting), donde el modelo se especializa demasiado en los datos de entrenamiento, perdiendo capacidad para generalizar a datos no vistos. Para mitigar este problema, se podrían aplicar técnicas como la regularización, el dropout o la parada temprana (early stopping), que ayudarían a equilibrar el aprendizaje y mejorar la generalización del modelo. En general, la tendencia descendente de ambas curvas es positiva, pero es crucial monitorear y ajustar el entrenamiento para evitar el sobreajuste y garantizar un rendimiento óptimo en datos reales.

La curva ROC según la Fig. 8, con un área bajo la curva (AUC) del 97% (0.97) indica que el modelo tiene un excelente desempeño en la clasificación binaria, es decir, es muy bueno distinguiendo entre las dos clases (por ejemplo, positivo y negativo).

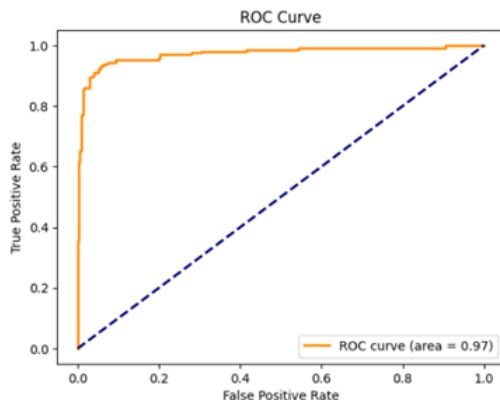


Fig. 8 Curva ROC

Modelo de Aprendizaje Profundo Fantasma: Se seleccionaron de forma automática, las 6 columnas más relevantes ['Embarazos', 'Glucosa', 'Insulina', 'IndiceDMC', 'Pedigri', 'Edad'] y que son las que aportan de forma significativa al proceso de aprendizaje, las cuales fueron usadas como características de entrada, del total de registros se tomó el 80% para train y el 20% para test, los datos fueron estandarizados y se aplicó el proceso de estratificación de los datos.

En este modelo también se aplicaron técnicas de regularización, dropout y la parada temprana (early stopping), que ayudaron a equilibrar el aprendizaje y mejorar la generalización del modelo. A partir de la época 20, el comportamiento de las dos curvas es hacia arriba y se mantienen tanto en entrenamiento como en validación evidenciando que el sobre ajuste se eliminó. Según la Fig. 9 la curva azul que aumenta representa el accuracy en el conjunto de entrenamiento, lo que indica que el modelo sigue mejorando su capacidad para predecir correctamente los datos con los que está siendo entrenado; durante el entrenamiento se produjo un stop learning en la época 92. En cuanto la curva naranja, que corresponde al accuracy en el conjunto de validación, muestra que el modelo es capaz de generalizar adecuadamente. Esto significa que, mientras el modelo se vuelve más preciso con los datos de entrenamiento, su capacidad para predecir correctamente datos nuevos o no vistos se cumple. Este comportamiento es una señal de que el modelo está aprendiendo patrones específicos del conjunto de entrenamiento que son generalizables a otros datos.

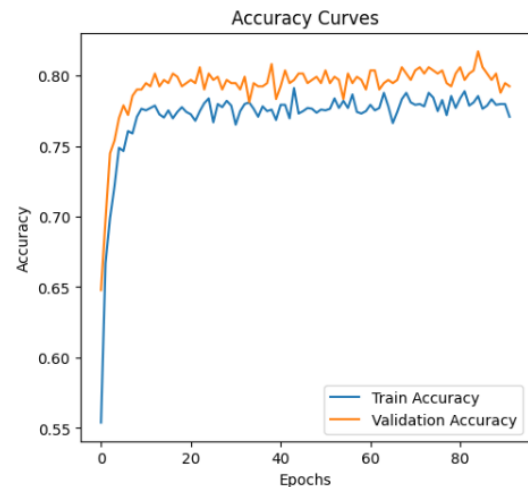
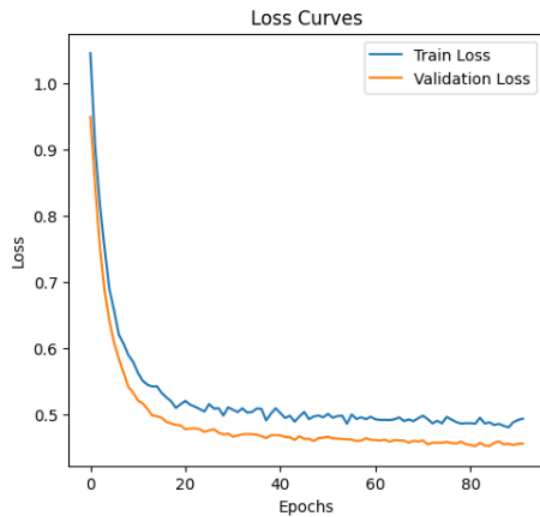


Fig. 9 Accuracy Ghost Deep Learning Model

Según la Fig. 10 de curvas de pérdida (loss) muestra el comportamiento del modelo de aprendizaje profundo fantasma a lo largo de 100 épocas, diferenciando entre la pérdida de entrenamiento (Total Loss) y la pérdida de validación (Validation Loss). Ambas curvas presentan una tendencia descendente, lo que indica que el modelo está aprendiendo de manera efectiva y reduciendo el error en sus predicciones a medida que avanza el entrenamiento.



. Fig. 10 Loss Ghost Deep Learning Model

En las primeras épocas, se observa una disminución rápida en ambas curvas, lo que refleja que el modelo está capturando los patrones principales de los datos y ajustando sus parámetros para minimizar la función de pérdida. Este comportamiento inicial es esperado y deseable en un proceso de entrenamiento bien configurado. Sin embargo, a partir de la época 20, se nota una ligera divergencia entre las dos curvas: mientras que la pérdida de entrenamiento continúa descendiendo de manera más pronunciada, la pérdida de validación disminuye al mismo ritmo incluso se estabiliza. En general, la tendencia descendente de ambas curvas es positiva, puesto que después de haberse ajustado el entrenamiento se evitó el sobreajuste y se garantizó un rendimiento óptimo en datos reales.

La curva ROC con un área bajo la curva (AUC) del 85% (0.85) según la Fig. 11, en esta se indica que el modelo tiene un excelente desempeño en la clasificación binaria, es decir, es muy bueno distinguiendo entre las dos clases indicando que el modelo tiene un rendimiento excepcional en la tarea de clasificación binaria.

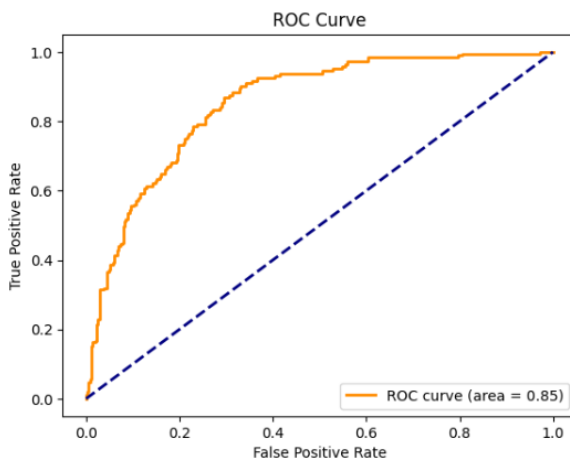


Fig. 11 Curva ROC Modelo Fantasma

IV. DISCUSIÓN Y CONCLUSIONES

El modelo de aprendizaje profundo fantasma presentado en el artículo destaca por su capacidad para optimizar la precisión y eficiencia en el diagnóstico temprano de la diabetes, un desafío crítico en el ámbito de la salud. Su arquitectura ligera y eficiente no solo reduce los costes computacionales, sino que también mantiene un alto rendimiento predictivo, lo que lo diferencia de los métodos convencionales. Esta combinación de eficiencia y precisión es particularmente relevante en entornos médicos, donde los recursos pueden ser limitados y la necesidad de diagnósticos rápidos y precisos es primordial. Además, su habilidad para manejar conjuntos de datos desequilibrados, una característica común en los datos clínicos, lo posiciona como una herramienta versátil y adaptable a diversos contextos, lo que amplía su potencial de aplicación en diferentes sistemas de salud.

Los resultados experimentales demuestran que el modelo fantasma supera a los métodos tradicionales, como las redes neuronales artificiales (RNA), en términos de precisión y estabilidad. Aunque el accuracy del modelo fantasma (79.24%) es ligeramente inferior al de las RNA (88.57%), este último presenta sobreajuste, lo que limita su capacidad de generalización. Por el contrario, el modelo fantasma muestra un menor error (loss de 0.4750 frente a 0.389 de las RNA), lo que sugiere una mayor capacidad para realizar predicciones precisas en datos no vistos. Esta diferencia resalta la importancia de priorizar la generalización sobre la precisión aparente en entornos clínicos, donde la capacidad de adaptación a nuevos casos es crucial. Además, la escalabilidad y los bajos requerimientos de recursos del modelo fantasma lo convierten en una opción viable para su implementación en sistemas de salud, tanto públicos como privados.

El estudio demuestra que el modelo de aprendizaje profundo fantasma representa un avance significativo en la detección temprana de diabetes, destacando su precisión, eficiencia computacional y escalabilidad. El modelo muestra especial robustez ante conjuntos de datos desbalanceados y requiere recursos mínimos para su implementación, lo que lo hace viable incluso en entornos clínicos con limitaciones tecnológicas. Además, su capacidad de generalización y estabilidad predictiva se mantiene consistente en diversos contextos médicos, lo que sugiere su potencial para mejorar estrategias de prevención y manejo de esta enfermedad crónica.

Más allá de la diabetes, el modelo fantasma tiene el potencial de transformar la forma en que se abordan otras enfermedades crónicas, proporcionando una herramienta accesible y eficaz para la prevención y el diagnóstico temprano. Su arquitectura ligera y su adaptabilidad lo hacen ideal para su implementación en entornos médicos diversos, lo que podría impulsar la adopción de la inteligencia artificial en la salud pública y privada. Este enfoque no solo mejora los resultados clínicos, sino que también contribuye a reducir los

costes asociados con el manejo de enfermedades crónicas, marcando un hito en la aplicación de la inteligencia artificial en el sector de la salud.

REFERENCES

- [1] M. E. Baldeón *et al.*, "Prevalence of metabolic syndrome and diabetes mellitus type-2 and their association with intake of dairy and legume in Andean communities of Ecuador," *PLoS One*, vol. 16, no. 7 July, 2021.
- [2] A. M. López Vaesken, A. B. Rodríguez Tercero, and P. C. Velázquez Comelli, "Conocimientos de diabetes y alimentación y control glucémico en pacientes diabéticos de un hospital de Asunción," *Rev. científica ciencias la salud*, vol. 3, no. 1, 2021.
- [3] J. Vlaški and I. Vorgučin, "Diabetes mellitus type 2 in children and adolescents," *Paediatr. Croat. Suppl.*, vol. 63, 2019.
- [4] R. N. Fatimah, "Diabetes Melitus Tipe 2," *Fak. Kedokt. Univ. Lampung*, vol. 4, 2015.
- [5] Y. C. Cujilan *et al.*, "Supervised Learning Techniques for the Optimization of Diagnosis Processes of Diabetes in Public Health Centers," in *Proceedings of the LACCEI international Multi-conference for Engineering, Education and Technology*, 2022, vol. 2022-July.
- [6] M. Zambrano-Vizueté *et al.*, "Segmentation of Medical Image Using Novel Dilated Ghost Deep Learning Model," *Comput. Intell. Neurosci.*, vol. 2022, 2022.
- [7] Y. Li, S. Bai, Y. Zhou, C. Xie, Z. Zhang, and A. Yuille, "Learning transferable adversarial examples via ghost networks," in *AAAI 2020 - 34th AAAI Conference on Artificial Intelligence*, 2020.
- [8] S. Liu, S. Huang, H. Cheng, J. Shen, and S. Chen, "A deep residual network with spatial depthwise convolution for large-scale point cloud semantic segmentation," *J. Image Graph.*, vol. 26, no. 12, 2021.
- [9] S. Wiedemann, K. R. Muller, and W. Samek, "Compact and Computationally Efficient Representation of Deep Neural Networks," *IEEE Trans. Neural Networks Learn. Syst.*, vol. 31, no. 3, 2020.
- [10] T. A. Davis, M. Aznavah, and S. Kolodziej, "Write quick, run fast: Sparse deep neural network in 20 minutes of development time via SuiteSparse: GraphBLAS," in *2019 IEEE High Performance Extreme Computing Conference, HPEC 2019*, 2019.
- [11] D. Patiño-Pérez *et al.*, "Stratification for the Improvement of the Performance of an ANN in Diabetes Detection," in *Proceedings of the LACCEI international Multi-conference for Engineering, Education and Technology*, 2023, vol. 2023-July.
- [12] A. Abraham *et al.*, "Scikit-learn: Machine Learning in Python," *J. Mach. Learn. Res.*, vol. 12, 2011.
- [13] S. Raschka and V. Mirjalili, *Python Machine Learning: Machine Learning & Deep Learning with Python, Scikit-Learn and TensorFlow 2, Third Edition*, no. January 2010. 2019.
- [14] C. Russo, H. Ramón, N. Alonso, B. Cicerchia, L. Esnaola, and J. P. Tessore, "Tratamiento Masivo de Datos Utilizando Técnicas de Machine Learning," *XVIII Work. Investig. en Ciencias la Comput.*, 2016.
- [15] M. M. Calderón *et al.*, "Application of Machine Learning for the Prediction of Covid19 through Classification Techniques and Supervised Learning," in *Proceedings of the LACCEI international Multi-conference for Engineering, Education and Technology*, 2022, vol. 2022-July.
- [16] Q. A. Hathaway *et al.*, "Machine-learning to stratify diabetic patients using novel cardiac biomarkers and integrative genomics," *Cardiovasc. Diabetol.*, vol. 18, no. 1, 2019.
- [17] K. Nakano *et al.*, "Machine learning approach to stratify complex heterogeneity of chronic heart failure: A report from the CHART-2 study," *ESC Hear. Fail.*, vol. 10, no. 3, 2023.
- [18] A. Kurani, P. Doshi, A. Vakharia, and M. Shah, "A Comprehensive Comparative Study of Artificial Neural Network (ANN) and Support Vector Machines (SVM) on Stock Forecasting," *Annals of Data Science*, vol. 10, no. 1, 2023.
- [19] M. G. Anderson D, "Artificial Neural Networks Technology," *Kaman Sci. Corp.*, vol. 258, no. 6, 1992.
- [20] V. Dohnal, K. Kuca, and D. Jun, "What are artificial neural networks and what they can do?," *Biomed. Pap. Med. Fac. Univ. Palacky. Olomouc. Czech. Repub.*, vol. 149, no. 2, 2005.
- [21] Y. He *et al.*, "Ghost imaging based on deep learning," *Sci. Rep.*, vol. 8, no. 1, 2018.
- [22] D. Navinchandra, D. Sriram, and R. D. Logcher, "Ghost: Project Network Generator," *J. Comput. Civ. Eng.*, vol. 2, no. 3, 1988.
- [23] W. Wang and J. Wang, "Double ghost convolution attention mechanism network: A framework for hyperspectral reconstruction of a single rgb image," *Sensors (Switzerland)*, vol. 21, no. 2, 2021.
- [24] X. Zhang, H. Yin, R. Li, J. Hong, Q. Li, and P. Xue, "Ghost network analyzer," *New J. Phys.*, vol. 22, no. 1, 2020.
- [25] Q. X. He, L. Zu, and F. A. Li, "Design and optimization of winding parameters for filament-wound toroidal pressure vessels," *Yuhang Xuebao/Journal Astronaut.*, vol. 27, no. 6, 2006.
- [26] D. E. Bambauer, "Ghost in the network," *University of Pennsylvania Law Review*, vol. 162, no. 5, 2014.
- [27] X. Jia, S. Du, Y. Guo, Y. Huang, and B. Zhao, "Multi-Attention Ghost Residual Fusion Network for Image Classification," *IEEE Access*, vol. 9, 2021.
- [28] C. Liu *et al.*, "Comprehensive Graph Gradual Pruning for Sparse Training in Graph Neural Networks," *IEEE Trans. Neural Networks Learn. Syst.*, 2023.
- [29] U. Michelucci, *Applied deep learning: A case-based approach to understanding deep neural networks*. 2018.
- [30] R. Borja-Robalino, A. Monleón-Getino, and J. Rodellar, "Estandarización de métricas de rendimiento para clasificadores Machine y Deep Learning," *Rev. Ibérica Sist. y Tecnol. la Inf.*, 2022.
- [31] P. de los Santos, "Machine Learning a tu alcance: La matriz de confusión," *Think Big Empresas*, 2018. .
- [32] A. J. Hung, J. Chen, and I. S. Gill, "Automated performance metrics and machine learning algorithms to measure surgeon performance and anticipate clinical outcomes in robotic surgery," *JAMA Surgery*, vol. 153, no. 8, 2018.