

Aplicaciones de la **inteligencia artificial** en la **solución de** **problemas reales**



DrC. Yailé Caballero Mota y otros

Aplicaciones de la **inteligencia artificial** en la **solución de problemas reales**

Yailé Caballero Mota y otros



Ediciones
Universidad
de Camagüey

2024

Edición: Vilda Rodríguez Méndez y Norbisley Fernández Ramírez
Corrección: Ernesto Piñero de Laosa
Diagramación y diseño de cubierta: Adrián Eduardo Cancino Gutiérrez
Imágenes para cubierta y separadores: Inteligencia Artificial

© Yailé Caballero Mota, 2024

© Sobre la presente edición:
Ediciones Universidad de Camagüey, 2024

ISBN: 978-959-7222-33-0

Ediciones Universidad de Camagüey
Dirección de Información Científico Técnica

Universidad de Camagüey *Ignacio Agramonte Loynaz*
Carretera Circunvalación Norte km 5 ½
Camagüey, Cuba (CP 74650)
E-mail: ediciones.uc@reduc.edu.cu
Web: <https://edicionesuc.reduc.edu.cu>

Índice

Nota de la editorial / 7

Prólogo / 11

Capítulo 1. Elementos básicos y aplicaciones de la
Inteligencia Artificial / 16

Capítulo 2. Aplicaciones de la inteligencia artificial para la
solución de problemas reales en el sector de la salud / 62

Capítulo 3. Redes neuronales: herramienta imprescindible en
estudios epidemiológicos de Fasciola hepatica / 109

Capítulo 4. Aplicaciones en la Ingeniería Civil / 137

Capítulo 5. Aplicaciones en la optimización del transporte. / 182

Capítulo 6. Aplicaciones en la Educación Superior / 204

Capítulo 7. Aplicaciones en los sistemas de generación,
transmisión, distribución y uso final de la energía / 231

Índice de Acrónimos y términos especializados / 267



Nota de la Editorial

Nota de la Editorial

El presente texto es el preludio de una serie de libros sobre Inteligencia Artificial (IA), que se propone publicar Ediciones Universidad de Camagüey. Este, es el resultado de investigaciones realizadas en diferentes momentos por profesionales de universidades de Cuba y México, con predominio de la participación de la Universidad de Camagüey *Ignacio Agramonte Loynaz*. Su relevancia está avalada por proyectos nacionales y de colaboración internacional, a la vez que, por la experiencia y el prestigio de los autores, merecedores de múltiples premios y reconocimientos por estos resultados. Tales circunstancias, junto a la actualidad del tema, tuvieron mucho peso en la selección de la temática y la obra en su conjunto.

El volumen fue concebido como una compilación de siete capítulos, cada uno enfocado desde una disciplina científica diferente. Lo que conecta entre sí a todas esas perspectivas disciplinares es el uso de la IA para la solución de los problemas específicos que les atañen a las disciplinas originales, demostrándose así la universalidad de esta rama de la tecnología, que combina aspectos de la ciencia de la computación, la psicología cognitiva, la matemática y la lógica, para desarrollar sistemas que puedan realizar tareas facilitadoras a la inteligencia humana del aprendizaje, el razonamiento, la percepción, la comprensión del lenguaje y la interacción social, lo que la convierte en un campo multidisciplinar, una herramienta cada vez más presente y útil en la vida de hoy.

Los trabajos que componen el libro se basan en investigaciones disponibles en el ecosistema digital de la publicación científica, ya sea publicadas en diferentes revistas científicas indexadas o alojadas en repositorios. No

obstante, consideramos que eso no desvaloriza la novedad de los resultados que estamos presentando, pues en todos los casos, se han renovado y atemperado a la publicación actual. El capítulo 1, sin embargo, no está enfocado a la solución de problemas reales específicos, como sucede con el resto. En él se realiza un recorrido teórico general por las diferentes aplicaciones que tiene la inteligencia artificial. Tal perspectiva general e informativa, lo convierte en un capítulo introductorio, importante para la aproximación a la temática y a su terminología.

Este no es un texto dedicado solo a lectores especializados, porque los problemas reales trascienden esa estrecha franja de involucrados. Hemos preparado un producto para una base amplia de consumidores: lectores muy especializados, que podrían —si lo desearan— replicar en otros contextos los resultados que se muestran; profesionales, especialistas en diferentes ramas, que no han incursionado aún en la IA pueden motivarse, replantearse sus propias perspectivas al entrar en contacto con el área inagotable que se nos está develando y el lector común, no especializado, que podrá encontrar aquí una vía para acercarse a tan fascinante mundo y sus potencialidades.

Esa heterogeneidad de receptores plantea el dilema de echar a navegar un texto científico en un mar en que podría confundirse con un libro de divulgación. Para sortear este desafío sin bajar los estándares requeridos por la ciencia y teniendo en cuenta la esencialidad del uso de metalenguaje internacional y mayormente en inglés, se optó por utilizar —siempre que fuera posible y aceptado— los términos en español. Solo en los casos en que no se cumplía esta condición, se mantuvieron los anglicismos, haciendo la aclaración correspondiente.

Otro elemento disruptivo de la rápida y eficaz comprensión de los contenidos por lectores diversos es el uso frecuente de acrónimos en las ciencias de la computación, por lo general

referidos a palabras en inglés, con respecto tanto a términos especializados, como a algoritmos de uso común, plenamente reconocidos por esa comunidad científica específica. Para esquivar los inconvenientes de tal hábito, se incluyó un Índice de acrónimos y términos especializados, al que puede recurrirse si se necesita esclarecer alguna sigla o concepto. Esto constituye un valor añadido del texto, especialmente útil para estudiantes y personas no especializadas.

Las imágenes de cubierta y los separadores fueron contruidos con Inteligencia Artificial, como muestra de los usos artísticos que puede tener esta tecnología; otra de las aristas que también cubre y que esperamos mostrar en algún texto futuro.

Las editoras



Prólogo

Prólogo

Este libro es resultado de más de 20 años de desarrollo de aplicaciones inteligentes por la Universidad de Camagüey, Cuba, en diversas áreas del conocimiento. Decidimos mostrar algunas de estas soluciones en la salud, la producción de alimentos, la Ingeniería Civil, el transporte, la educación, así como el manejo y uso eficiente de energía. Queremos significar que se muestran aplicaciones a problemas reales; no son soluciones a problemas simulados, ni se trata solamente del estudio experimental con bases de datos creadas para la verificación de los algoritmos: uno de los principales valores del libro.

En estas investigaciones hay fuertes alianzas con instituciones nacionales como la Universidad Central de Las Villas, La Universidad Tecnológica de La Habana, la Academia de Ciencias de Cuba y varios hospitales del país. Reflejo, además, de sólidos vínculos y participación en proyectos internacionales con Bélgica, España, China, México, entre otros. Los resultados combinan la utilización de herramientas probadas universalmente con propuestas ajustadas a necesidades locales.

El presente texto pretende familiarizar al lector con conocimientos que contribuyan a su cultura informática general y le permitan comprender el potencial transdisciplinar de la IA para resolver problemas del entorno. Cada uno de los contenidos tratados va acompañado de bibliografía actualizada, que permitirá al interesado ampliar la búsqueda de información, si así lo desea.

El libro se estructura en 7 capítulos:

El Capítulo 1 Elementos básicos y aplicaciones de la inteligencia artificial introduce el concepto de inteligencia artificial y las bases de sus métodos. Se muestra —a través de ejemplos de aplicaciones— cómo la IA ha trascendido de

los ámbitos académicos y científicos para llegar a ser parte de nuestras vidas. El capítulo finaliza con algunas reflexiones sobre el uso responsable de la IA: requisito para su aplicación exitosa. El contenido de este primer capítulo es una panorámica general, que es ampliada en los capítulos siguientes donde se muestran diversas técnicas y métodos de la IA en la solución de problemas específicos de diversas ramas.

En el Capítulo 2 Aplicaciones de la inteligencia artificial para la solución de problemas reales en el sector de la salud, se solucionan diversos problemas reales relacionados con la rama de la salud, a través de métodos de inteligencia artificial. En este capítulo se describen las soluciones a los problemas haciendo uso de métodos propuestos por los autores y se realiza la validación de los resultados a través del estudio experimental en cada una de las aplicaciones: la letalidad de la Covid19, los factores de riesgo para la aparición de mediastinitis postoperatoria en cirugía cardíaca, la predicción de la progresión de la ataxia de la marcha y la predicción de la mortalidad en pacientes con enfermedad renal crónica. Los resultados alcanzados permiten ratificar conocimiento existente sobre estas enfermedades y también pudieran ofrecer a los especialistas algunas aristas novedosas sobre su letalidad.

El Capítulo 3 está dedicado a ejemplificar el empleo de la inteligencia artificial en la soberanía alimentaria. Las redes neuronales: herramienta imprescindible en los análisis epidemiológicos de Fasciola hepatica muestra el empleo de las redes neuronales en estudios epidemiológicos de las provincias de Camagüey y Las Tunas, Cuba. Se pone de manifiesto que las redes neuronales son una herramienta poderosa para predecir la prevalencia de Fasciola hepatica y con ellos poder realizar una intervención de forma temprana con el fin de evitar la propagación del parásito.

El Capítulo 4 Aplicaciones en la Ingeniería Civil aborda la utilización de técnicas de inteligencia artificial en la predicción de respuestas térmicas y estructurales en

tipologías diversas: vigas compuestas de acero y hormigón y sus conexiones, columnas de acero restringidas y muros de mampostería confinada. Además, se ilustra la integración de la experimentación, la modelación numérica, y las técnicas de inteligencia artificial y estadísticas en la generación de datos para el desarrollo de formulaciones o métodos de diseño, y se reconocen las bondades de cada técnica y las potencialidades que ofrece su utilización conjunta.

El Capítulo 5 Operaciones de optimización del transporte en la provincia de Camagüey, Cuba, frente a la Covid-19, ofrece varios ejemplos de la aplicación de la optimización al transporte como la aplicación de la investigación de operaciones a la distribución de recursos relacionados con la COVID-19 y la aplicación de metaheurísticas en el ordenamiento del transporte urbano en Camagüey. En este capítulo se ejemplifican dos aplicaciones de la optimización en el transporte urbano y de carga, tomando como caso de estudio la provincia de Camagüey.

En el Capítulo 6 Aplicaciones en la Educación Superior se aborda la utilización de la inteligencia artificial en función de la atención personalizada de la experiencia educativa con el objetivo de describir contribuciones orientadas al desarrollo de metodologías y/o algoritmos para la predicción de la deserción escolar en instituciones universitarias, y al desarrollo de los sistemas expertos para la preparación de los estudiantes universitarios en la asignatura Álgebra Lineal.

Finalmente, el Capítulo 7 Aplicaciones en los sistemas de generación, transmisión, distribución y uso final de la energía muestra la utilización de sistemas expertos para la toma de decisiones basados en técnicas de inteligencia artificial —lógica difusa, algoritmos genéticos y redes neuronales artificiales— para el mejoramiento del desempeño energético y la calidad de la energía en sistemas eléctricos de potencia e industriales y la gestión de mantenimiento predictivo y proactivo en la industria de la energía eléctrica. Estos temas se presentan

mediante ejemplos desarrollados para la solución de diferentes problemas del campo de la Ingeniería Eléctrica.

El propósito es que este libro muestre a un público general —no solo miembros de la comunidad científica de las ciencias computacionales— las bondades de la inteligencia artificial en la solución de problemas reales para el desarrollo de la sociedad: su versatilidad e irregular desarrollo en diferentes ramas de la ciencia. Y al mismo tiempo pueda ser usado también con fines docentes e investigativos.

Al sumergirse en estas páginas el lector descubrirá cómo la inteligencia artificial no solo desafía a repensar la forma en que abordamos los problemas del mundo real, sino también a reflexionar sobre el papel de los seres humanos como arquitectos del futuro en una era digital en constante evolución. Si bien la necesidad de este conocimiento alcanza todos los campos de la ciencia y de la vida, su implementación no se produce al mismo ritmo en todos los casos; los avances no han sido homogéneos, aunque todas las áreas tienen resultados.

Este texto deja la puerta abierta a preguntas como ¿en qué medida la inteligencia artificial puede contribuir a la resolución de problemas globales urgentes —el cambio climático o la pobreza— y cuáles son sus limitaciones? ¿cómo podemos asegurar que las soluciones de inteligencia artificial sean accesibles para todos y no contribuyan a aumentar la brecha digital o social? ¿qué papel juega la colaboración entre humanos y sistemas de IA en la solución efectiva de problemas complejos desde una perspectiva ética?

En una nueva era de innovación impulsada por la inteligencia artificial, este libro invita al lector a explorar cómo la colaboración entre la mente humana y la potencia de los algoritmos puede allanar el camino hacia la resolución de los desafíos más apremiantes de la sociedad: IA responsable.

Dr. C. Yailé Caballero Mota



Capítulo 1. Elementos básicos y aplicaciones de la Inteligencia Artificial

Capítulo 1. Elementos básicos y aplicaciones de la Inteligencia Artificial

Rafael Esteban Bello Pérez

María Matilde García Lorenzo

1. Inteligencia Artificial. Elementos básicos

El término “Inteligencia Artificial” (IA) fue acuñado por el científico estadounidense John McCarthy en la década de 1950 (Moor, 2006). Durante muchos años, esta disciplina de la informática estuvo principalmente relacionada con la comunidad científica y la ciencia ficción. Sin embargo, en tiempos recientes, se ha vuelto tan popular que prácticamente cualquier persona puede estar familiarizada con el término y discutirlo. En la actualidad, los artículos sobre IA no se limitan a revistas y sitios científicos. Periódicos, revistas, portales web, blogs y otros medios de comunicación populares publican sistemáticamente artículos y noticias sobre esta disciplina.

La victoria de una computadora en el juego de GO (Chouard, 2016), el desarrollo de nuevos antibióticos mediante técnicas de IA (Trafton, 2020) y, más recientemente, la aparición de ChatGPT (Fernández, 2023) —que ha sido utilizado, evaluado y analizado desde diversas perspectivas— son solo algunos ejemplos de cómo la IA está presente en nuestra vida cotidiana. Además, numerosos sistemas y dispositivos que utilizamos en nuestra rutina diaria están respaldados por sistemas inteligentes, lo que ha llevado a que algunos hablen de la “IA invisible”. Otros incluso sugieren que la IA es la “electricidad del siglo xxi” (González, 2021), estableciendo un paralelo entre la revolución que representó el surgimiento de la electricidad en el pasado y el impacto que la IA ya tiene en nuestras vidas en el siglo actual.

¿Pero qué ha posibilitado este cambio sustancial en el desarrollo de la IA? El auge actual de esa tecnología avanzada de procesamiento de datos es el resultado de diversos factores. En primer lugar, se debe a la informatización de la sociedad, lo que ha permitido la generación y almacenamiento masivo de información de diversa procedencia y tipo en prácticamente todos los campos de actividad humana. Por otro lado, el desarrollo de la infraestructura informática, que posibilita la generación, captura, almacenamiento y procesamiento de esta información, así como la implementación de sistemas de análisis de información y algoritmos de alta complejidad. Finalmente, el avance de nuevos y potentes métodos en el campo de la IA, que han mejorado considerablemente su eficiencia y eficacia.

Aunque muchos consideran que estamos en las etapas iniciales de ese progreso, lo que plantea preguntas fundamentales, como hasta dónde llegará el hombre en el desarrollo de la llamada IA fuerte o general (Montes y Goertzel, 2019) o si alcanzará la singularidad tecnológica (Braga y Logan, 2019); es decir, el momento de la historia en el que el desarrollo tecnológico llegará a un punto sin precedentes, en el que las máquinas igualarán y superarán a la inteligencia humana.

El término “IA” se ha definido de diversas maneras. Una de las definiciones más conocidas se basa en el Test de Turing (Boden, 2017; Matthew, 2023), que establece que existirá Inteligencia Artificial cuando no pueda distinguirse entre un ser humano y un programa de computadora en una conversación a ciegas. Cabe destacar que en los últimos cinco años solo algunos sistemas han logrado, aunque parcialmente, superar esta prueba. El ejemplo actual más conocido es ChatGPT4.

Más recientemente, una comisión de la Unión Europea presentó una definición, en la cual la IA se asocia a sistemas de software —y posiblemente también de hardware— diseñados

por humanos que, frente a objetivos complejos, operan en el ámbito físico o digital. Estos sistemas son capaces de percibir su entorno mediante la adquisición e interpretación de datos estructurados o no estructurados, razonar sobre el conocimiento, procesar información derivada de estos datos y tomar decisiones para lograr el objetivo especificado (High-Level Expert Group on Artificial Intelligence, 2018).

En la práctica —como se pretende demostrar en este capítulo— la IA tiene el propósito de desarrollar métodos de resolución para problemas que carecen de algoritmos específicos o que, debido a su alta complejidad computacional, no pueden resolverse mediante enfoques convencionales. Esto se ejemplifica en casos como el juego de GO o el diseño de antibióticos mencionados previamente. También se aplica al problema del viajero vendedor, ayudándolo a encontrar la ruta más corta y eficiente para llegar a un destino después de recorrer múltiples ciudades (Salazar, 2019). A medida que el número de ciudades aumenta, la resolución se vuelve impracticable con los recursos computacionales disponibles en la actualidad.

En la actualidad, la IA se utiliza para abordar una amplia variedad de problemas, que incluyen búsqueda, optimización, planificación, programación (*scheduling*). Un aspecto distintivo de la IA, en contraste con otras técnicas computacionales, es su capacidad para manejar la incertidumbre inherente a la resolución de problemas del mundo real. Un elemento particularmente relevante y quizás una de las principales razones de su importancia actual son los métodos de aprendizaje, que permiten extraer conocimiento de diversos dominios de aplicación contenidos en la información digital disponible en la actualidad. Se mencionan habitualmente enfoques como el aprendizaje supervisado y no supervisado, el aprendizaje profundo y el aprendizaje por refuerzo.

La IA se ha convertido en un componente esencial de los sistemas ciberfísicos (Esterle y Grosu, 2016). En este contexto, emergen áreas de investigación innovadoras, como

los vehículos autónomos, la Industria 4.0 (I4.0) y las ciudades inteligentes, que serán revisadas en este capítulo.

Si la informatización fue uno de los pilares que ha permitido el desarrollo vertiginoso de la Inteligencia Artificial, esta última se convierte en un componente esencial y trascendental en la transformación digital de la sociedad. Tal transformación implica combinar diversas tecnologías, que incluyen la IA, la computación en la nube, la Analítica de Macrodatos (*Big Data*), la Internet de las cosas (IoT), dispositivos conectados, robots y drones, así como plataformas y redes sociales. A medida que los costos de estas tecnologías avanzadas disminuyen, se abren nuevas posibilidades y oportunidades para armonizar de maneras innovadoras, dando lugar a efectos combinatorios, en los cuales la capacidad de las tecnologías que trabajan en conjunto supera las capacidades individuales de cada una cuando se aplican por separado.

El uso de técnicas de IA está presente en prácticamente todas las actividades socioeconómicas humanas (Stone, 2016). Se encuentra en el desarrollo de nuevas tecnologías de transporte, como la planificación y la creación de vehículos autónomos. También se aplica en la educación mediante el uso de sistemas tutoriales inteligentes y la analítica del aprendizaje. Además, se utiliza para mejorar los sistemas de entretenimiento en línea, con el objetivo de proporcionar productos más personalizados y atractivos, generar contenido de mayor calidad y variedad, y crear sistemas interactivos basados en la empatía y la emoción. Además, la IA tiene aplicaciones en campos como la salud, la agricultura de precisión, la gestión de la energía y la ciberseguridad. Algunos de estos resultados se analizarán en la segunda sección de este capítulo. Sin embargo, todas estas aplicaciones requieren el desarrollo de un marco ético y legal que permita trabajar de manera responsable con la IA, como se examinará en la sección 3.

2. Revisión de algunas aplicaciones de la IA

A continuación, se presentan algunas aplicaciones de las técnicas de IA en diversos campos que ilustran la amplia difusión actual de esta disciplina.

2.1. Aplicaciones en el ámbito de la salud

Indudablemente, uno de los campos con mayor aplicación de las técnicas de IA es la medicina y la salud en general; en esa rama es conocida su tradición, actualidad y proyección hacia el futuro (Stone, 2016). Entre los ejemplos de tales aplicaciones se encuentran los asistentes clínicos que brindan apoyo en el diagnóstico, los análisis de salud, que abarcan desde la medicina personalizada hasta estudios de poblaciones, la robótica médica, como el sistema Robodoc-IBM en ortopedia, el sistema Da Vinci en cirugía cardíaca, nano robots antibacterianos en hematología y robots endoscópicos en el aparato digestivo; también el cuidado de adultos mayores, que incluye dispositivos de locomoción asistida y sistemas de apoyo emocional, así como la facilitación de la comunicación con la familia.

Una de las primeras aplicaciones de la IA en el sector de la salud fue MYCIN (Shortliffe, 1976) —un sistema experto diseñado para la detección de enfermedades infecciosas en la sangre— que empleaba razonamiento lógico en su motor de inferencias. Posteriormente, con el desarrollo del razonamiento basado en casos como vía de resolución de problemas, se introdujeron aplicaciones en el ámbito médico que adoptaron esta forma de razonar. Estas aplicaciones ajustaban las soluciones a problemas en función de experiencias previas y casos similares recuperados (Choudhury & Begum (2016).

Otros ejemplos notables incluyen CASEY (Koton, 1989) para el diagnóstico de enfermedades coronarias, PROTOS (Bareiss, Porter & Wier, 1987) para la clasificación y diagnóstico

de trastornos auditivos, MacRad (Macura & Macura, 1995) para la clasificación de imágenes radiológicas, SCINA (Haddad, Moertl & Porenta, 1995) para la interpretación de la perfusión miocárdica, AUGUSTE (Marling & Whitehouse, 2001) para el cuidado de pacientes con Alzheimer, SAPRIM (Rodríguez, *et al.*, 2008) en la predicción del riesgo pediátrico y CEDS (Chakraborty *et al.*, 2015) para el diagnóstico del cólera, entre otros.

Buch, Ahmed y Maruthappu (2018) destacan el crecimiento de la IA en medicina, señalando que, en 2016, los proyectos de salud atrajeron más inversión que cualquier otro sector de la economía global. Por su parte, IBM (2018) propone cuatro modos de utilización de la IA en este campo:

1. Como anotador para datos clínicos
2. Como generador de un resumen cognitivo de los registros de pacientes
3. Como medida de similitud clínica entre pacientes, lo que permite la creación de cohortes dinámicas de pacientes
4. Como buscador de información en la literatura médica no estructurada para respaldar hipótesis

La IA promete transformar la práctica de la medicina en diversas formas, incluyendo el diagnóstico de enfermedades, el desarrollo de medicamentos, la personalización de tratamientos, la digitalización de registros médicos, la planificación de citas en línea y recordatorios para citas de seguimiento. Como se mencionó anteriormente, también se utiliza para buscar datos en la bibliografía médica no estructurada con el fin de respaldar hipótesis, lo que contribuye al hallazgo de nuevos conocimientos. Además, se aplican algoritmos de dosificación y advertencias sobre efectos adversos al prescribir combinaciones de múltiples fármacos, entre otras aplicaciones.

También se está trabajando en el desarrollo de asistentes virtuales para el primer contacto con los pacientes, utilizando tecnología de reconocimiento de voz. Esto permite iniciar

un diálogo con el paciente para recopilar información inicial, lo que a su vez libera más tiempo para que los médicos mantengan conversaciones más específicas y brinden atención personalizada.

Tales avances permiten proponer fármacos antes de llevar a cabo la síntesis química y la evaluación experimental. En el trabajo de Yu-Chen, Ren Gui, Honda y Davis (2019), se exploran aspectos relacionados con el diseño de fármacos y la utilidad del Aprendizaje Automático (AA) y profundo en la creación de nuevos medicamentos, el análisis QSAR, la reutilización de fármacos y la visualización del espacio químico. Hessler & Baringhaus (2018) destacan el uso de redes neuronales artificiales en la predicción de propiedades fisicoquímicas y resaltan su eficacia en la predicción de relaciones cuantitativas estructura-propiedad o relaciones cuantitativas estructura-actividad.

En Schneider *et al.* (2019), se plantean que la IA está presentando nuevos desafíos tanto para los científicos involucrados como para la industria biofarmacéutica, así como para los procesos establecidos de descubrimiento y desarrollo de nuevos medicamentos. Se hacen referencias a los “grandes desafíos” en el descubrimiento de fármacos de molécula pequeña con el uso de IA y se abordan enfoques para enfrentarlos.

Entre otros casos de estudio, destaca Watson for Oncology (2017), entrenado por especialistas para analizar historiales médicos y ayudar a los médicos a definir tratamientos más personalizados basados en estudios científicos y en los historiales de los pacientes. Este sistema ha sido alimentado con una vasta cantidad de información, incluyendo 300 revistas médicas, cientos de libros y más de 15 millones de textos. En el trabajo de Suo, Ma, Yuan, Huai, Zhong, Gao & Zhang (2018), se aborda el problema de la progresión de pacientes similares mediante el uso del aprendizaje de funciones de distancias.

En el contexto de la clasificación multietiqueta, la predicción simultánea del riesgo de varias enfermedades, el análisis de las diversas mutaciones del VIH en respuesta a distintos tipos de retrovirales o la categorización de las hojas clínicas electrónicas, han mostrado un avance significativo en el campo de la IA. Los resultados teóricos obtenidos en estos estudios terminan siendo utilizados en problemas de gran impacto en las aplicaciones médicas, como se evidencia en los estudios de Heider, Senge, Cheng & Hüllermeier (2013); Zufferey, Hofer, Hennebert, Schumacher, Ingold & Bromuri (2015); Riemenschneider, Senge, Neumann, Hüllermeier & Heider (2016); y Li, Liu, Lin, Zhao & Zhang (2017).

2.2. Aplicaciones en la agricultura de precisión

“La agricultura de precisión es una especie de enfoque de gestión agrícola que utiliza la tecnología de la información (TI) para garantizar que los cultivos y la superficie reciban exactamente lo que necesitan para una salud y productividad óptimas” (Fouquet, 2021, pp. 5-6).

La revolución tecnológica de los últimos tiempos ha dejado su huella en el sector agrícola de manera permanente. Se están innovando constantemente formas de recopilar información que, desde un punto de vista agronómico, resultan de gran interés para comprender la condición de las fincas y los cultivos. La implementación de la Inteligencia Artificial en la agricultura amplía la capacidad de recopilar información detallada y optimizar el uso de herramientas agrícolas. Esta combinación de factores facilita un ajuste más preciso en aspectos como la cantidad de productos fitosanitarios, fertilizantes, el momento ideal de cosecha, el manejo de plagas y las necesidades de riego.

Para Pérez *et al.* (2019), la aplicación de la IA en la agricultura se refiere a todas las habilidades que una máquina, sensor, dispositivo de monitoreo o computadora poseen para

ejecutar tareas con una precisión excepcional, al recopilar datos específicos que permiten mejorar y optimizar diversas actividades agrícolas y cultivos.

Se trata de una herramienta innovadora que se está implementando en beneficio de agricultores de todo el mundo, especialmente aquellos que se dedican a la agricultura de precisión.

Ese enfoque utiliza diversas tecnologías en el campo con el propósito de recopilar la información esencial para que los agricultores puedan tomar decisiones anticipadas. Esta tecnología, a menudo referida como “4.0”, puede ser especialmente útil para determinar qué cultivos plantar, su ubicación, el momento adecuado y, en algunos casos, incluso predecir el volumen de la cosecha.

Dado que la agricultura es una industria con una larga historia, es crucial adoptar un enfoque más inteligente y eficiente para satisfacer las crecientes demandas. La aplicación de tecnologías de Inteligencia Artificial (IA) ofrece la oportunidad de mejorar una amplia variedad de tareas relacionadas con la agricultura, como la producción de cultivos más saludables con seguimiento en tiempo real, el control de plagas y la supervisión continua del estado del suelo (Observatorio para la Innovación Silvoagropecuaria y la Cadena Agroalimentaria, 2023).

La industria agrícola ha adoptado la automatización y la precisión como parte fundamental de sus prácticas para lograr una producción más eficiente. La incorporación de tecnologías como la orientación y la detección, así como el análisis de datos, ha llevado a la agricultura basada en datos. Estas tecnologías se aplican en todas las etapas del ciclo agrícola; abarcan desde la gestión de datos hasta la cartografía de cultivos, la cosecha y la planificación, entre otras áreas. Esto ha permitido la predicción y la respuesta rápida a situaciones aparentemente impredecibles en la agricultura.

Uno de los casos más extendidos del uso de la IA en la agricultura es el análisis predictivo, que contribuye a mantener una producción agrícola eficiente. Por ejemplo, utilizando técnicas de reconocimiento de patrones, es posible aprender del comportamiento de hojas de cultivos saludables y enfermas para determinar dónde aplicar herbicidas y dónde no. Estos algoritmos también ayudan a identificar brotes de plagas o malezas. Además, mediante el procesamiento y análisis de imágenes captadas por drones o satélites en fotos aéreas, las técnicas de IA permiten evaluar el estado de los cultivos, la presencia de malezas y crear mapas que identifican áreas con diferentes niveles de humedad, lo que facilita la optimización de la cantidad de riego necesaria para cada parcela (Rambauth, 2021).

Como se aborda en la investigación de Fouquet (2021), es esencial contar con métodos efectivos para identificar, gestionar y detectar enfermedades en la agricultura con el fin de garantizar una producción sostenible y productiva. Tradicionalmente, el diagnóstico de enfermedades se ha realizado mediante inspección visual, un proceso que es tanto laborioso como subjetivo. Con el propósito de automatizar este procedimiento, se ha trabajado en el desarrollo de técnicas de visión artificial respaldadas por IA.

Ese enfoque implica la detección de enfermedades utilizando imágenes y la identificación de características distintivas en estas. No obstante, en este tipo de aplicaciones, se ha observado que se requiere una cantidad significativa de ejemplos, tanto positivos como negativos, para crear un clasificador de reconocimiento de hojas eficaz. Sin embargo, la necesidad de un conjunto de datos extenso puede dificultar la aplicación de esas técnicas en la identificación y el conteo de plagas en otros tipos de cultivos.

De acuerdo con INTEC (2023), la agricultura y la ganadería de precisión, respaldadas por la IA, son tendencias líderes

en innovación en el sector agroalimentario a nivel global. Se ha observado un aumento del 30 % en la inversión en este campo durante el año 2023. Además, la financiación de nuevas empresas o *startups* en el ámbito de las tecnologías agroalimentarias experimentó un crecimiento del 45 % en 2021 en comparación con el año anterior.

En Agronet (2021, 16 de noviembre) se muestra que, la combinación de la IoT y la IA ha posibilitado la creación de una aplicación destinada a mejorar la productividad en la agricultura, aumentar la eficiencia durante la cosecha, rastrear el origen de los productos alimentarios y reducir el consumo de combustible de las máquinas, idea que en alguna medida fue referida antes por Montoya *et al.* (2017).

Esa herramienta de vanguardia tiene como objetivos la optimización de los rendimientos agrícolas, la eficiencia en la cosecha, la trazabilidad de los productos alimentarios y la reducción del consumo de combustible de las máquinas agrícolas. Igualmente proporciona información detallada sobre diversos indicadores, como la humedad de los cultivos, el rendimiento por metro cuadrado y la calidad del grano en términos de valor nutricional; incorpora funciones de seguridad, alertando al conductor cuando una persona se acerca al vehículo. Una vez que la cosechadora inicia su trabajo, la IA toma el control. Un modelo algorítmico —entrenado con miles de imágenes de granos que se etiquetaron manualmente desde el inicio del proceso de aprendizaje supervisado— cuenta los granos que quedan en el suelo a partir de una foto tomada desde la parte trasera del vehículo. El conductor tiene la opción de aceptar o rechazar los resultados, que también se utilizan para mejorar el modelo de IA.

La aplicación rastrea las ubicaciones en el campo mediante geolocalización de los vehículos y modelización espacial, lo que mejora significativamente la trazabilidad de los productos alimentarios. Además, proporciona recomendaciones al

conductor, como ajustes de velocidad o configuraciones de la cosechadora, basadas en datos en tiempo real.

En la literatura especializada se refiere la integración de la IA y la IoT, resaltándose los notables avances de su uso con el objetivo de mejorar diversos aspectos de la agricultura: Zha (2020); Ramón (2020); Segovia, Rojas & Quishpe (2021). Se destaca el papel crucial que han tenido los sensores en esta área en los últimos años; estos se utilizan en todo tipo de equipos, coches, máquinas, etc. También, a nivel climatológico, con estaciones meteorológicas situadas en puntos clave de las explotaciones, y para indicar tanto la humedad como otros factores claves en la fertilidad del suelo.

Los sensores se han empleado para optimizar la conducción de tractores, herramientas agrícolas y equipos utilizados en aplicaciones específicas. Ellos permiten un control más preciso y eficiente de las operaciones agrícolas, lo que se traduce en un uso más efectivo de los recursos y una mayor productividad.

Asimismo, se utilizan en aplicaciones relacionadas con la climatología. Se han instalado estaciones meteorológicas en ubicaciones estratégicas dentro de las explotaciones agrícolas, junto con sensores de suelo que monitorean la humedad y otros factores decisivos para evaluar la fertilidad del suelo. Estos datos son esenciales para tomar decisiones informadas sobre la gestión agrícola, como la programación de riego y la aplicación de fertilizantes de manera más precisa.

A la aplicación de la IA en la agricultura se le reconocen tanto ventajas como desventajas. Entre las primeras, se ubican las siguientes:

- **Datos Mejorados:** Permite la recopilación y análisis de datos detallados, lo que proporciona a los agricultores una base sólida para tomar decisiones informadas en situaciones específicas. Cuanta más información esté disponible, mayor control y capacidad de toma de decisiones tendrán.

- **Sostenibilidad Ambiental y Económica:** Permite ajustar las dosis de siembra, riego y productos fitosanitarios, lo que conduce a un uso más eficiente de estos insumos. Además, no solo tiene un impacto económico positivo en las explotaciones agrícolas, sino que también reduce el impacto ambiental al minimizar el uso excesivo de recursos.

- **Seguridad Alimentaria:** La aplicación de la IA a lo largo de todos los procesos agrícolas, desde la siembra hasta la post cosecha, permite un control exhaustivo de la trazabilidad de los cultivos y su seguridad alimentaria, garantizando la calidad de los productos.

Mientras que las desventajas, se centran en los siguientes aspectos:

- **Costo Económico:** La implementación de la IA en la agricultura puede ser costosa, y en algunos casos, los costos pueden superar los beneficios económicos que genera, lo cual puede tener un impacto negativo en las finanzas de las empresas agrícolas.

- **Mantenimiento Tecnológico:** Los sistemas de IA requieren un mantenimiento regular, y en el caso de la agricultura, están expuestos a condiciones climáticas adversas. Eso significa que necesitan una atención y mantenimiento constantes, lo que puede resultar costoso y requerir recursos adicionales.

- **Requisito de Formación:** El uso efectivo de la IA en la agricultura requiere que los agricultores adquieran conocimientos básicos de tecnología, condición que puede representar un desafío para aquellos que no están familiarizados con estas herramientas y tecnologías.

2.3. Aplicaciones en la Industria 4.0

La Industria 4.0 está liderando una nueva revolución que, como se explicó antes, algunos expertos ya consideran la cuarta revolución industrial. Este movimiento implica la creación de fábricas y productos inteligentes que están

conectados directamente con los consumidores a través de Internet y las nuevas tecnologías móviles, incluyendo la red 5G. Además, se caracteriza por un nivel de automatización y robotización masiva que no tiene precedentes. Esos cambios están teniendo un impacto transformador en la economía, la sociedad y el mercado laboral, afectando a múltiples sectores de manera transversal.

Esa nueva industria está influyendo en cómo los negocios operan (Asadollahi, Couzon, Nguyen, Ouazene & Yalaoui 2020; Ibiza, s.f.); ella implica la combinación de técnicas de producción y operación con tecnologías que se pueden integrar de manera fluida en los activos, el personal y los procesos de una empresa, lo que significa obtención de datos en tiempo real y capacidad de transformar productos, cadenas de suministro y expectativas de los clientes para mejorar las operaciones comerciales y el crecimiento de los ingresos. En esta cuarta revolución, se logra una conexión directa con el cliente a través de Internet o tecnologías móviles, con un alto nivel de automatización y robotización que involucra dispositivos, maquinaria e instalaciones, todo con el propósito de crear una industria inteligente.

La aplicación de IA se centra, en este caso, en la automatización de procesos industriales, el mejoramiento de las habilidades de la fuerza laboral y la innovación en el desarrollo de productos. Con ello, se logra una mayor productividad de los trabajadores, liberándolos para tareas más creativas que aporten un valor único. Además, esta tecnología ayuda a identificar y aprovechar el potencial del capital humano, mejorando su rendimiento y rentabilidad. En conjunto, contribuye a eliminar costos innecesarios y acelera la innovación en la creación de nuevos productos.

Lo anterior ha supuesto un cambio de patrón en el funcionamiento de la industria, impulsado por una nueva forma de interacción entre hombre y máquina, que se caracteriza

por fábricas inteligentes donde humanos y sistemas ciberfísicos interactúan en la nube (Nexus Integra, s.f.). Las fábricas inteligentes absorben estructuras automatizadas e incluyen habilitadores digitales que permiten a la maquinaria comunicarse entre sí y con los sistemas de la fábrica en su totalidad, a través de una configuración de IoT. Estas competencias están cada vez más demandadas por las fábricas de todos los sectores, que buscan garantizar la competitividad de sus plantas de producción en un contexto cada vez más tecnológico. El cambio de la automatización tradicional basada en robots industriales independientes hacia los “sistemas ciberfísicos” en red ha revolucionado la forma de trabajar de las plantas de producción y ha impuesto nuevos estándares de competitividad en el mercado, trayendo consigo a su vez, una serie de beneficios para productores y consumidores, como la fabricación *Just in time* (JIT) –Justo a Tiempo–, la introducción de nuevos productos, los cambios en el consumo y la evolución del mercado laboral. Ver también Uribe (2022).

Asadollahi *et al.* (2020) postulan también que, la influencia de IA repercutirá principalmente en tres ámbitos de la industria: aumento de la productividad debido a la automatización de procesos; ahorro de tiempo y costos de producción; mayor disponibilidad de productos y de servicios mejorados de IA y su consecuente incremento de la demanda. Entre las aplicaciones que se mencionan (Uribe, 2022) están las siguientes: el mantenimiento y la reparación predictiva, lo cual se basa en usar algoritmos para anticipar posibles incidencias o errores de un componente, máquina o proceso de producción; predecir y prevenir fallos de manera temprana, que no sólo asegura una intervención inmediata con la consiguiente reducción de costes, sino también genera una mayor eficiencia y un aumento de la productividad. La denominada calidad 4.0, referida a la alineación de la gestión de calidad con las capacidades emergentes de la I4.0 y de las nuevas tecnologías, con la IA a la

vanguardia, con el objetivo de alcanzar la excelencia operativa; la calidad 4.0 permite a las industrias mejorar continuamente la calidad de su producción al mismo tiempo que recopilan datos de uso y rendimiento. El diseño generativo es una tecnología que utiliza como herramienta base algoritmos de IA y la automatización de procesos, que permiten generar simultáneamente múltiples soluciones de diseño válidas de un mismo objeto, delimitadas únicamente por parámetros como peso, fortaleza, presupuesto o materiales; los algoritmos de diseño generativo producen, de forma autónoma, numerosas alternativas de diseño para un mismo producto y en base a un conjunto de requisitos previamente definidos por el ingeniero.

Uno de los campos de aplicación más conocidos es la robótica, pues, aunque todavía quedan algunos años para alcanzar la etapa de esplendor de esta rama, la adopción de máquinas robotizadas por parte del sector industrial es ineludible. Cuando llegue el momento, la IA desempeñará un papel crucial debido a su capacidad para hacer más flexibles, aplicaciones previamente rígidas. Por ejemplo, los robots colaborativos (cobots) se han convertido en el aliado perfecto de los trabajadores en la línea de producción; son herramientas de apoyo a los operarios con la finalidad de liberarles de las tareas más repetitivas y aburridas; ellos son ya una realidad. De esta manera, las personas podrán aprovechar el tiempo del que disponen para realizar tareas que los robots no son capaces de hacer, como la supervisión y el mantenimiento de la maquinaria, labores de programación y la optimización de procesos industriales; la robótica colaborativa es el futuro de la industria.

Según Nexus Integra (s.f.), el aprendizaje automático eleva significativamente la calidad de los productos al incorporar sistemas de mantenimiento predictivo en los procesos de producción. Esto implica la sustitución de inspecciones visuales por robots o cobots, que realizan los controles de calidad de

manera mucho más precisa y eficiente; esto da lugar a la denominada *Smart Manufacturing* (manufactura inteligente), donde se evalúan los datos recopilados durante la producción y se adaptan los cambios de forma automática.

Otro término relevante que ha surgido en este contexto son las *Smart Operations* (Operaciones Inteligentes), que se basan en la conectividad y la inteligencia automatizada para proporcionar de manera efectiva un mayor entendimiento y previsibilidad a través de análisis avanzados. El objetivo principal es lograr que las operaciones se realicen de manera cada vez más autónoma (Ang, Goh, Saldivar & Li, 2017).

Se pueden definir cinco características de las operaciones inteligentes: conexión y digitalización, capacidades analíticas, datos en tiempo real, procesos inteligentes y seguridad en los datos.

Según el Fondo de Recuperación Europeo (Ibiza, s.f.), se proyecta que para el año 2030, la Inteligencia Artificial tendrá un impacto significativo en el Producto Interno Bruto (PIB) mundial, llegando al 14 %, lo que equivale a aproximadamente 15,7 billones de dólares. Debido a este impacto, la IA se ha convertido en un motor de cambio fundamental en todos los sectores, independientemente del tamaño de las empresas o industrias.

2.4. Aplicaciones en energía

La informatización en el sector energético, como en el resto, está soportada por una creciente disponibilidad de datos, la potencia de cómputo y las tecnologías digitales, además de la necesidad de mejorar los patrones de análisis y planificación (IEA, 2017). El *white paper* del World Economic Forum (2021), titulado *Harnessing artificial intelligence to accelerate the energy transition*, explora el uso de la Inteligencia Artificial para acelerar la transición energética; con el objetivo general de mejorar la eficiencia del sistema energético, el movimiento

hacia los sistemas digitales es también soportado por las aplicaciones de la IA, las cuales pueden ayudar a disminuir los costos, reducir las pérdidas de energía y acelerar la integración de las energías renovables a las redes eléctricas. Purón (2022), en el Congreso Internacional de Inteligencia Artificial celebrado en la ciudad de Alicante, España, expuso que: “La IA es el mayor catalizador de oportunidades en el sector energético y cuenta con innumerables aplicaciones: optimizar la producción y almacenamiento, mejorar posiciones para comercializar en los mercados de compra-venta o ahorrar consumo y mejorar la huella de carbono, entre otras”; a lo que se añade que aunque es difícil estimar el valor añadido de las aplicaciones de la IA, pueden ayudar a minimizar las inversiones en activos, reducir los picos en la demanda de energía o incrementar la flexibilidad en los sistemas energéticos, lo cual podría conducir a un ahorro de miles de millones de dólares.

La cantidad masiva de datos generada por la infraestructura de las tecnologías de la información y las comunicaciones (TIC) instalada, requiere desarrollar vías automatizadas para analizar los datos. Adicionalmente, el cambio hacia sistemas de potencias más complejos, activos y descentralizados, genera tareas las cuales pueden rápidamente volverse inmanejables para los operadores humanos. La IA puede usarse para predecir la demanda de potencia y el volumen de generación, optimizar el mantenimiento y uso de los activos energéticos, comprender mejores patrones de uso de energía, así como proveer una mejor estabilidad y eficiencia del sistema energético. Puede también aliviar la carga sobre los humanos ayudando y parcialmente automatizando el proceso de toma de decisiones, así como automatizar la programación y el control de la multitud de dispositivos utilizados (Antonopoulos *et al.*, 2020).

De acuerdo a Energyavm (2020), algunas de las aplicaciones que ya se están utilizando en el sector son: optimización de las

redes eléctricas; desarrollo de redes inteligentes, las cuales transportan electricidad, pero también datos; y en el diseño, planeación y control de la movilidad basada en vehículos eléctricos.

El uso de la IA y el AA en el modelado del consumo de energía de los edificios (Platon, Dehkordi y Martel, 2015; Amasyali y El-Gohary, 2018; Ghoshchi, Zahedi, Pour y Ahmadi, 2022) puede tener también una incidencia relevante. El consumo de energía de las edificaciones es una tarea desafiante ya que muchos factores, como las propiedades físicas del edificio, las condiciones del tiempo, el equipamiento existente en los interiores de estas, y el comportamiento energético de los ocupantes, es difícil de predecir. Antonopoulos *et al.* (2020) hacen una revisión de aplicaciones de la IA al subcampo de respuesta a la demanda.

Otra implementación de la IA está relacionada con las energías renovables, las cuales han ganado un gran interés en los últimos años (Zahedi, Ahmadi, Eskandarpanah & Akbari, 2022). Actualmente, se está produciendo una transición global, que involucra a muchos países en todo el mundo, hacia un sistema energético sostenible. La mayoría de los países no están reduciendo sus emisiones lo suficientemente rápido, por lo que se vuelve esencial la integración de las energías renovables con las redes inteligentes (Martin, 2022).

Mediante el uso de la Inteligencia Artificial, es posible anticipar la disminución de la producción eléctrica en relación con la demanda de consumo para mantener un equilibrio en la disponibilidad de energía en el sistema, evitando así posibles cortes o déficits de suministro eléctrico en áreas específicas donde las energías renovables por sí solas no pueden cubrirlo por completo. Además, el análisis de datos y el aprovechamiento de los macrodatos permiten prever, con hasta 36 horas de anticipación, las regiones geográficas donde se producirá una reducción en la generación de energías renovables y cuándo se esperan aumentos en la demanda eléctrica.

La IA jugará un papel vital en el rediseño de los sistemas energéticos en el futuro, y en las cadenas de transmisión de la energía, a partir de una predicción y planificación más precisa de las fuentes de energía renovables no programables, la optimización y operación de las redes, la gestión de la demanda, y el empleo de sistemas inteligentes para la detección de fallas (Quest, Cauz, Heymann, Rod, Perret, Ballif, Virtuani & Wyrsh, 2022). En este contexto, la generación de fuentes fotovoltaicas es una de las áreas más promisorias.

El pronóstico de generación de energía fotovoltaica, su modelado y optimización (Das *et al.*, 2018) puede mejorar significativamente la potencia de generación, prediciendo las diversas variables que influyen y ayudando a los diseñadores a mejorar la configuración o ubicación del sitio de instalación. Otro tipo de energía renovable que puede modelarse usando las técnicas de IA es la energía eólica, para la cual se han llevado a cabo estudios para pronosticar y optimizar la generación de este tipo de energía (Lei *et al.*, 2009). También la IA podría ser útil en la gestión y suministro de la electricidad por fuentes renovables para los vehículos eléctricos (Entezari, Aslani, Zahedi & Noorollahi, 2023).

Dado el relevante papel que tendrán las energías renovables para atenuar el cambio climático, estas cada vez abarcan una parte más importante de los sistemas energéticos, y su desarrollo marca muchos de los retos de este sector; por eso, el empleo de las técnicas de IA despierta mucho interés y han conducido a numerosos estudios (Jha, Bilalovic, Jha, Patel & Zhang, 2017).

Según Milidonis *et al.* (2021), actualmente la mayoría de las aplicaciones de la IA en la energía renovable se centran en mejorar la previsión de recursos energéticos, en el mantenimiento predictivo de componentes operativos y la predicción de la demanda. En ese trabajo se revisan diversas aplicaciones de las técnicas de IA en el campo de la energía

solar, entre ellas: el diseño del campo de helióstatos, el sistema receptor y el sistema de almacenamiento térmico; la optimización del sistema receptor, el mantenimiento y diagnóstico de tales sistemas, la optimización de torres solares híbridas y sistemas fotovoltaicos; y el diseño y la optimización del sistema general.

Un ejemplo de aplicación es la predicción a corto plazo de la generación de electricidad requerida por las plantas de energía de combustibles fósiles y su uso de combustible, a partir de pronosticar las fuentes de energía renovable disponibles, como la demanda solar o eólica y de la red. Estos pronósticos a corto plazo no pueden lograrse con modelos estáticos; se necesitan modelos dinámicos con capacidad de toma de decisiones. La IA puede ser una herramienta poderosa para simular la elección de acciones humanas y operar sistemas de energía inteligentes sin ninguna interferencia de los operadores (Entezari, Aslani, Zahedi & Noorollahi, 2023).

Como un área de aplicación en pleno desarrollo, no está libre de retos y barreras que las investigaciones deben enfrentar para facilitar el trabajo en la industria (Quest *et al.*, 2022), entre ellas las siguientes:

- Falta de infraestructura de macrodatos: las industrias no siempre están tecnológicamente preparadas para manejar y acceder a macrodatos, y también deben abordarse las preocupaciones sobre la propiedad de los datos, la privacidad y la ciberseguridad.
- Planificación de proyectos y expectativas de producción contradictorias: las industrias a menudo confían en proyectos y resultados a corto plazo para mantener el liderazgo en el mercado.
- Discrepancias de recursos: la financiación y el intercambio de recursos son un aspecto clave para las transiciones exitosas de la investigación a la industria.

El riesgo general relacionado con las soluciones aportadas por la IA yace en el modelo usado, y más específicamente en la capacidad para verificar las decisiones propuestas por

los algoritmos. A largo plazo, su uso tiende a generalizarse, lo que lleva a un probable aumento del nivel de riesgo de las aplicaciones de IA (Quest *et al.*, 2022); lo que incrementa la necesidad de emplear una IA responsable, como se verá en la última sección de este capítulo.

2.5. Aplicaciones en ciudades inteligentes

A medida que el ser humano va despoblando las zonas rurales para concentrarse en núcleos urbanos, crece la necesidad de aprovechar el desarrollo tecnológico para mejorar la eficiencia de las ciudades. Esto pone a la humanidad ante un gran reto. A mayor población, más tendencia al caos y al desorden, de forma que las autoridades son conscientes del importante papel de la tecnología como nexo de unión y como solución para brindar a los ciudadanos entornos equilibrados, seguros, eficientes y saludables.

Las ciudades están experimentando una evolución hacia ecosistemas inteligentes (Garófalo, 2022). Este cambio se está produciendo mediante la convergencia de diversas tecnologías que incorporan elementos de la Inteligencia Artificial (IA), tales como el aprendizaje automático, las capacidades de redes neuronales, la inteligencia geoespacial, el análisis y la visualización de datos, así como sensores y objetos conectados inteligentes. Estos avances recientes en IA están acercándonos cada vez más al desarrollo de sistemas operativos urbanos que pueden simular patrones relacionados con aspectos humanos, mecánicos y ambientales, abarcando desde la infraestructura de transporte hasta las redes de comunicación.

En este contexto, las metodologías tradicionales de planificación y diseño urbano pueden ser mejoradas y enriquecidas con nuevas herramientas y procesos innovadores habilitados por la IA y las tecnologías inteligentes.

Actualmente una tendencia en el desarrollo de las ciudades es usar tecnologías de alto impacto para crear áreas más

eficientes y seguras en materia de seguridad electrónica, movilidad, energía, entre otras. Entre ellas, la Inteligencia Artificial, el análisis de los macrodatos y el rastreo de datos han llegado para transformar las urbes en ciudades inteligentes (Rebato, 2020).

Las ciudades inteligentes emplean la Inteligencia Artificial para recopilar datos de sus residentes y gestionar de manera eficiente los recursos. A partir de estas bases de datos, anticipan eventos o patrones que pueden prevenir en el futuro. Además, estas áreas hacen uso de la IoT con el fin de mejorar la calidad de vida de sus habitantes.

La infraestructura de las ciudades inteligentes está diseñada para evolucionar junto con el crecimiento de las áreas urbanas. Esto implica que la infraestructura debe ser adaptable a las transformaciones medioambientales, contribuir a la generación de energía y estar alineada con las últimas tecnologías disponibles.

Nuevos términos aparecen ligados a estos avances, entre ellos, las construcciones con domótica, arquitectura optimizada y autosuficiente (*smart buildings*) y las redes de distribución eléctrica eficientes (*smart grids*).

Se puede definir el nuevo concepto de ciudades inteligentes o *smart cities* como aquellas ciudades que han pasado a ser inteligentes gracias a la aplicación de tecnologías de la información y la comunicación, valiéndose para ello de una infraestructura surgida a raíz de la llamada Internet de las cosas. O sea que, las ciudades inteligentes emplean la IA para recopilar información de los ciudadanos y las infraestructuras urbanas, procesarla y gestionar eficientemente los recursos y servicios urbanos. Su objetivo principal es mejorar la calidad de vida de los ciudadanos, aumentar la eficiencia en el uso de recursos y lograr beneficios económicos. Por lo tanto, es esencial implementar adecuadamente esa tecnología en el funcionamiento urbano, adaptándola a las necesidades de la ciudad y su entorno.

A medida que la tecnología de la IA se vuelve más madura, más económica y más accesible, aparecen en la ciudad las aplicaciones impulsadas por la IA (Power, 2018). Las primeras ciudades adoptantes de esa tecnología preparan el camino para aplicarla mediante la recopilación de datos y desarrollando proyectos de exploración o previendo su aplicación en ámbitos como el transporte, la atención sanitaria, la educación, la seguridad, la cultura y la asistencia a grupos desfavorecidos, donde los ciudadanos notarán con más fuerza su impacto (Grosz, 2016).

La adopción de la IA genera nuevas oportunidades para los gobiernos; puede ayudar a poner a las personas en el centro de los servicios del gobierno electrónico, aumentar la velocidad y la calidad de los servicios y optimizar los procesos liberando al personal de tareas repetitivas. La IA es la clave para liberar el valor oculto en el desarrollo de la ciudad inteligente en general y será interesante en el contexto del desarrollo de un gobierno cognitivo.

La IA aporta a los procesos de la ciudad inteligente varios elementos clave, como la capacidad de aprendizaje automático, la automatización de procesos y la toma de decisiones autónoma, contribuyendo así a los cuatro componentes clave de las ciudades inteligentes identificados por del Cerro Santamaría (2023): conjuntos de datos, plataformas de comunicación, toma de decisiones oportunas y acción efectiva.

El concepto (y la práctica) de ciudades inteligentes son dinámicos y en constante evolución. Su definición no es fija, ni tiene un punto final, sino que se trata de un proceso en desarrollo. La implementación de la IA dentro de este paradigma de ciudad inteligente representa la maduración de la economía de la información, donde se busca maximizar el valor económico a través de la recopilación y el análisis de datos, y también refleja la expansión de la cuarta revolución industrial.

La evolución de las ciudades inteligentes hacia la adopción generalizada de técnicas, procesos y dispositivos de IA sugiere

que el nuevo régimen urbano basado en la Inteligencia Artificial continuará creciendo en los próximos años. Por lo tanto, es esencial analizar tanto los beneficios como los riesgos asociados con esta estrategia en constante desarrollo.

Según el autor referido antes, desde la gobernanza urbana, la tecnología en las ciudades inteligentes tiene como objetivo facilitar la interacción directa de los administradores urbanos con la infraestructura de la ciudad y la comunidad, permitiéndoles monitorear en tiempo real lo que ocurre en las áreas urbanas y su evolución.

En el análisis de Rebato (2020), también se aborda la transformación de la infraestructura de las ciudades, que va más allá de las construcciones y su funcionalidad. En este contexto, se destaca la importancia de la planificación urbana en relación con el progreso y la evolución de la ciudad. El enfoque principal es que esta infraestructura sea de fácil acceso y simplifique la vida de los ciudadanos. La infraestructura urbana se convierte en la puerta de entrada a nuevas formas de desarrollo social y cultural.

En una ciudad preparada para abordar sus problemas de movilidad utilizando la IA se recopila información del tráfico en tiempo real, lo que permite predecir embotellamientos y mejorar la movilidad. Además, se implementan soluciones como semáforos inteligentes programados con aprendizaje automático, capaces de adaptarse a las condiciones de tráfico para optimizar el servicio. En una ciudad inteligente, la red de transporte público facilita la conexión entre distintas áreas de manera eficiente, lo que ahorra tiempo a los ciudadanos. Además, se promueven medios de transporte que se ajusten a las necesidades energéticas y medioambientales de la ciudad, como los vehículos eléctricos.

Una de las características distintivas de las ciudades inteligentes es la implementación de soluciones inteligentes para optimizar la prestación de servicios públicos. Esta perspectiva tiene como pilar central la consideración del medioambiente, con el objetivo de reducir al máximo el impacto ecológico. La IA desempeña un papel fundamental en este proceso, ya que

permite a la ciudad procesar información sobre el uso y el consumo de servicios como el agua y la energía para mejorar su eficiencia.

Además, se promueven métodos efectivos para gestionar los residuos y fomentar la reutilización de estos. Otro ámbito en el que se puede potenciar la IA es la seguridad ciudadana. Las ciudades inteligentes se esfuerzan por garantizar la protección y tranquilidad de sus habitantes mediante la implementación de nuevas tecnologías. Esto incluye la reducción de la delincuencia a través de la innovación tecnológica.

En una ciudad segura, se pueden encontrar cámaras de seguridad, sistemas de vigilancia constante y un sistema de iluminación adecuado y sostenible. Además, se utilizan sistemas inteligentes de alerta temprana, biometría e identificación facial para mantener el control sobre las situaciones y prevenir delitos (Rebato, 2020).

Otro ejemplo está en la gestión de la energía (Garófalo, 2022). Pueden instalarse sensores en ubicaciones estratégicas dentro de edificios, lo que contribuirá a recopilar información sobre el consumo de energía y a prever el comportamiento de los consumidores. En este sentido, la Inteligencia Artificial y las ciudades inteligentes tienen el potencial de fortalecer la seguridad de las redes eléctricas y mejorar la gestión del rendimiento. Las redes eléctricas inteligentes tienen la capacidad de procesar grandes cantidades de datos provenientes de medidores inteligentes, lo que les permite evaluar y anticipar la respuesta de la demanda y la agrupación de cargas.

En *Action Pills* (2023) se plantea que la tecnología debe aplicarse en núcleos urbanos a través de tres capas para que pueda ser realmente útil a sus ciudadanos:

- La primera capa, basada en tecnologías de recepción como teléfonos inteligentes (smartphones) y sensores para recopilar la información y redireccionarla. Deben estar conectados a Internet a través de redes, destacando especialmente el 5G de reciente aparición.

- La segunda capa, a partir de aplicaciones específicas que gracias a su programación puedan procesar de forma inteligente todos estos datos, para convertirlos en información útil y aprovechable.

- La tercera capa, apoyada en una estructura altamente organizada y compleja que interactúe con esa información y la coordine para responder a las necesidades en tiempo real y para ofrecer soluciones a los ciudadanos, de forma que luego estos puedan tomar las mejores decisiones posibles para cada situación.

Indudablemente, construir una ciudad inteligente capaz de coordinar todos los procesos y satisfacer las necesidades de millones de personas a corto plazo no será una tarea sencilla. No obstante, este se perfila como uno de los principales desafíos para la humanidad en los próximos años. No obstante, los potenciales beneficios de la Inteligencia Artificial superan con creces el esfuerzo requerido para llevar a cabo esta empresa. De hecho, hay varios aspectos en los que los ciudadanos experimentarán mejoras significativas en su calidad de vida:

- Acceso y rentabilidad de los servicios: La incorporación de la domótica en las ciudades futuras tiene como objetivo simplificar la vida en los hogares, al mejorar la provisión de servicios como gas, agua y electricidad, al tiempo que reducirá el impacto ambiental, favorecerá el reciclaje y minimizará los desechos, a un menor costo económico y con un uso más eficiente de los recursos.

- Mayor seguridad: Se trata de prevenir y controlar los crímenes mediante una red de cámaras de seguridad y sistemas de vigilancia, sistemas de iluminación, alarmas o identificación, así como biometría, identificación facial con el fin de registrar movimientos y accesos de los ciudadanos.

- Mejora de las redes de transporte: Está dada por la capacidad de descongestionar las principales vías urbanas al procesar datos de tráfico en tiempo real y optimizar la organización de los sistemas de transporte público. Semáforos inteligentes, respaldados por sistemas de aprendizaje

automático y vehículos de emergencia, como ambulancias, conectados a la nube pueden perfeccionar sus rutas para evitar obstáculos.

- Fortalecimiento del mercado laboral y soluciones: La IA lejos de ser una amenaza para ciertos empleos, en realidad obligará al mercado laboral a adaptarse. Si bien algunas profesiones tradicionales podrían desaparecer debido a la automatización, surgirán nuevas oportunidades altamente calificadas para impacto del desarrollo de la IA sobre el mercado laboral; se plantea que como en gestionar estos sistemas avanzados.

Sobre este último aspecto en Sánchez (2023, 25 de marzo) se analiza los anteriores saltos tecnológicos, sí se prevé una masiva destrucción de tareas que componen las profesiones actuales, pero no de puestos de trabajo, este cambio exigirá un impulso en la formación de trabajadores cuyas funciones queden obsoletas. Allí se dice que, según algunas estimaciones, entre 2020 y 2025 se destruirán 85 millones de empleos obsoletos en el mundo y se crearán 97 millones por los nuevos oficios.

Respecto a esto, según el blog de (Acction Pills, 2023), se afirma que “la llegada a la mayoría de edad de las denominadas generaciones Y y Z está favoreciendo cada vez más la penetración de la tecnología y de la IoT en la sociedad, por lo que está surgiendo un nuevo tipo de ciudadano: el *smart citizen*, quien no sólo utiliza la IA en su propio beneficio, sino que también la emplea para mejorar su propia ciudad”.

La evolución de las ciudades inteligentes hacia la adopción generalizada de técnicas, procesos y dispositivos de IA sugiere que en los próximos años se verá una expansión de un nuevo paradigma urbano basado en esta. A pesar de que el desarrollo de las ciudades inteligentes es gradual, su naturaleza disruptiva es innegable. No obstante, es significativo señalar que este avance conlleva riesgos y amenazas, como se discute en el artículo de del Cerro Santamaría (2023).

El mismo autor también analiza que, la inteligencia artificial podría ser vulnerable a posibles ataques cibernéticos, lo que

podría permitir que ciberdelincuentes interfirieran con sistemas críticos como energía, transporte o alerta temprana. Dado que los sistemas de IA operan de manera autónoma, sus resultados pueden ser impredecibles, pues una IA sin control representa un peligro potencial para las personas. Además, existen otros riesgos y dilemas éticos asociados a la estrategia de ciudades inteligentes, como cuestiones de control, privacidad y seguridad. Estos problemas surgen debido a la alta recopilación, análisis, escaneo e identificación de grandes conjuntos de datos para la monitorización predictiva, lo que puede llevar a una situación de “vigilancia”.

2.6. Aplicaciones en ciberseguridad

Dado el desarrollo constante de la informatización de la sociedad, y la transformación digital que ocurre en todos los procesos socio-económicos de esta, los ciberataques están a la orden del día. El incremento de los ataques cibernéticos que reciben las instituciones y empresas ha provocado que las personas no puedan hacer frente a todas las amenazas de manera eficaz.

En la sección dedicada al desarrollo de las ciudades inteligentes utilizando la IA ya se adelantó sobre este otro campo de aplicación de las técnicas de la Inteligencia Artificial.

La aplicación de estas técnicas representa un avance significativo en el campo de la ciberseguridad. Esto habilita a las soluciones de ciberseguridad para analizar y aprender de los datos de manera más eficiente y precisa. Los profesionales en este ámbito a menudo se ven abrumados por una gran cantidad de tareas, una sobrecarga de datos, limitaciones de tiempo y una escasez de habilidades disponibles. Por lo tanto, la inteligencia artificial puede tener un impacto considerable en varias áreas clave, como la detección de amenazas con precisión; la automatización de la respuesta y la agilización de

la investigación de ataques (Ansari, Dash, Sharma & Yathiraju, 2022).

Tales beneficios son especialmente importantes dado que, en muchos casos, la demanda de recursos humanos en ciberseguridad supera la disponibilidad de profesionales capacitados.

Entre las técnicas de la IA más relevantes en este contexto están las relacionadas con el aprendizaje automático, dada la gran cantidad de datos disponibles y la continua evolución de los tipos de ataques a los sistemas informáticos; una característica de esta tecnología, que la sitúa como una aliada en la lucha contra este tipo de ataques, es que aprende continuamente.

La relación entre la ciberseguridad y la inteligencia artificial es estrecha, como señala Ayerbe (2020), a partir de tres enfoques distintos:

1. Enfoque de Defensa: Las técnicas de inteligencia artificial se utilizan para mejorar la ciberseguridad y la resiliencia de productos, servicios y sistemas, lo que beneficia a las empresas y la sociedad al fortalecer sus medidas de seguridad.

2. Enfoque de Ataque: Los cibercriminales y otros atacantes cibernéticos están empezando a utilizar la inteligencia artificial para comprometer la ciberseguridad y llevar a cabo diversos tipos de ataques, incluida la difusión de noticias falsas, lo que supone nuevos desafíos en términos de seguridad cibernética.

3. Enfoque de Confianza: Los sistemas de inteligencia artificial en sí mismos pueden ser vulnerables a ciberataques. Por lo tanto, es esencial desarrollar sistemas de IA seguros que protejan la privacidad, sean confiables y sean aceptados por sus usuarios.

Esta triple representación destaca la importancia de la inteligencia artificial en el ámbito de la ciberseguridad y subraya la necesidad de abordar tanto la defensa como la detección y prevención de ataques en este contexto.

La IA desempeña un papel fundamental en el campo de la ciberseguridad, ya que aborda los crecientes desafíos asociados con la complejidad de los sistemas informáticos

modernos (I4.0, infraestructura de la IoT, etc.) y la gran cantidad de datos generados por ellos. La ciberseguridad enfrenta diversas dificultades, como la detección de intrusiones, la protección de la privacidad, la defensa proactiva y la identificación de comportamientos anómalos, además de la constante evolución de las amenazas cibernéticas. Por esta razón, se han explorado vías basadas en la IA para agilizar el análisis y la toma de decisiones en tiempo real, permitiendo una detección y respuesta más rápida ante ciberataques.

También se emplea en múltiples etapas de un enfoque integral de seguridad inteligente, que incluye la identificación, protección, detección, respuesta y recuperación de incidentes. De esta manera, la ciberseguridad se convierte en otro campo de aplicación de la IA, al igual que lo son la energía, el transporte, la industria o la salud, pues ella proporciona herramientas y métodos para fortalecer la seguridad cibernética y hacer frente a las cambiantes amenazas que surgen constantemente en un entorno digital en constante evolución.

Algunas de las aplicaciones ventajosas que trae consigo la aplicación de esta tecnología para combatir ciberataques son, según Montoya (2023):

1. **Análisis predictivos:** A través de los análisis predictivos, la IA puede prevenir futuras amenazas. Desde qué dirección de IP tiene más probabilidades de ser atacada, hasta los tipos de amenazas que pueden producirse. Esta parte es clave porque la organización puede conocer dónde están sus brechas de seguridad para poder solucionarlas antes de que se ejecute un verdadero ataque.

2. **Detección automática de amenazas:** Los algoritmos de aprendizaje automático son entrenados para detectar, de manera automática, aquellas conductas o irregularidades en el rendimiento del sistema que indican un posible ataque a la seguridad.

3. **Prevención de potenciales ataques:** Esta tecnología puede detectar anomalías en el tráfico de la red o en el flujo

de datos a través del análisis de los patrones. Esto ayuda a detectar comportamientos anormales antes de que se produzca un ataque.

4. Automatización de la seguridad: Este apartado es muy útil para proteger a las empresas de amenazas típicas, como el *phishing*. Al automatizar los protocolos de seguridad se consigue reducir el error humano, ahorrar tiempo y esfuerzo.

5. Vulnerabilidades día cero: La IA en ciberseguridad elimina vulnerabilidades de día cero, es decir, cuando un hacker vulnera la seguridad antes de que los responsables del software puedan encontrar una solución.

En Kaspersky (2023), se destaca cómo las técnicas de inteligencia artificial resultan valiosas en el análisis de perfiles de los usuarios que interactúan con un sistema informático. Al crear perfiles personalizados del personal de la red a partir de sus comportamientos, es posible adaptar la seguridad de manera personalizada a las necesidades de la organización. Este modelo puede detectar a usuarios no autorizados al identificar desviaciones en su comportamiento. Incluso detalles sutiles, como la forma en que escriben en el teclado, pueden contribuir a construir un modelo predictivo de amenazas. Al anticipar posibles escenarios de manipulación, un sistema de seguridad basado en IA puede ofrecer soluciones para reducir la superficie de ataque potencial.

De manera análoga al perfil de comportamiento del usuario, se puede crear un perfil de diagnóstico personalizado para el rendimiento general de la computadora cuando está funcionando correctamente. Supervisar el uso de recursos como el procesador y la memoria, junto con indicadores como el consumo de datos de Internet, puede ser útil para identificar actividades maliciosas. No obstante, es importante considerar que algunos usuarios legítimos pueden requerir un alto uso de datos, por ejemplo, durante videoconferencias o descargas de archivos multimedia voluminosos. En estos casos, el algoritmo

debe ser capaz de distinguir estas prácticas habituales de las desviaciones sospechosas.

Las tecnologías de aprendizaje automático también pueden desempeñar un papel en la identificación y bloqueo de actividades de bots, que pueden afectar negativamente el rendimiento de los sitios web al consumir ancho de banda entrante, lo que resulta en una pérdida de tráfico y clientes potenciales. Incluso cuando los bots utilizan herramientas de anonimización como las redes privadas virtuales, los algoritmos pueden generar modelos predictivos basados en datos sobre comportamiento cibercriminal y tomar medidas proactivas, como el bloqueo de nuevas direcciones web que exhiben patrones de actividad similares.

Por su parte, en Ethical Hacking Consultores (2021) se identifican las siguientes aplicaciones de la IA en ciberseguridad: identificación de amenazas y neutralización de ciberataques, gestión de vulnerabilidades; optimización y monitorización de centros de datos, aprehensión de los patrones de comportamiento del tráfico en la red y recomendaciones a la agrupación de cargas de trabajo, así como la aplicación de políticas de seguridad, identificación segura de usuarios, ayuda para clasificar de manera automática la información por su nivel de criticidad de cara a distintas regulaciones, bloqueo de bots a partir de su comportamiento, entre otras.

Pero el uso de la IA en la ciberseguridad también tiene un lado negativo (Ayerbe, 2020). La IA ya se está utilizando en aplicaciones del mercado para conocer los patrones de comportamiento de usuarios y diseñar campañas comerciales utilizando software de IA a disposición de todo el mundo, por lo que sería muy ingenuo pensar que los ciberdelincuentes no lo estén utilizando también, en el caso más sencillo, para conocer mejor a sus víctimas e identificar el mejor momento en el que realizar una acción delictiva con las mayores garantías de éxito.

La utilización de la IA depende del perfil de los ciberatacantes, que van desde los más inofensivos, asociados a la cibermaleicia, a los más peligrosos, como pueden ser los relacionados con el

ciberterrorismo, el ciberespionaje o la ciberguerra. La misma variedad puede encontrarse en el nivel de sofisticación y complejidad de los ciberataques, que varía mucho de unos a otros. Detrás de los ciberataques más peligrosos y sofisticados susceptibles de utilizar la IA pueden estar grupos muy especializados, financiados por determinados estados y cuyos ataques pueden estar dirigidos hacia infraestructuras críticas de otro país o a generar campañas de desinformación. Si se analiza a los potenciales atacantes desde otra perspectiva, la de su forma de operar, pueden utilizar la IA tanto para explotar vulnerabilidades conocidas como para encontrar otras desconocidas o crearlas.

Además de la utilización del uso ofensivo de la IA por parte de los ciberatacantes para conocer patrones de comportamiento de las futuras víctimas, también puede utilizarse para romper más rápidamente contraseñas, construir malware que evite la detección, esconderse donde no puedan ser encontrados, adaptarse lo antes posible a las contramedidas que puedan tomarse, así como para la obtención automática de información utilizando métodos de procesamiento del lenguaje natural, la suplantación y generación de audios, vídeos y textos falsos. Los atacantes también están utilizando redes generativas conflictivas (*Generative Adversarial Networks*, GAN) para imitar patrones de tráfico de comunicaciones normales con el objetivo de distraer la atención de un ataque, encontrar y extraer rápidamente datos sensibles. Dada la relevancia de la IA en este contexto, el asunto es analizado con énfasis en Kissinger, Schmitdt & Huttenlocher (2021).

En Ayerbe (2020) se identifican cuatro tipos de ataques: acceso a los datos, envenenamiento, evasión y de oráculo (tipo *oracle*). El primero implica que el atacante pueda acceder a la totalidad o parte de los datos de entrenamiento para crear un modelo sustituto. En los de tipo “envenenamiento” se pueden alterar los datos o el modelo de manera indirecta o

directa: bien los atacantes, que no tienen acceso a los datos preprocesados utilizados por el modelo, modifican los datos antes del preprocesamiento (indirecto), o bien realizan cambios deliberados en los datos de entrenamiento, sin necesidad de acceder a los datos preprocesados utilizados por el modelo (directo), para influir en el comportamiento de este de manera perjudicial.

Las tácticas de ataque en este último tipo de envenenamiento son diversas; ellas son la Inyección de datos, consistente en la inserción de muestras conflictivas en los datos de entrenamiento originales para cambiar la distribución subyacente sin modificar las características o etiquetas iniciales; la Manipulación de datos de etiquetas, que es la modificación de las etiquetas de salida de los datos de entrenamiento originales; la Manipulación de datos de entrada, que radica en la alteración de los datos originales; y la Corrupción de la lógica, centrada en la manipulación del algoritmo de AA para afectar el proceso de aprendizaje y el modelo en sí.

En los ataques de tipo “evasión”, el atacante resuelve un problema de optimización restringida para encontrar una perturbación en la entrada que cause la clasificación errónea deseada. Finalmente, en los ataques de tipo “oráculo”, el atacante utiliza una interfaz o API del modelo para presentar entradas, observar las salidas y obtener la información que busca.

La información precedente permite coincidir con la conclusión de Prosegur (2023), acerca de que la Inteligencia Artificial es un arma de doble filo y reclama una política de seguridad al nivel de su potencial destructor. La IA con fines maliciosos podría potencialmente comprometer la integridad de la IA diseñada para proteger sistemas y datos. Los algoritmos inteligentes requieren ser entrenados con grandes volúmenes de información, por lo tanto, si un atacante logra introducir datos falsos durante esta fase, podría comprometer toda la

evaluación posterior, sin importar para qué se utilice, ya sea en aplicaciones como la conducción autónoma, el reconocimiento de voz o incluso en el campo de la telemedicina.

Otro aspecto muy importante a analizar es el siguiente. No se debe olvidar que la ciberseguridad es un dominio más de aplicación de la IA, pero que, al mismo tiempo, los sistemas de IA también pueden ser sistemas vulnerables, por lo que se tienen que desarrollar con privacidad y seguridad desde el diseño y hay que velar por su ciberseguridad a lo largo de todo el ciclo de vida del sistema y considerando toda la cadena de suministro, pues la seguridad de los sistemas inteligentes es uno de los aspectos que permiten incrementar la confianza de los usuarios en estos sistemas. Se debe trabajar en la IA para crear técnicas, métodos y herramientas que faciliten el diseño, desarrollo, validación y despliegue de sistemas basados en IA con un enfoque multicriterio que considere la ciberseguridad del dato, del modelo y del resultado (Ayerbe 2020). Finalmente, se pueden enumerar alertas y consejos sugeridos por los autores:

Alertas (o limitaciones):

1. El empleo de la IA requiere conjuntos de datos, pero usarlos puede entrar en conflicto con las leyes de privacidad de datos
2. El sector necesita más expertos en IA y ciberseguridad que los que están disponibles
3. Los equipos humanos seguirán siendo esenciales

Consejos (para abordar el futuro de la ciberseguridad):

1. Se debe invertir en tecnología que esté siempre orientada al futuro.
2. Complementar a los equipos con IA y AA, en lugar de sustituirlos.
3. Actualizar periódicamente las políticas de datos para cumplir con los cambios en la legislación.
3. Una IA responsable: requisito para su aplicación exitosa.

A pesar del potencial de la IA, hay muchas inquietudes ¿Cuáles son los límites, los riesgos, las limitaciones y los problemas de un uso prematuro o incorrecto? ¿Qué acciones se tienen que utilizar para garantizar que la IA se haga de una manera segura y fiable? ¿Qué uso de la IA por parte de todo el sector público puede proporcionar un resultado sostenible y deseado para todo el mundo? Es evidente la necesidad de tener una mirada ética en el desarrollo de servicios digitales ciudadanos basados en IA para hacer frente a los diversos problemas que se presentan, como son: la violación de la privacidad y de la seguridad digital de los ciudadanos, la desigualdad y la degradación ambiental, sólo por mencionar algunos.

Un ejemplo de la necesidad de que junto al desarrollo de la IA como disciplina científica y sus aplicaciones hace falta una IA responsable son el desarrollo de las ciudades inteligentes, o la gestión de la salud basada en IA, analizada antes. Se trata de un proceso complicado que viene acompañado de dilemas éticos y de trabas legales, asociados principalmente a la recopilación y tratamiento de los datos de los ciudadanos.

De aquí que se esté planteando con seriedad la necesidad de un marco normativo y regulatorio que se encargue de establecer los límites que no pueden traspasarse; como ha expresado el Jefe de IoT, Robótica y Ciudades Inteligentes para el Foro Económico Mundial, Jeff Merri: la clave es que la tecnología deje de ir por delante de la legislación y que se tenga siempre presente que, sin ley, las tecnologías se convierten en un riesgo masivo, pero con unas reglas claras y definidas no tienen que ser algo malicioso (Action Pills, 2023). Esto no tiene que representar un alto riesgo para el hombre, por eso, es fundamental que, aunque la Inteligencia Artificial venga definida por procesos automatizados, el ser humano siga estando siempre ahí para controlarla, regularla y velar por su correcto funcionamiento.

De acuerdo con Silva y González (2022), los proveedores que trabajen en el desarrollo de tecnologías futuras que aprovechen la IA, así como las organizaciones gubernamentales a nivel nacional y local que adquieran estas tecnologías para sus ciudades, deberán tener en cuenta cómo sortear los marcos legales y regulaciones actuales y futuras que supervisen el desarrollo e implementación de sistemas de IA. El uso de IA puede plantear diversas preocupaciones legítimas, entre las que se incluyen riesgos de privacidad de la información, especialmente cuando los datos procesados por estos sistemas involucran datos personales, como reconocimiento facial y sistemas biométricos utilizados para vigilancia y seguridad. Cumplir con los requisitos de transparencia representa un desafío significativo en el contexto de las ciudades inteligentes. Es fundamental establecer una comunicación efectiva con los ciudadanos que interactúan con los sistemas de IA mientras se desplazan por una ciudad inteligente.

Los sistemas de IA tendrían que someterse a pruebas rigurosas para comprobar la fiabilidad y la seguridad. Hay diferencias significativas entre un conjunto de datos de pruebas de IA y su despliegue en el mundo real. Las aplicaciones de IA pueden potenciar o esconder el sesgo intrínseco o basado en datos. No tienen conciencia, reflejan los prejuicios y las opiniones humanas y repiten y perpetúan datos incorrectos, lo cual provoca interpretaciones erróneas y errores. La policía predictiva conduce a la estigmatización, la discriminación, la desconfianza en esos cuerpos y tiene un efecto que se retroalimenta (Schrage, 2018).

La relevancia de este aspecto sobre el uso de la IA ha llevado a que la UNESCO elabore un documento donde se hace un análisis de la problemática y se dan recomendaciones a los gobiernos de cómo enfrentarla Unesco (2021). El 23 de noviembre de 2021, en la 41 Conferencia General de la UNESCO, los 193 Estados miembros aprobaron por aclamación el texto

de la Recomendación. La directora general de la UNESCO, Audrey Azoulay, planteó que el mundo necesita reglas para que la Inteligencia Artificial beneficie a la humanidad. Añadió que la Recomendación sobre la ética de la IA es una respuesta importante; establece el primer marco normativo mundial, al tiempo que otorga a los Estados la responsabilidad de aplicarlo a su nivel. La UNESCO apoyará a sus 193 Estados miembros en su aplicación y les pedirá que informen periódicamente sobre sus progresos y prácticas.

Referencias del capítulo

ActionPills (2023). Así transformará la inteligencia artificial las ciudades del futuro. <https://actions.es/asi-transformara-la-inteligencia-artificial-las-ciudades-del-futuro/>

Agronet (2021). El Internet de las cosas y la inteligencia artificial, al servicio de una agricultura de alta precisión. <https://www.theagilityeffect.com/es/article/el-internet-de-las-cosas-y-la-inteligencia-artificial-al-servicio-de-una-agricultura-de-alta-precision/>

Amasyali, K. y El-Gohary, N. M. (2018). A review of data-driven building energy consumption prediction studies. *Renewable and Sustainable Energy Reviews*, 81(1), 1192–1205. <https://doi.org-10.1016/j.rser.2017.04.095>

Ang, J. H., Goh, C., Saldivar, A. A. F., & Li, Y. (2017). Energy-efficient through-life smart design, manufacturing and operation of ships in an industry 4.0 environment. *Energies*, 10(5), 61-13. <https://doi.org/10.3390/en10050610>

Ansari, M. F., Dash, B., Sharma, P. & Yathiraju, N. (2022). The Impact and Limitations of Artificial Intelligence in Cybersecurity: A Literature Review. *International Journal of Advanced Research in Computer and Communication Engineering*. <https://www.researchgate.net/profile/Bibhu->

Dash-.5/publication/364122631_The_Impact_and_Limitations_of_Artificial_Intelligence_in_Cybersecurity_A_Literature_Review/links/633b425d76e39959d6959aad/The-Impact-and-Limitations-of-Artificial-Intelligence-in-Cybersecurity-A-Literature-Review.pdf

Antonopoulos, I., Robu, V., Couraud, B., Kirli, D., Norbu, S., Kiprakis, A., Flynn, D., Elizondo-Gonzalez, S. y Wattam, S. (2020). Artificial intelligence and machine learning approaches to energy demand-side response: *A systematic review. Renewable and Sustainable Energy Reviews*, 130, 109899. <http://dx.doi.org/10.1016/j.rser.2020.109899>

Asadollahi-Yazdi, E., Couzon, P., Nguyen, N. Q., Ouazene, Y. & Yalaoui, F. (2020). Industry 4.0: Revolution or Evolution? *American Journal of Operations Research*, 10(06), 241. https://www.scirp.org/html/1-1040746_103411.htm

Ayerbe, A. (2020). La ciberseguridad y su relación con la inteligencia artificial. <https://www.realinstitutoelcano.org/analisis/la-ciberseguridad-y-su-relacion-con-la-inteligencia-artificial/>

Bareiss, E. R., Porter, B. W. & Wier, C. (1987). Protos: An Exemplar-Based Learning Apprentice. *International Journal of Man-Machine Studies*, 29(5),12-23. [https://doi.org/10.1016/S0020-7373\(88\)80012-9](https://doi.org/10.1016/S0020-7373(88)80012-9)

Boden, M. A., (2017). Inteligencia Artificial, Turner Noema. <https://es.scribd.com/document/661576369/2017-Boden-Margaret-Inteligencia-Artificial>

Braga, A. & Logan, R. K. (2019). AI and the Singularity: A Fallacy or a Great Opportunity? *Information*, 10(2),1-4. <https://doi.org/10.3390/info10020073>

Brenes, J. A., Martínez, A., Quesada, C. U. & Jenkins, M. (2020). Sistemas de apoyo a la toma de decisiones que usan inteligencia artificial en la agricultura de precisión. <https://www.kerwa.ucr.ac.cr/bitstream/handle/10669/81549/Brenes%2C%20Martinez%2C%20Quesada%20y%20>

Jenkins%20-%20RISTI_E28.pdf?sequence=1

Buch, V. H., Ahmed, I. & Maruthappu, M. (2018). Intelligence in Medicine: Current Trends and Future Possibilities. *British Journal of General Practice*, 68(668), 143-144. <https://doi.org/10.3399/bjgp18X695213>

Chakraborty, S., Pal, C., Chatterjee, S., Chakraborty, B. & Ghoshal, N. (2015). Knowledge-Based System Architecture on CBR for Detection of Cholera Disease. In: Mansal, D., Kar, R., Das, S., Paiigrahi, B. (eds). *Intelligent Computing and Applications. Advances in Intelligent Systems and Computing*, 343. Springer, New Dali. https://doi.org/10.1007/978-81-322-2268-2_17

Chouard, T. (2016) The Go Files: AI computer clinches victory against Go champion. *Nature*. <https://doi.org/10.1038/nature.2016.19553>

Choudhury, N. & Begum, S. A. (2016). A Survey on Case-Based Reasoning in Medicine. *International Journal of Advanced in Computer Science and Applications*, 7(8), 136-144. <https://doi.org/10.14569/IJACSA.2016.070820>

Del Cerro Santamaría, G. (2023). Sobre la inteligencia de las ciudades. <https://telos.fundaciontelefonica.com/sobre-la-inteligencia-de-las-ciudades/>

Energyavm. (2020). La inteligencia artificial y su relación con el sector energético. [<https://www.editores.com.ar/inteligencia%20artificial%20y%20su%20relaci%C3%B3n%20con%20el%20sector%20energ%C3%A9tico>]

Entezari, A., Aslani, A., Zahedi, R. & Noorollahi, Y. (2023). Artificial intelligence and machine learning in energy systems: A bibliographic perspective. *Energy Strategy Reviews*, 45, 101017. <https://doi.org/10.1016/j.esr.2022.101017>

Ethical Hacking Consultores (2021). Inteligencia Artificial: un nuevo hito para la Ciberseguridad. Medium <https://ehcgroup.medium.com/inteligencia-artificial-un-nuevo-hito-para-la-ciberseguridad-57410236cfcf>

Fernández, Y. (2023). ChatGPT: qué es, cómo usarlo y qué puedes hacer con este chat de inteligencia artificial GPT-3.

Xataka. <https://www.xataka.com/basics/chatgpt-que-como-usarlo-que-puedes-hacer-este-chat-inteligencia-artificial>

Fouquet, F. (2021). Inteligencia Artificial aplicada a la agricultura de precisión. Control de hongos en la planta de tomate. [Tesis de grado no publicada, Universidad de Málaga]. <https://riuma.uma.es/xmlui/bitstream/handle/10630/23510/Fouquet%20Calder%C3%B3n%20Fabrice%20Memoria.pdf?sequence=1>

Garófalo, F. (15 de noviembre de 2022). Las Ciudades Inteligentes y la Inteligencia Artificial (IA). Neurona. <https://neurona-ba.com/las-ciudades-inteligentes-y-la-inteligencia-artificial-ia/>

Ghoshchi, A., Zahedi, R., Pour, Z. & Ahmadi, A. (2022). Machine learning theory in building energy modeling and optimization: a bibliometric analysis. *J Mod Green Energy*, 1(4), 1-16 <https://www.doi.org/10.53964/jmge.2022>

González, A. (12 de mayo de 2021). Inteligencia Artificial es la nueva electricidad para las industrias: IBM. dplnews. <https://dplnews.com/inteligencia-artificial-es-la-nueva-electricidad-para-las-industrias-ibm/>

Haddad, M., Moertl, D. & Porénta, G. (1995). SCINA: A Case-Based Reasoning System for the Interpretation of Myocardial Perfusion Scintigrams. *Artificial Intelligence in Medicine*, 9(1), 61-78. [https://doi.org/10.1016/S0933-3657\(96\)00361-2](https://doi.org/10.1016/S0933-3657(96)00361-2)

Heider, D., Senge, R., Cheng, W., & Hüllermeier, E. (2013). Multilabel Classification for Exploiting Cross-Resistance Information in HIV-1 Drug Resistance Prediction. *Bioinformatics*, 29(16), 1946-1952. <https://doi.org/10.1093/bioinformatics/btt331>

Hessler, G. & Baringhaus, K. H. (2018). Artificial Intelligence in Drug Design. *Molecules*, 23(10), 1-13. <https://doi.org/10.3390/molecules23102520>

High-Level Expert Group on Artificial Intelligence. (18 de December de 2018). A definition of AI: Main capabilities and

disciplines Independent. European Commission. https://ec.europa.eu/futurium/en/system/files/ged/ai_hleg_definition_of_ai_18_december_1.pdf

Ibiza. (s.f). Inteligencia artificial y su efecto en la industria 4.0. Ibiza. <https://ibisa.co/inteligencia-artificial-y-su-efecto-en-la-industria-4-0/>

International Business Machines. (2018). What is Artificial Intelligence in Medicine? <http://www.ibm.com/watson-health/learn/artificial-intelligence-medicine>

International Energy Agency (2017). Digitalization and energy – Analysis. <https://iea.blob.core.windows.net/assets/b1e6600c-4e40-4d9c-809d-1d1724c763d5/DigitalizationandEnergy3.pdf>

Instituto tecnológico de Santo Domingo. (2023). Agricultura de precisión e Inteligencia Artificial, las tendencias globales del sector agroalimentario en 2023. <https://www.interempresas.net/Grandes-cultivos/Articulos/465242-agricultura-precision-Inteligencia-Artificial-tendencias-sector-agroalimentario-2023.html>

Jha, S. K., Bilalovic, J., Jha, A., Patel, N. & Zhang, H. (2017). Renewable energy: Present research and future scope of artificial intelligence. *Renewable and Sustainable Energy Reviews*, 77, 297–317. <https://doi.org/10.1016/j.rser.2017.04.018>

Kaspersky (2023). La IA y el aprendizaje automático en la ciberseguridad: cómo determinarán el futuro. <https://latam.kaspersky.com/resource-center/definitions/ai-cybersecurity>

Kissinger, H., Schmitdt, E. & Huttenlocher, D. (2021). The Age of AI and our human future. John Murray Publishers.

Koton, P. (1989). A Medical Reasoning Program that Improves with Experience. *Computer Methods and Programs in Biomedicine*, 30(2), 177-184. <https://www.sciencedirect.com/science/article/abs/pii/0169260789900709?via%3Dihub>

Lei, M., Shiyan, L., Chuanwen, J., Hongling, L., & Yan, Z. (2009). A review on the forecasting of wind speed and generated power. *Renewable and Sustainable Energy Reviews*, 13(4), 915 - 920. <http://www.sciencedirect.com/science/article/pii/>

S1364-0321(8)00028-2

Li, R., Liu, W., Lin, Y., Zhao, H. & Zhang, C. (2017). An Ensemble Multilabel Classification for Disease Risk Prediction. *Journal of Healthcare Engineering*, 1-10. <https://doi.org/10.1155/2017/8051673>

Lo, Y. C., Ren, G., Honda, H. & Davis, K. L. (2019). Artificial Intelligence-Based Drug Design and Discovery. <https://doi.org/10.5772/intechopen.89012>

Macura, R. T., & Macura, K. J. (1995). MacRad: Case-Based Retrieval System for Radiology Image Resource. AAAI Technical Report FS-95-03. <https://www.aaai.org/Papers/Symposia/Fall/1995/FS-95-03/FS95-03-031.pdf>

Marling, C. & Whitehouse, P. (2001). Case-Based Reasoning in the Care of Alzheimer's Disease Patients. In D. W. Aha & I. Watson (Eds.), *Case-Based Reasoning Research and Development*. Lecture Notes in Computer Science, (pp. 702-715). Macura

Martin, E. (2022). La relación entre la Inteligencia artificial y las energías renovables. <https://www.narasolar.com/la-inteligencia-artificial-y-las-energias-renovables/>

Matthew S Lindia (2023). Phenomenology of the Turing test: a Levinasian perspective. *Journal of Communication*. jqad026, <https://doi.org/10.1093/joc/jqad026>

Milidonis, K., Blanco, M. J., Grigoriev, V., Panagiotou, C. F., Bonanos, A. M., Constantinou, M., Pye, J. y Asselineau, C.A. (2021). Review of application of AI techniques to Solar Tower Systems. *Solar Energy*, 224, 500–515. https://www.researchgate.net/publication/352538835_Review_of_application_of_AI_techniques_to_Solar_Tower_Systems

Montes, G.A. y Goertzel, B. (2019). Distributed, decentralized, and democratized artificial intelligence. *Technological Forecasting and Social Sciences*, 141, 354-358. <https://ideas.repec.org/a/eee/tefoso/v141y2019icp354-358.html>

Montoya, B. (2023). El papel de la inteligencia artificial para

combatir ciberataques. Think Big . <https://blogthinkbig.com/inteligencia-artificial-en-la-ciberseguridad>

Montoya, E. A. Q., Colorado, S. F. J., Muñoz, W. Y. C. & Golondrino, G. E. C. (2017). Propuesta de una arquitectura para agricultura de precisión soportada en IoT. *Revista Ibérica de Sistemas e Tecnologías de Informação*, (24), 39-56. <https://pdfs.semanticscholar.org/4c6f/37715e4314c1089a009cb5899cbac9a67ce0.pdf>

Moor, J. (2006). The Dartmouth College Artificial Intelligence Conference: The Next Fifty Years. *AI Magazine*, 27(4), 87-91. <https://ojs.aaai.org/aimagazine/index.php/aimagazine/article/view/1911>

Nexus Integra (s.f). Inteligencia Artificial: el motor detrás de la industria 4.0. <https://nexusintegra.io/es/inteligencia-artificial-el-motor-detras-de-la-industria-4-0/>

Nexus Integra (s.f.b). Smart Operations: el futuro de la industria inteligente. <https://nexusintegra.io/es/smart-operations/>

Observatorio para la Innovación Silvoagropecuaria y la Cadena Agroalimentaria (2023). Innovaciones de la inteligencia artificial en la agricultura de precisión. <https://opia.fia.cl/601/w3-article-116120.html>

Pérez, M. R., Mendoza, M. A. & Suarez, M. J. (2019). Paradigma IoT: desde su conceptualización hacia su aplicación en la agricultura. *Paradigma*, 40(18), 1-8. <https://scholar.google.com/scholar?hl=es&as>

Rambauth-Ibarra, G. (2021). Agricultura de Precisión: La integración de las TIC en la producción Agrícola, *J. Comput. Electron. Sci. Theory Appl*, 38(1), pp. 37–38. <https://doi.org/10.17981/cesta.03.01.2022.04>



Capítulo 2. Aplicaciones de la inteligencia artificial para la solución de problemas reales en el sector de la salud

Capítulo 2. Aplicaciones de la inteligencia artificial para la solución de problemas reales en el sector de la salud

Yanela Rodríguez Alvarez

Yailé Caballero Mota

María Matilde García Lorenzo

Sergio Orlando Escalona González

Resumen

En este capítulo se usaron técnicas de inteligencia artificial, especialmente métodos de aprendizaje automático para adquirir conocimiento sobre enfermedades como la COVID-19, la mediastinitis, la ataxia y la enfermedad renal crónica, que ayude a la predicción del peligro de muerte para pacientes de estas enfermedades. Como resultado, se han determinado los rasgos que resultan relevantes para pronosticar un mayor riesgo de muerte en el paciente (riesgo de letalidad alta), para el caso de la COVID-19 y la enfermedad renal crónica, o pronosticar el desarrollo o progresión de la enfermedad, como en el caso de la mediastinitis y la ataxia, respectivamente. Además, se ha generado conocimiento para predecir estas clases de diagnóstico.

Abstract

In this chapter, various real problems related to the health branch are solved, through artificial intelligence (AI) methods. In this area there are problems with continuous domains and with mixed data. The solutions to the problems are described using the methods proposed by the authors and the validation of the results is carried out through the experimental study in each one of the applications.

Introducción

En el contexto de los sistemas sanitarios y médicos, la gestión de la incertidumbre y la complejidad de las decisiones son retos muy importantes que pueden abordarse mediante técnicas de IA y *Soft Computing* (Corona, 2018; Di Vaio, Palladino *et al.*, 2020). Los estudios sobre aplicaciones médicas de las técnicas inteligentes pueden ser agrupados en cuatro ejes temáticos principales: aprendizaje automático, robot inteligente, tecnología de reconocimiento de imágenes y sistemas expertos (Liu, Rong *et al.*, 2020). Esas técnicas se han identificado como componentes importantes para la gestión de las infecciones asociadas a la asistencia sanitaria con el fin de apoyar la vigilancia, comprender los factores de riesgo, mejorar la estratificación del riesgo de los pacientes e identificar a tiempo las vías de transmisión (Scardoni, Balzarini *et al.*, 2020).

En este capítulo se presenta la aplicación de estas técnicas en diferentes áreas de la salud. El procedimiento seguido para resolver estos problemas deviene de Álvarez (2023):

- Construir el sistema de decisión para el dominio de aplicación.
- Determinar las funciones de comparación locales para los rasgos y la función de semejanza global.
- Establecer los umbrales de similaridad deseados (valores recomendados entre 0,75 y 0,95).
- Calcular la importancia de los rasgos (usando PSO+RST/ PSO+RST+FUZZY).
- Aplicar el método GP/SP.

Adquisición de conocimiento sobre la letalidad de la COVID-19

El propósito de esta investigación fue gestionar la adquisición de conocimiento útil para el pronóstico de un posible desarrollo fatal de la enfermedad utilizando técnicas de IA. Los datos utilizados fueron publicados sobre la situación de

la enfermedad en México (Epidemiología, 2020). Se aplicaron técnicas de aprendizaje automático (AA) para identificar los rasgos determinantes en el pronóstico y extraer conocimiento para esta problemática. Esto puede ser útil para confirmar patrones de comportamiento ya conocidos o eventualmente generar algunos novedosos.

En la situación extrema provocada por el brote de COVID-19, varios ejemplos de estas técnicas demostraron su potencial para la lucha contra la pandemia (Vaishya *et al.*, 2020); (Vinod & Prabakaran, 2020) y se consideran componentes clave en la era posterior a COVID-19 (Bullock, Luccioni *et al.*, 2020; Madurai & Pugazhendhi, 2020). En particular, la IA y el AA se incluyen entre los componentes clave en el contexto de COVID-19, junto con las cadenas de suministro, las comunicaciones, la vigilancia, el Internet de las Cosas, la robótica, los sistemas basados en drones y las aplicaciones móviles (Madurai & Pugazhendhi, 2020). Este interés se ha puesto de manifiesto en la publicación de revisiones, tanto con una perspectiva general (Lalmuanawma *et al.*, 2020; Sipior, 2020; Swapnarekha, Behera *et al.* 2020), como en temas particulares como el análisis de imágenes médicas (Suri, Puvvula *et al.*, 2020; Shaikh *et al.*, 2021) o el diseño de fármacos (Kaushik & Raj, 2020). Algunos usos potenciales de las técnicas inteligentes en el contexto de COVID-19 incluyen el diagnóstico médico, la predicción de la evolución de los pacientes, el diseño de fármacos y procedimientos médicos, la predicción de brotes y la gestión de pandemias (Bullock *et al.*, 2020).

En cuanto a diagnóstico y/o clasificación una aplicación interesante es la clasificación precoz de los pacientes en función de la gravedad de la enfermedad para optimizar el uso del personal médico (Suri, Puvvula *et al.* 2020; Shaikh, Andersen *et al.*, 2021). En la misma línea, otro enfoque singular se basa en el uso de la Lógica Difusa para decidir el procedimiento correcto para pacientes que no se encuentran

en situación grave (Ren *et al.*, 2020). Asimismo, las técnicas de *soft computing* para la gestión de la incertidumbre se han utilizado con éxito para modelar relaciones causales entre protocolos y estrategias de relajación (Ocampo & Yamagishi, 2020).

En general, algunas de las técnicas más utilizadas incluyen Aprendizaje Profundo (Suri, *et al.* 2020, Yan *et al.* 2020), Máquinas de Soporte Vectorial y Bosque Aleatorio (Lalmuanawma, 2020), combinaciones de Lógica Difusa y Metaheurísticas (Yan *et al.*, 2020), etc. En algunos contextos, la combinación de técnicas ha demostrado una ventaja notable (Da Silva *et al.*, 2020). Por ejemplo, las combinaciones de conjuntos difusos y metaheurísticas se han utilizado con éxito en la detección del cáncer (Salmeron & Arévalo, 2020). Otra técnica inteligente que ha demostrado un gran potencial en los sistemas de decisión médica es la Teoría de los Conjuntos Difusos (Yao & Azam, 2014).

Un aspecto importante a tener en cuenta es permitir estrategias de diagnóstico que minimicen la necesidad de pruebas de laboratorio (Oyelade & Ezugwu, 2020). También se ha señalado la importancia de considerar en el diagnóstico aspectos como la edad, el sexo y otras condiciones (diabetes, hipertensión, condiciones previas hepáticas y renales) (Swapnarekha *et al.*, 2020).

Como se ha mencionado antes uno de los aspectos investigados sobre la COVID-19 es su posible tratamiento. Dado el desconocimiento que existía sobre la enfermedad al propagarse a inicios del 2020 para su estudio fue necesario aplicar diferentes técnicas desde perspectivas clínicas, matemáticas y computacionales. Referido a su diagnóstico uno de los aspectos de interés fueron los factores de riesgo para la enfermedad, es decir, cuáles podían incidir en una evolución más agresiva y riesgosa. La COVID1-9 se caracteriza por tener un comportamiento muy variable, desde ser totalmente

asintomática hasta provocar la muerte muy rápidamente. De modo que predecir una posible evolución letal de la enfermedad resulta de mucha relevancia.

Entre los estudios desarrollados con este propósito está Yan, Zhang *et al.*, (2020), que desarrolla un modelo para clasificar la severidad y predecir una tendencia de desarrollo fatal de la enfermedad. En esta investigación se estudiaron varios centenares de casos, descritos inicialmente por más de 300 atributos y aplicó XGBoost (*Extreme Gradient Boosting*), un algoritmo *ensemble*¹ de AA, con un gran potencial de interpretación debido a su sistema de decisión basado en árboles recursivos; e identificó tres atributos relevantes, así como reglas para predecir un desenlace negativo para un paciente. Es decir, con el conocimiento que se genera pueden construirse sistemas inteligentes, y puede utilizarse para incrementar el conocimiento existente sobre el dominio de aplicación; este último aspecto es especialmente útil en dominios donde el conocimiento todavía es limitado, como es el caso de la epidemia de COVID-19.

García (2020) trabajó la adquisición de conocimiento con los datos sobre desarrollo de la COVID-19 en México hasta el 20 de abril de 2020 (disponible en Epidemiología, 2020). Esta base de casos estaba constituida por 8 775 casos positivos a la COVID-19, con dos clases, 8 060 pacientes con No fallecieron y 712 pacientes con Sí fallecieron. Lo cual indica que los datos tienen alto desbalance, es decir, de las dos clases de decisión la segunda está mucho menos representada. Cada caso queda conformado por 16 atributos o rasgos predictores y el rasgo de decisión o clase Letalidad, como se muestra en la Tabla 2.1. Con esa base de casos se realizó el proceso de descubrimiento de conocimiento, en el cual se ejecutaron tres tareas principales: análisis de la consistencia de los datos, selección de rasgos importantes para predecir una evolución desfavorable del paciente y adquisición de conocimiento para la predicción.

¹ Nota: Un ensemble es una combinación de modelos individuales simples que juntos crean uno más potente.

Tabla 2. 1. Descripción de rasgos predictores en los datos sobre el desarrollo de la Covid19 en México

Nº	Nombre del rasgo	Descripción	Formato o fuente
	Origen	La vigilancia centinela se realiza a través del Sistema de Unidades de Salud Monitoras de Enfermedades Respiratorias (USMER).	1 USMER 2 NO USMER
	Sexo	Identifica al sexo del paciente	1 MUJER 2 HOMBRE
	Tipo_paciente	Identifica el tipo de atención que recibió el paciente en la unidad. Se denomina como ambulatorio si regresó a su casa o se denomina como hospitalizado si fue ingresado a hospitalización	1 Ambulatorio 2 Hospitalizado
	Edad	Identifica la edad del paciente	Númérico
	Embarazo	Identifica si la paciente está embarazada	SI_NO
	Diabetes	Identifica si el paciente tiene un diagnóstico de diabetes	SI_NO
	Epoc	Identifica si el paciente tiene un diagnóstico de EPOC	SI_NO
	Asma	Identifica si el paciente tiene un diagnóstico de asma	SI_NO
	Inmupsupr	Identifica si el paciente presenta inmunosupresión.	SI_NO
	Hipertension	Identifica si el paciente tiene un diagnóstico de hipertensión	SI_NO
	Otras_com	Identifica si el paciente tiene diagnóstico de otras enfermedades	SI_NO
	Cardiovascular	Identifica si el paciente tiene un diagnóstico de enfermedades cardiovasculares	SI_NO
	Obesidad	Identifica si el paciente tiene diagnóstico de obesidad	SI_NO
	Renal_cronica	Identifica si el paciente tiene diagnóstico de insuficiencia renal crónica	SI_NO
	Tabaquismo	Identifica si el paciente tiene hábito de tabaquismo	SI_NO
	Otro_caso	Identifica si el paciente tuvo contacto con algún otro caso diagnosticado con SARS CoV-2	SI_NO

Para el análisis de la consistencia se utiliza la medida calidad de la clasificación de la *Rough Sets Theory* (RST) (Caballero *et al.*, 2010), la cual mide el grado de separación entre los objetos de las clases, en este problema las clases {Letal=Sí, Letal=No}. Utilizando la herramienta de análisis de datos basada en conjuntos aproximados (Bello, 2019), al aplicar las medidas a la base de casos estudiada se obtienen los resultados siguientes: Calidad de la aproximación de la clase Letal=No= 0,76, Calidad de la aproximación de la clase Letal= Sí = 0,37 y Consistencia de los datos (calidad de la clasificación) = 0,73.

Estos valores indican que la clase Letal= Sí, además de ser una clase minoritaria en la base de casos, no tiene un valor alto de calidad, en este caso la causa es que los casos que tienen valor de decisión Letal=Sí son similares a otros casos de la clase Letal=No. Por su parte, el valor de la consistencia de los datos se puede considerar como satisfactorio (ni alto, ni bajo) y esto incidirá posiblemente en que la eficacia de los métodos de aprendizaje automático (García *et al.*, 2020).

Para realizar el proceso de aprendizaje se utilizaron formas de representación del conocimiento que permitan a los especialistas comprender el conocimiento descubierto. Una forma de representación del conocimiento que preserva la interpretabilidad son los prototipos. Un prototipo es un caso que sirve de patrón o representante de un conjunto de los casos que están en la base de casos. Los métodos basados en prototipos usualmente forman grupos de casos similares y para cada grupo se determina un prototipo; este puede ser uno de los elementos del grupo o ser construido a partir de los casos que forman el grupo. Se ha utilizado el método de construcción de prototipos basado en la granulación de la base de casos usando una relación de similaridad, en la que se forman clases de similaridad de los elementos de la base de casos y para cada clase se construye un prototipo.

El conocimiento adquirido se formalizó como prototipos para la clase Letalidad=Sí. El proceso de adquisición del conocimiento incluyó la construcción de los prototipos que mejor representan las clases del problema por ser una forma interpretable de formalización del conocimiento para los especialistas. Atendiendo a esto, se empleó el método de construcción de prototipos según el método IMBNPBASIR SEL-CLASS para la clasificación de conjuntos de datos desbalanceados (Rodríguez *et al.*, 2019). En (García *et al.*, 2020) se presentan los mejores prototipos construidos para la clase Letal=Sí

Al realizar la clasificación basada en prototipos es necesario tener presente que estos son representantes de un grupo de casos, de modo que el cálculo de la precisión del conocimiento descubierto —es decir, del conjunto de prototipos— se realiza buscando para cada caso cuál prototipo resulta más similar. Los prototipos seleccionados por el método IMBNPBASIR SEL-CLASS permiten realizar una clasificación con una exactitud de la clasificación general de 79,14. Teniendo en cuenta que la clase relevante en este estudio es la que denota una posible alta letalidad de la enfermedad, solo se presentan los mejores prototipos construidos para la clase Letal=Sí.

Para evaluar la calidad de los prototipos (CNF) se determina el conjunto de casos similares a él, y de ellos cuántos tienen el valor de decisión sí. La calidad o certeza (CNF) del prototipo seleccionado por gránulo se calcula en función del rasgo de decisión del mismo y con respecto a todos los casos del gránulo. Cuando todos los casos son de la misma clase que el prototipo seleccionado, la CNF= 1; en caso contrario y dependiendo de la cantidad de casos que sean de clases diferentes, el valor de CNF será $0 < \text{CNF} < 1$. A mayor CNF, mayor calidad o certeza del conocimiento descubierto, en forma de prototipo.

La CNF de un prototipo es alta si clasifica correctamente las instancias situadas en su gránulo. La precisión se define en (1), donde es el conjunto de instancias situadas en el gránulo de con la misma clase que . Entonces, representa todas las instancias situadas en el gránulo de :

$$\text{CNF}(p_i) = \frac{|V_{ii}|}{|V_i|}$$

El conjunto de prototipos representantes de los diferentes gránulos o grupos es utilizado posteriormente en la clasificación o predicción de un nuevo caso.

Luego la exactitud general del clasificador es calculada como el promedio de las precisiones obtenidas con todos los subconjuntos de prueba.

En el caso del rasgo Edad, el valor representa un promedio de los valores de la edad para los casos comprendidos en el prototipo; por lo tanto, debe verse como un valor aproximado para ese rasgo. Por ejemplo, el prototipo 2 puede describirse en general como “pacientes fallecidos cuya edad promedio es 76 años y tienen diabetes, inmunosupresión, hipertensión, obesidad e insuficiencia renal crónica”.

Además, como resultado de la construcción de la relación de similaridad se obtiene también el peso de los rasgos del dominio de aplicación, mostrando total coincidencia con los métodos de selección y ordenamiento de rasgos anteriormente empleados.

Como un resultado de la construcción de la relación de similaridad se obtiene también el peso de los rasgos del dominio de aplicación; para este problema se calcularon los pesos que se muestran en la Tabla 2.3. Nótese la coincidencia, en general, con los resultados de la selección de rasgos presentada en la Tabla 2.2.

Tabla 2.2 Prototipos obtenidos (García et al., 2020)

No	Prototipo	CNF
P1	EDAD=60, DIABETES=Si, EPOC=No, ASMA=Si, INMUSUPR=No, 0,63 HIPERTENSION=Si, OTRA_COM=No, CARDIOVASCULAR=Si, OBESIDAD=Si, RENAL_CRONICA=No, TABAQUISMO=No, LETAL=Si	
P2	EDAD=76, DIABETES=Si, EPOC=No, ASMA=No, INMUSUPR=Si, 1,00 HIPERTENSION=Si, OTRA_COM=No, CARDIOVASCULAR=No, OBESIDAD=Si, RENAL_CRONICA=Si, TABAQUISMO=No, LETAL=Si	
P3	EDAD=65, DIABETES=Si, EPOC=No, ASMA=No, INMUSUPR=Si, 0,61 HIPERTENSION=Si, OTRA_COM=No, CARDIOVASCULAR=No, OBESIDAD=No, RENAL_CRONICA=No, TABAQUISMO=No, LETAL=Si	
P4	EDAD=65, DIABETES=Si, EPOC=No, ASMA=No, INMUSUPR=No, 0,57 HIPERTENSION=No, OTRA_COM=No, CARDIOVASCULAR=No, OBESIDAD=No, RENAL_CRONICA=Si, TABAQUISMO=No, LETAL=Si	
P5	EDAD=60, DIABETES=Si, EPOC=No, ASMA=No, INMUSUPR=No, 0,55 HIPERTENSION=Si, OTRA_COM=No, CARDIOVASCULAR=No, OBESIDAD=No, RENAL_CRONICA=No, TABAQUISMO=No, LETAL=Si	
P6	EDAD=72, DIABETES=Si, EPOC=No, ASMA=No, INMUSUPR=Si, 0,55 HIPERTENSION=No, OTRA_COM=No, CARDIOVASCULAR=No, OBESIDAD=No, RENAL_CRONICA=Si, TABAQUISMO=No, LETAL=Si	

Tabla 2.3 Peso de los rasgos según el método IMBNPBASIR SEL-CLASS

No. orden	Rasgo	Peso	No. orden	Rasgo	Peso
1	EPOC	0,19	8	Obesidad	0,04
2	Cardiovascular	0,17	9	Origen	0,04
3	Renal_cronica	0,14	10	Tabaquismo	0,02
4	Diabetes	0,11	11	Hipertension	0,02
5	Inmusupr	0,08	12	Asma	0,01
6	Embarazo	0,08	13	Otro_caso	0,00
7	Edad	0,05	14	Otra_com	0,00

Utilización del enfoque de prototipos en los estudios COVID-19 con el conjunto de datos de pacientes cubanos

Apartir de la experiencia mexicana anteriormente comentada se quiso evaluar cómo impactaron —para el caso de Cuba— los antecedentes patológicos personales en la letalidad, así como en los pacientes que desarrollan un estadio grave y crítico de la enfermedad. Para realizar este análisis, se tomaron los datos de pacientes graves, críticos y fallecidos en Cuba, desde la fecha 1/11/2020 al 11/03/2021, reportados en los partes diarios del Ministerio de Salud Pública de Cuba (MINSAP). Se obtuvo una base de casos constituida por 4321 casos, descritos por 30 rasgos, que se describen en la Tabla 2.4.

Fue necesario realizar un filtrado de la información, en el caso de los pacientes que fallecieron y estuvieron reportados como graves y críticos, para eliminar de la base de casos, los registros asociados a este paciente y solo dejar el registro referido a Fallecido. Luego de realizar el filtrado, se obtuvieron 1151 registros, 230 fallecidos, 291 críticos y 630 graves.

Se decidió discretizar el valor del rasgo edad según los grupos etarios que considera el MINSAP en los partes diarios que emitió a la población sobre el comportamiento de la COVID-19 en el país.

Los prototipos obtenidos que describen a los pacientes fallecidos o a aquellos que su estado pudiera conducir a la

Tabla 2.4 Descripción de rasgos predictores para el caso de estudio de Cuba, desde 1/11/2020 al 11/03/2021

No.	Nombre del rasgo	Descripción	Formato fuente
1	Fecha_actualización	Fecha del parte de cierre del día a las 12 de la noche	DD/MM/AAAA
2	Nacionalidad	Identifica la nacionalidad del paciente	Texto
3	Sexo	Identifica el sexo del paciente	F_M
4	Edad	Identifica la edad del paciente	N Numérico
5	Provincia_res	Identifica la provincia de residencia del paciente	Texto
6	Municipio_res	Identifica el municipio de residencia del paciente	Texto
7	Embarazo	Identifica si la paciente está embarazada	SI_NO
8	Diabetes	Identifica si el paciente tiene un diagnóstico de diabetes	SI_NO
9	EPOC	Identifica si el paciente tiene un diagnóstico de EPOC	SI_NO
10	Asma	Identifica si el paciente tiene un diagnóstico de asma	SI_NO
11	Inmunsupr	Identifica si el paciente presenta inmunosupresión	SI_NO
12	Hipertension	Identifica si el paciente tiene un diagnóstico de hipertensión	SI_NO
13	Cardiovascular	Identifica si el paciente tiene un diagnóstico de enfermedades cardiovasculares	SI_NO
14	Obesidad	Identifica si el paciente tiene diagnóstico de obesidad	SI_NO
15	Renal_cronica	Identifica si el paciente tiene diagnóstico de insuficiencia renal crónica	SI_NO
16	Tabaquismo	Identifica si el paciente tiene hábito de tabaquismo	SI_NO
17	Otra_app	Identifica si el paciente tiene diagnóstico de otras enfermedades	SI_NO
18	Estadia_hospitalaria	Identifica la cantidad de días que el paciente ha estado ingresado hospitalariamente	N Numérico
19	Uci	Identifica si el paciente ha estado ingresado en la Unidad de Cuidados Intensivos	SI_NO
20	Fiebre	Identifica si el paciente ha tenido fiebre	SI_NO
21	Arritmia	Identifica si el paciente ha tenido arritmia	SI_NO
22	Intubado	Identifica si el paciente ha estado intubado	SI_NO
23	Distress_resp	Identifica si el paciente ha tenido estrés respiratorio	NO_Ligero_Mod erado_Severo
24	Hemodinámica	Identifica los valores de hemodinámica del paciente	Estable_Inestable
25	Ritmo_diuretico	Identifica el estado del ritmo diurético del paciente ha tenido	Anuria_Conservado_Normal_oligúrico_Oligoanúrico_Poligúrica_Piúrica
26	Gasometría	Identifica los valores de gasometría del paciente	Texto
27	Neumonía	Identifica si el paciente ha tenido neumonía	SI_NO
28	Pcr_evol	Identifica si al paciente se le ha realizado un PCR evolutivo	NO_Positivo_Negativo
29	Otras_comp	Identifica si el paciente ha presentado otras complicaciones	Texto
30	Estado	Identifica el estado del paciente al cierre del parte.	GRAVE_CRITICO_FALLECIDO

letalidad (pacientes graves o críticos), se muestran en la Tabla 2.5, estos tienen una exactitud de la clasificación de 78,79.

La Tabla 2.6 muestra las características más importantes y el peso de cada una de ellas en el estudio de caso de Cuba,

Tabla 2.5 Prototipos obtenidos para el caso de estudio de Cuba

No	Prototipo	CNF
P1	EDAD=entre 20-39, DIABETES=No, EPOC=No, ASMA=Sí, INMUSUPR=No, HIPERTENSION=Sí, OTRA_COM=Sí, CARDIOVASCULAR=No, OBESIDAD=No, RENAL_CRÓNICA=No, TABAQUISMO=No, CLASE=grave	0,75
P2	EDAD= entre 60-69, DIABETES=Sí, EPOC=No, ASMA=No, INMUSUPR=No, HIPERTENSION=Sí, OTRA_COM=No, CARDIOVASCULAR=Sí, OBESIDAD=No, RENAL_CRÓNICA=No, TABAQUISMO=No, CLASE=grave	0,84
P3	EDAD= ≥70, DIABETES=No, EPOC=No, ASMA=Sí, INMUSUPR=No, HIPERTENSION=Sí, OTRA_COM=No, CARDIOVASCULAR=Sí, OBESIDAD=No, RENAL_CRÓNICA=No, TABAQUISMO=No, CLASE=fallecido	0,63
P4	EDAD= ≥70, DIABETES=No, EPOC=No, ASMA=No, INMUSUPR=No, HIPERTENSION=Sí, OTRA_COM=No, CARDIOVASCULAR=Sí, OBESIDAD=No, RENAL_CRÓNICA=No, TABAQUISMO=No, CLASE=grave	0,82
P5	EDAD=entre 40-59, DIABETES=Sí, EPOC=No, ASMA=Sí, INMUSUPR=No, HIPERTENSION=Sí, OTRA_COM=No, CARDIOVASCULAR=Sí, OBESIDAD=Sí, RENAL_CRÓNICA=No, TABAQUISMO=No, CLASE=fallecido	0,67
P6	EDAD= ≥70, DIABETES=Sí, EPOC=Sí, ASMA=No, INMUSUPR=No, HIPERTENSION=Sí, OTRA_COM=Sí, CARDIOVASCULAR=Sí, OBESIDAD=Sí, RENAL_CRÓNICA=No, TABAQUISMO=No, CLASE=crítico	1

Tabla 2.6 Las características más importantes y el peso de cada una de ellas para el estudio de caso de Cuba

Peso	Rasgos
0,18	Edad
0,17	Hipertensión
0,16	Otras_comp
0,14	Cardiovascular
0,09	Diabetes
0,08	Epoc
0,07	Renal_crónica
0,03	Obesidad
0,03	Asma
0,00	Inmusupr
0,00	Tabaquismo

coincidentes con las características más importantes para el estudio de caso de México, mostradas en la Tabla 2.3.

Confirmándose así que las comorbilidades se asocian directamente a mayor riesgo de tener una evolución desfavorable y superior probabilidad de muerte. Tal como lo corrobora el estudio (León Álvarez, Calderón Martínez *et al.*,

2021), desde el punto de vista médico al respecto en Cuba, donde se reflejó además que la hipertensión arterial (HTA) constituye la comorbilidad más frecuente vinculada a la mortalidad por el virus. Tras un análisis de la totalidad de los pacientes confirmados con el nuevo coronavirus en la nación, con 18 años o más, diagnosticados entre el 11 de marzo de 2020 (cuando se detectaron los primeros casos en el territorio) y el 15 de octubre de ese año, se observó que más del 47 por ciento de los fallecidos presentaba HTA. En León, Calderón Martínez et al., 2021) los investigadores precisaron que esos datos están en sintonía con los reportados en el resto del mundo, y el 80 por ciento de los decesos en ese período en el país vinculados al COVID-19 fueron personas de más de 60 años, para una edad promedio de 73. Asimismo, gran parte de ellos tenían dos o tres comorbilidades, donde otras suelen ser la diabetes mellitus y complicaciones vasculares derivadas de ese padecimiento y de la HTA, seguidas de la obesidad y la enfermedad cardiovascular.

Entre otros Antecedentes Patológicos Personales en los pacientes cubanos, se destacan enfermedades cerebrovasculares, cáncer, demencia, enfermedades hepáticas, hipertiroidismo, colitis, Parkinson y hepatopatía crónica. Basándose en esta evidencia inicial, parece que la predicción de la letalidad también está influenciada por las condiciones nacionales o los procedimientos médicos utilizados. La letalidad inferior de la pandemia en Cuba puede provocar una proporción diferente de afecciones en cada clase. Este tipo de comparación entre los prototipos obtenidos mediante técnicas de inteligencia computacional puede ser útil para mejorar nuestra comprensión de esta compleja enfermedad.

Factores de riesgo para la aparición de mediastinitis postoperatoria en cirugía cardíaca

En cirugía cardiororácica la esternotomía media es una de las vías de acceso a los órganos mediastinales más utilizada,

donde la prevención y el tratamiento de sus complicaciones continúa retando a los cirujanos de esta especialidad (Belarj, Dahraoui *et al.*, 2018; Chello, Lusini *et al.*, 2020). La mediastinitis es una complicación severa —aunque poco común— de la esternotomía media, que posee gran importancia por su elevada mortalidad, morbilidad y costes asociados a la asistencia médica (Chello *et al.*, 2020; Pastene, Cassir *et al.*, 2020 y Lin & Jimenez, 2022). Se informa principalmente después de una intervención de revascularización coronaria (asociada a la disección uni o bilateral de la arteria mamaria interna) y se acompaña de un incremento en los costos; los pacientes requieren hospitalizaciones prolongadas, así como diversos procedimientos quirúrgicos complementarios (Pastene *et al.*, 2020).

Una mediastinitis posquirúrgica es la infección profunda de cualquier punto del espacio que esté limitado: delante por el esternón, detrás por la columna vertebral y lateralmente por las pleuras, con el antecedente reciente de cirugía cardíaca o torácica. A pesar de los recientes avances en el tratamiento médico, su incidencia oscila entre 1 % y el 5 %, y la mortalidad asociada a esta complicación, en general, oscila entre 20 % y el 50 % (Chello *et al.*, 2020).

Entre los múltiples factores predisponentes reconocidos se hallan: la obesidad, la diabetes mellitus, las enfermedades pulmonares, las reintervenciones, el bajo gasto cardíaco, la ventilación mecánica, el tiempo quirúrgico y de derivación cardiopulmonar prolongados, entre otros (Pastene *et al.*, 2020 y Mleyhi *et al.*, 2021).

Como se plantea en (Bermúdez-Yera, Naranjo-Ugalde *et al.*, 2019) se conoce que la esternotomía media es la vía de acceso más utilizada en el Hospital Universitario Cardiocentro *Ernesto Guevara* de Santa Clara, Cuba, donde se decide analizar los factores de riesgo presentes en una serie de pacientes, de forma tal que su presencia oriente hacia una

profilaxis correcta y un diagnóstico y tratamiento adecuados de esta grave complicación de la cirugía cardíaca abierta; con el objetivo general de diseñar un modelo predictivo de mediastinitis postoperatoria a través de factores de riesgo.

La incidencia de la mediastinitis postoperatoria, de acuerdo al total de pacientes intervenidos quirúrgicamente del corazón o grandes vasos mediante la esternotomía mediana longitudinal como incisión de abordaje, en el Cardiocentro Ernesto Guevara de Santa Clara, es de 0,98 % desde enero de 2000 hasta mayo de 2019. Valga aclarar que desde 2000 hasta 2008 la incidencia fue superior a 1,60, pero desde 2009 a la actualidad decreció a 0,60.

El estudio de los factores de riesgo de la mediastinitis postoperatoria es muy útil, pues, aunque muchos de ellos son ya conocidos, hay particularidades en cada serie de pacientes. Además, se puede predecir la aparición de esta severa complicación en aquellos que presentan determinados factores de riesgo o combinación de estos, estudiando estadísticamente la relación entre ellos y la presencia de mediastinitis. (Bermúdez *et al.*, 2019 y Mleyhi *et al.*, 2021).

En Bermúdez *et al.* (2019) se presenta modelo predictivo de mediastinitis postoperatoria en cirugía cardiovascular que utiliza regresión logística con ajustes para obtener el modelo. Los autores presentan como resultados de la investigación una incidencia de mediastinitis postoperatoria de 0,98%, con múltiples factores de riesgo predictores, pero los de mayor valor según el modelo obtenido son: enfermedad pulmonar obstructiva crónica, hiperglucemia postoperatoria, tiempo de ventilación artificial mecánica mayor de 24 horas, transfusión de más de 2 unidades de hemoderivados, neumotórax y sepsis endovascular que fueron los que arrojó el modelo de regresión logística.

Por su parte en el estudio de (Konishi *et al.*, 2019) se relacionan variables comunes al análisis realizado por

(Bermúdez *et al.*, 2019) como el tiempo de ventilación artificial mecánica y la estadía en cuidados intensivos, pero asocia otras como el tiempo de cirugía y pinzamiento aórtico.

La problemática se encuentra en adquirir conocimiento cada vez más certero y que acerque a una profilaxis correcta y un diagnóstico y tratamiento oportunos y adecuados de esta letal complicación postquirúrgica. Con el objetivo de resolver dicha problemática se emplean los métodos de selección de prototipos basados en relaciones difusas por las bondades, ventajas y buen desempeño mostrado en secciones anteriores.

Se realizó un estudio analítico retrospectivo de casos y controles en torno a la presencia de mediastinitis postoperatoria en el Hospital Universitario Cardiocentro *Ernesto Guevara* de Santa Clara desde el año 2000 hasta mayo de 2019, con el objetivo de diseñar un modelo predictivo de esta afección mediante factores de riesgo. El universo de estudio lo constituyeron todos los pacientes intervenidos quirúrgicamente por esternotomía mediana longitudinal en dicho centro y periodo, y la muestra quedó constituida por: 45 casos (clase positiva), que fueron la totalidad de los que presentaron el diagnóstico de mediastinitis postoperatoria en el periodo de estudio y 180 controles, 4 por cada caso, pacientes que no tuvieron diagnóstico de mediastinitis postoperatoria (clase negativa), que se constituyeron por muestreo estratificado según tipo de intervención, cercanía en fecha con la aparición de los casos, edad cercana (± 5 años) y el mismo sexo del caso.

El conjunto de datos cuenta con 225 ejemplos de entrenamiento y 40 atributos. El índice de desbalance del conjunto de datos es $IR=4$. Dentro de los 40 atributos se tienen 28 atributos ordinales con valores de “Sí” o “No”. Además, existen 4 atributos nominales, siete atributos numéricos y el atributo clase que toma valores *positive* y *negative* en dependencia de si tiene o no diagnóstico de mediastinitis

postoperatoria. La descripción del conjunto de atributos que caracterizan a los pacientes se muestra en la Tabla 2.7.

Para el atributo edad se decidió discretizar el rango de valores posibles según la escala de la OMS.

Como un resultado de la construcción de la relación de similitud difusa se obtiene también el peso de los rasgos del dominio de aplicación; para este problema se calcularon los pesos que se muestran en la Tabla 2.8.

Tabla 2.7 Descripción del conjunto de datos factores de riesgo para la aparición de mediastinitis postoperatoria en cirugía cardíaca

No.	Atributo	Descripción	Tipo	Rango de valores
	Edad	Edad	nominal	{A,B,C,D,E}
	Genero	Genero	nominal	{F,M}
	Diabetes	Diabetes	nominal	{si,no}
	Obesidad	Obesidad	nominal	{si,no}
	EPOC	EPOC	nominal	{si,no}
	Fuma	Fuma	nominal	{si,no}
	Anemia	Anemia	nominal	{si,no}
	Estadio_preop	Estadía preoperatoria prolongada	Número	[1,36]
	Infecç_preop	Infecciones periféricas sobre todo en el mes previo a la cirugía	nominal	{si,no}
	Tipo_de_cirugía	Tipo de cirugía	nominal	{urgente,electiva}
	Op_inicial	Tipo de Operación Inicial	nominal	{valvular, coronaria}
	Doble_AMI	Empleo de doble arteria mamaria interna en la revascularización	nominal	{si,no}
	Unid_hem_tr	Cantidad de unidades de Transfusión de hemoderivados	número	[0,6]
	Prof_antib	Profilaxis antimicrobiana	nominal	{cefazolina, vancomicina}
	Arritm_trans	Arritmia Transoperatoria	nominal	{si,no}
	Hemor_trans	Hemorragia transoperatoria	nominal	{si,no}
	Bajo_g_trans	Bajo gasto cardíaco transoperatorio	nominal	{si,no}
	Uso_BCPIA	Uso de balón de contrapulsación intraaórtico	nominal	{si,no}
	Estad_pos_Uci	Estadía postoperatoria	número	[1,120]
	Tiemp_VAM	Tiempo de ventilación artificial mecánica	número	[0,92]
	Tiemp_snd_V	Tiempo con sondas vesical	número	[1,120]
	Reint_pos_in	Reintervención en el postoperatorio inmediato	nominal	{si,no}
	Der_pericard	Derrame pericárdico	nominal	{si,no}
	Tapon_car	Taponamiento cardíaco	nominal	{si,no}
	Bajo_gast_post	Bajo gasto cardíaco postoperatorio	nominal	{si,no}
	Arrit_postop	Arritmia Postoperatoria	nominal	{si,no}
	TQ	Traqueostomía	nominal	{si,no}
	Uso_MPP	Uso medicamentos potencialmente peligrosos	nominal	{si,no}
	Disf_VI_post	Disfunción ventricular izquierda postoperatoria	nominal	{si,no}
	IMA_periop	Infarto Agudo de Miocardio perioperatorio	nominal	{si,no}

Tabla 2.8 Peso de los rasgos según el método Imb-SPBASIR-Fuzzy_V2

Atributo	Peso	Atributo	Peso
Estad_pos_Uci	0,054619119	Bajo_gast_post	0,025071247
Obesidad	0,053515022	Der_pericard	0,024946895
Sepsis_HQ	0,048689201	Sepsis_a_ot_niv	0,023911216
Tapon_car	0,048639418	Diabetes	0,017394201
IMA_periop	0,047859918	Arrit_postop	0,015968844
Tiemp_snd_V	0,046080131	Reint_pos_in	0,015777652
Doble_AMI	0,045652113	Edad	0,015045475
Derram_pl	0,043234129	Prof_antib	0,014925308
Uso_BCPA	0,042181076	Hiperglicemia	0,013745137
Cultivo_+	0,03831815	Ins_renal_ag	0,01253217
Arritm_trans	0,035368362	Op_inicial	0,012070088
Genero	0,033904313	Bajo_g_trans	0,01207001
Hemor_trans	0,031995213	Estadio_preop	0,010496395
Fuma	0,030533747	Anemia	0,007306968
TQ	0,029764896	Infec_preop	0,004497214
Sepsis_resp	0,02800097	Neumotorax	0,003545909
Tipo_de_cirugía	0,027533017	EPOC	0,002727766
Disf_VI_post	0,026999231	Uso_MPP	0,001863429
Unid_hem_tr	0,026020886	Tiemp_VAM	0,001761816
Estad_post	0,025433349		

El método de selección de prototipos Imb-SPBASIR-Fuzzy_V2 para este conjunto de datos de factores de riesgo para la aparición de mediastinitis postoperatoria en cirugía cardíaca obtiene un AUC (*Area under the curve*): 89,44 y 99%. Los prototipos obtenidos y su calidad puede observarse en la Tabla 2.9. Para una mejor interpretación del conocimiento descubierto no se muestran los atributos con valor “no”, que en este caso particular significa que no está presente el factor que describe.

Como puede verse los atributos Tipo_de_cirugía, Operación_inicial y Prof_antib no permiten distinguir para los prototipos encontrados si el paciente padece o no de mediastinitis. Esto puede significar que son atributos redundantes y/o irrelevantes por lo que sería conveniente aplicar alguna técnica de reducción de atributos o selección de atributos relevantes. Por otro lado, el conocimiento descubierto muestra que atributos

Tabla 2.9 Prototipos seleccionados con el clasificador *Imb-SPBASIR-Fuzzy_V2* y el conjunto reducido de rasgos

No.	Prototipo	CNF
P1	Edad="entre 61 y 79 años", EPOC="si", Fuma="si", Estadio_preoperatorio="+-5", Tipo_de_cirugía="electiva", Operación_inicial="valvular", Unid_hem_tr= "+4", Prof_antib="Cefazolina", Estad_pos_UCI= "+5", Tiemp_VAM= "+2", Tiemp_snd_V="+6", Sepsis_HQ="si", Estad_post="+32", Hiperglicemia="si", Mediastinitis="positive".	1,0
P2	Edad="entre 18 y 60 años", Genero="Masculino", Estadio_preop="+7", Tipo_de_cirugía="electiva", Operación_inicial="valvular", Unid_hem_tr= "+2", Prof_antib="Cefazolina", Estad_pos_UCI= "+3", Tiemp_VAM= "+1", Tiemp_snd_V="+4", Estad_post="+8", Mediastinitis="negative".	1,0

como: edad, EPOC, tabaquismo, estancia preoperatoria y estancia preoperatoria y postoperatoria, unidades de transfusión sanguínea, tiempo de ventilación mecánica, tiempo de sonda vesical, sepsis de la herida quirúrgica, diabetes mellitus, son factores de riesgo determinantes de mediastinitis postoperatoria.

La clasificación basada en prototipos selecciona los prototipos que son representantes de un grupo de casos, de modo que el cálculo de la precisión del conocimiento descubierto, es decir, del conjunto de prototipos, se realiza buscando para cada caso a cuál prototipo resulta más similar. Para evaluar la calidad de los prototipos (CNF) se determina el conjunto de casos similares a él, y de ellos cuántos tienen el valor de decisión igual al suyo. Estos prototipos pueden interpretarse como los pacientes tipo que representan a cada conjunto (centroide de cada grupo) y a partir de los cuales se puede predecir la clase de un nuevo paciente sin clasificar.

Predicción de la progresión de la ataxia de la marcha

Una de las enfermedades neurodegenerativas es la ataxia, causada principalmente por el daño a la parte del cerebro que se ocupa del movimiento, de la coordinación y del balance. Desde las Ciencias Médicas existen numerosos estudios del tema a nivel nacional e internacional, que tuvieron sus inicios en la segunda mitad del siglo XIX. Recientemente, en la literatura se registran avances significativos desde la inteligencia artificial

orientada al desarrollo de herramientas hardware y software para la detección temprana y evaluación de enfermedades neurodegenerativas (Yang, Zhong *et al.*, 2014, Peng, Chen *et al.* 2021, Tartarisco, Bruschetta *et al.*, 2021, Kaya, Karakuş *et al.*, 2022, Ru, Li *et al.* 2022). Estos estudios han sido motivados por la importancia de establecer procedimientos que permitan detectar síntomas tempranos de dicha enfermedad en aras de propiciar la intervención terapéutica mediante un oportuno proceso rehabilitador (Caparrós *et al.*, 2021).

Con la colaboración de la Academia de Ciencias de Cuba, el Centro de Neurociencias de Cuba y el Centro para la Investigación y Rehabilitación de las Ataxias Hereditarias (CIRAH), que radica en Holguín, se ha emprendido el estudio de técnicas inteligentes en la adquisición de conocimiento sobre las ataxias, por la gravedad que presenta esta enfermedad y su prevalencia de la misma en el país, sobre todo en esta región oriental.

La Ataxia Espinocerebelosa tipo 2 (SCA2) constituye una enfermedad hereditaria, con carácter degenerativo del sistema nervioso, presentación clínica progresiva e invalidante y perteneciente a la clase de las ataxias de tipo autosómico dominante. Se caracteriza por una amplia gama de rasgos cerebelosos y no cerebelosos debidos a una patología que afecta al cerebelo, el tronco encefálico, la médula espinal, los nervios periféricos y otras estructuras. Las manifestaciones cerebelosas más frecuentes son ataxia de la marcha, inestabilidad postural, dismetría, disartria y disdiadococinesia. Estas características suelen ir acompañadas de neuropatía periférica, calambres musculares, trastornos cognitivos, anomalías del sueño, disautonomía y otras anomalías menos frecuentes (Gonzalez *et al.*, 2021)

Cuba tiene las mayores tasas de incidencia y prevalencia en el mundo de la SCA2. La provincia de Holguín exhibe la más alta prevalencia con valor de 47,86 por cada 100000 habitantes (Gonzalez *et al.*, 2021). Dicha región muestra los mayores porcentajes tanto de enfermos como de familias con

ataxias pertenecientes a la forma molecular SCA2 (Hersheson *et al.*, 2012).

Medir la progresión de la ataxia de la marcha —un signo clínico determinante— es complejo usando las técnicas médicas y cuantitativas habituales, no obstante, los avances en ciencias de datos ponen a disposición un grupo de técnicas que pueden ser fácilmente empleadas para el análisis de la marcha.

El estudio presentado en este capítulo, es el primer acercamiento, en Cuba, a la resolución de este problema empleando minería de datos y aprendizaje automático sobre los datos en bruto. Otros trabajos en el área han utilizado técnicas inteligentes, pero orientadas al procesamiento de imágenes y señales como las resonancias magnéticas y los electrooculogramas fundamentalmente (Becerra *et al.*, 2013; Cabeza *et al.*, 2021; Becerra, 2022).

En esta investigación se utilizan técnicas inteligentes para la predicción de la progresión de la SCA2 analizando datos de la marcha de un grupo de pacientes y controles, a través de los métodos basados en prototipos.

Con los resultados obtenidos se logra revelar entre todas las variables que caracterizan la marcha de una persona, aquellas que se afectan proporcionalmente al grado de avance y progresión de la SCA2, es decir, aquellos rasgos que permiten una mejor separabilidad entre los grupos enfermos y controles; reduciendo la dimensión del problema y abriendo el camino a futuros trabajos. Adicionalmente, se emplearon técnicas de visualización de información para mostrar la distribución de los datos según los casos analizados. Luego se utilizaron técnicas de pesado de rasgos y selección de atributos relevantes y finalmente se procedió a presentar los prototipos de pacientes que mejor representan a los enfermos que padecen SCA2.

Las evaluaciones instrumentales de la marcha se realizaron utilizando un conjunto de 6 unidades de medición inercial corporales (The Opal TM, Mobility Lab (APDM Inc.), versión 4.6.3.v 20170301-0400, APDM Inc., Portland, OR, EE.UU.) situadas en ambas manos, ambos pies, el esternón y la región lumbar (L5). (Figura 1).

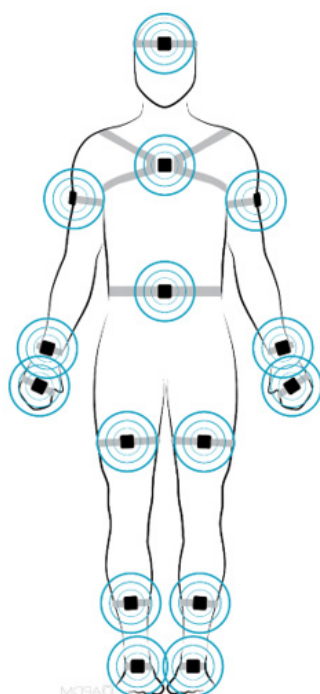


Figura 1.

Los protocolos experimentales para la evaluación de la marcha incluían dos tareas: (1) marcha natural y (2) marcha en tándem hacia delante. En la tarea de marcha natural, se pidió a los sujetos que caminaran normalmente (como lo hacen en su vida diaria) una distancia de 10 m en una trayectoria rectilínea y luego giraran y volvieran a caminar la misma distancia. En la tarea de marcha en tándem formal, se indicó a los pacientes que caminaran colocando un pie tocando delante del otro en una trayectoria rectilínea para completar 10 pasos sin apoyo.

Las medidas de la marcha, Tabla 2.10, se calcularon utilizando el software *Mobility Lab*.

Las medidas de la marcha incluían las métricas de las extremidades inferiores (velocidad de la marcha, longitud de la zancada, elevación del pie a mitad del balanceo, ángulo

de la punta del pie hacia fuera, ángulo de la punta del pie hacia fuera, periodo de doble apoyo y periodo de balanceo), métricas lumbares/troncales (rangos de movimiento en los planos coronal, sagital y transversal en las regiones lumbar y del tronco) y el rango de movimiento del brazo (ROM). Para todas estas medidas, se obtuvieron las medias intrasujeto y la variabilidad paso a paso para múltiples zancadas. Las medias reflejan las características cinemáticas de la marcha, mientras que las variabilidades paso a paso consisten en las desviaciones estándar intrasujeto de cada medida cinemática a lo largo de los ciclos de la marcha durante la tarea experimental (Sangeux, Passmore *et al.*, 2016; Rennie, Löfgren *et al.*, 2018).

Tabla 2.10 Medidas de la marcha calculadas con el software Mobility Lab

No.	Atributos
Marcha - Miembro inferior -	
1	Cadencia (pasos/min) [media]
2	Cadencia (pasos/min) [std]
3	Doble apoyo (%GCT) [mean]
4	Doble apoyo (%GCT) [std]
5	Elevación a medio recorrido (cm) [mean]
6	Elevación a medio recorrido (cm) [std]
7	Duración del ciclo de la marcha (s) [mean]
8	Duración del ciclo de la marcha (s) [std]
9	Velocidad de la marcha (m/s) [mean]
10	Velocidad de la marcha (m/s) [std]
11	Variabilidad del paso lateral (cm)
12	Circunducción (cm) [mean]
13	Circunducción (cm) [std]
14	Número de ciclos de la marcha (#)
15	Ángulo R de impacto del pie (grados) [mean]
16	Ángulo R de impacto del pie (grados) [std]
17	Ángulo R de desviación de la punta del pie (grados) [mean]
18	Ángulo R de desviación de la punta del pie (grados) [std]
19	Apoyo de un solo miembro (%GCT) [mean]
20	Apoyo de un solo miembro (%GCT) [std]
21	Postura (%GCT) [mean]

22	Postura (%GCT) [std]
23	Duración del paso (s) [mean]
24	Duración del paso (s) [std]
25	Longitud de zancada (m) [mean]
26	Longitud de zancada (m) [std]
27	Balanceo (%GCT) [mean]
28	Balanceo (%GCT) [std]
29	Apoyo doble terminal R (%GCT) [mean]
30	Apoyo doble terminal R (%GCT) [std]
31	Ángulo de salida de la punta del pie (grados) [mean]
32	Ángulo de salida de la punta del pie (grados) [std]
	Marcha – Lumbar
33	Rango de movimiento coronal (grados) [mean]
34	Rango de movimiento coronal (grados) [std]
35	Rango de movimiento sagital (grados) [mean]
36	Rango de movimiento sagital (grados) [std]
37	Rango de movimiento transversal (grados) [mean]
38	Rango de movimiento transversal (grados) [std]
	Marcha – Tronco
39	Rango de movimiento coronal (grados) [mean]
40	Rango de movimiento coronal (grados) [std]
41	Rango de movimiento sagital (grados) [mean]
42	Rango de movimiento sagital (grados) [std]
43	Rango de movimiento transversal (grados) [mean]
44	Rango de movimiento transversal (grados) [std]
	Marcha - Miembro superior -
45	Velocidad de balanceo del brazo (grados/s) [mean]
46	Velocidad de balanceo del brazo (grados/s) [std]
47	Rango de movimiento del brazo (grados) [mean]
48	Rango de movimiento del brazo (grados) [std]
	Giros
49	Ángulo (grados) [mean]
50	Ángulo (grados) [std]
51	Duración (s) [mean]
52	Duración (s) [std]
53	NG (#)
54	Velocidad de giro (grados/s) [mean]
55	Velocidad de giro (grados/s) [std]
56	Pasos en el giro (#) [mean]
57	Pasos en el giro (#) [std]
58	Clase

Tabla 2.11 Medidas de la marcha calculadas con el software Mobility Lab tras la combinación de medidas

No.	Medidas de la marcha
1	Marcha - Miembro inferior
2	Variabilidad de la cadencia (pasos/min)
3	Variabilidad del doble apoyo (%GCT)
4	Variabilidad de la elevación a media oscilación (cm)]
5	Variabilidad de la duración del ciclo de la marcha (s)
6	Variabilidad de la velocidad de la marcha (m/s)
7	Variabilidad del paso lateral (cm)
8	Variabilidad de la circunducción (cm)
9	Número de ciclos de la marcha (#)
10	Variabilidad del Ángulo R de impacto del pie (grados)
11	Variabilidad del Ángulo R de salida de la punta del pie (grados)
12	Variabilidad de apoyo de un miembro (%GCT)
13	Variabilidad de la postura (%GCT)
14	Variabilidad de la duración del paso (s)
15	Variabilidad de la longitud de zancada (m)
16	Variabilidad del balanceo (%GCT)
17	Variabilidad del doble apoyo terminal (%GCT)
18	Variabilidad del ángulo de salida de la punta del pie (grados)
19	Marcha – Lumbar
20	Variabilidad del rango de movimiento coronal (grados)
21	Variabilidad del rango de movimiento sagital (grados)
22	Variabilidad del rango de movimiento transversal (grados)
23	Marcha – Tronco
24	Variabilidad del rango de movimiento coronal (grados)
25	Variabilidad del rango de movimiento sagital (grados)
26	Variabilidad del rango de movimiento transversal (grados)
27	Marcha - Miembro superior
28	Variabilidad de la Velocidad L de balanceo del brazo (grados/s)
29	Variabilidad del rango de movimiento del brazo (grados)
30	Giros
31	Variabilidad del ángulo (grados)
32	Variabilidad de la duración (s)
33	NG (#)
34	Variabilidad de la velocidad de giro (grados/s)
35	Variabilidad de los Pasos en los giros (#)
36	Clase

El conjunto de datos cuenta con la información de las 122 personas (52 pacientes sanos o ejemplos de control y 70 pacientes enfermos que padecen SCA2) incluidas en el

estudio longitudinal caracterizadas por 58 medidas de la marcha.

La variabilidad de la marcha fue reportada en (Galna, Lord *et al.*, 2013; Buckley, Mazzà *et al.*, 2018) como Coeficiente de Variación (CV) definido como Desviación Estándar (SD)/media (%), o Desviación Estándar combinada (cSD) definida como la raíz cuadrada de la varianza media de los pasos izquierdo y derecho (cm). En este estudio la variabilidad de la marcha que utilizan para reducir significativamente la dimensión del problema. También permite aumentar su interpretación y mejorar la precisión y exactitud de los métodos utilizados. Cada medida de la marcha representada por las medias dentro de un sujeto y la variabilidad paso a paso para múltiples zancadas (SD) se sustituye por su variabilidad, Tabla 2.13, por ejemplo, para las medidas Giros - Pasos en Giro (#) [media] y Giros - Pasos en Giro (#) [SD] se calcula y cambia por Giros - Pasos en Giro Variabilidad (#) (Giros - Pasos en Giro (#) [SD]/ Pasos en Giro (#) [media]).

Para dar solución a este problema se utiliza el método propuesto SPBASIRFuzzy, que es un clasificador basado en selección de prototipos, presentado anteriormente. Este es un método de reducción de datos que combina de la teoría de los conjuntos aproximados extendida, los conjuntos difusos y la metaheurística PSO para obtener mejores resultados en la clasificación de nuevos casos o ejemplos. Con este método se obtiene un AUC: 91,02 y : 99%.

Donde Red (Bermejo and Cabestany 2000) es el coeficiente de reducción; esta medida, indica el número de objetos que es generado/seleccionado por los algoritmos en estudio (relación entre el tamaño del conjunto de prototipos (P) generados/seleccionados y la cantidad de ejemplos (X) del conjunto original). Se calcula utilizando la expresión .

Como un resultado de la construcción de la relación de similitud se obtiene también el peso de los rasgos del dominio de aplicación; para este problema se calcularon los pesos que se muestran en la Tabla 2.12.

Tabla 2.12 Peso de los rasgos según el método SPBASIRFuzzy

Ranking	Atributo	Peso	Ranking	Atributo	Peso
1	27	0,069	16	22	0,030
2	29	0,069	17	21	0,027
3	6	0,066	18	25	0,022
4	15	0,064	19	20	0,021
5	10	0,061	20	7	0,021
6	30	0,059	21	23	0,017
7	19	0,058	22	18	0,014
8	3	0,053	23	11	0,013
9	8	0,051	24	5	0,011
10	2	0,050	25	12	0,011
11	14	0,048	26	1	0,008
12	17	0,043	27	4	0,005
13	26	0,041	28	28	0,002
14	16	0,032	29	9	0,001
15	24	0,031	30	13	1,34E-04

Es posible utilizar estos pesos, siguiendo la lógica del estudio realizado, para reducir la dimensión del problema que se analiza, eliminar características redundantes y mejorar la interpretación de los resultados. Se utiliza como función de evaluación en este caso el peso de los atributos, quedando seleccionados aquellos con un peso $> 0,06$. Luego de realizar la selección de atributos SPBASIRFuzzy con el conjunto de datos reducido se obtiene un AUC: 92,60 y : 99%. Los prototipos obtenidos mantienen su calidad, como puede observarse en la Tabla 2.13.

Tabla 2.13 Prototipos seleccionados con el clasificador SPBASIRFuzzy y el conjunto reducido de rasgos

No.	Prototipo	CNF
P1	AtributoNo.6=4.575, AtributoNo.10=0.078, AtributoNo.15=0.034 , Atributo No.27=0.159, AtributoNo.29=0.109, Clase=SCA2	1
P2	AtributoNo.6= 2.3, AtributoNo.10=0.034, AtributoNo.15=0.017 AtributoNo.27=0.126, AtributoNo.29=0.083, Clase=CON	1

Como los atributos son de tipo numérico el valor dado representa un promedio de los valores para los casos comprendidos en el prototipo; por lo tanto, debe verse como

un valor aproximado para ese rasgo. Por ejemplo, el prototipo 1 puede describirse en general como “pacientes con SCA2 cuya variabilidad del paso lateral promedio es 4,575 cm, la variabilidad del Ángulo R de salida de la punta del pie promedio es 0,078 grados, la variabilidad del balanceo promedio es 0,034 grados, la variabilidad de la velocidad de giro promedio es 0,126 s y la variabilidad de los pasos en un giro promedio es de 0,083”.

Si se analiza la Tabla 2.11 Medidas de la marcha calculadas con el software *Mobility Lab* tras la combinación de medidas se comprende que los atributos evalúan parámetros relacionados con la marcha a partir del comportamiento de las diferentes partes del cuerpo y se pueden agrupar de la siguiente forma: del 1 al 17, extremidades inferiores, del 18 al 20 lumbar, del 21 al 23 tronco, 24 y 25 extremidades superiores y del 26 al 30 analizando los giros.

Como resultado de la aplicación de las técnicas inteligentes de selección difusa de prototipos y la importancia de los atributos obtenida también con estas técnicas y como se muestran en las tablas, Tabla 2.12 y la Tabla 2.13 respectivamente, de las 17 medidas de la marcha teniendo en cuenta las extremidades inferiores, las que mayor información nos brindan en la predicción de progresión de la SCA2 son: Variabilidad del paso lateral, Variabilidad del Ángulo R de salida de la punta del pie y Variabilidad del balanceo.

No resalta significativamente ninguna medida para este conjunto de entrenamiento y estos algoritmos inteligentes, teniendo en cuenta las medidas que reflejan los movimientos de las partes del cuerpo lumbar, tronco y extremidades superiores. Finalmente, de las 5 medidas que se estudian de los giros sobresalen: Variabilidad de la velocidad de giro y Variabilidad de los pasos en un giro. Por todo lo anterior podemos concluir que las técnicas inteligentes utilizadas permiten reducir la dimensionalidad del problema y mejorar la interpretabilidad del mismo. Es posible con solo seis variables realizar el pronóstico de la progresión de la ataxia de la marcha que dependía inicialmente de 58 variables, con altos valores de precisión

(aproximadamente 93%, AUC= 92,60 y) y un porciento de los ejemplos totales de la base de casos (R= 99).

Predicción de la mortalidad en pacientes con enfermedad renal crónica

La enfermedad renal crónica se ha convertido en la tercera causa de muerte de más rápido crecimiento y se estima que para el 2040 sea la quinta causa de años de vida potencialmente perdidos. A nivel mundial, alrededor de 850 millones de personas conviven con la enfermedad renal, de estos, 4 millones reciben terapias de reemplazo renal para una tasa bruta de 298,4 pacientes por millón de habitantes (Bello *et al.*, 2022).

En Cuba, durante el año 2021 la enfermedad renal crónica fue la causa número catorce de mayor mortalidad. Se reportaron 1 576 muertes, 398 más que el año 2020, para una tasa bruta de 14,1 y ajustada de 7,0 (Salud. 2022). La mortalidad por enfermedad renal crónica en el territorio nacional ha aumentado considerablemente con respecto a años anteriores por lo que constituye una necesidad científica la realización de investigaciones que contribuyan a identificar precozmente factores pronósticos. El envejecimiento poblacional y el aumento de enfermedades como la hipertensión arterial y la diabetes mellitus son desafíos importantes para nuestro país en aras de obtener mejores resultados en la salud renal de los individuos.

Las Tunas, Cuba —según reportes de la Dirección Provincial de Salud— se sitúa entre las cinco primeras del país de mayor prevalencia con 1 797 pacientes para una tasa de 2,98 por cada mil pacientes. En cuanto a la mortalidad, durante el año 2021 se reportaron 29 muertes.

En el Hospital General Docente *Dr. Ernesto Guevara de la Serna* existe un Servicio de Nefrología que brinda atención a los pacientes con enfermedad renal crónica en hemodiálisis.

Por las particularidades del paciente nefrópata, en todos los servicios de hemodiálisis se almacenan con sistematicidad grandes cantidades de datos de diferentes variables.

El análisis de estas variables con la implementación de métodos de la inteligencia artificial generaría un sistema de alerta temprana que podría contribuir a la aplicación de intervenciones futuras que logren modificar el escenario en la mortalidad de los pacientes.

Los factores causales de la enfermedad renal crónica varían mucho a nivel internacional. La diabetes mellitus y la hipertensión arterial se perfilan como las principales causas. Aunque la enfermedad renal poliquística, las glomerulopatías primarias y secundarias, la uropatía obstructiva, las enfermedades hereditarias y otras etiologías de origen desconocido también influyen en el daño renal (Rahman, 2022).

Para el diagnóstico de la enfermedad renal crónica tradicionalmente el funcionamiento renal se ha estimado mediante la tasa de filtración glomerular y la albuminuria. Estos dos parámetros presentan numerosas ventajas y pocas limitaciones, por lo que han sido determinaciones analíticas vanguardias en el diagnóstico y pronóstico de la enfermedad renal crónica (Landler *et al.*, 2022).

Cuando la enfermedad renal crónica progresa hacia el estadio final, generalmente el paciente es sometido a terapias dialíticas. La hemodiálisis es la principal modalidad que se realiza. La eficacia de la purificación de sangre por hemodiálisis depende de la condición del acceso vascular, el tipo de membrana de diálisis, el tiempo de tratamiento y la tasa de flujo sanguíneo. A pesar de los beneficios que ofrece el proceder dialítico, se suman al paciente con enfermedad renal factores pronósticos de mortalidad dependientes de la hemodiálisis (Hadaad *et al.*, 2022).

Varios factores de riesgo se asocian a la mortalidad de los pacientes en hemodiálisis, entre ellos: la edad avanzada, la presencia de enfermedades cardiovasculares y otros factores modificables como la presencia de catéter centrovenoso

central, la anemia y la hipervolemia. Además, el diagnóstico tardío de la enfermedad renal y los años de evolución en la terapia dialítica pueden influir potencialmente en el resultado (Barra *et al.*, 2022).

La adherencia a la hemodiálisis es un elemento de gran importancia para evaluar de forma rigurosa las intervenciones médicas destinadas a mejorar los resultados en diálisis. Un meta-análisis realizado por Sousa y colaboradores (Sousa *et al.*, 2022) con el objetivo de explorar la validez de criterio de las variables influyentes en la adherencia a la diálisis mostró que existía una fuerza de asociación mediana-alta en la ganancia de peso interdialítica, el fósforo sérico y los niveles de potasio.

La naturaleza asintomática de la enfermedad renal crónica provoca que muchos pacientes evolucionen a etapas finales de la enfermedad. Es indispensable contar con herramientas que puedan predecir la progresión para que los pacientes se beneficien de estrategias terapéuticas oportunas para evitar complicaciones y la mortalidad prematura.

En los pacientes con enfermedad renal crónica avanzada se instalan factores pronósticos de progresión de la enfermedad. El control de la presión arterial, cifras adecuadas de glucemia, el tratamiento de la albuminuria, evitar toxinas nefrotóxicas, modificar la obesidad, tratar la dislipidemia, ajustar los fármacos y disminuir el riesgo cardiovascular se consideran medidas de protección.

Los factores de progresión de la enfermedad renal crónica demandan detección temprana. Una encuesta realizada para determinar tendencias actuales del progreso de la enfermedad renal demostró que los métodos de la inteligencia artificial predecían correctamente la evolución de la enfermedad. Se analizaron modelos de imputación de valores faltantes, reducción de características, selección de características, combinación de clasificadores, lógica difusa y enfoques de aprendizaje automático y profundo (Malakar *et al.*, 2022).

En los últimos 30 años la inteligencia artificial en el campo de la Nefrología ha sido poco estudiada. Entre los años 1990

y 2019 solamente se visualizaron 218 artículos científicos que hacían referencia a las enfermedades renales y los métodos de la inteligencia artificial. Del total de estudios mencionados, 8 artículos abordaban la mortalidad en la enfermedad renal (Kanda *et al.*, 2020, Garcia *et al.*, 2021, Park *et al.*, 2021, Kanda *et al.*, 2022).

En Cuba no se ha implementado hasta el momento herramientas de inteligencia artificial en la enfermedad renal crónica. En esta investigación se propone la implementación de la inteligencia artificial en la predicción de mortalidad en pacientes con enfermedad renal crónica, con vistas a automatizar y optimizar el pronóstico del paciente nefrópata y estratificar el riesgo de mortalidad.

Análisis de los factores pronósticos de mortalidad en los pacientes con enfermedad renal crónica en hemodiálisis en el Hospital General Docente *Dr. Ernesto Guevara de la Serna*

Existe un estudio no experimental, longitudinal de cohorte retrospectivo en los pacientes con enfermedad renal crónica en hemodiálisis en el Hospital General Docente *Dr. Ernesto Guevara de la Serna* durante el período del 1 de enero de 2017 al 31 de diciembre de 2021. La población estuvo constituida por 169 pacientes con enfermedad renal crónica que recibió hemodiálisis en la institución y período declarados anteriormente. De los 169 pacientes 62 pertenecen a la clase minoritaria Mortalidad=Sí y 107 pertenecen a la clase mayoritaria Mortalidad=No, para un IR= 1,73. Se analizaron un total de 30 atributos divididos en variables epidemiológicas, clínicas, paraclínicas y dependientes del tratamiento como pueden observarse en la Tabla 2.14.

Los prototipos obtenidos con el método IMBNPBASIR SEL-CLASS-V2 y que describen a los pacientes fallecidos, se muestran en la Tabla 2.15, los mismos tienen un AUC de 77,51.

La Tabla 2.16 muestra las características más importantes y el peso de cada una de ellas en el estudio de caso de los pacientes con enfermedad renal crónica.

Tabla 2.14 Descripción de rasgos predictores en los datos de los pacientes con enfermedad renal crónica en hemodiálisis

No.	Atributo	Descripción	Tipo	Rango Valores
Epidemiológicas				
	Sexo	Identifica el sexo del paciente	Nominal	{F,M}
	Edad	Identifica la edad del paciente	Numérico	[23,91]
	ECV	Identifica si el paciente padece de enfermedad cardiovascular	Nominal	{si, no}
	ERPAD	Identifica si el paciente padece de enfermedad renal poliquística	Nominal	{si, no}
	DM	Identifica si el paciente padece de diabetes mellitus	Nominal	{si, no}
	Gx	Identifica si el paciente padece de Glomerulopatía	Nominal	{si, no}
	HTA	Identifica si el paciente padece de Hipertensión arterial	Nominal	{si, no}
	Tabaquismo	Identifica si el paciente fuma	Nominal	{si, no}
	Dislipidemia	Identifica si el paciente padece de Dislipidemia	Nominal	{si, no}
	HVI	Identifica si el paciente padece de hipertrofia ventricular izquierda	Nominal	{si, no}
	Obesidad	Identifica si el paciente padece de Obesidad	Nominal	{si, no}
	Desnutrición	Identifica si el paciente padece de Desnutrición	Nominal	{si, no}
	Depresión	Identifica si el paciente padece de Depresión	Nominal	{si, no}
Clínicas y paraclínicas				
	FC	Frecuencia cardíaca del paciente según valor determinado al examen físico.	Numérico	[44,96]
	Hb	Hemoglobina del paciente según valor determinado al examen físico.	Numérico	[6,19]
	Albúmina	Albúmina del paciente según valor determinado al examen físico.	Numérico	[19,55]
	Urea	Urea del paciente según valor determinado al examen físico.	Numérico	[9,27]
	Creatinina	Creatinina del paciente según valor determinado al examen físico.	Numérico	[396,2075]

Filtrado	Filtrado glomerular (MDRD-4) del paciente según valor determinado al examen físico.	N Numérico	[3,17]
Ácido úrico	Ácido úrico del paciente según valor determinado al examen físico.	N Numérico	[41,726]
Glucemia	Glucemia del paciente según valor determinado al examen físico.	N Numérico	[1,10]
Colesterol	Colesterol del paciente según valor determinado al examen físico.	N Numérico	[1,7]
Calcio	Calcio del paciente según valor determinado al examen físico.	N Numérico	[1,3]
Dx_ESTADIO_5	Diagnóstico de la enfermedad renal crónica en el estadio 5	N Nominal	{si, no}
Dependientes del tratamiento			
Atención_nefro_previa	Identifica si el paciente tuvo atención nefrológica previa al inicio del programa en hemodiálisis.	N Nominal	{si, no}
Tiempo_HD	Identifica el tiempo en meses de cada paciente desde que comenzó el programa en hemodiálisis.	N Numérico	
Inicio_tardío	Identifica si el paciente, independientemente de las causas, comenzó el programa en hemodiálisis posterior a seis meses de ser prescrito por el especialista.	N Nominal	{si, no}
CCV	Identifica la presencia de catéter centrovenoso para el acceso vascular en la hemodiálisis.	N Nominal	{si, no}
GPI_Excesiva	Identifica si la ganancia de peso obtenida por cada paciente después del inicio de la hemodiálisis (interdialítica) es excesiva	N Nominal	{si, no}
Hepatitis	Identifica el contagio de hepatitis secundario al proceder de hemodiálisis en el paciente.	N Nominal	{si, no}
Mortalidad	Estado del paciente al culminar la investigación	N Nominal	{si, no}

Tabla 2.15 Prototipos obtenidos para el caso de los pacientes con enfermedad renal crónica

No	Prototipo	CNF
P1	Sexo=F, Edad=65, ECV=sí, DM=Sí, HVI=Sí, FC=67, Hb=12, Albúmina=38, Urea=11, Creatinina=952, Filtrado=5, Ácido_úrico=266, Glucosa=3, Colesterol=5, Dx_estadio_5=Sí, Atención_Nefro_previa=Sí, Calcio=2, Tiempo_HD=17, CCV=Sí, Mortalidad=Sí	0,74
P2	Sexo=F, Edad=54, DM=Sí, FC=72, Hb=11, Albúmina=31, Urea=13, Creatinina=975, Filtrado=5, Ácido_úrico=360, Glucosa=3, Colesterol=4, Atención_Nefro_previa=Sí, Calcio=2, Tiempo_HD=10, CCV=Sí, Mortalidad=Sí	0,58
P3	Sexo=F, Edad=64, ECV=Sí, DM=Sí, HTA=Sí, Tabaquismo=Sí, FC=44, Hb=8, Albúmina=25, Urea=16, Creatinina=965, Filtrado=5, Ácido_úrico=352, Glucosa=3, Colesterol=4, Dx_estadio_5=Sí, Atención_Nefro_previa=No, Calcio=2, Tiempo_HD=24, Inicio_tardío=sí, GPI_excesiva=Sí, Hepatitis=Sí, Mortalidad=Sí	1,00
P4	Sexo=F, Edad=65, ECV=Sí, DM=Sí, HVI=Sí, FC=67, Hb=12, Albúmina=38, Urea=11, Creatinina=952, Filtrado=5, Ácido_úrico=266, Glucosa=3, Colesterol=5, Dx_estadio_5=sí, Atención_Nefro_previa=Sí, Calcio=2, Tiempo_HD=17, CCV=Sí, Mortalidad=Sí	0,70
P5	Sexo=M, Edad=72, DM=Sí, FC=75, Hb=11, Albúmina=23, Urea=17, Creatinina=1011, Filtrado=6, Ácido_úrico=295, Glucosa=5, Colesterol=4, Atención_Nefro_previa=Sí, Calcio=2, Tiempo_HD=8, Inicio_tardío=Sí, GPI_excesiva=Sí, Mortalidad=Sí	1,00
P6	Sexo=M, Edad=82, HVI=Sí, FC=58, Hb=12, Albúmina=28, Urea=16, Creatinina=963, Filtrado=6, Ácido_úrico=286, Glucosa=4, Colesterol=5, Atención_Nefro_previa=No, Calcio=2, Tiempo_HD=50, Hepatitis=Sí, Mortalidad=Sí	0,59
P7	Sexo=M, Sexo=39, DM=Sí, Desnutrición=Sí, FC=69, Hb=8, Albúmina=20, Urea=18, Creatinina=1236, Filtrado=5, Ácido_úrico=386, Glucosa=4, Colesterol=4, Atención_Nefro_previa=Sí, Calcio=2, Tiempo_HD=2, CCV=Sí, Mortalidad=Sí	0,57

Teniendo en cuenta los pesos obtenidos para los atributos, los prototipos obtenidos y el criterio de los expertos y especialistas médicos sobre la capacidad de algunos atributos

Tabla 2.16 Las características más importantes y el peso de cada una de ellas para el estudio de caso de Cuba

Peso	Rasgos	Peso	Rasgos
SEXO	0,4143	OBESIDAD	0,0230
EDAD	0,0332	COLESTEROL	0,0228
ECV	0,0325	CREATININA	0,0205
ALBÚMINA	0,0322	UREA	0,0174
HEPATITIS	0,0321	CALCIO	0,0153
GLUCOSA	0,0307	DEPRESIÓN	0,0151
Gx	0,0288	Dx_ESTADIO_5	0,0123
INICIO_TARDÍO	0,0278	Atención_Nefro_previa	0,0121
TIEMPO_HD	0,0277	CCV	0,0120
Hb	0,0277	TABAQ	0,0113
FILTRADO	0,0267	HVI	0,0097
DM	0,0245	DISLIP	0,0092
GPI_EXCESIVA	0,0244	HTA	0,0040
ACIDO_URICO	0,0239	DESNUTR	0,0032
ERPAD	0,0230	FC	0,0026

de discriminar en el desenlace mortal o no de la enfermedad, —ya sea por sus valores de normalidad o por ser una constante que en los pacientes con ERC estén alterados estos valores— los factores de progresión de la enfermedad renal crónica más influyentes según este estudio realizado son:

- Epidemiológicas: el sexo, la edad, enfermedad cardiovascular, diabetes, hipertensión arterial, tabaquismo, desnutrición, hipertrofia ventricular izquierda

- Clínicas y paraclínicas: frecuencia cardíaca, hemoglobina, albúmina, urea, creatinina, filtrado glomerular (MDRD-4), diagnóstico de la enfermedad renal crónica en el estadio V, inicio en el programa en hemodiálisis posterior a seis meses de ser prescrito por el especialista.

- Dependientes del tratamiento: si el paciente tuvo atención nefrológica previa al inicio del programa en hemodiálisis, tiempo en hemodiálisis, ganancia de peso después del inicio de la hemodiálisis (interdialítica) excesiva, contagio de hepatitis secundario al proceder de hemodiálisis.

Conclusiones

La extracción de conocimiento a partir de datos es una línea de trabajo muy utilizada en la actualidad gracias a la cantidad de información en formato digital disponible. El conocimiento descubierto puede ser utilizado para construir sistemas inteligentes que ayuden a resolver problemas de un dominio de aplicación o puede ser útil para enriquecer el conocimiento sobre el dominio que tienen los especialistas.

La COVID-19 resulta una enfermedad de reciente desarrollo. Su inicio se reporta en diciembre del 2019, de modo que todavía el conocimiento sobre la misma es limitado y tiene asociado un grado de incertidumbre mayor que otras enfermedades estudiadas por más tiempo a pesar de que por ser una pandemia, se tiene una gran cantidad de datos recogidos. Por su parte, la mediastinitis, la ataxia y la enfermedad renal crónica, tienen gran relevancia por su elevada mortalidad, morbilidad y costes asociados a la asistencia médica. De este modo, la adquisición de conocimiento para las mismas resulta de interés y es valioso para los especialistas en el campo de la salud.

Para representar el conocimiento se han utilizado un formalismo que permite una mejor interpretación de los resultados por los especialistas, en este caso los prototipos. En el caso del conocimiento inferido sobre la letalidad de la enfermedad, los prototipos representan patrones para la predicción del riesgo de muerte para el paciente, lo que permite identificar diferentes alternativas que pueden elevar el riesgo de muerte. Los resultados alcanzados, basados en el empleo de las técnicas de IA, permiten ratificar conocimiento existente sobre estas enfermedades y también pudieran ofrecer a los especialistas algunas aristas novedosas sobre su letalidad.

Trabajos futuros

Se pretende extender de forma inmediata los resultados obtenidos a la predicción de patrones corporales y enfermedades dermatológicas asociadas a otras afecciones. Por su parte, se desea realizar un modelo geoespacial de predicción de enfermedades transmisibles. De igual forma se aplicarán los métodos propuestos por los autores en la reducción de los paros cardiaco-respiratorios a través del pronóstico para la detección precoz, así como en la reducción de la estadía hospitalaria y el costo de esta en los pacientes del Centro de Atención Cardiovascular del Hospital Universitario Manuel Ascunce Domenech de Camagüey.

Referencias del capítulo

Barra, A. B. L., Roque-da-Silva, A. P., Canziani, M. E. F., Lugon, J. R. & Strogoff-de-Matos, J. P. (2022). Characteristics and predictors of mortality on haemodialysis in Brazil: a cohort of 5,081 incident patients. *BMC nephrology*, 23(1), 1-8.

Becerra, R. A. (2022). Technology for processing saccadic electrooculograms in people with Spinocerebellar Ataxia type 2 [Tesis de doctorado, Universidad de Málaga]. <https://riuma.uma.es>

Becerra, R., Joya, G., García Bermúdez, R. V., Velázquez, L., Rodríguez, R. & Pino, C. (2013). *Saccadic points classification using multilayer perceptron and random forest classifiers in EOG recordings of patients with ataxia SCA2* [Ponencia]. Paper presented at the International Work-Conference on Artificial Neural Networks.

Belarj, B., Dahraoui, S., Rar, L., Atmani, N., Frikh, M., Ben, Y. y Boulahya, A. (2018). Exceptional association of two species of bacteria causing mediastinitis: *Haemophilus influenzae* (H. influenzae) and *Aggregatibacter aphrophilus* (A. aphrophilus). *BMC Infectious Diseases*, 18(1), 1-4.

Bello-García, B. e. a. (2019). *Implementación de métodos para el pre-procesamiento de datos usando Teoría de los conjuntos aproximados (RST) en Python* [Ponencia]. Paper presented at the Conferencia Internacional de Procesamiento de la Información (CIPi2019), Cuba.

Bello, A. K., Okpechi, I. G., Osman, M. A., Cho, Y., Htay, H., Jha, V. y Johnson, D. W. (2022). Epidemiology of haemodialysis outcomes. *Nature Reviews Nephrology*, 18(6), 378-395.

Bermejo, S. & Cabestany, J. (2000). A batch learning vector quantization algorithm for nearest neighbour classification. *Neural Processing Letters*, 11(3), 173-184.

Bermúdez-Yera, G. d. J., Naranjo-Ugalde, A. M., Rabassa-LópezCallejas, M. A., Lagomasino-Hidalgo, Á. L., Chaljub-Bravo, E. & Barreto-Fiu, E. E. (2019). Modelo predictivo de mediastinitis postoperatoria en cirugía cardiovascular. *Cirugía Cardiovascular*, 26(6), 277-282.

Buckley, E., Mazzà, C. & McNeill, A. (2018). A systematic review of the gait characteristics associated with Cerebellar Ataxia. *Gait & Posture*, 60, 154-163.

Bullock, J., Luccioni, A., Pham, K. H., Lam, C. S. N., & Luengo-Oroz, M. (2020). Mapping the landscape of artificial intelligence applications against COVID-19. *Journal of Artificial Intelligence Research*, 69, 807-845.

Caballero, Y., Bello, R., Arco, L., García, M. & Ramentol, E. (2010). Knowledge discovery using rough set theory. In *Advances in Machine Learning I* (pp. 367-383). Springer.

Cabeza-Ruiz, R., Velázquez-Pérez, L. & Pérez-Rodríguez, R. (2021). *Convolutional Neural Networks as Support Tools for Spinocerebellar Ataxia Detection from Magnetic Resonances* [Ponencia]. Paper presented at the International Workshop on Artificial Intelligence and Pattern Recognition. Habana, Cuba.

Caparrós, G. J., Ávila, R. M. Á., Bermúdez, R. V. G., & Montero, J. M. (2021). Grupo de investigación: Laboratorio interuniversitario para el procesado de señales biomédicas y

Garcia-Montemayor, V., Martin-Malo, A., Barbieri, C., Bellocchio, F., Soriano, S., Pendon-Ruiz de Mier, V. y Rodriguez, M. (2021). Predicting mortality in hemodialysis patients using machine learning analysis. *Clinical kidney journal*, 14(5), 1388-1395.

Gatti, G., Dell'Angela, L., Barbati, G., Benussi, B., Forti, G., Gabrielli, M. y Pappalardo, A. (2016). A predictive scoring system for deep sternal wound infection after bilateral internal thoracic artery grafting. *European Journal of Cardio-Thoracic Surgery*, 49(3), 910-917.

Gonzalez-Garces, Y., Dominguez-Barrios, Y., Zayas-Hernandez, A., Sigler-Villanueva, A. A., Canales-Ochoa, N., Hernandez Oliver, M. O. y Carrillo-Rodes, F. J. (2021). Impacts of the COVID-19 pandemic on the mental health and motor deficits in cuban patients with cerebellar ataxias. *The Cerebellum*, 20(6), 896-903.

Gutiérrez-Urbón, J. M., Pereira-Rodríguez, M. J. & Cuenca-Castillo, J. J. (2013). Estudio de casos y controles de los factores de riesgo de mediastinitis en cirugía de revascularización miocárdica. *Cirugía Cardiovascular*, 20(1), 13-17.

Hadad-Arrascue, F., Nilsson, L. G., Rivera, A. S., Bernardo, A. A. & Cabezuelo Romero, J. B. (2022). Expanded hemodialysis as effective alternative to on-line hemodiafiltration: A randomized mid-term clinical trial. *Therapeutic Apheresis and Dialysis*, 26(1), 37-44.

Hershenson, J., Haworth, A. & Houlden, H. (2012). The inherited ataxias: genetic heterogeneity, mutation databases, and future directions in research and clinical diagnostics. *Human mutation*, 33(9), 1324-1332.

Kanda, E., Epureanu, B. I., Adachi, T., Tsuruta, Y., Kikuchi, K., Kashiara, N. y Nitta, K. (2020). Application of explainable ensemble artificial intelligence model to categorization of hemodialysis-patient and treatment using nationwide-real-world data in Japan. *Plos one*, 15(5), e0233491.

Kanda, E., Okami, S., Kohsaka, S., Okada, M., Ma, X., Kimura, T. y Yajima, T. (2022). Machine Learning Models Predicting Cardiovascular and Renal Outcomes and Mortality in Patients with Hyperkalemia. *Nutrients*, 14(21), 4614.

Kaushik, A. C. y Raj, U. (2020). AI-driven drug discovery: A boon against COVID-19? *AI Open*, 1, 1-4.

Kaya, M., Karakuş, S., & Tuncer, S. A. (2022). Detection of ataxia with hybrid convolutional neural network using static plantar pressure distribution model in patients with multiple sclerosis. *Computer Methods and Programs in Biomedicine*, 214, 106525.

Konishi, Y., Fukunaga, N., Abe, T., Nakamura, K., Usui, A. & Koyama, T. (2019). Efficacy of new multimodal preventive measures for post-operative deep sternal wound infection. *General Thoracic and Cardiovascular Surgery*, 67(11), 934-940.

Lalmuanawma, S., Hussain, J. & Chhakchhuak, L. (2020). Applications of machine learning and artificial intelligence for COVID-19 (SARS-CoV-2) pandemic: A review. *Chaos, Solitons & Fractals*, 139, 110059.

Landler, N. E., Olsen, F. J., Christensen, J., Bro, S., Feldt-Rasmussen, B., Hansen, D. y Sørensen, I. M. H. (2022). Associations Between Albuminuria, Estimated GFR and Cardiac Phenotype in a Cohort with Chronic Kidney Disease: The CPH-CKD ECHO Study. *Journal of Cardiac Failure*, 28(11), 1615-1627.

Lemaigren, A., Birgand, G., Ghodhbane, W., Alkhoder, S., Lolom, I., Belorgey, S. y Dilly, M.-P. (2015). Sternal wound infection after cardiac surgery: incidence and risk factors according to clinical presentation. *Clinical Microbiology and Infection*, 21(7), 674. e611-674. e618.

León, J. L., Calderón, M. & Gutiérrez, A. R. (2021). Análisis de mortalidad y comorbilidad por COVID-19 en Cuba. *Revista Cubana de Medicina*, 60(2). e2117.

de fijación de la mirada. *Revista de la Asociación Interacción Persona Ordenador (AIPO)*, 2(2), 90-93.

Chello, C., Lusini, M., Nenna, A., Nappi, F., Spadaccio, C., Satriano, U. M. y Chello, M. (2020). Deep Sternal Wound Infection (DSWI) and Mediastinitis After Cardiac Surgery: Current Approaches and Future Trends in Prevention and Management. *Surgical Technology International*, 36, 212-216.

Corona, C. C. (2018). Soft computing techniques and sustainability science, an introduction. In *Soft Computing for Sustainability Science* (pp. 1-10). Springer.

Cutrell, J. B., Barros, N., McBroom, M., Luby, J., Minhajuddin, A., Ring, W. S. & Greilich, P. E. (2016). Risk factors for deep sternal wound infection after cardiac surgery: influence of red blood cell transfusions and chronic infection. *American Journal of Infection Control*, 44(11), 1302-1309.

Da Silva, R. G., Ribeiro, M. H. D. M., Mariani, V. C., & dos Santos Coelho, L. (2020). Forecasting Brazilian and American COVID-19 cases based on artificial intelligence coupled with climatic exogenous variables. *Chaos, Solitons & Fractals*, 139, 110027.

Di Vaio, A., Palladino, R., Hassan, R. & Escobar, O. (2020). Artificial intelligence and business models in the sustainable development goals perspective: A systematic literature review. *Journal of Business Research*, 121, 283-314.

Epidemiología, D. G. d. (2020). *BC COVID-19 México 20 de abril 2020*. <https://www.gob.mx/salud/documentos/datos-abiertos-152127>

Galna, B., Lord, S. & Rochester, L. (2013). Is gait variability reliable in older adults and Parkinson's disease? Towards an optimal testing protocol. *Gait & Posture*, 37(4), 580-585.

García-Lorenzo, M. M., Rodríguez, Y., Hernández, A. R., Bello-García, B., Filiberto, Y., Rosete, A. y Bello, R. (2020). Adquisición de conocimiento sobre la letalidad de la COVID-19 mediante técnicas de inteligencia artificial. *Anales de la Academia de Ciencias de Cuba*, 10(3), 891.

Lin, J. & Jimenez, C. A. (2022). *Acute mediastinitis, mediastinal granuloma, and chronic fibrosing mediastinitis: A review* [Ponencia]. Paper presented at the Seminars in Diagnostic Pathology. www.sciencedirect.com.

Liu, R., Rong, Y. & Peng, Z. (2020). A review of medical artificial intelligence. *Global Health Journal*, 4(2), 42-45.

Machín, J. C. (2011). Factores predictores de mediastinitis aguda en cirugía cardiovascular, protocolo de prevención y algoritmos diagnóstico y terapéutico. [Tesis de doctorado, Universidad de Ciencias Médicas de Santiago de Cuba]. <https://tesis.sld.cu/>

Madurai, R. & Pugazhendhi, R. (2020). Science of the Total Environment Restructured society and environment: A review on potential technological strategies to control the COVID-19 pandemic. *Science of the Total Environment-Elsevier BV*(725-P), 138858.

Malakar, S., Roy, S. D., Das, S., Sen, S., Velásquez, J. D. & Sarkar, R. (2022). Computer Based Diagnosis of Some Chronic Diseases: A Medical Journey of the Last Two Decades. *Archives of Computational Methods in Engineering*, 1-43.

Mleyhi, S., Miri, R., Messai, M., Ellouze, H., Tesnim, B., Ben Omrane, S. & Denguir, R. (2021). Post cardiac surgery Mediastinitis: Risk factors and therapeutic management. *European Journal of Cardiovascular Nursing*, 20 (Supplement_1), zvab060. 048.

Nieto, M. (2014). Diseño y validación de un modelo predictivo de mediastinitis en cirugía cardíaca [Tesis de doctorado, Universidad Complutense de Madrid]. dialnet.unirioja.es/

Ocampo, L. & Yamagishi, K. (2020). Modeling the lockdown relaxation protocols of the Philippine government in response to the COVID-19 pandemic: An intuitionistic fuzzy DEMATEL analysis. *Socio-Economic Planning Sciences*, 72, 100911.

Oyelade, O. N., & Ezugwu, A. E. (2020). A case-based reasoning framework for early detection and diagnosis of novel coronavirus. *Informatics in Medicine Unlocked*, 20, 100395.

Park, S., Park, B. S., Lee, Y. J., Kim, I. H., Park, J. H., Ko, J. y Park, K. M. (2021). Artificial intelligence with kidney disease: A scoping review with bibliometric analysis, PRISMA-ScR. *Medicine*, 100(14). e254422

Pastene, B., Cassir, N., Tankel, J., Einav, S., Fournier, P.-E., Thomas, P. & Leone, M. (2020). Mediastinitis in the intensive care unit patient: a narrative review. *Clinical Microbiology and Infection*, 26(1), 26-34.

Peng, L., Chen, Z., Chen, T., Lei, L., Long, Z., Liu, M. y Wan, L. (2021). Prediction of the age at onset of spinocerebellar ataxia type 3 with machine learning. *Movement Disorders*, 36(1), 216-224.

Rahman, S. (2022). Causative Factors of Chronic Kidney Disease in Patiens with Hemodialysis Therapy. *KEMAS: Jurnal Kesehatan Masyarakat*, 18(1). 114-121.

Ren, Z., Liao, H. & Liu, Y. (2020). Generalized Z-numbers with hesitant fuzzy linguistic information and its application to medicine selection for the patients with mild symptoms of the COVID-19. *Computers & Industrial Engineering*, 145, 106517.

Rennie, L., Löfgren, N., Moe-Nilssen, R., Opheim, A., Dietrichs, E. & Franzén, E. (2018). The reliability of gait variability measures for individuals with Parkinson's disease and healthy older adults—The effect of gait speed. *Gait & Posture*, 62, 505-509.

Rodríguez, Y., Caballero, Y., Filiberto, Y., García, I. M., Fernández, Y. & Frias, M. (2019). Similar prototype methods for class imbalanced data classification. In *Uncertainty Management with Fuzzy and Rough Sets* (pp. 193-209). Springer.

Ru, D., Li, J., Xie, O., Peng, L., Jiang, H. & Qiu, R. (2022). Explainable artificial intelligence based on feature optimization for age at onset prediction of spinocerebellar ataxia type 3. *Frontiers in Neuroinformatics*. 16. 978630

Salmeron, J. L., & Arévalo, I. (2020). *A Privacy-Preserving, Distributed and Cooperative FCM-Based Learning Approach for Cancer Research* [Ponencia]. Paper presented at the International Joint Conference on Rough Sets. Habana, Cuba.

salud., D. d. R. m. y. E. d. (2022). *Anuario Estadístico de Salud 2021*. <http://bvscuba.sld.cu/anuario-estadistico-de-cuba/>

Sangeux, M., Passmore, E., Graham, H. K. & Tirosh, O. (2016). The gait standard deviation, a single measure of kinematic variability. *Gait & Posture*, 46, 194-200.

Scardoni, A., Balzarini, F., Signorelli, C., Cabitza, F. & Odone, A. (2020). Artificial intelligence-based tools to control healthcare associated infections: a systematic review of the literature. *Journal of infection and public health*, 13(8), 1061-1077.

Shaikh, F., Andersen, M. B., Sohail, M. R., Mulero, F., Awan, O., Dupont-Roettger, D. y Bisdas, S. (2021). Current landscape of imaging and the potential role for artificial intelligence in the management of COVID-19. *Current Problems in Diagnostic Radiology*, 50(3), 430-435.

Sipior, J. C. (2020). Considerations for development and use of AI in response to COVID-19. *International Journal of Information Management*, 55, 102170.

Sousa, H., Ribeiro, O., Costa, E., Christensen, A. J. & Figueiredo, D. (2022). Establishing the criterion validity of self-report measures of adherence in hemodialysis through associations with clinical biomarkers: A systematic review and meta-analysis. *Plos one*, 17(10), e0276163.

Suri, J. S., Puvvula, A., Biswas, M., Majhail, M., Saba, L., Faa, G. y Chadha, P. S. (2020). COVID-19 pathways for brain and heart injury in comorbidity patients: A role of medical imaging and artificial intelligence-based COVID severity classification: A review. *Computers in Biology and Medicine*, 124, 103960.

Swapnarekha, H., Behera, H. S., Nayak, J. & Naik, B. (2020). Role of intelligent computing in COVID-19 prognosis: A state-

of-the-art review. *Chaos, Solitons & Fractals*, 138, 109947.

Tartarisco, G., Bruschetta, R., Summa, S., Ruta, L., Favetta, M., Busa, M. y Cerasa, A. (2021). Artificial intelligence for dysarthria assessment in children with ataxia: A hierarchical approach. *IEEE Access*, 9, 166720-166735.

Vaishya, R., Javaid, M. & Khan, I. H. (2020). Aplicaciones de Inteligencia Artificial (IA) para la pandemia de COVID-19. *Diabetes y síndrome metabólico: investigación clínica y revisiones*, 14(4), 337-339.

Vinod, D. N. & Prabakaran, S. (2020). Data science and the role of Artificial Intelligence in achieving the fast diagnosis of COVID-19. *Chaos, Solitons & Fractals*, 140, 110182.

Yan, L., Zhang, H.-T., Xiao, Y., Wang, M., Sun, C., Liang, J. y Xiao, Y. (2020). Prediction of criticality in patients with severe COVID-19 infection using three clinical features: a machine learning-based prognostic model with clinical data in Wuhan. *MedRxiv*, 27

Yan, T., Wong, P. K., Ren, H., Wang, H., Wang, J. & Li, Y. (2020). Automatic distinction between COVID-19 and common pneumonia using multi-scale convolutional neural network on chest CT scans. *Chaos, Solitons & Fractals*, 140, 110153.

Yang, Z., Zhong, S., Carass, A., Ying, S. H. & Prince, J. L. (2014). *Deep learning for cerebellar ataxia classification and functional score regression* [Ponerencia]. Paper presented at the International Workshop on Machine Learning in Medical Imaging. Boston, USA.

Yao, J. & Azam, N. (2014). Web-based medical decision support systems for three-way medical decision making with game-theoretic rough sets. *IEEE Transactions on Fuzzy Systems*, 23(1), 3-15.



**Capítulo 3. Las redes neuronales:
herramienta imprescindible en los análisis
epidemiológicos de *Fasciola hepatica***

Capítulo 3. Las redes neuronales: herramienta imprescindible en los análisis epidemiológicos de *Fasciola hepatica*

Amílcar Arenal Cruz

Lisney Lisandra González Rodríguez

Annelis Arteaga Campbell

Julio Madera Quintana

Damián Fuentes Milanés

Resumen

Uno de los modelos matemáticos prometedores es Redes Neuronales Artificiales (RNA), un método de inteligencia artificial para modelar funciones objetivas complejas. La adaptabilidad de la RNA permite la aplicación en casos en que es imposible crear un modelo matemático estricto, pero donde hay un conjunto suficientemente representativo de muestras. En el presente trabajo empleamos redes neuronales en estudios epidemiológicos de *Fasciola hepatica* de las provincias de Camagüey y Las Tunas, Cuba. Se emplearon variables meteorológicas (fundamentalmente precipitaciones y temperatura) para el estudio de ambas provincias. En el caso de Camagüey, el ajuste al modelo tuvo un coeficiente de determinación de 0,96, con las precipitaciones el de mayor importancia para la predicción. En el caso de las Tunas, con una red neuronal de series de tiempo, la relación entre el valor real y el que se predice por RNA, solo tuvo una diferencia con un residual entre 0,5 % y -0,5 %, y un coeficiente de determinación fue $R^2 = 0,91$. Las redes neuronales son una herramienta poderosa para predecir la prevalencia de *F. hepatica* y con ellos poder realizar una intervención de forma temprana con el fin de evitar la propagación del parásito.

Abstract

One of these promising mathematical models is Artificial Neural Networks (RNA), an artificial intelligence method for modeling complex objective functions. The adaptability of the ANN allows its application in cases where it is impossible to create a strict mathematical model but where there is a sufficiently representative set of samples. In the present work, we use neural networks in epidemiological studies of *Fasciola hepatica* in the provinces of Camagüey and Las Tunas. The study of both provinces used Meteorological variables, mainly precipitation, and temperature. In the case of Camagüey, the fit to the model had a coefficient of determination of 0.96, with rainfall being the most important for the prediction. In the case of Tunas, with a neural network of time series, the relationship between the real value and the one predicted by RNA only had a difference with a residual between 0.5% and -0.5%, and a coefficient of determination was $R^2=0,91$. Neural networks are a powerful tool to predict the prevalence of *F. hepatica* and can intervene early to prevent the spread of the parasite.

Introducción

Una de las enfermedades más comunes en la ganadería bovina es la producida por la *Fasciola hepatica*. Es el agente causal de una de las parasitosis de mayor impacto del ganado, la fascioliasis o fasciolosis; debido a las grandes pérdidas económicas que se producen en el mundo (Thanh, 2012b). Está considerada como una zoonosis emergente con millones de personas en riesgo de infección mayormente en el Sur de América y África (Recalde *et al.*, 2014).

A pesar de la existencia de bases de datos en los mataderos de las patologías detectadas, en Cuba solo existen evaluaciones y/o determinaciones de pérdidas económicas y prevalencias de enfermedades pertenecientes a la región del matadero en

cuestión: pueden consultarse trabajos como González *et al.* (2007a) sobre la evaluación de la prevalencia de *Fasciola* en una Empresa Ganadera de La Habana o Brito *et al.* (2010b) que evaluó la prevalencia de tres provincias centrales de nuestro país y estimó las pérdidas económicas por concepto de carne y leche basados en investigaciones donde se registró el porciento de pérdidas de peso y disminución de leche en rebaños bovinos parasitados por fasciolosis.

Sólo Arenal *et al.* (2014) y Palacio (2016) se enfocaron en describir el comportamiento de la prevalencia durante seis años en un matadero de la provincia de Camagüey, fuera de esto, al término de esta entrega, la región oriental de Cuba no presentaba trabajos sobre estos temas.

El uso de modelos Bayesianos para incrementar el valor epidemiológico de los análisis (Durr *et al.*, 2005); el empleo de Redes Neuronales para mejorar el valor epidemiológico en los mataderos estableciendo un comportamiento en la provincia de Camagüey con variables climatológicas (Arenal *et al.*, 2014); la estimación de las pérdidas por *Fasciola* en la provincia de Pinar del Río durante los años 2005-2013, la base en el método bayesiano clásico del diagnóstico *a priori* de la enfermedad y su pronóstico *a posteriori*, y con ellos la intervención de forma temprana con el fin de evitar la propagación del parásito (González, 2014) y el establecimiento de comportamientos estacionales de fasciolosis con el empleo de *Autoregressive Integrated Moving Average* (ARIMA) (Palacio, 2016) demuestran como los modelos matemáticos logran un gran auge en los últimos años y tienen suma importancia que tienen para que investigadores y médicos veterinarios puedan desarrollar métodos para predecir el comportamiento de la enfermedad con el fin de disminuir la prevalencia de fasciolosis (González, 2014).

Estos modelos son muy útiles porque capturan propiedades esenciales de la dispersión de una enfermedad de una forma simplificada. Además, al modificar los parámetros del modelo se pueden representar o descubrir situaciones que difícilmente se pueden obtener mediante experimentación. Por lo tanto, contribuyen a prevenir futuras situaciones patológicas, determinar la prevalencia e incidencia y contribuir a tomar decisiones objetivas para el control o supresión de las enfermedades.

Autores como da Cunha et al. (2010) refirieron la técnica de predicción o pronóstico en la epidemiología como principal objetivo en la proyección de las necesidades futuras de la salud pública y animal, o la detección de un comportamiento de una serie de tiempo que indique la ocurrencia de la incidencia y prevalencia de una enfermedad en cuestión. Algunas de las técnicas utilizadas para el pronóstico y la predicción incluyen funciones de ajuste como las RNA, los modelos de regresión, ARIMA, etc.

Uno de estos modelos matemáticos prometedores es Redes Neuronales Artificiales (RNA), un método de inteligencia artificial para modelar funciones objetivas complejas. Es un enfoque computacional inspirado en el sistema nervioso humano. Se basa en las teorías de la interconexión masiva y arquitectura de procesamiento paralelo de los sistemas neuronales biológicos (Almási *et al.*, 2015).

Se utilizan para reconocer patrones, incluyendo imágenes, manuscritos y secuencias de tiempo y tienen capacidad de aprender y mejorar su funcionamiento Aichouri *et al.* (2015).

La característica importante de RNA es la adaptabilidad, esto permite que se puedan aplicar en casos en que es imposible crear un modelo matemático estricto, pero donde hay un conjunto suficientemente representativo de muestras. La otra característica importante de las redes neuronales es su capacidad de generalizar la información de entrada y dar

respuestas correctas de los datos “desconocidos”, lo que las hace eficaces en la solución de problemas de clasificación complicados (Almási *et al.*, 2015).

Ventajas que ofrecen las redes neuronales:

- Aprendizaje: Las RNA tienen la habilidad de aprender mediante una etapa que se llama etapa de aprendizaje. Esta consiste en proporcionar a la RNA datos como entrada a su vez que se le indica cuál es la salida (respuesta) esperada.

- Autoorganización: Una RNA crea su propia representación de la información en su interior, quitándole esta tarea al usuario. Esta autoorganización provoca la generalización: facultad de las redes neuronales de responder apropiadamente cuando se les presentan datos o situaciones a las que no fue expuesta anteriormente. El sistema puede generalizar la entrada para obtener una respuesta. Esta característica es muy importante cuando se tiene que solucionar problemas en los cuales la información de entrada no es muy clara; además permite que el sistema dé una solución, incluso cuando la información de entrada está especificada de forma incompleta (Matich, 2001).

- Tolerancia a fallos: Debido a que una RNA almacena la información de forma redundante, esta puede seguir respondiendo de manera aceptable aun si se daña parcialmente.

- Tiempo real: La estructura de una RNA es paralela, por lo cual, si se implementa con computadoras o en dispositivos electrónicos especiales, se pueden obtener respuestas en tiempo real.

- Fácil inserción dentro de la tecnología existente. Se pueden obtener chips especializados para redes neuronales que mejoran su capacidad en ciertas tareas. Ello facilitará la integración modular en los sistemas existentes.

- Flexibilidad: Una RNA puede manejar cambios no importantes en la información de entrada, como señales con ruido u otros cambios en la entrada (por ejemplo, si la

información de entrada es la imagen de un objeto, la respuesta correspondiente no sufre cambios si la imagen cambia un poco su brillo o el objeto cambia ligeramente).

Aplicaciones de las RNA

Las redes neuronales artificiales (RNA) se emplean en los últimos años para realizar tareas complejas en diferentes áreas y como estrategia de problemas de modelización matemática, diseñadas como sistemas de entrada y salida. A diferencia de otras estrategias de modelado, no es necesario conocer la relación matemática entre entradas y salidas. Por lo tanto, a diferencia del modelo de regresión, no es necesario proponer una función para el modelo, ya que las RNA son funciones aproximadores universales (da Cunha *et al.*, 2010). Las RNA resolverán cualquier problema que pueda representar. Los ejemplos incluyen el reconocimiento de patrones de sonido e imagen, previsión de series temporales, simulación, control y optimización del sistema.

Tales características las hacen bastante apropiadas para aplicaciones en las que no se dispone *a priori* de un modelo identificable que pueda ser programado, pero se dispone de un conjunto básico de ejemplos de entrada (previamente clasificados o no) (Agrawal *et al.*, 2012 y Almási *et al.*, 2015). Esto incluye problemas de clasificación y reconocimiento de patrones de voz, imágenes, señales, etc. Asimismo, se utilizan para encontrar patrones de fraude económico, hacer predicciones en el mercado financiero, hacer predicciones de tiempo atmosférico, etc. Las RNA demostraron la capacidad de diagnosticar y predecir eventos clínicamente significativos, siendo el primero el infarto del miocardio; el cuál introduciéndole datos de pacientes enfermos, radiografías y electrocardiogramas aprendió a identificar la enfermedad y a diferenciarlos de los sanos. Luego de los ciclos de aprendizaje el programa fue capaz de identificar los enfermos con solo

introducirle datos de su historia clínica o pronosticar si era propenso padecer patologías cardíacas en el futuro (Atkov *et al.*, 2012).

De la misma forma, basado en la incidencia mensual de fiebre tifoidea de los años 2005 hasta 2009 y para el año 2010 desde el Centro de control y prevención de enfermedades en China. Las primeras se utilizaron como el modelado y las últimas para la previsión de muestras, respectivamente. Las predicciones se compararon con la predicción realizada con el modelo estadístico ARIMA. Los resultados mostraron más cercana a la real la predicción con RNA (Zhang *et al.*, 2013).

Almási *et al.* (2015) aplicaron redes neuronales artificiales al diagnóstico de cáncer a través de imágenes. Durante la fase de entrenamiento el sistema recibió imágenes de tejidos de pacientes positivos de cáncer y tejidos sanos; así como, las respectivas clasificaciones de dichas imágenes. Con el entrenamiento es el adecuado, una vez concluido, el sistema podría recibir imágenes de tejidos no clasificados y obtener su clasificación (sano/ no sano) con un buen grado de seguridad. Las variables de entrada podían ser desde los puntos individuales de cada imagen hasta un vector de características de las mismas que se puedan incorporar al sistema (por ejemplo, procedencia anatómica del tejido de la imagen o la edad del paciente al que se le extrajo la muestra).

En el campo de la parasitología, la investigación de Theodoropoulos *et al.* (2000) demostró como las RNA son capaces de aprender a identificar imágenes de estadio 3 de larvas del parásito *Strongylus* de los animales domésticos. Se entrenó la neurona con una base de datos cuantitativos obtenidos a partir de imágenes digitales procesadas de larvas. Se registraron un total de 255 imágenes de 57 larvas individuales en diversas formas pertenecientes a cinco géneros. Tras el procesamiento de imágenes, se midieron 16 características, de los cuales siete fueron seleccionados como

invariantes a la forma de larva. Este RNA junto con un centro de análisis de imágenes y una base de datos, se convirtió en la base para el desarrollo de un sistema de identificación de larvas por computadora cuyo rendimiento de identificación en general fue del 91,9 %. Las ventajas de este sistema son su velocidad y objetividad. La objetividad del sistema se basa en el hecho de que no está sujeto a la variabilidad inter e intra observador que surge de perfil del usuario de la competencia en la interpretación de las descripciones subjetivas y no cuantificables.

Resultados en la Universidad de Camagüey en el estudio epidemiológico de *Fasciola hepatica* con redes neuronales

Provincia Camagüey

En los años 2008, 2009 y 2010 los kilogramos de hígados afectados y desviados al consumo animal superaron la tonelada (1117,8; 1043,3 y 1070 respectivamente). Aunque el 2008 representó casi 100 hígados más con fasciolosis (508) que en el 2009 y 2010. No obstante, fue el año 2008 el de mayor cantidad de ganado sacrificado y el 2010 fue el de menor durante el período estudiado. El año 2012 dictaminó la menor cantidad de hígados decomisados por *F. hepatica* y por consiguiente el valor más bajo de kilogramos no obtenidos para el consumo humano (Tabla 1).

De forma general se manifiestan los seis años analizados con afectación de hígado para un total aproximado de 4,8 toneladas. El Instituto de Medicina Veterinaria, como institución rectora del trabajo veterinario en Cuba (IMV, 2007), reportó una infestación alta en bovinos, con una existencia de 1170 focos, 66 969 enfermos y 3059 muertos.

Esto demuestra que el nivel de afectación hepática + que produce este parásito es considerable (Schweizer *et al.*, 2005),

Tabla 1. Decomisos de hígados de bovinos por presencia de *Fasciola hepatica*, en el período 2008 – 2013, en el matadero Chacuba.

Años	Ganado Sacrificado	Hígados Afectados	Total kilogramos Decomisados Total				
			Parcial		Total		kilogramos Decomisados
			Uno	Kg	Uno	Kg	
2008	5390	508	374	542,1	134	575,7	1117,8
2009	2603	415	277	501,4	138	541,9	1043,3
2010	1956	462	347	625,4	115	444,6	1070
2011	2022	254	195	397,74	59	289,05	686,79
2012	2740	162	130	235,2	32	137,7	372,9
2013	3423	293	256	366	37	154,8	520,8
Total	18134	2094	1579	2667,84	515	2143,75	4811,59

pero un programa de vigilancia epizootológica, muestreos coprológicos sistemáticos y tratamientos antiparasitarios frecuentes, sin duda, los daños por decomiso pudieran reducirse ostensiblemente, en coincidencia con otros trabajos realizados en Cuba. Son sumamente mayores las pérdidas aún sin contabilizar, teniendo en cuenta lo relacionado con la afectación de la producción de leche y conversión en carne (González *et al.*, 2007b). Un animal afectado por fasciolosis puede disminuir hasta un 28 % su producción de carne, pues consume como promedio un 15 % menos de alimentos, además de la cantidad y calidad de la leche que deja de producir (Fredes *et al.*, 1997).

Las prevalencias observadas por cada año y en los trece municipios estudiados debido a *F. hepatica* son inferiores a las obtenidas en investigaciones realizadas en Cuba en otros períodos (Figura 1). En nuestro caso en el año 2010 fue el de mayor prevalencia, influenciado fundamentalmente por el aporte de animales afectados de Vertientes. En La Habana, se informaron prevalencias superiores al 30 % en cada año, en matanzas normales solo en una Empresa Ganadera (González *et al.*, 2007b). Esto obedece a que en el matadero estudiado el promedio de matanzas es de 25 animales diariamente por

lo general y en época de bajas compras del ganado (mayo a julio), los mismos se ven disminuidos, puesto que el centro no es priorizado como entidad perteneciente al Ministerio de la Industria Alimenticia. Cifras superiores a las investigadas se obtuvieron en el centro del país con 37,2 % en Villa Clara, 50,9 % en Sancti Spíritus y similares a Cienfuegos con 20,1 % en los años 2000-2004 (Brito *et al.*, 2010a).

A nivel internacional en el estado de Trujillo, Venezuela se informó una prevalencia de 2,2 %, baja en comparación con nuestra investigación (Hernandez y Rojas, 2009). Similares a las observadas en Suiza en el año 2006 con un 18,0 %, (Rapsch *et al.*, 2006), de igual forma en Reino Unido, reportaron resultados del 22 % en 2011 y una disminución a 19,43 % en 2012 y en Escocia un 28,9 % de prevalencia en el mismo año (Gracia, 2012). También hay reportes de valores muy por encima de los obtenidos, como en diversos estados de Vietnam donde alcanzaron prevalencias desde un 20 % hasta 90 %, (Thanh, 2012a).

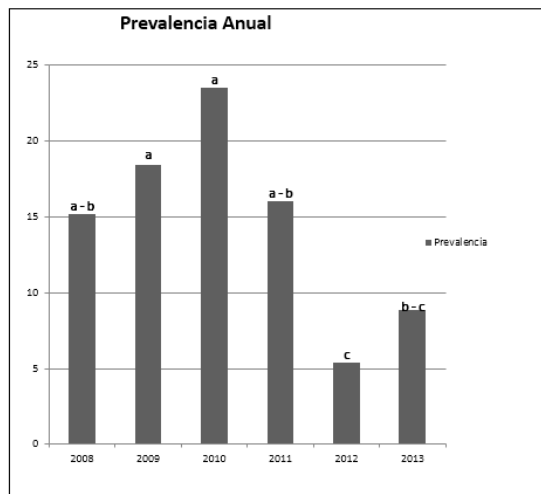


Figura 1. Prevalencia de *Fasciola hepatica* por año en el matadero Chacuba, Camagüey, Cuba (2008-2013). Letras diferentes significan, diferencias significativas ($P < 0,05$). ANOVA de un factor seguido de un Student Newman Keuls

En zonas más cercanas a nuestro país, en el Sur de América, Brasil: varias investigaciones evidenciaron en estados como Río Grande del Sur un 18,7 % y Santa Catarina un 10,1 % en un período de cinco años (Dutra *et al.*, 2010a). Río Grande del Sur fue inferior al año 2010 y de igual forma para los años 2008, 2009, 2010 y 2011 con respecto a Santa Catarina.

Rozo y Margarita (2009), en Latinoamérica informaron una prevalencia en Perú y México de fasciolosis bovina de 95,5 % y 96,5 % respectivamente, superiores a los resultados obtenidos por año y municipio. Se informa prevalencias de Brasil 3,32 %, en Bolivia de 16,59 % y por último Colombia con 25 % valores aproximados a los vistos en esta investigación con excepción de Brasil, ya que nuestros menores resultados fueron de 5,4 % para el año 2012 y 5,5 % en el municipio de Najasa.

La figura 2, muestra la relación entre el valor real y el que se predice por red neuronal artificial. Se observa una relación lineal entre ambos valores. Con un residual menor del 10% en casi todos los casos (Figura 3).

Interesante que las variables de mayor importancia para el modelo fueron las relacionadas con el clima (precipitaciones y las temperaturas). Coincidentemente, varios modelos colocan

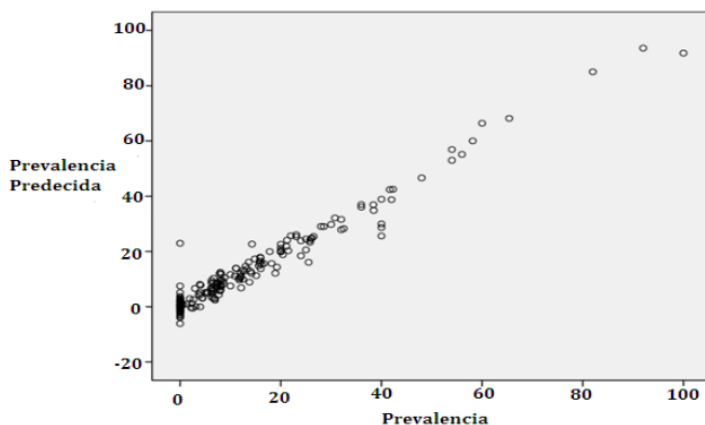


Figura 2. Relación de la prevalencia de *Fasciola hepatica* en el matadero Chacuba, Camagüey, Cuba. (2008-2013) y su predicción por red neuronal artificial MLP

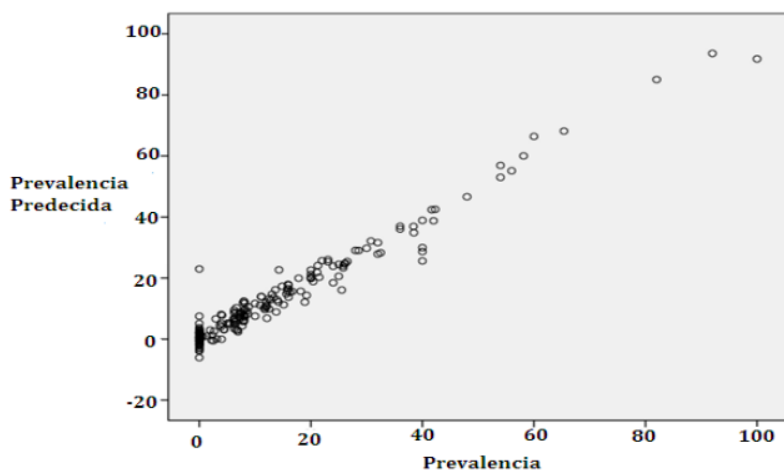


Figura 3. Valor residual de cada prevalencia predicha por el modelo de red neuronal artificial

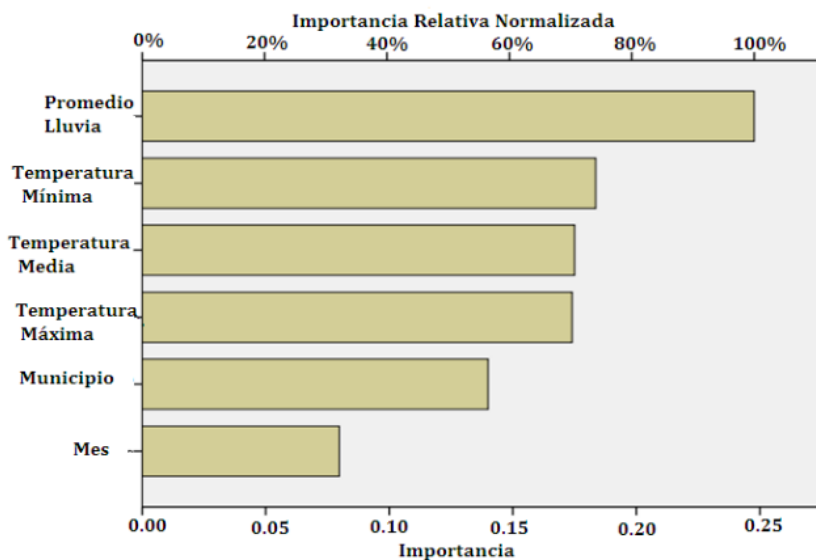


Figura 4. Importancia relativa de las variables para el modelo de red neuronal artificial.

a estos factores como los primeros determinantes para la predicción de *F. hepatica* (De Waal *et al.*, 2007).

Sin embargo, los meses presentaron la menor importancia (Figura 4).

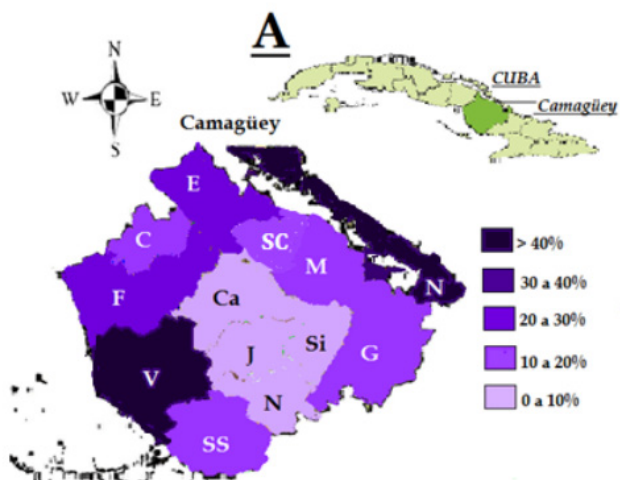
El promedio de precipitaciones fue el de mayor relevancia debido a su nivel de influencia para la realización del ciclo biológico de la *F. hepatica* que incluye el desarrollo del hospedero intermediario fundamental para un ciclo completo. Se necesita alta humedad tanto para la *F. hepatica* como para *Lymnaea* spp. (Fox *et al.*, 2011, Relf *et al.*, 2011). Otros modelos, indican a la temperatura como el de mayor importancia en la predicción de fasciolosis bovina (Fox *et al.*, 2011).

Situación para cada uno de los municipios

Al realizar una representación en el mapa de la provincia en análisis (Camagüey), sobre prevalencia de *F. hepatica*, según la intensidad de los colores que se muestra respecto a los valores obtenidos durante el período 2008-2013 en estudio.

La disminución en la estimación de Vertientes, con una diferencia del 9,9%, puede deberse a que no teníamos lotes de este municipio en el último año y el penúltimo solo tenía un lote con baja prevalencia. En el caso del Sta. Cruz del Sur se utilizó en la estimación siete lotes menos que en lo observado, lo cual influye en la prevalencia predicha por la red neuronal. Nuestro modelo tiene coeficiente de determinación de $r^2=0,96$; superior al planteado en Reino Unido donde se utilizaron modelos de regresión lineal que explicaban 70 % (McCann *et al.*, 2010).

En redes neuronales para otras enfermedades se informan altos r^2 (Eg: 0,97 para predecir nódulos en pacientes con cáncer (Batuello *et al.*, 2001). Recientemente, se usó una red neuronal en series de tiempo para la estimación de malaria en el Amazona, aunque los autores informan un error cuadrático medio más bajo que la regresión logística, no describen el



Ca – Camagüey C – Carlos Manuel de Céspedes E- Esmeralda
 F – Florida J – Jimaguayú S - Sibanicú SC – Sierra de Cubitas
 G – Guáimaro M - Minas N – Najasa Nv - Nuevitas SS – Santa Cruz del
 Sur V - Vertientes

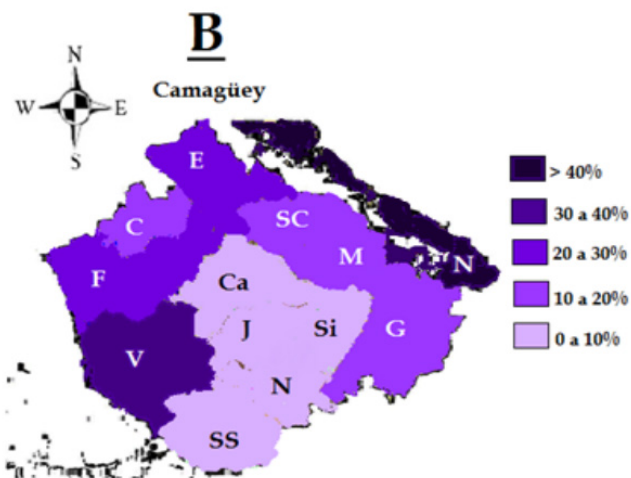


Figura 5. En el mapa A: Mapa de prevalencia de la fasciolosis observada por municipio de la provincia Camagüey, Cuba. En el matadero Chacuba (2008-2013). En el mapa B: Mapa de prevalencia de la fasciolosis estimada de acuerdo a red neuronal MLP, por municipio de la provincia Camagüey, Cuba. En el matadero Chacuba (2008-2013)

coeficiente de determinación (Cunha *et al.*, 2010). En un análisis de riesgo de Schistosomiasis el coeficiente de regresión fue de 0,65, más alto que el de regresión logística con 0,59 (Xu *et al.*, 2013).

Sin embargo, las redes neuronales no están exentas de inconvenientes. La principal desventaja es su calidad de “caja negra”, es decir, sin esfuerzo adicional, es difícil, si no imposible, obtener una perspectiva de un problema basado en un modelo (Sargent, 2001). Aunque cada vez más se comprende la interpretación de las sinapsis neuronales.

Provincia Las Tunas

Durante el periodo de 2006-2015 fueron faenados un total de 213 593 bovinos en los mataderos y losas de la provincia de Las Tunas. Se retiraron 17 543 hígados que estaban afectados por la *F. hepática*. La prevalencia global en estos diez años fue de 12,5 %. La prevalencia anual presentó diferencias significativas en las proporciones ($P < 0,0001$), se destacó el 2010 con el valor más alto. Los años 2006 y 2007 presentaron las prevalencias más bajas.

La mayor cantidad de hígados decomisados, ocurrió en el 2009; a pesar de sacrificar menos animales y tener un menor número de animales afectados en comparación con los años posteriores.

Sin embargo, mayor porcentaje de hígados decomisados en el año 2008 (79,98 %) de los animales afectados, con una prevalencia inferior, pudieran indicar una mayor severidad en la infestación de la *F. hepatica* (Tablas 2 y 3).

Los reportes de prevalencia en bovinos de *F. hepatica* en países con climas templados son variables. En el este de Europa, se estimaron prevalencias de 37 %, 50 % y 61 % en Bélgica, Alemania y España respectivamente (Howell *et al.*, 2015). La prevalencia de fasciolosis en Reino Unido fue de 77,5 %, se incrementó en los años recientes, y se reportó en

Tabla 2. Prevalencia de fasciolosis y decomisos de hígados en la provincia de Las Tunas 2006-2015

Año	Animales Sacrificados	Animales Afectados		Hígados Decomisados		
		Cant.	%	Cant.	%	kg
2006	15 326	1 373	8,95a	883	5,76	-
2007	16 272	1 056	6,48a	538	3,3	-
2008	22 538	2 803	9,24a	1 666	7,39	-
2009	18 696	2 518	13,46a	1 875	10,02	-
2010	19 172	2 873	14,98a	1 889	9,85	-
2011	20 040	2 706	13,50 a	1 668	8,32	6 710
2012	21 207	2 937	13,84 b	1 617	7,62	7 589,10
2013	25 074	3 222	12,84 b	2 259	9,00	8 477,86
2014	27 439	3 370	12,28 a	2 357	8,58	8 673,66
2015	27 829	3 872	13,91 b	2 791	9,21	9 901,27
Total	213 593	26 730	12,5	17 543	8,21	41 351,89

Tabla 3. Prevalencia de fasciolosis y decomisos de hígado en los municipios de la provincia Las Tunas en el periodo 2006-2015

Municipios	Animales Sacrificados	Animales Afectados		Hígados Decomisados		
		Cant.	%	Cant.	%	Kg
Puerto Padre	42 969	12 383	28,8	10 418	24,2	23 046,5
Jesús Menéndez	1 186	435	36,6	370	31,2	716
Majibacoa	3 675	12	0,3	12	0,32	48,5
Manatí	10 905	7 328	67,2	2 699	24,7	9 165,6
Las Tunas	81 985	4 367	5,3	3 097	3,7	5 691,4
Colombia	63 050	712	1,13	452	0,7	1 033,5
Amancio R.	6 511	535	8,22	429	6,5	1 639
Jobabo	3 312	25	0,75	25	1,0	11,4

áreas nuevas, y se cree que fue resultado de veranos más lluviosos e inviernos más cálidos. Uno de los motivos de variación se explicó por la combinación de factores de manejo y lluvia promedio por cinco años durante mayo, junio y julio (Howell *et al.*, 2015).

En Suiza en el año 2006 se informó 18,0%, después del análisis anatomopatológico de 1131 animales que llegaron al matadero (Rapsch *et al.*, 2006). En cinco distritos de Punjab,

Pakistán reportaron afectaciones elevadas (25,46 %) de *F. hepatica* de 1 190 animales en 200 granjas (Khan *et al.*, 2009).

Thanh (2012b) en su estudio sobre fasciolosis en la zona centro-sur de Vietnam obtuvo resultados entre 60-70 % de prevalencias en cinco distritos y 27-32 % en otros tres y refirió los rangos de infección bovina entre un 20 y 90 % en las regiones del norte de Vietnam. Además, valoró la ubicación de los animales en pastoreo en zonas de riego o ríos, y en montañas

Por otro lado, en los climas tropicales como en América del Sur, aparecen informes en el matadero del estado de Trujillo, Venezuela donde la incidencia global de los nueve años estudiados fue 22,25 por cada 1000 animales sacrificados (Hernandez y Rojas, 2009). En Brasil estudiaron 12,3 millones de ganado sacrificado en cinco años y resultaron en afectaciones de 18,7 % y 10,1 % en los estados de Río Grande del Sul y Santa Catarina (Dutra *et al.*, 2010b). En Chile, la fasciolosis tiene una amplia distribución, y afecta a animales de abasto y silvestres, con una prevalencia de 87,45 % en bovinos.

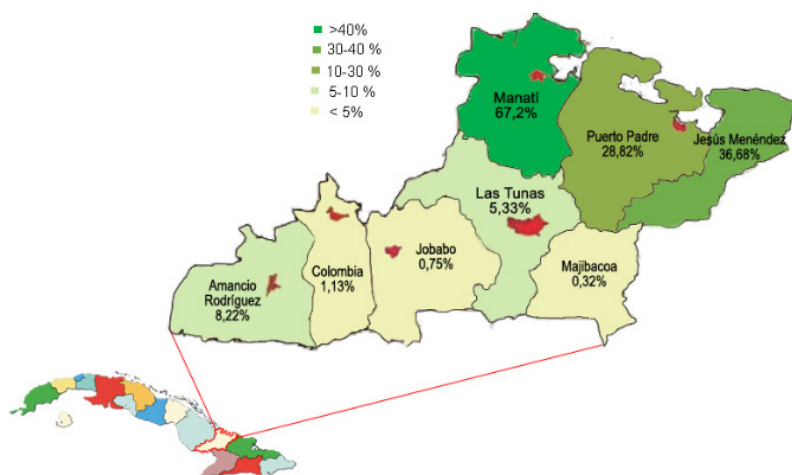
En Quindío, uno de los departamentos o estados de Colombia, se encontraron valores de hasta 3,74 % (Recalde Reyes *et al.*, 2014); y en la región lechera altoandina del departamento de Antioquia registraron hatos con prevalencias mayores al 30% y municipios con prevalencias que superan el 40 % (Morales *et al.*, 2013). En el distrito de Yabebyry, departamento de Misiones, Paraguay informaron la prevalencia en 80 % (Giménez *et al.*, 2014).

En La Habana, se encontraron prevalencias superiores al 30 % por año en una Empresa Pecuaria (González *et al.*, 2007a). Se concluye que esta parasitosis provocó una afectación en 1 de cada 3 bovinos que se sacrificaron en el matadero. En el centro del país se obtuvieron afectaciones de 37,2% en Villa Clara, 20,1 % en Cienfuegos y 50,9 % en Sancti Spíritus en los años 2000-2004 (Brito *et al.*, 2010b). En un matadero de

Camagüey sacrificaron 18 134 animales con afectaciones de 2 094 animales, en seis años (Arteaga, 2014).

Los datos por municipios mostraron diferencias en la prevalencia, con valores por debajo del 5 % en los municipios del Sur de la provincia (Figura 6). Los municipios de Jesús Menéndez, Majibacoa, Amancio y Jobabo deben su menor número de sacrificios a que proceden solo de losas.

Figura 6. Prevalencia de fasciolosis y decomisos de hígado en los municipios de la provincia Las Tunas en el periodo 2006-2015



Los municipios de Camagüey presentaron prevalencias que oscilaban entre 30-80 % en el año 2009 y 2010 en un promedio de matanzas de 25 animales diarios, además el número de animales procedentes de cada municipio fluctúa entre 50-150 en solo un año (Arteaga, 2014).

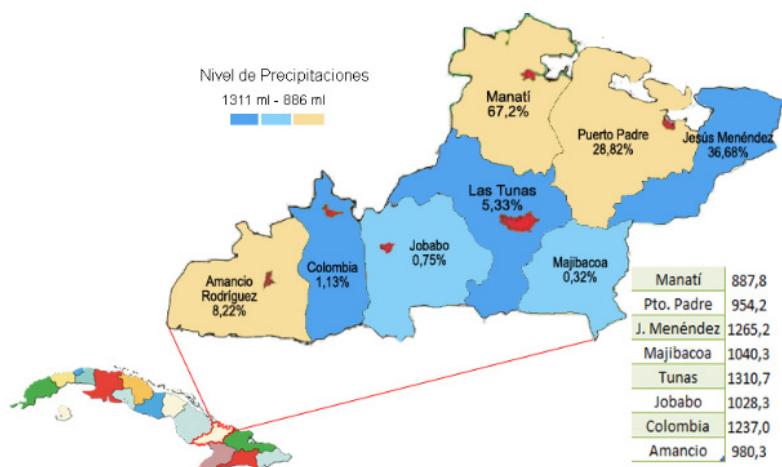
Los municipios del Norte de la provincia de Las Tunas (Manatí, Puerto Padre y Jesús Menéndez) mantienen constante y con poca variabilidad la alta afectación por la enfermedad (Tabla 2) y entre ellos el municipio de Manatí (67,2%), cuya proporción se elevó a más de la mitad de los animales sacrificados. Por el contrario de Camagüey presenta a Nuevitas (84%) municipio

del Norte y del Sur a Vertiente (58%), con mayor afectación por *Fasciola*.

El comportamiento estacional en los municipios del Norte podría favorecerse de factores climáticos. Un factor en común de ellos es su posición geográfica, estando ubicados en la zona costera por lo que el clima es más fresco por las brizas del mar y la humedad relativa es más alta. Coincidiendo con las situaciones geográficas de trabajos donde la alta prevalencia es en zonas costeras (Dutra *et al.*, 2010b, Howell *et al.*, 2015, Thanh, 2012b).

Al analizar las precipitaciones de la provincia en este periodo obtuvimos que solo Jesús Menéndez se comportó con valores superiores de lluvia y altas afectaciones (Figura 7).

Figura 7. Comportamiento de las precipitaciones en la provincia Las Tunas en el período 2006-2015



En una investigación realizada en Australia, emplearon el método de cálculo Bayesiano para relacionar prevalencia con precipitaciones. Se obtuvo una evidente relación negativa, pero la probabilidad de uno emergió cuando áreas regadas fueron

identificadas (Durr *et al.*, 2005). También resultaron negativos los cálculos por el método de Spearman que usaran para relacionar el nivel de infestación con niveles de lluvia (Bennema *et al.*, 2011). Hasta ahora, numerosos estudios se enfocan en describir la distribución espacial de la presencia de infección por *Fasciola hepatica*, pero los resultados no permiten concluir si la presencia de la infección esta también asociada con pérdidas significantes de producción (Bennema *et al.*, 2011).

El decomiso de hígados por animales sacrificados fue de 8,21% (Tabla 2) y a más de la mitad de los animales afectados se le decomiso el hígado totalmente (65,63 %), un reflejo del grado de alteración de este órgano. En Kenia se decomisan anualmente entre el 10 y el 20% de los hígados. En Brasil el grado de decomiso 50,6 %, 61,2 % y 42 % en los estados de Grande do Sul, Santa Catarina y Paraná respectivamente. En Trujillo Venezuela se retiraron de la cadena de producción 3 037 hígados, por presentar alta infestación por *Fasciola* en nueve años.

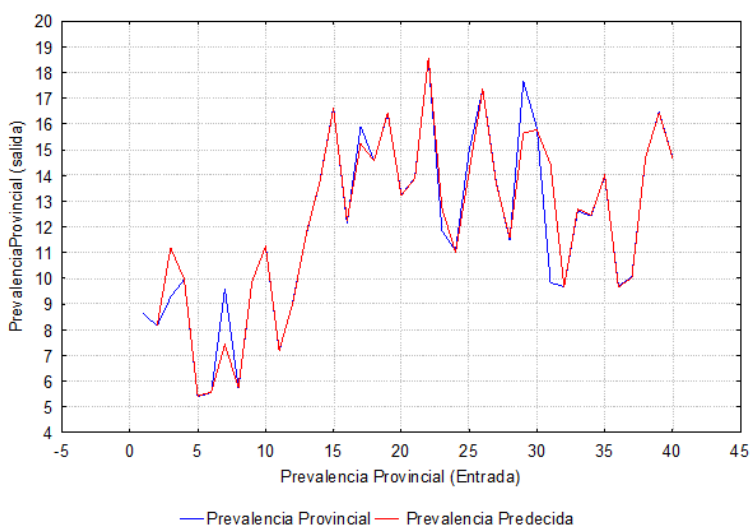
En Cuba se informan decomisos de 15,76 % en un periodo de 4 años en una empresa ganadera en La Habana (González *et al.*, 2007a); y en el centro del país un 18% promedio en tres provincias en un periodo de cinco años (Brito *et al.*, 2010b). En Pinar del Rio se decomisaron en un periodo de diez años 21 439 hígados en los mataderos, un total de 85 756 kg.

Realizada la evaluación de la prevalencia de *Fasciola hepatica* en la provincia de las Tunas, se procedió a emplear la base de datos obtenida para modelar la Red Neuronal Artificial contribuyendo este modelo a la predicción de fasciolosis en la provincia. Para ello se tomó en cuenta los factores favorecedores de la prevalencia. Al entrenar la red se asignó la capa de entrada la prevalencia que sería la predicha, de capas intermedias los datos de prevalencia de los municipios, precipitaciones de los municipios y los meses y años. Se probaron los datos mensualmente y trimestral. De los modelos obtenidos se escogió el que mostro validación mayor.

El resultado se obtuvo por MLP con función de activación Identidad-Exponencial. El valor de $R=0,67$ es una fuerte indicación que el modelo es efectivo en la predicción de *F. hepatica* en las series de tiempo de Redes Neuronales. Santana (2010) explicó que cuando una validación responde con $R>0,5$ los valores de aprendizaje se aproximan correctamente. Podemos observarlo en la Figura 6 donde se presenta los valores predichos y la prevalencia real, comportándose similar. En la investigación de Haydock *et al.* (2016) con regresión lineal obtuvieron un $R=0,71$ con este resultado logró una buena predicción de la infestación de fasciolosis en Nueva Zelanda.

En la Figura 8 muestra la relación entre el valor real y el que se predice por Red Neuronal Artificial, con un residual entre 0,5 y -0,5, solo ocho valores están fuera de este rango (Figura 9). El coeficiente de determinación fue $R^2= 0,91$ y el de una sola neurona fue $R^2=0,34$; lo que demuestra que mientras mayor información tenga la red neuronal resulta en mejor predicción.

Figura 8. Predicción por Redes Neuronales en series de tiempo (regresión) de la prevalencia de *Fasciola hepatica*.



En el informe de predicción del riesgo de infestación por fasciolosis con modelos de matemáticos en Nueva Zelanda obtuvieron una $R^2=0,71$ y con ello encontraron que el modelo fue efectivo (Haydock *et al.*, 2016). Otro informe, en Río Grande do Sul, Brasil emplearon regresión lineal en la predicción de las aéreas de riesgo usando prevalencias de ocho años y variables climáticas. Obtuvieron coeficientes de correlación de 0,32 y demostraron la importancia del uso de variables del clima como temperatura y lluvias para predecir el comportamiento de la *Fasciola* (Pereira *et al.*, 2016).

En la predicción de la incidencia de la malaria en el estado de Roraima en Brasil, las redes neuronales mostraron menor errores cuadráticos absolutos de los valores predictivos comparados con el modelo de regresión lineal.

Las redes neuronales son una herramienta poderosa para predecir la prevalencia de *Fasciola hepatica*.

Referencias del capítulo

Agrawal, A., Kumar, V., Pandey, A. & Khan, I. (2012). An application of time series analysis for weather forecasting. *International Journal Of Engineering Research And Applications (IJERA)*, 2, 974-980.

Aichouri, I., Hani, A., Bougherira, N., Djabri, L., Chaffai, H. & Lallahem, S. (2015). River flow model using artificial neural networks. *Energy Procedia*, 74, 1007-1014.

Almásí, A.-D., Woźniak, S., Cristea, V., Leblebici, Y. & Engbersen, T. (2015). Review of advances in neural networks: neural design technology stack. *Neurocomputing*.

Arenal, A., Arteaga, A., Rodríguez, I., Vázquez, R., Filiberto, Y., Casaert, S., López, V., Charlier, J. & Vercruysse, J. (2014). Red neuronal artificial para mejorar el valor epidemiológico de los estudios de mataderos para la fasciolosis bovina. 8th cuban congress on microbiology and parasitology, 5th national Congress On Tropical Medicine and 5th International Symposium on Hiv/Aids Infection in Cuba.

Arteaga, A. (2014). *Redes neuronales en la predicción de prevalencia de Fasciola hepatica en el matadero chacuba de camagüey*. Tesis de maestría en Ciencias del Diagnóstico Veterinario, Camagüey, Universidad Ignacio Agramonte Loynaz.

Atkov, O. Y., Gorokhova, S. G., Sboev, A. G., Generozov, E. V., Muraseyeva, E. V., Moroshkina, S. Y. & Cherniy, N. N. (2012). Coronary heart disease diagnosis by artificial neural networks including genetic polymorphisms and clinical parameters. *Journal of Cardiology*, 59, 190-194.

Batuello, J. T., Gamito, E. J., Crawford, E. D., Han, M., Partin, A. W., Mcleod, D. G. & O'donnell, C. (2001). Artificial neural network model for the assessment of lymph node spread in patients with clinically localized prostate cancer. *Urology*, 57, 481-485.

Bennema, S. C., Ducheyne, E., Vercruysse, J., Claerebout, E., Hendrickx, G. & Charlier, J. (2011). Relative importance of management, meteorological and environmental factors in the spatial distribution of *Fasciola hepatica* in dairy cattle in a temperate climate zone. *International Journal for Parasitology*, 41, 225-233.

Brito, A. E., Barreto, M. A. H., Rodríguez, P. D. L. F. & Prado, E. A. S. (2010^a). Prevalencia, decomisos de hígado y pérdidas económicas por *Fasciola hepatica* en mataderos bovinos de tres provincias de la región central de Cuba . *Abril/2010*, vol. 11.

Brito, A. E., Hernández Barreto, M. A., De La Fé Rodríguez, P. & Silveira Prado, E. A. (2010b). *Prevalencia, decomisos de hígado y pérdidas económicas por Fasciola hepatica en mataderos bovinos de tres provincias de la región central de Cuba*. *Redvet. Revista electrónica de veterinaria*, 11, 1-7.

Cunha, G., Luitgards-Moura, J. F., Naves, E. L. M., Andrade, A. O., Pereira, A. A. & Milagre, S. T. (2010). A utilização de uma rede neural artificial para previsão da incidência da malária no

município de cantá, estado de roraima. *Revista da sociedade brasileira de medicina tropical*, 43, 567-570.

Da Cunha, G. B., Luitgards-Moura, J. F., Martins Naves, E. L., Oliveira Andrade, A., Alves Pereira, A. & Milagre, S. T. (2010). A utilização de uma rede neural artificial para previsão da incidência da malária no município de cantá, estado de roraima. *Revista da sociedade brasileira de medicina tropical*, 43, 567-570.

De Waal, T., Relf, V., Good, B., Gray, J., Murphy, T., Forbes, A., Mulcahy, G., Holden, N., Hochstrasser, T. & Schulte, R. (2007). Developing models for the prediction of fasciolosis in ireland. Making science work on the farm: a workshop on decision support systems for irish agriculture, 60-63.

Durr, P. A., Tait, N. & Lawson, A. B. (2005). *Bayesian hierarchical modelling to enhance the epidemiological value of abattoir surveys for bovine fasciolosis*.

Dutra, L. H., Molento, M. B., Naumann, C. R. C., Biondo, A. W., Fortes, F. S., Savio, D. & Malone, J. B. (2010a). *Mapping risk of bovine fasciolosis in the south of brazil using geographic information systems*.

Dutra, L. H., Molento, M. B., Naumann, C. R. C., Biondo, A. W., Fortes, F. S., Savio, D. & Malone, J. B. (2010a). Mapping risk of bovine fasciolosis in the south of brazil using geographic information systems. *Veterinary parasitology* 169 76-81.

Fox, N. J., White, P. C., McClean, C. J., Marion, G., Evans, A. & Hutchings, M. R. (2011). Predicting impacts of climate change on *Fasciola hepatica* risk. *Plos One*, 6, e16126.

Fredes, F., Gorman, T., Silva, M. & Alcaíno, H. (1997). Evaluación diagnóstica de fracciones cromatográficas de *Fasciola hepatica* mediante western blot y elisa en animales infectados. *Archivos de Medicina Veterinaria*, 29, 283-294.

Giménez, T., Núñez, A., Chamorro, N. & Alarcón, G. (2014). Compendio de estudio de la infección natural por *Fasciola hepatica* en *lymnaea spp.* En el distrito de Yabebry,

departamento de misiones–Paraguay. Compendio de ciencias veterinarias.

González, N. (2014). Modelo matemático para el diagnóstico y pronóstico de fasciola hepática en el ganado bovino. *Redvet. Revista Electrónica de Veterinaria*, 15, 1-10.

González, R., Pérez, M. & Brito, S. (2007^a). Fasciolosis bovina. Evaluación de las principales pérdidas provocadas en una empresa ganadera. *Revista Salud Animal*, 29, 167-175.

González, R., Pérez Ruano, M. & Brito, S. (2007^b). Fasciolosis bovina. Evaluación de las principales pérdidas provocadas en una empresa ganadera. *Revista Salud Animal*, 29, 167-175.

Gracia, B. G. (2012). Cómo reducir las pérdidas causadas por la fasciolosis hepática: grave impacto económico en la industria.

Haydock, L. A. J., Pomroy, W. E., Stevenson, M. A. & Lawrence, K. E. (2016). A growing degree-day model for determination of *Fasciola hepatica* infection risk in new zealand with future predictions using climate change models. *Veterinary parasitology*, s0304-4017, 30197-2.

Hernandez, E. & Rojas, E. (2009). *Fasciola hepatica* en el matadero industrial trujillo, venezuela. *Instituto experimental “Jose Witremundo Torrealba”, U.L.A-N.U.R.R, Trujillo-Venezuela*, 3.

Howell, A., Baylis, M., Smith, R., Pinchbeck, G. & Williams, D. (2015). Epidemiology and impact of *Fasciola hepatica* exposure in high-yielding dairy herds. *Preventive Veterinary Medicine*, 121, 41-48.

IMV (2007). Reporte anual de fasciolosis bovina en cuba. *Instituto de medicina veterinaria de Cuba*.

Khan, M. K., Sajid, M. S., Khan, M. N., Iqbal, Z. & Iqbal, M. U. (2009). *Bovine fasciolosis: prevalence, effects of treatment on productivity and cost benefit analysis in five districts of punjab, pakistan*. Elsevier.

Matich, D. J. (2001). *Redes neuronales: conceptos básicos y aplicaciones*. Tesis de Maestría, Universidad Tecnológica Nacional.

Mccann, C. M., Baylis, M. & Williams, D. J. (2010). The development of linear regression models using environmental variables to explain the spatial distribution of *Fasciola hepatica* infection in dairy herds in England and Wales. *International Journal for Parasitology*, 40, 1021-1028.

Morales, C. G., Sánchez, G. A., Jiménez, C. C. C., Carmona, C. G., Pérez, F. M. & Trujillo, L. E. V. (2013). Control de *Fasciola hepatica* en el agua de consumo animal a través de filtración rápida y lenta. *Revista EIA*, 133-141.

Palacio, D. (2016). *Comportamiento estacional de la infestación por Fasciola hepatica en bovinos sacrificados en el matadero Chacuba de la provincia de Camagüey en el período 2008-2014*. Tesis de maestría en Diagnóstico veterinario, Universidad de Camagüey, Ignacio Agramonte Loynaz.

Pereira, A. E., Da Costa, C., Vieira, L. & Beltrão, M. (2016). Assessing the risk of bovine fasciolosis using linear regression analysis for the state of Rio Grande do Sul, Brazil. *Veterinary Parasitology* 217, 7-13.

Rapsch, C., Schweizer, G., Grimm, F., Kohler, L., Bauer, C., Deplazes, P., Braun, U. & Torgerson, P. R. (2006). *Estimating the true prevalence of Fasciola hepatica in cattle slaughtered in Switzerland in the absence of an absolute diagnostic test*. *International Journal for Parasitology*.

Recalde Reyes, D. P., Padilla Sanabria, L., Giraldo Giraldo, M. I., Toro Segovia, L. J., Gonzalez, M. M. & Castaño Osorio, J. C. (2014). Prevalencia de *Fasciola hepatica*, en humanos y bovinos en el departamento del Quindío-Colombia 2012-2013. *Infectio*, 18, 153-157.

Relf, V., Good, B., Hanrahan, J., McCarthy, E., Forbes, A. & Dewaal, T. (2011). Temporal studies on *Fasciola hepatica* in *Galba truncatula* in the west of Ireland. *Veterinary parasitology*, 175, 287-292.

Rozo, B. & Margarita, W. (2009). Consideraciones sobre estrategias sostenibles para el control de *Fasciola hepatica* en Latinoamérica. *Revista Colombiana de Ciencias Pecuarias*, 14, 28-35.

Santana, J. C. (2010). Predicción de series temporales con redes neuronales: una aplicación a la inflación colombiana. *Revista Colombiana de Estadística*, 29, 77-92.

Sargent, D. J. (2001). Comparison of artificial neural networks with other statistical approaches. *Cancer*, 91, 1636-1642.

Schweizer, G., Braun, U., Deplazes, P. & Torgerson, P. (2005). Papers & articles. *The Veterinary Record*, 157, 188-193.

Thanh, N. T. G. (2012a). *Zoonotic fasciolosis in vietnam: molecular identification and geographical distribution*.

Thanh, N. T. G. (2012b). *Zoonotic fasciolosis in vietnam: molecular identification and geographical distribution*. Doctorado (phd), ghent

Theodoropoulos, G., Loumos, V., Anagnostopoulos, C., Kayafas, E. & Martinez, B. (2000). A digital image analysis and neural network based system for identification of third-stage parasitic *strongyle larvae* from domestic animals. *Computer methods and programs in biomedicine*, 62, 69–76.

Xu, J.-F., Xu, J., Li, S.-Z., Jia, T.-W., Huang, X.-B., Zhang, H.-M., Chen, M., Yang, G.-J., Gao, S.-J. & Wang, Q.-Y. (2013). Transmission risks of schistosomiasis japonica: extraction from back-propagation artificial neural network and logistic regression model. *Plos Neglected Tropical Diseases*, 7, e2123.

Zhang, X., Liu, Y., Yang, M., Zhang, T., Young, A. A. & Li, X. (2013). Comparative study of four time series methods in forecasting typhoid fever incidence in china. *Plos One*, 8, 6-11.



Capítulo 4. Aplicaciones en la Ingeniería Civil

Capítulo 4. Aplicaciones en la Ingeniería Civil

Rafael Larrúa Quevedo

Yaima Filiberto Cabrera

Yisel Larrúa Pardo

Ernesto Treviño Treviño

Resumen

En este capítulo se aborda la utilización de técnicas de inteligencia artificial, como el aproximador de funciones de los k -vecinos más cercanos (k -NN), la meta heurística *Particle Swan Optimization* (PSO) y los algoritmos genéticos (AG), así como las redes neuronales artificiales, en la predicción de respuestas térmicas y estructurales en tipologías diversas, tales como las vigas compuestas de acero y hormigón y sus conexiones, las columnas de acero restringidas y los muros de mampostería confinada. Además, se ilustra la integración de la experimentación, la modelación numérica y las técnicas de inteligencia artificial y estadísticas en la generación de datos para el desarrollo de formulaciones o métodos de diseño, y se reconocen las bondades de cada técnica y las potencialidades que ofrece su utilización de manera conjunta.

Abstract

This chapter deals with the use of artificial intelligence techniques, such as the k -nearest neighbor (k -NN) function approximator, the Particle Swan Optimization (PSO) metaheuristic, and genetic algorithms (GA), as well as artificial neural networks, in the prediction of thermal and structural responses in diverse typologies, such as steel and concrete composite beams and their connections, restrained steel columns and confined masonry walls. In addition, the integration of experimentation, numerical modeling, and artificial intelligence and statistical techniques in the generation

of data for the development of formulations or design methods is illustrated, and the benefits of each of them and the potential offered by their joint use are recognized.

Introducción

Para el estudio del comportamiento de las estructuras en Ingeniería Civil la experimentación ha sido el elemento clave para el fundamento y desarrollo de los métodos de diseño hasta nuestros días. A finales del siglo xx, con el avance y desarrollo tecnológico en el campo de la informática, se han podido desarrollar herramientas computacionales basadas en métodos numéricos que permiten simular de manera virtual experimentos reales.

Se hace necesario comentar que existe un nexo indisoluble entre simulación y experimentación, pues los modelos deben someterse a un proceso de calibración matemática y física, que garantice una aproximación adecuada y estabilidad de la solución obtenida, y la correspondencia con el fenómeno físico que se estudia. Una vez calibrado adecuadamente un modelo numérico, éste debe ser validado a partir de su verificación con respecto a ensayos similares con características relativamente distintas a los ensayos con los que se calibró inicialmente. Sólo entonces puede ser utilizado para su extrapolación, sustituyendo la experimentación física de nuevos casos, con el consecuente ahorro de recursos materiales y humanos, y ofreciendo la posibilidad de realizar estudios más amplios y representativos, lo que incluye el aprovechamiento de las amplias posibilidades de visualización que ofrecen esas técnicas.

Sin embargo, en ocasiones los modelos numéricos introducen complejidad, laboriosidad y costos computacionales, que pueden ser aminorados por medio de la introducción de las técnicas de inteligencia artificial en la generación de nuevos resultados de interés, como complemento de la experimentación

y la modelación, técnicas más convencionales en el ámbito de la Ingeniería Civil.

A lo anterior se añade, la relevancia que conceden los autores del presente trabajo a la utilización del diseño estadístico de experimentos en la planeación de los estudios y a la valoración por medio de las técnicas estadísticas de los resultados, como otro elemento clave del enfoque integrado que se expone en el presente capítulo por medio de cinco casos de estudio relacionados con diferentes tipologías estructurales y técnicas de inteligencia artificial o sus combinaciones utilizadas.

En los tres primeros casos (4.1 a 4.3) se aborda la demostración de la efectividad de las técnicas de IA en la predicción de respuestas térmicas y estructurales en tres tipologías estructurales diferentes, en tanto los restantes casos (4.4 y 4.5), unido a lo anterior, ilustran como pueden integrarse la experimentación, la modelación numérica, las técnicas de inteligencia artificial y estadísticas en la generación de datos para el desarrollo de un método gráfico de diseño y de una nueva formulación de diseño, respectivamente.

Todos los estudios que componen este artículo forman parte de investigaciones previamente publicadas por separado, en diferentes revistas científicas especializadas, como se indica al inicio de cada apartado.

4.1. Predicción de la capacidad resistente última de las conexiones tipo perno en estructuras compuestas de acero y hormigón, mediante “k- vecinos más cercanos” (k-NN) y “Optimización por Enjambre de Partículas”

En el presente caso de estudio se demuestra la efectividad del algoritmo k-NN de conjunto con la meta heurística *Particle Swarm Optimization* (PSO), en la predicción de la capacidad resistente última de la conexión tipo perno en presencia de lámina plegada perpendicular, a partir de una base de datos de resultados experimentales, y se procede a la comparación de los resultados alcanzados con los que se obtienen al evaluar

los datos considerados utilizando las formulaciones de las normativas *American Institution of Steel Construction*. (2010). y *European Committee for Standardization* (2004).

Descripción del objeto de estudio

Las estructuras compuestas acero-hormigón han sido utilizadas ampliamente en la construcción de puentes y edificios desde mediados del pasado siglo hasta la fecha; en las mismas se aprovechan las bondades de cada material, lo que las hace muy eficientes y económicas. Desde 1960 ha sido común el empleo de sistemas de pisos compuestos que incorporan el empleo de una lámina plegada de acero conformada en frío que actúa no solo como encofrado permanente de la losa de hormigón colocada *in situ* sino también como refuerzo en tracción de la losa.

Un componente esencial de una viga compuesta es la conexión entre la sección de acero y la losa de hormigón armado. Esta conexión se asegura mediante conectores que se instalan en el ala superior de la viga de acero, usualmente mediante soldadura, antes del hormigonado de la losa. Los conectores aseguran que los diferentes materiales que constituyen la sección compuesta actúen de manera conjunta. Una amplia variedad de tipos han sido utilizados como conectores, en tanto diversas consideraciones de comportamiento y economía continúan motivando el desarrollo de nuevos productos. Actualmente, el conector tipo perno con cabeza (*stud*) es el más ampliamente difundido debido a su probado desempeño y a la facilidad en su colocación con tecnología específicamente diseñada a tal efecto (Ver Figura 1).

Puede afirmarse que han sido desarrollados un número notable de estudios experimentales para profundizar en el comportamiento de las conexiones. Específicamente, los ensayos de conectores del tipo *push-outh* han sido una importante vía para la evaluación de la influencia de diferentes parámetros en el comportamiento de estas, así como para la obtención de formulaciones que permitan predecir su capacidad resistente última, lo que ha marcado significativamente la evolución de

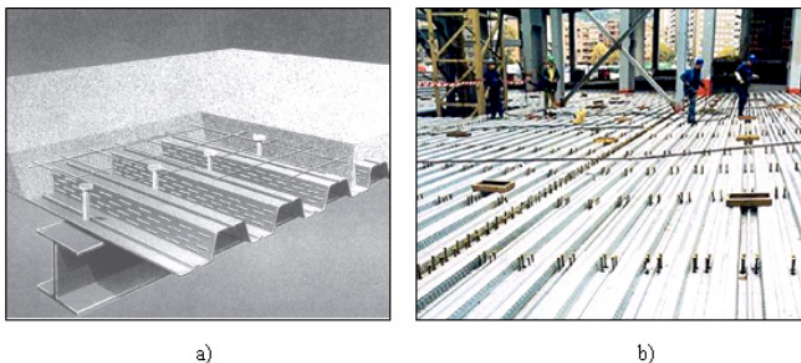


Figura 1. Conectores tipo perno con cabeza (*stud*) colocados sobre vigas laminadas y láminas plegadas con nervaduras perpendiculares. a) Esquema general. b) Entrepiso en construcción antes del hormigonado. Fuentes: Elaboración propia/Catálogo Canan Group Inc.

los métodos de cálculo. Como valor agregado a lo anterior, es posible contar con la valiosa información derivada del conjunto de programas experimentales internacionales realizados en el ámbito de las conexiones en construcción compuesta acero-hormigón, en la que prevalecen los estudios dedicados a la tipología perno con cabeza (*stud*).

El espécimen para el ensayo de conectores está compuesto por un perfil de acero de sección I o dos perfiles T que se unen para formar una sección I (Lyons, 1992; Rambo-Roddenberry, 2002), al cual son soldados simétricamente en las caras exteriores de sus alas los conectores a estudiar. Posteriormente se hormigonan losas a ambos lados, quedando embebidos en ellas los conectores. Estas losas pueden ser rectangulares macizas de hormigón o losas compuestas con lámina metálica plegada. En este segundo caso la losa puede estar ubicada con las nervaduras paralelas o perpendiculares a la viga. Al espécimen se le aplica una carga centrada controlada hasta que ocurre el fallo de la conexión, determinándose así la capacidad resistente última de la misma (ver Figura 2).

Para la tipología considerada los conectores pueden estar ubicados en cantidades de uno o dos y en diversas posiciones

en el valle de la lámina y no está reglamentada una probeta estándar (ver Figura 4.3 (d, e, f)). Se destacan por su relevancia las investigaciones de Robinson (1988), Sublett *et al.*(1992), Lyons (1994), Díaz *et al.* (1998), Johnson y Yuan (1998), y Rambo-Roddenberry (2002).



Figura 2. Ensayo de conexiones con lámina plegada perpendicular. a) Vista general del ensayo. b) Detalle del espécimen. Fuentes: Elaboración propia.

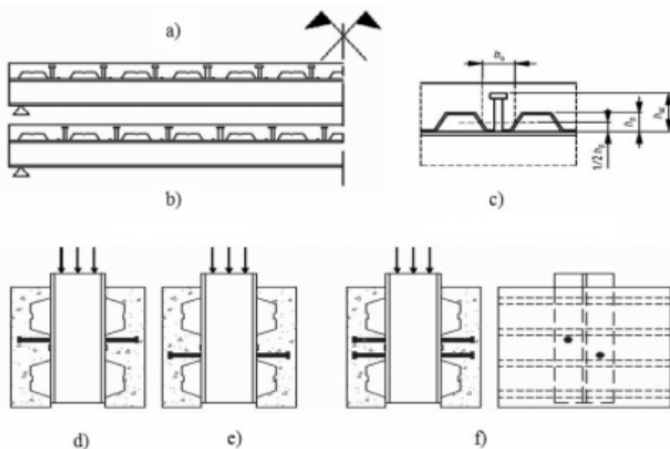


Figura 3. Posiciones de los conectores en el valle de la lámina: a) y d) Posición favorable, b) y e) Posición desfavorable, c) Posición central. Dimensiones generales, f) Posición escalonada. Fuente: Elaboración propia.

Concepción de la aplicación de las técnicas de inteligencia artificial

El objetivo central es demostrar la efectividad del algoritmo *k-NN*, como aproximador de funciones conjuntamente con la meta heurística Particle Swan Optimization (*PSO*) y una función objetivo basada en lograr la mayor similaridad posible entre la semejanza entre objetos según los rasgos predictores y la que se obtiene a partir del rasgo objetivo, que permiten realizar la asignación de pesos a los atributos considerados por la función de semejanza del método *k-NN*, en la predicción de la capacidad resistente última de la conexión seleccionada, para las múltiples condiciones de diseño que pueden presentarse, a partir de una base de datos de resultados experimentales, de cara al perfeccionamiento de los métodos de cálculo.

El procedimiento seguido es el siguiente:

1. Obtención de la base de datos de resultados experimentales

Se define una base de datos única a partir de los datos experimentales relativos a un total de 66 ensayos internacionales, que definen el total de registros o instancias, con las variables de entrada (8) descritas en la Tabla 1 y como variable de respuesta única la capacidad resistente última de la conexión.

2. Definición de las técnicas de inteligencia artificial a utilizar y de la herramienta computacional a desarrollar

Se decide la utilización del algoritmo *k-NN*. Para definir el peso de las variables de entrada, se acude a la técnica de optimización *PSO*, basada en el comportamiento de una población como los enjambres de abejas, cardúmenes de peces o bandadas de pájaros.

Para facilitar la predicción de la capacidad resistente última de las conexiones fue desarrollado un sistema automatizado, denominado PROCON (Pronósticos en la Construcción), en

Tabla 1 Variables de entrada. Fuente: Elaboración propia.

N°	Descripción	Valores
1	Área de la sección transversal del vástago del perno (A)	Áreas (m ²) para pernos con diámetros: 9,53; 12,70; 15,88; 19,05; 22,23 (mm)
2	Número de conectores en el valle de la lámina (n)	1;2
3	Posición del conector en el valle de la lámina (p) (ver Figura 2)	Favorable (1); desfavorable (2); central (3); escalonada (4)
4	Ancho medio del valle de la lámina (bo) (ver Figura 2c)	0,113; 0,140; 0,153; 0,229 (m)
5	Altura de la lámina (hP) (ver Figura 2c)	0,025; 0,051; 0,060; 0,076; 0,080; 0,114; 0,152 (m)
6	Altura del conector (hSC) (ver Figura 2c)	,064; 0,089; 0,095; 0,0102; 0,114; 0,127; 0,152; 0,197 (m)
7	Resistencia del hormigón a la compresión (f'c)	Entre 18,41 y 48,82 (MPa)
8	Resistencia última a la tracción del acero del perno (Fy)	Entre 430,94 y 546,08 (MPa)

el que fueron implementados los algoritmos requeridos. El sistema sigue los lineamientos de software libre y utiliza el patrón arquitectónico Modelo-Vista-Controlador (MVC), cuenta con Java como lenguaje de programación y Eclipse como plataforma y está diseñado para establecer conexiones con el gestor de bases de datos MySQL.

3. Diseño y ejecución del experimento numérico. Selección de las mejores combinaciones

Se diseñó y ejecutó un experimento con el objetivo de determinar los resultados de la precisión del estimador de funciones k-NN, tomando como casos almacenados los recogidos en la base de datos (ver Tabla 1) y considerando dos variables: k (número de vecinos cercanos) y “variantes de pesos”. Fueron combinados cuatro valores de k (1, 3, 5 y k óptimo) y cinco variantes para establecer los pesos (utilizando

el algoritmo PSO, variante estándar con pesos uniformes y las restantes propuestas en base a criterios cualitativos y cuantitativos basados en estudios precedentes), configurando una matriz de 5 x 4 para un total de 20 combinaciones. Para cada combinación se utilizó Validación Cruzada de K-Particiones o Validación Cruzada de K-Fold (*K-Folds Cross-Validation*), subdividiendo el conjunto de datos en 10 subconjuntos de igual tamaño que fueron utilizados una vez como conjunto de prueba mientras los demás formaban el conjunto de entrenamiento. Finalmente, para cada combinación se calculó la eficacia del algoritmo considerando las precisiones obtenidas con todos los subconjuntos de prueba, tomando en cuenta las medidas estadísticas: Error Porcentual Absoluto Medio (MAPE)-, Raíz del Error Cuadrático Medio (RMSE), el promedio del módulo de las diferencias entre los valores experimentales y los que se obtienen por la predicción (PMD), el coeficiente de correlación, el coeficiente de determinación R^2 , así como la relación entre los valores experimentales y las predicciones (Q_{exp} / Q_m).

4. Comparación entre resultados experimentales, de inteligencia artificial y normativas internacionales

Una vez determinada la variante de mayor efectividad, se procedió a comparar los resultados alcanzados con los que se obtienen al evaluar los datos considerados utilizando las formulaciones de las normativas AISC (2005) y *European Committee for Standardization* (2004).

Principales resultados

Se observa que siempre se alcanza la mayor efectividad en la predicción, con base en las medidas estadísticas utilizadas, cuando se combina la determinación de los pesos según el algoritmo *PSO* con el valor de k igual a 1.

Además, tal como muestra la Figura 4, que ilustra la comparación entre las predicciones con Inteligencia Artificial y

normas internacionales (*Eurocode-4* y *AISC*); la relación Q_{exp} / Q_m para el método propuesto (IA) se aproxima de manera certera a la unidad, en tanto los valores mínimo y máximo se separan moderadamente de 1, mientras la propia relación presenta resultados más desfavorables para las normas internacionales valoradas.

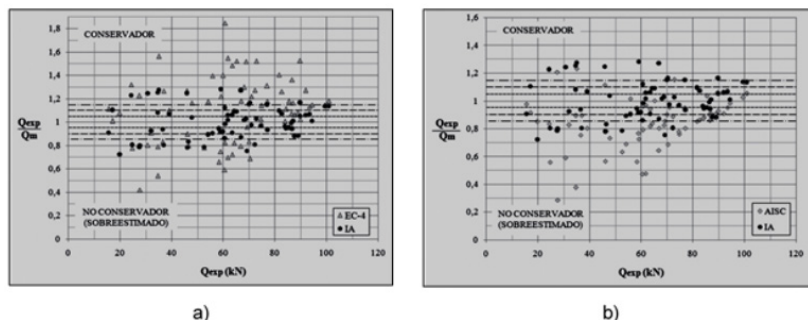


Figura 4. Comparación entre las predicciones con inteligencia artificial y normas internacionales. a) EC-4. b) AISC.

En resumen, puede afirmarse que se obtienen resultados promisorios que demuestran que el algoritmo $k - NN$ y el método de cálculo de pesos de los atributos (*PSO*) que se propone constituyen una técnica eficaz que favorece la creación de nuevos datos para la generación de un conjunto representativo de las posibles situaciones de diseño, de cara al perfeccionamiento de los métodos de cálculo, de una manera rápida y simple, complementando a la experimentación clásica y la simulación numérica. A su vez, libera al investigador del área de Ingeniería Civil de la definición de los pesos de las variables de entrada con base en otros criterios cualitativos o cuantitativos menos fundamentados. Poder contar con la herramienta computacional *PROCON* favorece notablemente la amplitud de la experimentación desarrollada. La concepción general de este instrumento y sus facilidades probadas permiten su uso generalizado en la predicción de la respuesta de otras

conexiones acero-hormigón y de tipologías estructurales diversas.

4.2. Predicción de la respuesta de muros de mampostería confinada, mediante k-NN y PSO

En el presente caso de estudio se demuestra, como en el anterior, la efectividad del algoritmo *k-NN* de conjunto con la meta heurística *Particle Swan Optimization (PSO)*, implementados en PROCON, herramienta computacional propia, en la predicción de la respuesta de muros de mampostería confinada ante cargas laterales cíclicas reversibles en su plano, que simulan la acción sísmica, a partir de una base de datos de resultados experimentales, y se procede a la comparación de los resultados alcanzados con los pronósticos realizados mediante las expresiones de las “Normas Técnicas Complementarias para el Diseño y Construcción de Estructuras de Mampostería” (NTC-DF-2004) (GD, 2004) y mediante las expresiones propuestas por Riahi *et al.* (2008). El caso cuenta con las particularidades, respecto al descrito en 4.1, dadas por la mayor complejidad de la experimentación que sirve de base, la mayor cantidad de variables de entrada y de respuestas a evaluar, así como por la mayor dispersión de los resultados de estas últimas.

Descripción del objeto de estudio

La mampostería es uno de los materiales de construcción con mayor diversidad de usos tanto en el pasado como en el presente. Es quizás el más antiguo que se mantiene en uso generalizado alrededor del mundo en la actualidad, y fue sin dudas el más importante material de construcción, hasta el advenimiento del hormigón armado y del acero estructural hace aproximadamente una centuria.

En la mampostería confinada, objeto de estudio del presente apartado, el acero de refuerzo se concentra dentro de

elementos de hormigón, de sección transversal relativamente pequeña, de dimensiones típicamente iguales al espesor del muro, que lo rodean o confinan; y cuyas funciones primordiales son proporcionar continuidad a los muros, tanto entre sí como con las losas de entrepiso y cubierta, y mejorar las capacidades de resistencia, deformación y disipación de energía, ante sollicitaciones laterales cíclicas reversibles como las inducidas por los sismos.

En la Figura 5 se presenta la imagen de una edificación en construcción que cuenta con muros de mampostería confinada debidamente concebidos. Además se muestra el dispositivo para la investigación experimental de este tipo de muros ante cargas laterales cíclicas reversibles en su plano, representativas de la acción sísmica.



Figura 5. Mampostería confinada. a) Edificación en construcción. b) Dispositivo para la investigación experimental ante cargas laterales cíclicas reversibles en su plano. Fuentes: Elaboración propia

Concepción de la aplicación de las técnicas de inteligencia artificial

El objetivo central es valorar la efectividad del enfoque de trabajo y de la herramienta computacional PROCON descritos en 4.1, en la predicción de dos de las principales respuestas experimentales de los muros estudiados, como son el cortante

de agrietamiento (V_{agr}) y el cortante máximo o resistencia (V_{max}) y comparar los resultados con respecto a pronósticos realizados mediante las expresiones de la norma mexicana (GD, 2004) y mediante las expresiones propuestas por Riahi *et al.* (2008).

A continuación se describe el procedimiento seguido que coincide, en términos generales, con el planteado en 4.1; se exponen las particularidades asociadas al nuevo objeto de estudio.

1. Obtención de las bases de datos de resultados experimentales

Se parte de una extensa base de datos conformada por los resultados de programas de investigación experimental realizados durante los últimos 45 años, provenientes de fuentes internacionales diversas y confiables, que incluye los datos de 428 ensayos, obtenidos de referencias reunidas de nueve países.

Los datos o variables de entrada incluyen, entre otros: a) dimensiones de las piezas; b) dimensiones de los especímenes, tanto del tablero de mampostería como de los elementos verticales de confinamiento; c) armado de los elementos verticales de confinamiento; d) refuerzo horizontal y vertical distribuido a lo largo del tablero de mampostería; e) características de los materiales; y f) tipo y modalidad del ensayo. Los tipos de variables, se clasifican en discretas (D) o continuas (C).

Entre las respuestas o variables de salida se incluyen, entre otras: a) el tipo de fallo predominante; b) fuerzas y deformaciones para varios estados del comportamiento como son el primer agrietamiento y la resistencia; c) fuerza y deformación para el estado de daño posterior a la resistencia asociado al 20 % de degradación de la misma; d) fuerza y deformación para el estado final alcanzado en el ensayo; e) la rigidez inicial o del primer ciclo; f) la rigidez al primer agrietamiento y a la resistencia; y g) la energía disipada equivalente a distorsiones de 0.004 y 0.006.

En correspondencia con los propósitos del trabajo, se analizaron cuatro variantes de bases de datos, las cuales se identificaron consecutivamente con los números romanos I, II, III y IV. El alcance definido para cada una de ellas se define tomando en cuenta las variables de entrada que intervienen en el comportamiento de los especímenes y en las restricciones que imponen las formulaciones de Riahi *et al.* (2008).

Un resumen de la composición final de las bases de datos I, II, III y IV, se muestra en la Tabla 2. Las variables de entrada de la base de datos en extenso se adecuaron en función de las características de las cuatro bases de datos para inteligencia artificial, pudiendo ajustarse a un total de 32 para cada una de ellas, de lo cual se muestra un fragmento en la Tabla 3.

C = variable continua, D = variable discreta

Tabla 2. Cantidad de instancias en las cuatro bases de datos para procesamiento por inteligencia artificial. Fuente: Elaboración propia

Clasificación	Base de datos			
	I	II	III	IV
Clase de confinamiento 2	112	112	104	31
Clase de confinamiento 1, 3, 4 y 5	146	0	0	0
Sin refuerzo adicional	131	67	59	31
Con refuerzo adicional	127	45	45	0
Acero en castillos $p < 1\%$	92	26	24	0
Acero en castillos $p \geq 1\%$	166	86	80	31
$\sigma > 0.12\text{fm}$	11	11	9	0
$\sigma \leq 0.12\text{fm}$	247	101	95	31
$0.7 \leq H/L \leq 1.2$	197	79	79	31
$H/L > 2$	8	8	0	0
$H/L < 0.7; 1.2 \leq H/L \leq 2$	53	25	25	0
Total de instancias	258	112	104	31

Tabla 3. Variables de entrada (fragmento). Fuente: Elaboración propia.

Título	Descripción / Rango de Valores / Índices	Tipo
Piezas	Tipo de piezas. (1) Sólidas; (2) Huecas.	D
An/Ab	Relación área neta y área bruta de las piezas de mampostería. 0.433-1	C
axc	Área de las piezas de mampostería. (36 - 800 cm ²)	D
Material	Material. (1) Barro artesanal; (2) Barro extruido; (3) Concreto no industrializado, tabicón; (4) Concreto industrializado; (5) Adobe; (6) Silico calcáreas.	D
Clase de confinamiento	Del muro. (1) No confinado; (2) Confinado con castillos exteriores; (3) Confinado con castillos ahogados; (4) Mal confinado con castillos exteriores; (5) Mal confinado con castillos ahogados.	D
Mod_Ensaye	Modalidad del ensayo. (1) Monótono; (2) Cíclico; (3) Dinámico; (4) Cíclico y dinámico; (5) otros	D
H/L	Relación de aspecto. (0.235-3.125)	C
L x t	Área del muro. (1110.4-12960 cm ²)	C
RefAdic	efuerzo adicional en el tablero de mampostería. (1) Ninguno; (2) Reforzado verticalmente; (3) Reforzado horizontalmente; (4) Mal reforzado verticalmente; (5) Mal reforzado horizontalmente; (6) Reforzado vertical y horizontalmente; (7) Reforzado vertical y mal reforzado horizontalmente; (8) Reforzado horizontal y mal reforzado verticalmente; (9) Mal reforzado vertical y mal reforzado horizontalmente; (10) Mortero por ambos lados; (11) Yeso por ambos lados; (12) Malla y mortero por ambos lados; (13) Malla y mortero por un lado; (14) Fibra de carbono y mal refuerzo horizontal; (15) Bandas de fibras de carbono; (16) Rehabilitado con bandas de fibras de carbono; (17) Mortero de alta adherencia con fibras a ambos lados y juntas secas; (18) Mortero con fibras empobrecidas a ambos lados y junta secas; (19) Rehabilitado con malla y mortero a ambos lados.	D
t_rec	Espesor de recubrimiento de mortero por una cara. (0 - 2.5 cm)	C
fb	Resistencia a compresión de los cubos del mortero de pega o junta. (0-438.17 kg/cm ²)	C

fb2	Resistencia a compresión de los cubos del mortero de recubrimiento. (0-165.7 kg/cm ²)	C
f'c	Resistencia del concreto en castillos y dalas. (0-540 kg/cm ²)	C
fp	Media de la resistencia a compresión de las piezas sobre el área bruta. (44.1-235 kg/cm ²)	C
fm	Media a la resistencia a compresión de las pilas. (8.704- 257.89 kg/cm ²)	C
vm	Media de la resistencia a compresión diagonal de los muretes. (0.3-16.4 kg/cm ²)	C
Em	Módulo de elasticidad de la mampostería. (2840- 205430.4 kg/cm ²)	C
Gm	Módulo de cortante de la mampostería. (1519- 82172 kg/cm ²)	
hc/bc	Razón de las dimensiones del castillo. (0-4.67 cm)	
Barra Cast	Tipo de acero utilizado en varillas longitudinales de castillos perimetrales o de confinamiento. (1) Ninguno; (2) Corrugada laminada en caliente; (3) Alambre corrugado estirado en frío; (4) Alambrón liso, (5) Malla soldada; (6) Malla de gallinero; (7) otro.	D

2. Definición de las técnicas de inteligencia artificial a utilizar y de la herramienta computacional a desarrollar

En correspondencia con 4.1 se decide la utilización del algoritmo *k-NN* de conjunto con la técnica de optimización *PSO*, implementados en *PROCON*.

3. Diseño y ejecución del experimento numérico. Selección de las mejores combinaciones

Se diseña y ejecuta un experimento con el objetivo de determinar los resultados de la precisión del estimador de funciones *k-NN*, tomando como casos almacenados los recogidos en las bases de datos consideradas (I, II, III y IV), y asignando a la variable *k* (número de vecinos cercanos)

los valores 1, 2 y 3. En cada caso se utilizó el método de Validación Cruzada usando K-Subconjuntos (*K-Fold Cross-Validation*), subdividiendo el conjunto de datos en diez subconjuntos de igual tamaño los cuales fueron utilizados una vez como conjunto de prueba, mientras los demás formaban el conjunto de entrenamiento. Para cada combinación se calculó la precisión del algoritmo tomando en cuenta las precisiones obtenidas con todos los subconjuntos de prueba, por medio, igualmente, de las medidas estadísticas descritas en 4.1.

4. Comparación entre resultados experimentales, de inteligencia artificial y normativas internacionales

Una vez determinada la variante de mayor efectividad para cada base de datos, se procede a comparar los resultados alcanzados con los que resultan al utilizar las expresiones de la norma NTC-DF (2004) y las propuestas por Riahi *et al.* (2008), respectivamente.

Principales resultados

Fue posible comprobar que, para todas las bases de datos, la mayor efectividad de la predicción se manifiesta con el valor de k igual a 2, tanto para el cortante de agrietamiento (V_{agr}) como para la resistencia (V_{max}). A los efectos de ilustrar las comparaciones entre los resultados para la base de datos I, la Figura 6 muestra gráficos de la relación (Q_{exp}/Q_m) entre los resultados experimentales del cortante de agrietamiento (V_{agr}) y los obtenidos con los diferentes métodos utilizados: inteligencia artificial ($k = 2$), expresiones de la norma NTC-DF (2004) y las propuestas por Riahi *et al.* (2008), respectivamente.

Se observa, en todos los casos, que la media de la relación Q_{EXP}/Q_M para las predicciones con inteligencia artificial se aproxima de manera certera a la unidad, en tanto la media de dicha relación se aleja de la unidad un 68% para la norma mexicana y un 12 % para las ecuaciones propuestas por Riahi (2008).

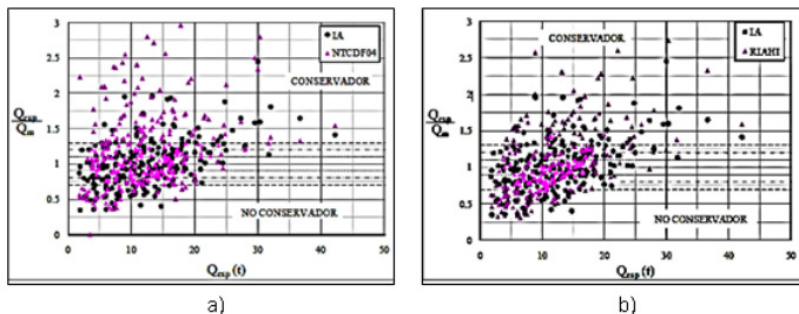


Figura 6. Relaciones Q_{exp}/Q_m para las instancias de la base de datos I en la predicción del cortante de agrietamiento (V_{agr}). a) Según IA y GDF (2004). b) Según IA y Riahi et al. (2008). Fuente: Elaboración propia

Por su parte, a los efectos de ilustrar las comparaciones entre los resultados para la base de datos IV Figura 7 muestra gráficos de la relación (Q_{exp}/Q_m) entre los resultados experimentales del cortante de agrietamiento (V_{agr}) y los obtenidos con los diferentes métodos utilizados: inteligencia artificial ($k = 2$), expresiones de la norma NTC-DF (2004) y las propuestas por Riahi et al. (2008), respectivamente.

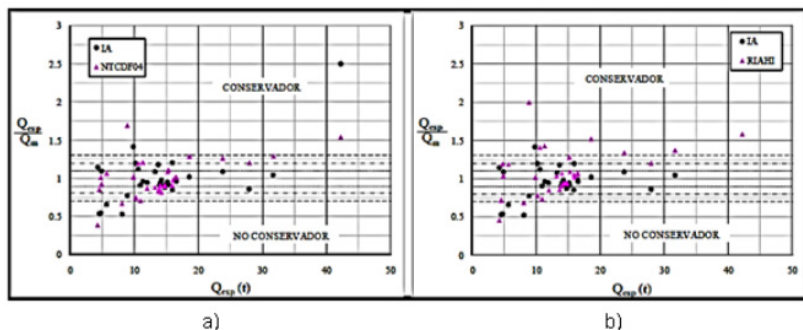


Figura 7. Relaciones Q_{exp}/Q_m para las instancias de la base de datos IV en la predicción del cortante de agrietamiento (V_{agr}). a) Según IA y GDF (2004). b) Según IA y Riahi (2008). Fuente: Elaboración propia

Se observa que la media de la relación $Q_{\text{EXP}}/Q_{\text{M}}$ para las predicciones con inteligencia artificial se aproxima, igualmente, de manera certera a la unidad. En el caso de las predicciones de la norma mexicana, la media de dicha relación se aleja de la unidad un 1 %, en tanto para las predicciones efectuadas con las ecuaciones propuestas por Riahi (2008), la media de la relación $Q_{\text{EXP}}/Q_{\text{M}}$ se aleja de la unidad un 9 %.

En resumen, se confirman las conclusiones expuestas en 4.1 respecto a la efectividad del algoritmo k -NN de conjunto con el método de cálculo de pesos de los atributos (PSO) y la utilidad de la herramienta computacional PROCON, para un objeto de estudio que se caracteriza por una mayor complejidad de la experimentación que sirve de base, mayor cantidad de variables de entrada y de respuestas a evaluar, así como, por la mayor dispersión de los resultados de estas últimas.

4.3. Predicción de respuestas térmicas y estructurales de columnas de acero restringidas en situación de incendio mediante redes neuronales.

En el presente caso de estudio se demuestra la efectividad de las redes neuronales, implementadas en MATLAB 2013, en la predicción del comportamiento térmico y estructural de columnas de acero restringidas en situación de incendio, a partir de bases de datos de origen experimental (Correia, 2011). Las particularidades del caso están dadas por la diversidad de respuestas valoradas, y su dependencia del tiempo de exposición al fuego. La información experimental disponible es amplia y suficiente para la realización de las predicciones y en consecuencia los resultados de los modelos numéricos, desarrollados por el propio autor de la investigación experimental Correia (2011), no se utilizan como insumos para las predicciones, si no para establecer comparaciones acerca de su cercanía al comportamiento experimental y a las predicciones por medio de redes neuronales.

Descripción del objeto de estudio

La columna es el elemento estructural vertical empleado para sostener la carga de las edificaciones. Es utilizada ampliamente por la posibilidad que proporciona al distribuir espacios y soportar el peso de la construcción. Es un elemento fundamental en el esquema de una estructura y la adecuada selección de su tamaño, forma, espaciamiento y composición influyen de manera directa en su capacidad de carga.

En el transcurso de la historia la resistencia al fuego de columnas de acero ha sido puesta a prueba. Con el objetivo de proporcionar un mejor comportamiento ante la influencia de las altas temperaturas en la degradación de la resistencia y la rigidez de los elementos, se han llevado a cabo programas experimentales para diferentes condiciones de apoyo, características geométricas, orientación y cargas aplicadas. Dentro de estos, ocupa un lugar relevante, por su amplitud y rigor científico, la investigación experimental y numérica, llevada a cabo por Correia (2011) en columnas de acero con restricciones.

La Figura 8 muestra el dispositivo de ensayo concebido especialmente por ese autor para realizar las pruebas experimentales en columnas cargadas con restricciones axiales y rotacionales restringidas, en situación de incendio. El sistema comprende un marco de acero de rigidez variable, con la función de simular y modificar la rigidez de la estructura circundante a la columna en situación de incendio. La columna queda en el interior de un horno vertical mediante el cual se simula un incendio estandar, según la curva ISO - 834 (Correia, 2011).

Como resultado del programa experimental se genera una abarcadora e invaluable información acerca de las respuestas obtenidas: temperaturas (a lo largo de la altura de la columna y en puntos seleccionados de la sección transversal), fuerzas de restricción, desplazamientos verticales, deflexión lateral (respecto a los ejes de mayor y menor inercia) y la rotación en la base de la columna.

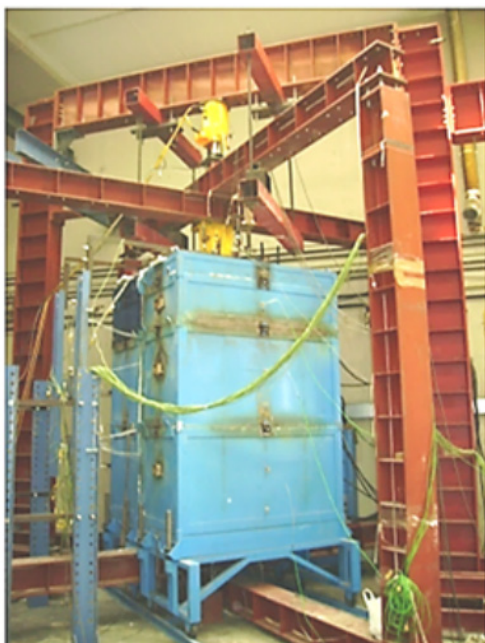


Figura 8. Dispositivo del ensayo de columna restringida en incendio. Fuente: Correia (2011)

Se consideraron diferentes variables de entrada como: características geométricas de los perfiles de acero (h , bf , t_f , t_w), restricción axial (α_A), restricción rotacional (β_R), esbeltez (λ), % de carga respecto a la máxima estimada (P), rigidez ($K_{a,s}$) y tiempo de exposición al fuego. Adicionalmente, el autor desarrolla modelos numéricos tridimensionales en el software ABAQUS 2013, calibrados frente a los resultados experimentales (ver Figura 9).

Concepción de la aplicación de las técnicas de inteligencia artificial

El objetivo central es demostrar la efectividad de las redes neuronales en la predicción de las diferentes variables de respuesta consideradas, a los efectos de comparar con los resultados que se obtienen de manera experimental o por medio de la modelación numérica tridimensional.

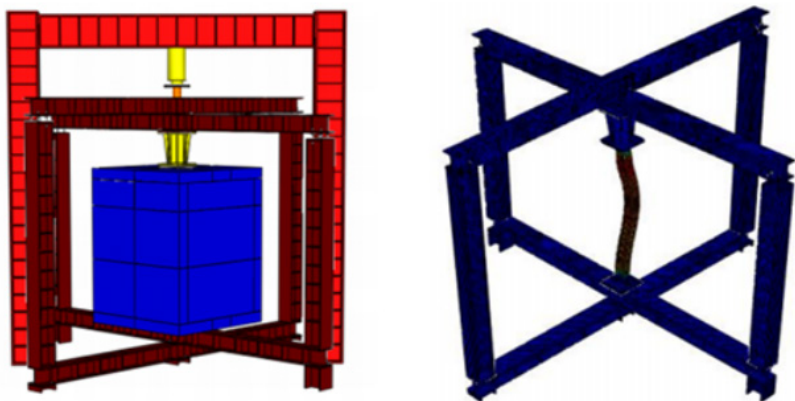


Figura 9. Modelos numéricos tridimensionales. Fuente: Correia (2011)

El procedimiento seguido es el siguiente:

1. Obtención de las bases de datos de resultados experimentales a considerar

Para obtener las bases de datos requeridas fue necesaria la conversión de la voluminosa información gráfica de origen experimental ofrecida por el autor en información numérica válida para realizar las predicciones de la evolución de las variables de respuesta antes citadas. La Tabla 4 muestra un resumen de las bases de datos concebidas, diferenciadas para cada tipo de respuesta.

2. Definición del tipo de red neuronal y de la herramienta computacional a utilizar

Se utilizan redes neuronales del tipo *Multilayer-Perceptron* (MLP), implementadas en MATLAB 2013, con conexiones hacia adelante (*feed-forward*), y aprendizaje supervisado y algoritmo de propagación de los errores hacia atrás (*backpropagation*).

Tabla 4 Resumen de las bases de datos. Fuente: Elaboración propia.

Variables de respuesta	Variables de entrada	Cantidad de instancias
Temperatura media en los extremos de la columna (TS3).	$h, bf, tf, tw, \alpha A, \beta R, \lambda$ L, P, KA, S, t	81
Temperatura media en la mitad de la altura de la columna (TS1)	$h, bf, tf, tw, \alpha A, \beta R, \lambda$ L, P, KA, S, t	81
Fuerzas de restricción (P/Po).	$h, bf, tf, tw, \alpha A, \beta R, \lambda$ L, P, KA, S, t	98
Desplazamientos verticales (DV).	$h, bf, tf, tw, \alpha A, \beta R, \lambda$ L, P, KA, S, t	98
Deflexión lateral en el eje de mayor inercia (Dmay.)	$h, bf, tf, tw, \alpha A, \beta R, \lambda$ L, P, KA, S, t	414
Deflexión lateral en el eje de menor inercia (Dmen.)	$h, bf, tf, tw, \alpha A, \beta R, \lambda$ L, P, KA, S, t	451
Rotaciones en la base (R).	$h, bf, tf, tw, \alpha A, \beta R, \lambda$ L, P, KA, S, t	84

Se definen redes con una o dos capas ocultas, con diferentes números de neuronas, con función de transferencia *tansig* para las capas ocultas, en tanto para la capa de salida se utiliza una función de transferencia *purelin*.

3. Diseño y desarrollo del experimento numérico. Selección de las mejores combinaciones

Se diseña un experimento que considera redes con una o dos capas ocultas, con un número variable de neuronas en las capas ocultas entre 10 y 50. Para determinar las combinaciones más efectivas se tienen en cuenta, igualmente, las medidas estadísticas descritas en 4.1.

4. Comparación entre resultados experimentales, de modelos numéricos y de inteligencia artificial

Se ilustra la comparación entre resultados por medios gráficos.

Principales resultados

Se evalúan 14 combinaciones para cada respuesta, para un total de 98. La Tabla 5 muestra las combinaciones más efectivas y se indica, por su representatividad, el valor de R^2 alcanzado.

Tabla 5. Selección de las mejores combinaciones para cada respuesta

Variable de respuesta	Arquitectura de la red	R ²
Temperatura media en los extremos de la columna.	50-20-1	0,998
Temperatura media en la mitad de la altura de la columna.	10-20-1	1,000
Fuerzas de restricción.	30-10-1	0,998
Desplazamientos verticales.	20-10-1	0,994
Deflexión lateral en el eje de mayor inercia.	10-40-1	0,989
Deflexión lateral en el eje de menor inercia.	40-50-1	0,981
Rotaciones en la base	40-40-1	0,993

Nótese como en todos los casos la mayor efectividad se alcanza con redes de dos capas ocultas; a su vez en las mejores combinaciones el número de neuronas en las capas ocultas es diferente en dependencia de la respuesta evaluada y en todos los casos se alcanzan valores de R^2 muy cercanos a 1, lo que es indicativo de la calidad lograda en las predicciones. En la

Figura 10, a modo de ilustración representativa de lo obtenido, se presentan los resultados de la combinación más efectiva para la respuesta “Fuerzas de restricción”. Pueden apreciarse los valores cercanos a 1 del coeficiente R.

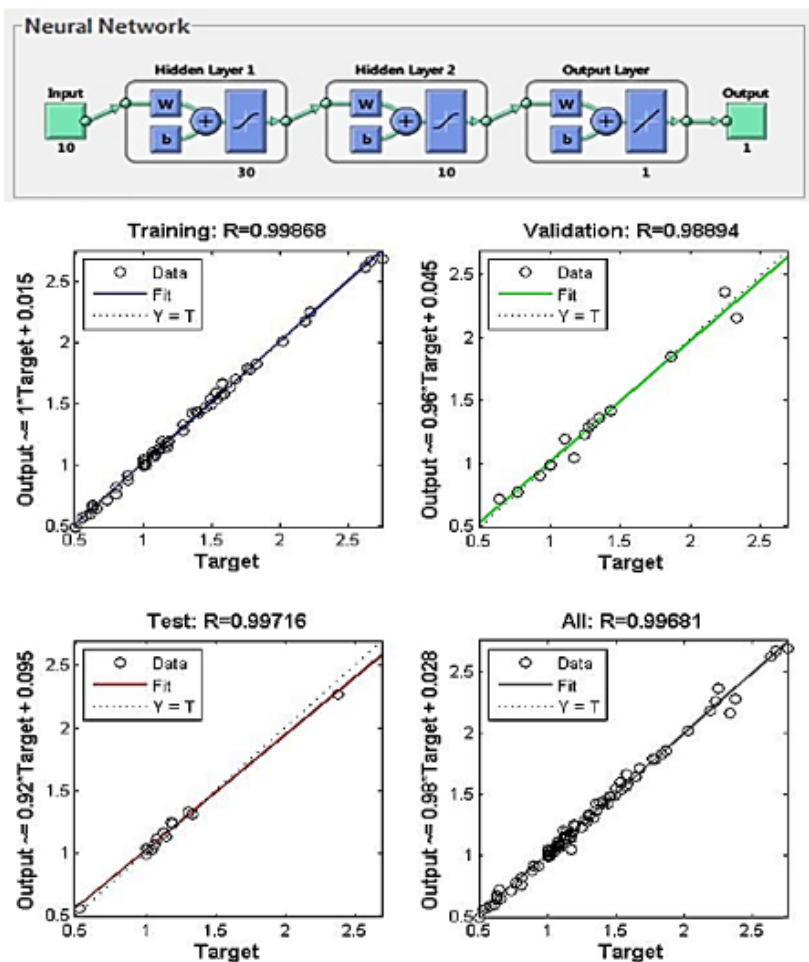


Figura 10. Resultados de la combinación más efectiva para la respuesta “fuerza de restricción”. MATLAB2013. Fuente: Elaboración propia. (con salidas de MATLAB 2013)

La Figura 11 muestra, de manera gráfica, la comparación entre las predicciones con la red neuronal 30-10-1, para la propia respuesta “Fuerzas de restricción”, con los resultados experimentales y los del modelo numérico tridimensional en ABAQUS, realizado por Correia (2011).

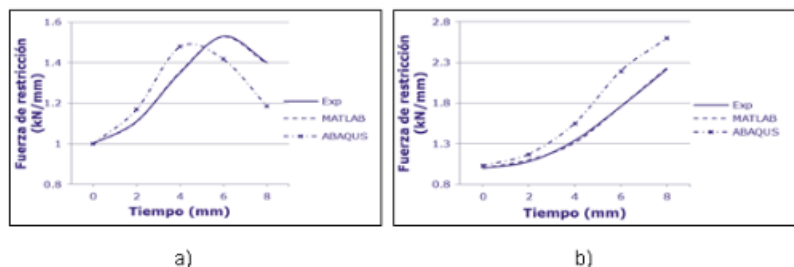


Figura 11. Comparación entre las predicciones con la red neuronal 30-10-1, resultados experimentales y de modelación numérica (ABAQUS). Respuesta: fuerzas de restricción. a) Especimen experimental HEA160- K128- L70. b) Especimen experimental HEA160- K45- L30. Fuente: Elaboración propia

Se ratifica la cercanía entre las predicciones de Inteligencia Artificial con los resultados experimentales y la mayor efectividad, en este caso, respecto a los resultados de los modelos numéricos. Se obtienen resultados similares para las restantes respuestas.

En resumen, puede afirmarse que resulta claramente efectiva la predicción de las respuestas térmicas y estructurales de columnas con restricciones sometidas a incendio por medio de redes neuronales del tipo *Multilayer-Perceptron* (MLP), con dos capas ocultas y las características específicas detalladas previamente. Se trata de un resultado especialmente relevante dada la complejidad que introducen las características de los ensayos realizados y la evolución según el tiempo de exposición al fuego de las respuestas.

4.4 Predicción de la degradación de las resistencias a momento flector y cortante de vigas compuestas de acero y hormigón en situación de incendio mediante redes neuronales, PSO y AG. Método de diseño gráfico alternativo

En el presente caso de estudio se demuestra la efectividad en la predicción de la degradación de las resistencias a momento flector y cortante de vigas compuestas en situación de incendio de un enfoque que integra la utilización de redes neuronales MLP de una capa oculta con el algoritmo de propagación de los errores hacia atrás (*backpropagation*) con un método para calcular los pesos iniciales basado en la calidad de la medida de la similaridad propuesto en el marco de la *Rough Set Theory* (RST) extendida. Se trabaja con bases de datos derivadas de modelaciones térmicas y determinaciones analíticas de las resistencias consideradas, ante el limitado número de datos experimentales disponibles, de cara al desarrollo de un método gráfico de diseño.

Descripción del objeto de estudio

Las vigas compuestas de acero y hormigón se caracterizan por presentar secciones transversales en las que ambos materiales trabajan de forma conjunta, tal como se expuso en 4.1. La exposición de estos materiales estructurales a las altas temperaturas que ocurren en situación de incendio provoca la degeneración de sus propiedades físicas y químicas, lo que conduce a una severa reducción de la resistencia y la rigidez de los elementos estructurales (ver Figura 12).

Dado el limitado número de datos experimentales se desarrollan modelos numéricos térmicos implementados en el módulo *SuperTempcalc* del programa *Temperature Calculation and Design* (TCD) desarrollado por FSD (*Fire Safety Design*, Suecia) de demostrada confiabilidad. Para la modelación de la geometría de la viga compuesta de acero y hormigón se

considera la losa maciza con anchos efectivos en un rango entre uno y tres metros; el espesor de losa se encuentra en un rango entre 0,1 m y 0,2 m. Los perfiles de acero se tomaron del surtido *American Wide Flange Beams (Universal Beams*, 2005), dada su amplia utilización internacional.

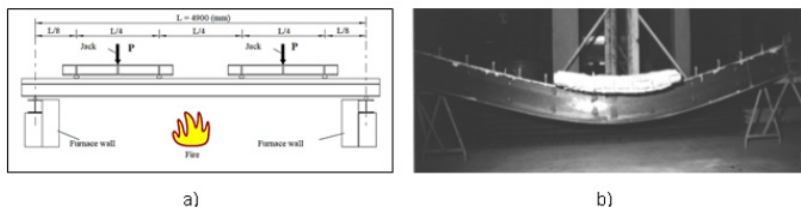


Figura 12. Viga compuesta de acero y hormigón. a) Esquema general de ensayo en situación de incendio. b) Viga deformada después del ensayo. Fuente: Kruppa y Zhao (1995)

La Figura 13 muestra el campo de temperatura y las isotermas de una sección transversal con revestimiento de 10 mm. para un tiempo de exposición al fuego de 120 minutos.

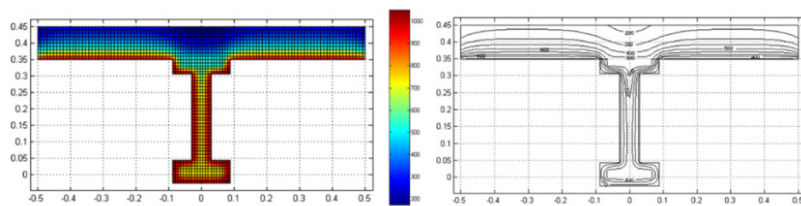


Figura 13. Campo de temperaturas e isotermas de la sección transversal de la viga compuesta para 120 minutos de exposición al fuego. Fuente: Elaboración propia (con imágenes de SuperTempcalc)

Además de los gráficos de campo de temperaturas e isotermas que ilustran el comportamiento térmico de las secciones modeladas, es de primordial utilidad la salida en formato Microsoft Excel que emite *SuperTempcalc* y que

proporciona la evolución de las temperaturas para cada nodo de la sección, lo que constituye el punto de partida para la herramienta computacional SCBEAM desarrollada en el presente trabajo, que permite determinar la evolución de las resistencias a momento flector y a cortante vertical de las secciones en estudio a partir del conocimiento de la evolución de las temperaturas.

Concepción de la aplicación de las técnicas de inteligencia artificial

Dada la necesidad de obtener un conjunto de resultados de las respuestas estructurales para desarrollar el método gráfico pretendido, se determinan inicialmente las técnicas de inteligencia artificial más efectivas para una base de datos menos compleja (vigas sin revestimiento) y se extrapolan los mismos al caso más complejo (vigas con revestimiento); en este último se generan los datos necesarios para construir las ayudas gráficas de diseño combinando resultados de inteligencia artificial, en lugar de una parte de los resultados provenientes de laboriosas modelaciones numéricas que hubieran sido requeridos.

A continuación se describe el procedimiento seguido:

1. Creación de la base de datos para realizar las predicciones de inteligencia artificial en vigas no revestidas

Se define una base de datos que cuenta con 9 variables de entrada relacionadas con la geometría de cada uno de los 32 perfiles de acero considerados, una variable de entrada asociada al tiempo de resistencia al fuego (TRF) y como variables de salida los valores relativos entre las resistencias ante incendio y en temperatura ambiente μ y ν , determinadas para cada tiempo de exposición al fuego (0, 5, 10, 15, 20, 25, 30, 35 y 40 minutos) para un total de 288 instancias, a las que se le determinan las variables de respuesta por medio

de modelación térmica numérica complementada por análisis estructural (SuperTempcalc + SCBEAM). La Tabla 6 muestra las características generales de la base de datos.

Tabla 6. Características generales de la base de datos para vigas no revestidas. Fuente: Elaboración propia

Variables de entrada	Rango de valores
Altura total	(303 ~ 612) mm
Altura total del alma	(276,6 ~ 573) mm
Altura parcial del alma	(245 ~ 547) mm
Espesor del alma	(5,1 ~ 13) mm
Ancho del ala	(101 ~ 324) mm
Espesor del ala	(5,7 ~ 21,7) mm
Área del alma	(1487,16 ~ 7277,1) mm ²
Área de las alas	(1151,4 ~ 12312,0) mm ²
Área total	(2638,56 ~ 19589,1) mm ²
TRF	(0; 5; 10; 15; 20; 25; 30; 35; 40) min

2. Selección de la combinación de técnicas de inteligencia artificial más efectiva

Se evalúan los resultados de las combinaciones MLP + PSO y MLP + AG para la base de datos de vigas no revestidas, por medio de medidas estadísticas y valoración de resultados gráficos.

3. Definición del número mínimo de instancias en el conjunto de entrenamiento (número máximo de instancias en el conjunto de prueba) para vigas no revestidas

Se define la efectividad que se alcanza con diferentes proporciones de instancias consideradas como conjunto de entrenamiento o prueba: a) (75 % / 25 %); b) (59,4 % / 40,6 %); c) (50 % / 50 %), con el propósito de establecer criterios acerca del número mínimo de instancias requeridas para realizar predicciones efectivas, a los efectos de reducir

el número de modelos térmicos a desarrollar. En esta etapa se toma en cuenta, además de las medidas estadísticas, la correspondencia entre los gráficos determinados con resultados de IA y los determinados anteriormente con resultados del SuperTempcalc + SCBEAM.

4. Consideración de los resultados de vigas no revestidas en el proceso de obtención de datos para la construcción de los gráficos de μ y σ en función del TRF en vigas con revestimiento contra incendio

En el caso de vigas revestidas, a las variables de entrada indicadas en la Tabla 4.6 se añaden el espesor y la conductividad del revestimiento, con tres valores cada una, lo que incrementa el número de instancias totales a 2 592, engorroso de enfrentar únicamente con modelación numérica y análisis estructural (SuperTempcalc + SCBEAM). Es por ello que se concibe integrar a la base de datos valores determinados por inteligencia artificial, lo que se realiza de la siguiente forma:

- Obtención de las respuestas para el % definido de las instancias por medio de (SuperTempcalc + SCBEAM).
- Predecir las restantes instancias por medio de la combinación de técnicas más efectiva, según los resultados en vigas sin revestimiento.
- Validación de resultados para un subconjunto del conjunto de prueba determinado por ambas vías.

5. Obtención de las ayudas de diseño de vigas no revestidas con datos numéricos (SuperTempcalc + SCBEAM), en tanto para las vigas revestidas se obtienen con la integración de los resultados numéricos y de IA (SuperTempcalc + SCBEAM) + IA)

Principales resultados

Las combinaciones (MLP + PSO) y (MLP + AG) para vigas no revestidas alcanzan resultados del mismo orden de magnitud,

con valores cercanos a 1 del coeficiente de determinación R^2 , en tanto las medidas estadísticas no se afectan sensiblemente por la reducción del número de instancias que forman parte del conjunto de entrenamiento, lo que significa que puede lograrse una excelente precisión considerando como conjunto de entrenamiento el 50 % de las instancias.

Con base en estos resultados, para vigas revestidas se consideran conjuntos de entrenamiento integrados por el 50 % de las instancias totales (2592) con valores de las respuestas determinados por medio de SuperTempcalc + SCBEAM y, en consecuencia, conjuntos de prueba integrados por el 50% restante a los efectos de determinar sus respuestas por medio de MLP + PSO.

En todos los casos se alcanza gran efectividad con valores bajos del error y coeficientes de correlación y de determinación R^2 cercanos a 1. Lo anterior valida la efectividad de los resultados y justifica la utilización de las predicciones realizadas por medio de MLP + PSO de manera combinada con los resultados de origen numérico (SuperTempcalc + SCBEAM) en la construcción de los gráficos de μ vs TRF y v vs TRF. En la Figura 14 se corrobora de manera gráfica la validez de las anteriores afirmaciones, para una selección de los casos de validación.

Finalmente, los gráficos que describen la degradación de las resistencias a momento flector y constante (μ y v contra TRF, respectivamente) pueden ser agrupados conveniente y utilizarse como ayudas de diseño, dos para vigas sin revestimiento y 18 para vigas revestidas, de un método gráfico alternativo, que permite determinar el tiempo de resistencia al fuego de una viga dada, de manera directa, a partir de las solicitaciones en situación de incendio y las resistencias a temperatura ambiente. La Figura 15 muestra una de las ayudas confeccionadas. G-1 a G-VII son los grupos de perfiles afines por cercanía de sus curvas de degradación.

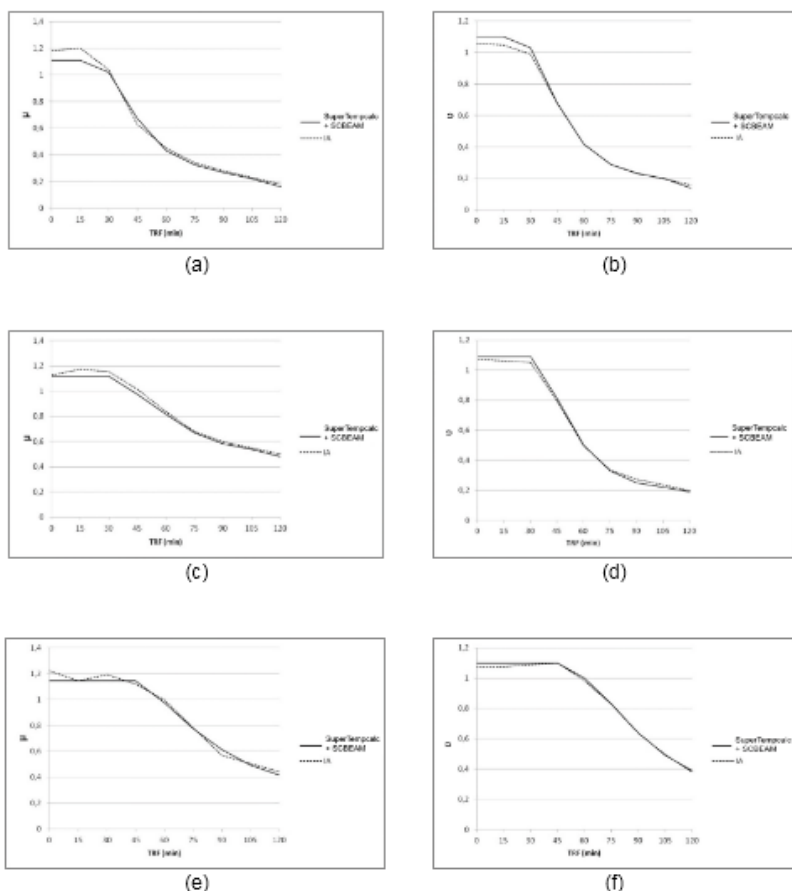


Figura 14. Comparación de curvas de degradación de la resistencia a momento flector y a cortante vertical en vigas con revestimiento contra incendio (Super Tempcalc + SCBEAM vs IA). Espesor de revestimiento contra incendio de 10 mm. Conductividad del material de revestimiento de 0,1 W/(K•m) (a) y (b) W 310 x 100 x 28,3; (c) y (d) W 460 x 150 x 60; (e) y (f) W 310 x 310 x 117

En resumen, puede afirmarse que resulta efectivo el enfoque de generar respuestas basadas en modelación térmica y análisis estructural, combinadas con predicciones

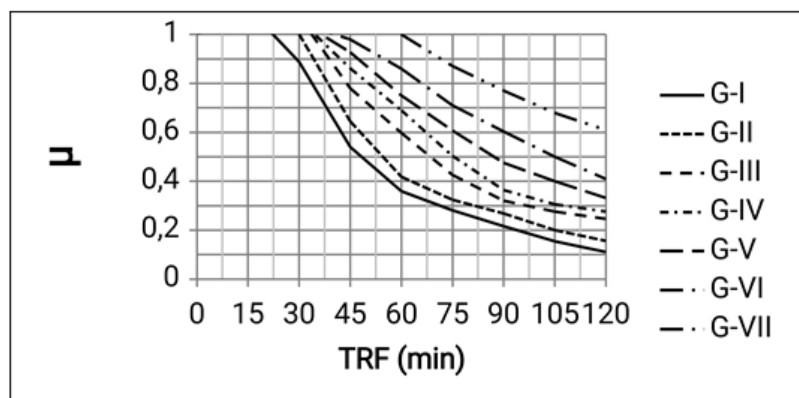


Figura 15. Ayuda de diseño μ vs TRF. Espesor del revestimiento: 10 mm y conductividad del material de revestimiento: 0,1 W/(K.m)

de inteligencia artificial por medio de redes neuronales del tipo Multilayer-Perceptron (MLP) de una capa oculta con el algoritmo de propagación de los errores hacia atrás (*backpropagation*) de conjunto con un método para calcular los pesos iniciales basado en la calidad de la medida de la similitud propuesto en el marco de la *Rough Set Theory* (RST) extendida, de cara a la propuesta de un método de diseño gráfico de vigas compuestas de acero y hormigón en situación de incendio.

4.5 Predicción de la resistencia de las conexiones tipo canal en estructuras compuestas de acero y hormigón

En el presente caso de estudio se abordan las conexiones tipo canal en estructuras compuestas de acero y hormigón, con el objetivo principal de obtener una nueva formulación para determinar la resistencia de este tipo de conexión de amplio uso en el área americana. Se recopilan, a partir de fuentes confiables, datos experimentales y numéricos de ensayos de conexiones tipo canal, a partir de lo cual se crea una base de datos válida para realizar predicciones por medio de diferentes combinaciones de técnicas de IA, a los efectos de determinar la

más efectiva. Se realiza un diseño estadístico de experimento que permite identificar las instancias inexistentes en las fuentes experimentales o numéricas, a ser generadas por la mejor técnica de IA identificada previamente; finalmente se propone una nueva formulación para el cálculo de la resistencia de la conexión aplicando técnicas estadísticas a una base de datos que integra resultados experimentales, numéricos y de inteligencia artificial, y se demuestra su superioridad respecto a las formulaciones existentes en la literatura internacional.

Descripción del objeto de estudio

El conector canal, objeto de estudio del presente trabajo, puede ser fabricado con facilidad a partir de perfiles de acero laminados en caliente y ser unido a la viga de acero mediante soldadura convencional; posee una resistencia considerablemente mayor que la de los conectores tipo perno con cabeza, estudiados en 4.1, la que a su vez puede ser regulada en el diseño a través de su longitud, que es directamente proporcional a la resistencia de la conexión. Por tales razones, el conector canal es considerado como un tipo viable y efectivo en el contexto latinoamericano.

Las conexiones canal, no han sido abordadas de manera extensiva en los programas experimentales desarrollados en comparación con los dedicados a las conexiones tipo perno con cabeza; sin embargo, pueden citarse los estudios de este corte realizados por Viest *et al.* (1952), Larrúa (1992), Pashan (2006), Maleki y Bagheri (2008), y Ramírez *et al.* (2010). Esto ha justificado la necesidad de su complementación con estudios con base en modelos numéricos, que logran reproducir con efectividad el comportamiento experimental (Ramal, 2011). La Figura 16 ilustra el comportamiento experimental de la conexión estudiada y su adecuada reproducción por medio de un modelo numérico implementado en ABAQUS.

Respecto al objeto de estudio, ha sido comprobado por los autores del presente trabajo que las formulaciones propuestas en las principales normativas (*Canadian Institute of Steel Construction (CSA), 2001; American Institution of Steel*

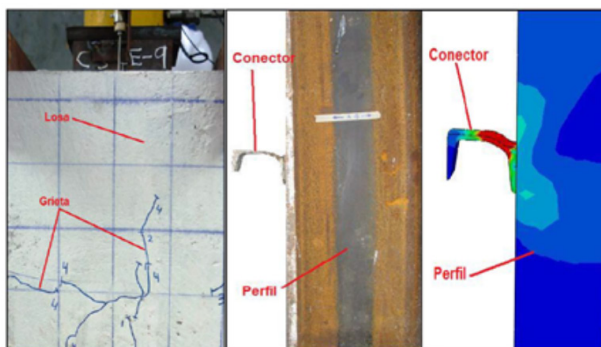


Figura 16. Comportamiento de las conexiones canal. a) Agrietamiento del hormigón en la losa. b) Deformaciones experimentales del conector. c) Reproducción de las deformaciones del conector en el modelo numérico. Fuente: Ramal (2011)

Construction (AISC), 2010) o las investigaciones realizadas por autores de resultados relevantes (Pashan, 2005), subestiman la resistencia de las conexiones determinadas por métodos experimentales, al tiempo que tanto estas como la información experimental disponible, no son representativas de la totalidad de situaciones de diseño posibles. Por tanto, existe la necesidad del perfeccionamiento de tales formulaciones.

Concepción de la aplicación de las técnicas de Inteligencia Artificial

El objetivo central es valorar la efectividad del enfoque de trabajo y de la herramienta computacional PROCON descritos en 4.1 y 4.2, en la predicción de la capacidad resistente de la conexión tipo canal en losa maciza, a los efectos de generar una base de datos representativa de las situaciones de diseño más frecuentes en el uso de ese tipo de conexión, válida para proponer una nueva formulación y evaluar su desempeño en comparación con las principales normativas y autores internacionales.

A continuación se describe el procedimiento seguido que

coincide en términos generales con el planteado en 4.1 y 4.2, con la incorporación de la combinación de técnicas k -NN + AG:

1. Obtención de la base de datos de resultados experimentales y de origen numérico

La recopilación de información experimental y numérica pertinente condujo a la conformación de una base de datos integrada por 46 instancias provenientes de resultados experimentales de los programas de Pashan (2005), con conectores de 102 mm y 127 mm de altura, y del Proyecto SENACYT 2010 (Ramírez *et al.*, 2010), con conectores de 76,2 mm de altura, así como 15 instancias derivadas de resultados numéricos de especímenes con alturas de canal de 76,2 mm (Ramal, 2011). Este último autor desarrolló la simulación virtual tridimensional del ensayo *push out* de conexiones tipo canal en losa maciza de hormigón armado, implementada en el programa computacional ABAQUS. La calibración de los modelos fue realizada tomando en cuenta los resultados experimentales del Proyecto SENACYT COL 06-007 (Ramírez *et al.*, 2010).

La base de datos cuenta con cinco variables de entrada continuas con datos de la conexión a estudiar, como son: altura del conector canal (H), espesor del ala (t) y del alma (w) del conector canal, longitud del conector (L) y la resistencia a compresión del hormigón (f'_c). Incluye además como variable de respuesta continua la resistencia última de cada instancia (Q_{con}).

2. Realizar predicciones por medio de diferentes combinaciones de técnicas de inteligencia artificial, y determinar la más efectiva

Se utiliza el estimador de funciones k -NN complementado con dos tipos de algoritmos de optimización, que permiten la asignación de pesos a los atributos de las bases de datos: algoritmo de optimización de enjambres de partículas (PSO)

y algoritmo genético (AG). Para predecir la resistencia de las conexiones por medio de técnicas de inteligencia artificial se utiliza el sistema automatizado PROCON antes utilizado en 4.1 y 4.2. Para evaluar los resultados de las predicciones fueron empleadas las siguientes medidas estadísticas: Q_{exp} / Q_m (media), coeficiente de correlación, coeficiente de determinación múltiple (R^2).

3. Proponer una nueva formulación para el cálculo de la resistencia de las conexiones tipo canal en losa maciza

Una vez seleccionada la combinación de técnicas de inteligencia artificial más efectiva, con base a los resultados de las medidas estadísticas antes expuestas, se realiza un diseño de experimentos que asegura la generación de combinaciones (en este caso un total de 90), que abarcan las situaciones de diseño más frecuentes en la práctica internacional; a tal efecto se combinan alturas de 70, 100 y 130 mm, espesores del alma de 4 y 8 mm, longitudes de 50, 100 y 150 mm, y resistencias a la compresión del hormigón de 20, 25, 30, 35 y 40 MPa. Las combinaciones que arroja el anterior diseño se comparan con la base de datos original (con datos experimentales y numéricos), con el propósito de seleccionar las instancias no representadas en la misma (un total de 39), y se procede a la predicción de la variable de salida de estas (resistencia de la conexión), por medio de la combinación de técnicas de inteligencia artificial más efectiva definida en el paso previo.

Finalmente, se propone una nueva formulación para el cálculo de la resistencia de las conexiones tipo canal en losa maciza, aplicando técnicas estadísticas a la base de datos que integra resultados experimentales, numéricos y de inteligencia artificial, según ha sido descrito previamente, y se realiza un cuidadoso análisis de su efectividad por medio de las medidas estadísticas antes indicadas.

Principales resultados

Se selecciona la combinación k-NN + AG con k=3, por la calidad de las medidas estadísticas que ofrece; haciendo uso de la misma, se completa la base de datos original (con datos experimentales y numéricos) con las predicciones de la variable de salida de las 39 instancias no representadas y se obtienen, por técnicas estadísticas, diferentes alternativas de formulaciones, entre las cuales se selecciona por su efectividad la expresión siguiente:

(4.1)

Donde:

Q_{sc} : capacidad resistente del conector (kN)

L: ancho del conector (mm)

w espesor del alma del conector (mm)

f'_c : resistencia a compresión del hormigón (MPa)

La Figura 17 ilustra las bondades de la nueva formulación propuesta respecto a dos importantes normas internacionales del ámbito americano (*Canadian Institute of Steel Construction, 2001* y *American Institution of Steel Construction, 2010*).

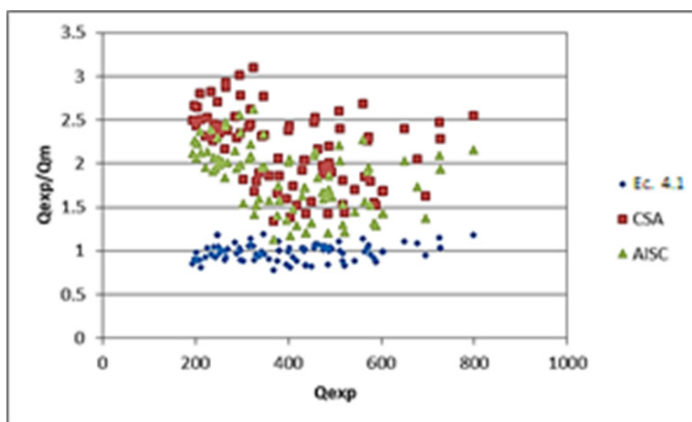


Figura 17. Relaciones Q_{exp}/Q_m en la predicción de la resistencia de la conexión canal según la ecuación 4.1, CSA y AISC. Fuente: Elaboración propia

En resumen, se demuestra la efectividad del procedimiento propuesto, que integra resultados experimentales, numéricos y de inteligencia artificial, de cara a obtener una nueva formulación para la determinación de la capacidad resistente de la conexión canal en losa maciza, más eficiente que las existentes internacionalmente. Como particularidad importante resalta la utilización de modelos numéricos y predicciones con IA para complementar la base de datos experimentales con los casos necesarios para representar las situaciones de diseño más frecuentes y lograr así una formulación efectiva y válida para todo el conjunto de las mismas. Por su parte, se destaca el rol del diseño estadístico de experimentos y las técnicas estadísticas en general, en la planeación del estudio, la generación de la formulación pretendida y la demostración de su efectividad.

Conclusiones

A través de los casos de estudio tratados se demuestra la efectividad de las técnicas de inteligencia artificial en la predicción de respuestas térmicas y estructurales en tipologías estructurales diversas, así como se ilustra cómo pueden integrarse la experimentación, la modelación numérica y las técnicas de inteligencia artificial y estadísticas en la generación de datos para el desarrollo de formulaciones o métodos de diseño.

Lo expuesto demuestra que las bases de datos de origen experimental o numéricas, o con orígenes combinados, disponibles en la literatura técnica internacional deben ser valoradas desde una perspectiva diferente cuando han de servir de base para realizar predicciones por medio de las herramientas que brinda la inteligencia artificial; a tal efecto es primordial el trabajo multidisciplinario de los especialistas de los ámbitos de la ingeniería civil y de la informática. Se identifica, a su vez, que la naturaleza de los fenómenos que dan sustento a las bases de datos indudablemente influye en las técnicas o

combinaciones de técnicas que resultan más efectivas.

A juicio de los autores, la integración de técnicas, que como enfoque de trabajo se potencia en este capítulo, parte de reconocer y aprovechar las bondades de cada una de las mismas, entre las que sobresalen la cercanía al fenómeno real de la experimentación, la capacidad de visualización la modelación numérica, el reducido costo computacional y versatilidad de las técnicas de inteligencia y la capacidad de organización de la investigación que aporta el diseño estadístico de experimentos unido al rigor científico que conceden las técnicas estadísticas en la evaluación de los resultados.

Trabajos futuros

Dada la generalidad de los enfoques de trabajo descritos es posible la extensión de su aplicación a otras tipologías estructurales tanto a temperatura ambiente como en situación de incendio. Además, ya se incursiona en la utilización de técnicas de inteligencia artificial en el no menos importante campo de la optimización del diseño estructural con criterios de costo e impacto ambiental.

Referencias del capítulo

AMERICAN INSTITUTION OF STEEL CONSTRUCTION. (2010). *Specification for Structural Steel Buildings*. AISC.

CANADIANINSTITUTEOFSTEELCONSTRUCTION.(2001).*Handbook of steel constructions*. (8va ed.). [s.n.]

CANAM GROUP INC. (1995). *Joist and Joist Girders Catalogue*. [s.n.]

CASTELLANOS, L. y CÉSPEDES E. (2016). *Comportamiento de columnas de acero en situación de incendio*. [Tesis de pregrado]. Universidad de Camagüey.

CORREIA A. (2011). *Fire resistance of steel and composite steel-concrete columns*. [Tesis de doctorado]. Universidad de Coimbra. Portugal.

DIAZ, B., EASTERLING W.S, y MURRAY T.M. (1998). *Behavior of Welded Shear Studs Used with 1.0 C Deck Internal report*. Blacksburg, VA.

ESPINO, J. (2012). *Predicción de la capacidad resistente de conexiones tipo canal en estructuras compuestas hormigón-acero*. [Tesis de pregrado]. Universidad de Camagüey.

EUROPEAN COMMITTEE FOR STANDARDIZATION (2004). *Eurocode 4 (EN 1994-1-1): Design of composite steel and concrete structures - Part 1-1: General rules and rules for buildings*. CEN.

México, Gobierno del Distrito Federal (2004). *Normas Técnicas Complementarias para el Diseño y Construcción de Estructuras de Mampostería*. México: GDF.

JAYAS, B. S. y HOSAIN, M.U. (1988). Behaviour of headed studs in composite beams: push-out tests. *Canadian Journal of Civil Engineering*, 15, pp. 240-253. <https://journals.sagepub.com/doi/abs/10.1260/1369-4332.15.9.1547>.

JOHNSON R.P, YUAN H. (1998). Existing Rules and New Tests for Stud Shear Connectors in Troughs of Profiled Sheeting. *Proc. Instn Civ. Engrs Structs & Bldgs*, 128(3). pp. 244-251. <https://doi.org/10.1680/istbu.1998.30458>.

KRUPPA, J. y ZHAO, B. (1995). Fire resistance of composite beams to Eurocode 4 part 1.2. *Journal of Constructional Steel Research*, 33, 51-69. [https://doi.org/10.1016/0143-974X\(94\)00014-9](https://doi.org/10.1016/0143-974X(94)00014-9).

LARRÚA, R., CABALLERO, Y., FILIBERTO, Y., OLIVERA, I., GUERRA, M., BELLO, R. y BONILLA J. (2009). Aplicación de la inteligencia artificial a la predicción de la capacidad resistente última de las conexiones en estructuras mixtas acero - hormigón. *Revista de la Construcción*, 8(2), 109-119. <https://www.redalyc.org/pdf/1276/127619798010.pdf>.

LARRÚA, R., LARRÚA, Y., SILVA, V., FILIBERTO, Y. y CABALLERO, Y. (2020). Using numerical modelling and artificial intelligence for predicting the degradation of the resistance to vertical shear in

steel – concrete composite beams under fire. *Applied Computer Sciences in Engineering*, 35-47. https://doi.org/10.1007/978-3-030-61834-6_1.

LARRÚA, R., PARDO, M., CASTELLANOS, L. y CÉSPEDES E. (2016). *Predicción de respuestas térmicas y estructurales de columnas de acero restringidas en situación de incendio*. [ponencia]. 2da Convención Internacional de Ciencias Técnicas. Universidad de Oriente. Santiago de Cuba, Cuba.

LARRÚA, R., VARGAS, R., ESPINO, J. (2017). Predicción de la resistencia de las conexiones tipo canal en estructuras compuestas de acero y hormigón. [ponencia]. 6th *Engineering, Science and Technology Conference*, Panamá.

LARRÚA, Y. (2016). *Comportamiento termo-estructural y diseño de vigas compuestas de acero y hormigón en situación de incendio*. [Tesis de doctorado]. UCLV. Cuba

LARRÚA, R. (1992). *Cálculo de estructuras compuestas de hormigón y acero sometidas a flexión bajo carga estática con fundamentación experimental de los conectores*. [Tesis de doctorado]. ISPJAE. Habana. Cuba.

LYONS, J. C., EASTERLING, W.S. y MURRAY T.M. (1994). *Strength of Welded Shear Studs. Report No. CE/VPI-ST 94/07*. Virginia Polytechnic Institute and State University, Blacksburg, VA.

MALEKI, S. y BAGHERI, S. (2008). Behavior of channel shear connectors. Part I: Experimental study". *Journal of Structural Engineering*, 131(1), 96-106. <https://es.scribd.com/document/488183572/01-Son-Behavior-of-channel-shear-connectors-Part-I-Experimental-study>.

PASHAN, A. (2006). *Behaviour of channel shear connectors: push-outs tests*. [Tesis de maestría]. Universidad de Saskatchewan, Canada.

RAMAL, CH. (2011). *Modelación numérica de conectores canal en secciones de viga y losa maciza de hormigón*. [Tesis de pregrado]. Universidad de Camagüey. Cuba.

RAMBO-RODDENBERRY, M. D. (2002). *Behaviour and Strength of Welded Stud Shear Connector*. PhD. [Tesis de doctorado].

University of Blacksburg, Virginia]. <https://vtechworks.lib.vt.edu/server/api/core/bitstreams/ec0fd0b4-6c4e-4e04-9cbe-22de00b85e6c/content>.

RAMÍREZ, O., LARRÚA, R., VARGAS, R., YEOMANS, F. y PINTO, M. (2010). *Caracterización experimental de conectores en estructuras compuestas de hormigón y acero*. [Proyecto SENACYT COL 006-007]. Universidad Tecnológica de Panamá.

RIAH, Z., ELWOOD, K. y ALCOCER, S. (2008). *Modelo Envolvente del Comportamiento de Muros de Mampostería Confinada para un Diseño Sísmico Basado en Desempeño* [ponencia]. XVI Congreso Nacional de Ingeniería Estructural, Sociedad Mexicana de Ingeniería Estructural, Veracruz, México.

ROBINSON, H. (1988). Multiple Stud Shear Connectors in Deep Ribbed Metal Deck. *Canadian Journal of Civil Engineering*, 15(4), 553-569. https://www.researchgate.net/publication/237190201_Multiple_stud_shear_connections_in_deep_ribbed_metal_deck.

SUBLETT, C.N., EASTERLING, W.S. y MURRAY, T.M. (1992). *Strength of Welded Headed Studs in Ribbed Metal Deck on Composite Joists*. [Report No. CE/VPI-ST 92-03]. Blacksburg, VA.

TREVIÑO, E., LARRÚA, R., FLORES, L., FILIBERTO, Y. y ACEVEDO, A. (2013). Pronóstico de la respuesta de muros de mampostería confinada mediante técnicas de inteligencia artificial. [Memorias] *X Simposio de estructuras, geotecnia y materiales de construcción*. Editorial FEIJO.

TREVIÑO, E., LARRÚA, R., ALCOCER, S., FLORES, L., PARDO, M. y FILIBERTO, Y. (2010). *Aplicación de la inteligencia artificial a la predicción de la respuesta de muros de mampostería confinada*. [ponencia]. XVII Congreso de la Sociedad Mexicana de Ingeniería Estructural. León, Guanajuato.

UNIVERSAL BEAMS (2005). www.steel.web.info.

VIEST, I.M., SIESS, C.P., APPLETON, J.H. y NEWMARK, N.M. (1952). Full scale tests on channel shear connectors in composite T-Beams. *University of Illinois Bulletin*, 405. <http://illinois.edu/>.



Capítulo 5. Operaciones de optimización del transporte en la provincia de Camagüey, Cuba, frente a la COVID-19

Capítulo 5. Operaciones de optimización del transporte en la provincia de Camagüey, Cuba, frente a la COVID-19²

Yoan Martínez López

Julio Madera Quintana

Resumen

Los esfuerzos para aplicar la Inteligencia Artificial al transporte de carga y humano han sido múltiples, varios ejemplos de la aplicación de la optimización al transporte. En este capítulo se ejemplifica la optimización en el transporte de carga, tomando como caso de estudio la provincia de Camagüey durante el año 2020. Se ejemplifica la aplicación de un modelo de enrutamiento de vehículos combinado con algoritmos de optimización para la toma de decisiones en la distribución de insumos relacionados con el servicio asistencial a pacientes hospitalizados y sospechosos de la COVID-19 en Camagüey, Cuba. El objetivo principal es determinar las rutas óptimas para una flota de vehículos que deben realizar entregas o recogidas en varios puntos, con el fin de minimizar el costo total, expresado en términos de distancia total recorrida, tiempo total de viaje, costo de combustible, o una combinación de estos y otros factores.

Abstract

Efforts to apply Artificial Intelligence to both cargo and human transportation have been numerous, with several

² Este capítulo está basado en el artículo “Aplicación de la investigación de operaciones a la distribución de recursos relacionados con la COVID-19” de Yoan Martínez-López, Hilda Oquendo Ferrer; Yailé Caballero Mota, Lenier E. Guerra-Rodríguez, Raudel Junco-Villegas, Isnel Benítez Cortés, Ansel Y. Rodríguez González, y Julio Madera Quintana (2020), publicado originalmente en *Retos de la Dirección*, 14(2), 86-105.

examples showcasing optimization in transportation. This chapter illustrates optimization in cargo transportation, using the province of Camagüey during the year 2020 as a case study. It exemplifies the application of a vehicle routing model combined with optimization algorithms for decision-making in the distribution of supplies related to healthcare services for hospitalized patients and COVID-19 suspects in Camagüey, Cuba. The main objective is to determine optimal routes for a fleet of vehicles that need to make deliveries or pickups at various points, aiming to minimize the total cost, expressed in terms of total distance traveled, total travel time, fuel cost, or a combination of these and other factors.

Introducción

El año 2020 inició con una pandemia que situó a la humanidad ante grandes retos. La aparición de la enfermedad COVID-19 (Ludvigson y Ng, 2020), causada por el coronavirus SARS-CoV-2 ha provocado que muchos de los países afectados realicen esfuerzos para detener esta pandemia y minimizar los daños de su impacto.

Muchas de las personas que enfermaron de COVID-19 en los primeros dos años de esa pandemia, experimentaron síntomas de leves a moderados y se recuperaron sin procedimientos médicos especiales. El protocolo cubano para el tratamiento de la enfermedad tuvo la particularidad de que, mediante la realización de una pesquisa en la población, se trataba de identificar la mayor cantidad posible de enfermos y sus contactos para aislarlos de la población sana, trasladándolos temporalmente hacia centros habilitados con ese propósito u hospitales, para su tratamiento oportuno. Esta circunstancia hizo necesario hacer llegar a estos lugares con la mayor celeridad todos los recursos necesarios, desde medicamentos hasta insumos médicos y no médicos.

La distribución de esos recursos se realizaba normalmente por la Empresa Comercializadora y Distribuidora de

Medicamentos (EMCOMED), la Empresa de Suministros Médicos (EMSUME) y la Empresa de Abastecimientos y Servicios a la Educación (EPASE), para lo cual contaban con medios de transporte propios. Dada la emergencia, era vital optimizar los recorridos, de modo que pudiera prestarse un servicio de calidad y, sobre todo, oportuno. De aquí que, una correcta planificación del transporte, con enfoque de ahorro y uso eficiente de los medios disponibles, resultara indispensable para hacer llegar a su destino los recursos requeridos para la atención de enfermos y sospechosos de haber contraído la enfermedad (López, González y Campos, 2020; Lozano, Marmolejo y Rodríguez, 2020).

La situación descrita puso de relieve la necesidad de llevar a cabo una investigación, cuyo objetivo principal fue la aplicación de un modelo de enrutamiento de vehículos combinado con algoritmos de optimización para la toma de decisiones en la distribución de insumos relacionados con el servicio asistencial a pacientes hospitalizados y sospechosos de la COVID-19 en la provincia de Camagüey. Además, se utilizaron los problemas de enrutamiento de vehículos heterogéneos con ventanas de tiempo, en combinación con algoritmos de optimización para solucionar la distribución de estos recursos.

Los principales resultados se enfocaron en la experimentación con un total de 15 modelos para el estudio del comportamiento de los algoritmos aplicados al problema. Para ello, se utilizó la biblioteca *Capacitated Vehicle Routing Problem* (CVRP), en español “Problema de Enrutamiento de Vehículos con Capacidad”, implementada en Matlab. Además, se desarrollaron tres de metaheurísticas: *Estimation of Distribution Algorithm* (EDA), en español Algoritmo de Estimación de Distribución (AED), *Simulated Annealing* o Recocido Simulado (SA), *Variable Neighborhood Search* (VNS), en español Búsqueda de Vecindad Variable. A partir de la información modelada se procedió a la solución del problema *Fleet Size and Mix*

Vehicle Routing Problem with Time Windows (FSMVRPTW), en español “Problema de Enrutamiento de Vehículos con Tamaño y Mezcla de Flota y Ventanas de Tiempo”, o “Problema de Enrutamiento de Vehículos de Flota Mixta y Tamaño con Ventanas de Tiempo” a través de algoritmos EDA y VNS. Este último se seleccionó porque cuenta con una implementación de código abierto en Excel. Finalmente, los estudios acerca del problema de enrutamiento de vehículos demostraron su utilidad en diferentes situaciones complejas —como la generada por esta pandemia— para optimizar la distribución de recursos.

En tiempos de COVID-19, contar con una organización del transporte óptima para distribuir los recursos médicos, constituyó una herramienta vital para la toma de decisiones en la provincia Camagüey.

Desarrollo

La planificación de las rutas de transporte *Vehicle Routing Problem* (VRP) es uno de los problemas de optimización combinatoria más importantes, el cual ha sido ampliamente estudiado, con aplicaciones en el mundo real, en la logística de distribución y del transporte (Toth & Vigo, 2002). Esto es significativo para la toma de decisiones por parte de los directivos.

Desde su aparición en 1959 por Dantzig y Ramser —quienes realizaron por primera vez una formulación del problema para una aplicación de distribución de combustible—, el estudio de VRP ha generado numerosas investigaciones y se han escrito artículos sobre diversas variantes del problema (Dantzig y Ramser, 1959).

El objetivo es minimizar la distancia total recorrida por un conjunto de vehículos ubicados en un almacén central, para satisfacer la demanda de un determinado conjunto de clientes. Cada cliente tiene una demanda conocida y cada vehículo sirve una única ruta durante el período de planificación, la

cual debe comenzar y finalizar en el almacén central. Esto es una generalización del problema del agente viajero *Travelling Salesman Problem* (TSP) (Hatamlou, 2018).

En la práctica, una flota de vehículos es raramente homogénea, ya que la necesidad de estar presente en diversos segmentos del mercado obliga a las empresas a disponer de vehículos que se adapten a la tipología de la mercancía transportada. De igual modo, el disponer de vehículos con diferentes unidades de carga, permite una mejor adaptación a la demanda (Yepes y Medina, 2002). Por este motivo, se optó por el estudio de esta variante del problema VRP, en el que, el número de vehículos disponible para satisfacer las demandas de los clientes era limitado.

En la literatura se han propuesto diferentes versiones del VRP clásico con el fin de acercarse a los contextos reales de los diferentes problemas planteados (Golden, Assad, Levy y Gheysens, 1984). Estas variantes son formuladas mediante la incorporación de variables y restricciones adicionales al problema original.

La variante VRP para flotas heterogéneas (HFVRP, *Heterogeneous Fleet VRP*), aparece cuando los diferentes vehículos que conforman la flota difieren en equipamiento, capacidad, antigüedad, estructura de costos o incluso nivel de emisiones, si estas son consideradas (Golden *et al.*, 1984).

Los primeros ejemplos del HFVRP se deben a los problemas FSM (*Fleet Size and Mix*), propuestos por Golden *et al.* (1984). Estos autores proponen dos heurísticas, donde la primera de ellas se basa en el algoritmo de ahorros de Clarke & Wright (1964), y la segunda utiliza un esquema de ruta gigante (Beasley, 1983); este autor formuló 20 problemas que posteriormente han sido tomados como referencia por muchos estudiosos para la presentación de los resultados de sus algoritmos en flotas heterogéneas.

Otras variantes del HFVRP fueron los problemas con un número limitado de vehículos HVRP, introducidos inicialmente por (Taillard, 1999) que presentó un método heurístico basado en la generación de columnas. Este comienza con la solución al problema VRP homogéneo para cada tipo de vehículo y utiliza un algoritmo de búsqueda tabú.

Los problemas HFVRP son un caso especial del VRP clásico, por lo tanto pertenecen igualmente a la categoría de NP-Completo, donde el tiempo de cómputo crece exponencialmente con el tamaño del problema; razón por la cual, varios autores proponen heurísticas o metaheurísticas para la resolución de este tipo de problemas, ya que son los métodos adecuados para obtener soluciones de calidad en un tiempo de cálculo razonable (Taillard, 1999).

En la literatura existen numerosas variantes del problema HFVRP; tienen en cuenta si la flota de vehículos es limitada o ilimitada y del tipo de coste considerado. Las diversas variantes que se pueden encontrar son las siguientes:

- VRP heterogéneo con flota fija y costes variables dependientes de las rutas (HVRPD).
- VRP heterogéneo con flota fija con costes fijos y variables (HVRPFD).
- VRP heterogéneo con flota ilimitada con costes variables (FSMD).
- VRP heterogéneo con flota ilimitada con costes fijos (FSMF).
- VRP heterogéneo con flota ilimitada con costes fijos y variables (FSMFD).

Por otra parte, a las variantes que consideran restricciones de ventanas de tiempo, se les añade TW, al acrónimo del problema específico.

Cuando la suma de las demandas de todos los clientes excede la capacidad del vehículo, aparece la variante del CVRP (Baldacci, Battarra, & Vigo, 2008). Al igual que la mayoría

de los problemas VRP, el problema CVRP es NP-completo. Esto es así porque el número de posibles soluciones crece exponencialmente con el número de nodos del grafo (clientes o puntos de paso) y rápidamente sobrepasa las capacidades de cálculo de los ordenadores más potentes.

En el CVRP se considera una flota fija y heterogénea de vehículos que se encuentra estacionada en un almacén central para atender la demanda de los clientes conocidos. El CVRP consiste en el diseño de un conjunto de rutas hamiltonianas de menor coste, de tal manera que cada cliente ha de ser visitado una única vez por un único vehículo y todas las rutas de los vehículos han de comenzar y finalizar en el almacén.

El problema CVRP básico trata de determinar los recorridos de k vehículos de capacidad C_k que parten de un origen común y deben pasar por un conjunto de lugares de interés (clientes) para distribuir o recoger mercancías según una demanda d_i , y volver de nuevo al origen de manera que la distancia total recorrida (el coste o el tiempo empleado) por el conjunto de vehículos sea mínima. En el tipo de problema más sencillo no se tiene en cuenta el horario de entrega o recogida en cada lugar de interés (ventanas horarias) (Baldacci *et al.*, 2008).

El Problema de Enrutamiento de Vehículos con Ventanas de tiempo (VRPTW, VRP *with time windows*) es la variante que ha recibido la mayor atención en la literatura, por la importancia práctica que posee. Las ventanas de tiempo se producen cuando los clientes requieren que el servicio de entrega o recogida se produzca dentro de una franja horaria especificada, la cual está determinada por una hora temprana y otra tardía de servicio. Igualmente, se puede incluir un límite en el tiempo total de conducción de los vehículos debido a las regulaciones existentes en los contratos de los conductores.

En la bibliografía consultada se hace una distinción entre las ventanas de tiempo duras y suaves. En el caso de las primeras, si un vehículo llega demasiado temprano a realizar la

entrega, se le permite esperar al cliente hasta que el vehículo esté listo para ser atendido, aunque no se permite llegar más tarde del horario en que puede ser atendido. Para el caso de ventanas de tiempo suaves, los horarios de los clientes pueden ser violados a costa de una penalización en la función objetivo

Problema de Enrutamiento de Vehículos con Flota Heterogénea y Ventanas de Tiempo

El FSMVRPTW puede definirse a partir de los siguientes elementos (Baldacci *et al.*, 2008):

Sea $G = (V, E)$ un grafo dirigido donde $V = N \cup \{0\}$ es el conjunto de nodos (clientes) $E = \{(i,j) : i,j \in V\}$ y es el conjunto de arcos, donde el nodo 0 denota el depósito.

Para cada arco $(i,j) \in E$, denotaremos con d_{ij} el tiempo mínimo de viajar del nodo i al nodo j .

Cada cliente $h \in H$ es asociado con una demanda no negativa q_i (por conveniencia en la notación, al depósito se le asigna una demanda $q_0 = 0$).

Para cargar/descargar la cantidad q_i se necesita un tiempo de servicio s_i y una ventana de tiempo $[a_i, b_i]$.

El servicio del cliente i , debe comenzar entre a_i y b_i , es decir, en cualquier solución factible el vehículo que atiende al cliente i debe arribar en un instante de tiempo $t \in [a_i, b_i]$, o en un instante de tiempo $t < a_i$, por lo que debe esperar $a_i - t$ unidades de tiempo antes de comenzar el servicio.

Para simplificar la formulación del problema, se asume que todas las ventanas de tiempo están dentro de un horizonte de tiempo dado (por ejemplo, un día). Notemos que esto implica que para cada cliente $h \in H$

Se tiene una flota compuesta por H diferentes tipos de vehículos, donde cada vehículo del tipo $h (h=1, \dots, H)$ tiene una capacidad Q^h y un costo fijo F^h .

El objetivo de FSMVRPTW es determinar la cantidad óptima de vehículos heterogéneos y sus rutas asociadas para

minimizar la suma de los costos fijos de los vehículos y de los costos de las rutas sujeto a las siguientes restricciones:

1. Cada ruta comienza y termina en el depósito.

2. Cada ruta es asignada a un solo vehículo.

3. La demanda total de los clientes servidos en una ruta no puede exceder la capacidad del vehículo asignado a dicha ruta.

4. Cada cliente es visitado exactamente una vez y su servicio comienza dentro de su ventana de tiempo.

Sin que se pierda la idea original, se asume que los tipos de vehículos son enumerados en orden no decreciente del valor del costo fijo F^h y que, para cada cliente u , existe (al menos) un tipo de vehículo h tal que $Q^{h^*} > Q^h$, de otra manera los vehículos de tipo h^* pueden ser eliminados de la instancia, en cualquier solución factible, los vehículos del tipo h^* pueden ser reemplazados por vehículos del tipo h sin incrementar el costo de la solución.

Formulación matemática de FSMVRPTW

Se delimita el conjunto K de vehículos distintos obtenidos por la definición de n vehículos de tipo h para cada $h \in H$. Para cada $k \in K$, sea $Q^{h^*} > Q^h$ denotan la capacidad y el costo del vehículo k , respectivamente. Esta es una generalización de la formulación de tres índices del VRP (Toth & Vigo, 2002).

La parte de enrutado del problema es modelada a través de dos conjuntos de variables binarias: (i) las variables ϵ_{ij} toman valor 1 si el arco (i, j) es atravesado por el vehículo k ; (ii) las variables ϵ_i toman el valor 1 si el cliente i es servido por el vehículo k . Para seleccionar un conjunto adecuado de vehículos, se introduce la variable binaria Z^k , que toma valor 1 si el vehículo k y i^k es utilizado, 0 en caso contrario.

El manejo de las ventanas de tiempo y la duración de las rutas requiere la definición de los siguientes conjuntos de variables: (i) las variables t_i^k indicarán el instante mínimo de tiempo en el que el vehículo k puede arribar a cada nodo

i $t_i^k \in V$; (ii) las variables τ_i para indicar el instante de tiempo mínimo en el que el servicio al cliente i puede comenzar; (iii) las variables π^k para determinar el instante de tiempo en el cual el vehículo k , si se utiliza, comienza su ruta. Obsérvese que, para cada vehículo k , el tiempo de inicio y finalización de su ruta es dado por las variables t_i^k .

A partir del uso de estas variables y una constante positiva muy grande (se le asigna el valor, t_i^k el problema FSMVRPTW puede formularse como sigue:

$$\sum_{k \in K} (F^k Z^k + t_0^k - \pi^k) - \sum_{i \in N} S_i \quad (1)$$

Sujeto a:

$$\sum_{k \in K} y_i^k = 1, \quad i \in N \quad (2)$$

$$y_i^k = \sum_{j \in V} (j, i) \in x_{ji}^k, \quad i \in N, k \in K \quad (3)$$

$$y_i^k = \sum_{j \in V} (i, j) \in x_{ij}^k, \quad i \in N, k \in K \quad (4)$$

$$\sum_{i \in V} q_i y_i^k \leq Q^k Z^k, \quad k \in K \quad (5)$$

$$\sum_{i \in S} \sum_{j \in S} x_{ij}^k \geq y_i^k, \quad S \subseteq N, \quad i \in S, \quad K \in K \quad (6)$$

$$t_j^k \geq \tau_i + s_i + d_{ij} - M(1 - x_{ij}^k), \quad (i, j) \in E, \quad N, k \in K \quad (7)$$

$$t_j^k \geq \pi^k + d_{0j} - M(1 - x_{0j}^k), \quad j \in N, \quad k \in K \quad (8)$$

$$t_0^k \geq \pi^k, \quad k \in K \quad (9)$$

$$\tau_i \geq t_j^k, \quad i \in N, \quad k \in K \quad (10)$$

$$a_i \leq \tau_i \leq b_i, \quad i \in N \quad (11)$$

$$x_{ij}^k \in \{0, 1\}, \quad (i, j) \in E, \quad k \in K \quad (12)$$

$$y_i^k \in \{0, 1\}, \quad (i, j) \in E, \quad k \in K \quad (13)$$

$$Z^k \in \{0, 1\}, \quad k \in K \quad (14)$$

$$t_j^k \geq 0, \quad i \in V, \quad k \in K \quad (15)$$

$$\tau_i \geq 0, \quad i \in N \quad (16)$$

$$\pi^k \geq 0, \quad k \in K \quad (17)$$

La restricción (2) propone que cada cliente debe ser visitado exactamente por un vehículo, mientras que la restricción (3) y (4) proponen que, si el cliente i es visitado por el vehículo k , este debe entrar y salir, respectivamente, del nodo asociado. La desigualdad (5) está relacionada con la restricción de capacidad de cada vehículo $k \in K$, y la desigualdad (6) tiene en cuenta los requerimientos de conectividad de cada ruta. Las restricciones (7), (10) y (11) definen el instante de tiempo para servir a cada cliente i , e impone las restricciones de las ventanas de tiempos. Por otro lado, la restricción (8) define el instante de tiempo en el cual cada vehículo k $t_0^k - \pi^k \geq 0$ comienza su recorrido. Nótese que cuando un vehículo k no es utilizado, las variables π^k y $t_0^k - \pi^k \geq 0$ no tienen restricciones, pero (9) establece que, $t_0^k - \pi^k \geq 0$ y la función objetivo los establecerá en un valor común para cualquier solución óptima.

La gran cantidad de variables en el modelo (2) - (17) y la presencia en las restricciones de valores de M muy grandes, hace imposible (o extremadamente complejo) obtener la solución exacta del problema.

Solución de VRP

Los EDA son un grupo de algoritmos de EA (*Evolutionary Algorithms*) que permiten ajustar el modelo a la estructura de un problema determinado, realizados por una estimación de distribuciones de probabilidades a partir de soluciones seleccionadas (Larrañaga, Lozano y Mühlenbein, 2003). El modelo es una distribución de probabilidad. Estos algoritmos se basan en sustituir los operadores de cruce y mutación de los algoritmos genéticos por la estimación y posterior muestreo de una distribución de probabilidad, aprendida a partir de los individuos seleccionados de una población (Martínez, Madera y Leguen, 2016; Martínez *et al.*, 2019; Martínez, Madera, Rodríguez y Barigye, 2019).

El SA (*Simulated Annealing*, recocido simulado, cristalización simulada o enfriamiento simulado) es un algoritmo de búsqueda

metaheurística para problemas de optimización global; el objetivo general de este tipo de algoritmo es encontrar una buena aproximación al valor óptimo de una función en un espacio de búsqueda grande. A este valor óptimo se lo denomina óptimo global. El nombre e inspiración viene del proceso de recocido del acero y cerámicas, una técnica que consiste en calentar y luego enfriar lentamente el material para variar sus propiedades físicas. El calor causa que los átomos aumenten su energía y puedan así desplazarse de sus posiciones iniciales (un mínimo local de energía); el enfriamiento lento les da mayores probabilidades de recrystalizar en configuraciones con menor energía que la inicial (mínimo global) (Van Laarhoven & Aarts, 1987).

La VNS, propuesta por Mladenović & Hansen (1997), es un método metaheurístico para solucionar un conjunto de problemas de optimización combinatoria. Explora vecindarios distantes de la solución actual, y movimientos de allí a nuevos valores si, y solo si, hubo mejora a la solución actual. El método de búsqueda local se aplica repetidamente para conseguir soluciones en el vecindario a local óptima. VNS se diseñó para aproximar soluciones de problemas de optimización discreta y continua, está enfocada a solucionar problemas de programación lineal, problemas de programación en enteros, problemas de programación en enteros mixtos y problemas de programación no lineal (Mladenović & Hansen, 1997).

Resultados experimentales

Para realizar el siguiente experimento se utilizó la biblioteca CVRP (Mostapha, 2015), implementada en MATLAB 19a. y las tres metaheurísticas anteriormente analizadas (EDA, SA, VNS). Se experimentó con un total de 15 modelos (cantidad de nodos x cantidad de recursos disponibles) para el estudio del comportamiento de estos algoritmos aplicados al problema.

Como se puede apreciar en la Tabla 1, las metaheurísticas tienen comportamientos similares en la eficiencia evaluativa para las diferentes configuraciones de los modelos (cantidad de nodos x cantidad de recursos disponibles), los cuales convergen a la solución óptima. Tanto los algoritmos VNS como EDA, se comportaron ligeramente superiores al SA, aunque la cantidad de las iteraciones para encontrar el óptimo es similar. Estos experimentos permiten analizar el comportamiento de las diferentes metaheurísticas para luego seleccionar cuál utilizar para resolver los diferentes problemas en la toma de decisiones.

Tomando como base el estudio anterior, se seleccionó una herramienta que usa la variante VNS (Erdogan, 2017), para realizar dos estudios de casos en la provincia de Camagüey, en la toma de decisión por parte de los directivos, en la organización del transporte en tiempos de COVID-19.

Aplicación del problema FSMVRPTW en la empresa EPASE de Camagüey

Apartir de la modelación del problema FSMVRPTW, se procedió a su aplicación en la provincia Camagüey. La mayoría de las empresas que transportan mercancías, relacionadas con la salud pública (EMCOMED, EMSUME y EPASE), cuentan con un almacén central en el municipio de Camagüey. Con salida de dicho almacén, las mercancías son transportadas/distribuidas al resto de los 12 municipios de la provincia. En contacto con algunas de estas empresas se caracterizó el transporte, las mercancías que transportan, los horarios de servicios, etc. Una vez caracterizada la información se procedió a simular una posible carga a distribuir con todos los parámetros para este tipo de problema. Las distancias entre los diferentes municipios de la provincia Camagüey fueron calculadas a partir del servicio de BingMaps de Microsoft. Nótese que las distancias son simétricas por lo que se representa como una matriz triangular.

Caso I

1. La empresa cuenta con una flota de cuatro vehículos, todos del mismo tipo ($H = 1$) y con la misma capacidad ($Q^h = 100$), interpretada como cajas con las mismas dimensiones).
2. Cada viaje de un vehículo tiene un costo fijo de 100 CUP y un costo por kilómetro de 10 CUP.
3. Cada vehículo puede recorrer, como máximo, 400 km y los choferes no pueden manejar más de 9 horas.
4. La velocidad promedio de los vehículos es de 50 km/h; esta se relaciona con la calidad de las carreteras y con las condiciones técnicas de los vehículos.
5. Número de clientes
6. $N = 12$, representa los 12 municipios de la provincia Camagüey.
7. Se tienen las distancias entre los municipios, las que representan el peso en los arcos del grafo.
8. La ventana de tiempo de carga de los vehículos en el almacén central es $[6:00, 8:00]$; indica que los vehículos deben iniciar su carga en ese horario y el tiempo de serviciado del vehículo es de 55 minutos.
9. La ventana de trabajo modelada para los clientes es $[8:00, 16:00]$, indica que todas las descargas de vehículos deben comenzar después de las 8:00 a.m. y antes de las 4:00 p.m.
10. El tiempo de descarga de la mercancía en los municipios es la misma para todos e igual a 20 minutos.

Además, se definieron las cantidades de envases a transportar a cada municipio, o sea, la demanda del municipio.

Tabla 1 Comparación de las tres metaheurísticas analizadas en los problemas de la biblioteca CVRP

Modelo	cEDA		SA		VNS		GRASP	
	Iteración	Valor del costo	Iteración	Valor del costo	Iteración	Valor del costo	Iteración	Valor del costo
Bajmaj_R3_0.35_2.13_2 x 12 min.mat	2	16.221	1	16.221	5	16.221	18	16.221
Bajmaj_R3_0.35_2.13_2 x 18 min.mat	1	16.221	1	16.221	15	16.221	6	16.221
Bajmaj_R3_0.35_2.13_2 x 24 min.mat	2	16.221	1	16.221	12	16.221	7	16.221
Bajmaj_R3_0.35_2.13_3 x 12 min.mat	1	16.221	1	16.221	51	16.221	46	16.221
Bajmaj_R3_0.35_2.13_3 x 18 min.mat	1	16.221	1	16.221	54	16.221	13	16.221
Bajmaj_R3_0.35_2.13_3 x 24 min.mat	2	16.221	2	16.221	10	16.221	48	16.221
Bajmaj_R3_0.35_2.13_4 x 12 min.mat	1	16.221	2	16.221	61	16.221	12	16.221
Bajmaj_R3_0.35_2.13_4 x 18 min.mat	1	16.221	1	16.221	38	16.221	16	16.221
Bajmaj_R3_0.35_2.13_4 x 24 min.mat	2	16.221	2	16.221	22	16.221	6	16.221
Bajmaj_R19_0.35_2.13_4 x 6 min.mat	2	22.85	2	22.85	34	22.85	90	22.85
Bajmaj_R19_0.35_2.13_4 x 12 min.mat	1	22.85	4	22.85	36	22.85	29	22.85

Bajmaj_ R19_0.35_2.13_4 x 18 min.mat	1	22.85	11	22.85	29	22.85	16	22.85
Bajmaj_ R19_0.35_2.13_5 x 6 min.mat	2	22.85	1	22.85	19	22.85	38	22.85
Bajmaj_ R19_0.35_2.13_5 x 12 min.mat	1	22.85	3	22.85	40	22.85	82	22.85
Bajmaj_ R19_0.35_2.13_5 x 18 min.mat	1	22.85	4	22.85	94	22.85	65	22.85
Bajmaj_ R19_0.35_2.13_6 x 6 min.mat	1	22.85	1	22.85	38	22.85	25	22.85
Bajmaj_ R19_0.35_2.13_6 x 12 min.mat	3	22.85	2	22.85	102	22.85	28	22.85
Bajmaj_ R19_0.35_2.13_6 x 18 min.mat	2	22.85	2	22.85	89	22.85	4	22.85

A partir de la información modelada se procedió a la solución del problema FSMVRPTW a través de un algoritmo VNS, obteniéndose la solución óptima que se muestra en la Figura 1. Obsérvese que la solución presenta cuatro rutas, lo que implica la utilización de cuatro vehículos.

Todos los vehículos parten del almacén central ubicado en la ciudad de Camagüey. El primer vehículo visita los municipios Sierra de Cubitas, Esmeralda, Céspedes y Florida por ese orden y regresa a la base. Este recorrido se realiza en 2 horas y 43 minutos, con una distancia de 180,6 km .

El segundo vehículo visita los municipios de Minas y Nuevitas en un tiempo de 2 horas y 34 minutos, con una distancia de 162,1 km .

El tercer vehículo visita Jimaguayú, Najasa, Santa Cruz y Vertientes en un tiempo de 3 horas y 15 minutos, recorre en total 226,25 km .

El último vehículo utilizado recorre los municipios Sibanicú y Guáimaro en un tiempo de 2 horas y 29 minutos, a lo largo de una distancia de 166,29 km .

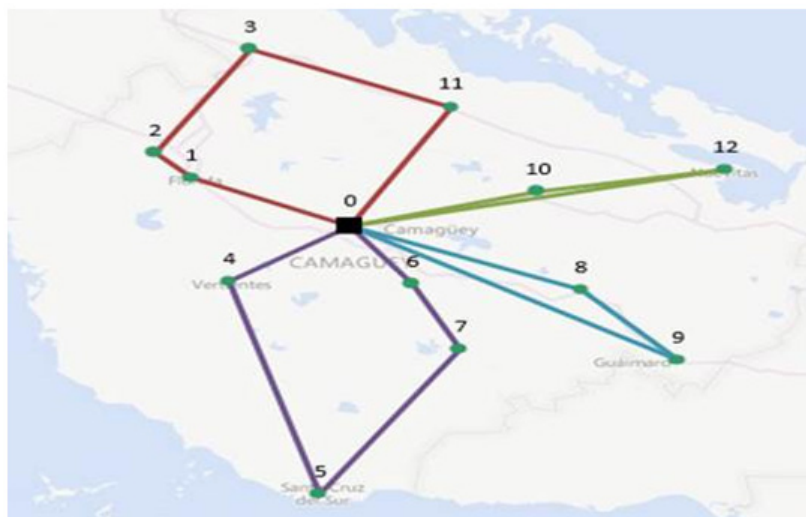


Fig. 1 Solución óptima al problema FSMVRTW para la transportación de insumos relacionados con la COVID-19 por la empresa EPASE en la provincia de Camagüey

Caso II

1. La empresa cuenta con una flota de cinco vehículos, todos del mismo tipo ($H = 1$) y con la misma capacidad ($Q^h = 100$), interpretada como cajas con las mismas dimensiones).
2. Cada viaje de un vehículo tiene un costo fijo de 100 CUP y un costo por kilómetro de 10 CUP.
3. Cada vehículo puede recorrer, como máximo, 400 km y los choferes no pueden manejar más de 9 horas.
4. La velocidad promedio de los vehículos es de 50 km/h, relacionada con la calidad de las carreteras y con las condiciones técnicas de los vehículos.

5. Número de clientes $N = 36$, representa los 12 municipios de la provincia Camagüey, los hospitales y centros de aislamiento.
6. Se tienen las distancias entre los municipios, las que representan el peso en los arcos del grafo.
7. La ventana de tiempo de carga de los vehículos en el almacén central es $[6:00, 8:00]$, indica que los vehículos deben iniciar su carga en ese horario y el tiempo de serviciado del vehículo es de 55 minutos.
8. La ventana de trabajo modelada para los clientes es $[8:00, 16:00]$, indica que todas las descargas de vehículos deben comenzar después de las 8:00 a.m. y antes de las 4:00 p.m.
9. El tiempo de descarga de la mercancía en los municipios es la misma para todos e igual a 20 min (ver Figura 1).

Conclusiones

La investigación de operaciones constituyó una poderosa herramienta para la planificación de los recursos en el combate de pandemias.

Para aplicar las herramientas de IO tiene que existir interés y voluntad política por parte de las autoridades gubernamentales, demostrado en el énfasis puesto por el gobierno cubano en la aplicación de la investigación matemática al combate de la COVID-19.

La programación lineal y los problemas de transporte afloran como herramientas útiles para la utilización eficiente de los recursos a disposición del combate de la COVID-19.

Con la aplicación de algoritmos de búsqueda variable y de estimación de distribuciones a la solución del problema VRP se obtuvieron soluciones óptimas, que minimizaron los recorridos

a realizar por los medios de transporte.

Al reducir al mínimo las distancias y la cantidad de vehículos asignados, se redujo la cantidad de combustible y por lo tanto el ahorro de este portador energético, así como su impacto sobre el medioambiente, en particular, sobre las emisiones de CO_2 .

Para el combate de la COVID-19 se logró tener una organización óptima del transporte para el traslado de recursos médicos, la cual influyó en la toma de decisiones de los directivos en la provincia.

Trabajo futuro

Como trabajo futuro se propone desarrollar otros modelos de optimización del transporte de carga para su aplicación en la provincia de Camagüey

Referencias del capítulo

Baldacci, R., Battarra, M., & Vigo, D. (2008). *Routing a heterogeneous fleet of vehicles*. En B. Golden y S. Raghavan (ed.). *The vehicle routing problem: latest advances and new challenge* (pp. 3-27). Springer.

Beasley, J. E. (1983). *Route first-cluster second methods for vehicle routing*. *Omega*, 11(4), 403-408. [https://doi.org/10.1016/0305-0483\(83\)900336](https://doi.org/10.1016/0305-0483(83)900336)

Clarke, G. & Wright, J. T. (1964). *Scheduling of Vehicles from a Central Depot to a Number of Delivery Points*. *Operations Research*, 12(4), 568-581. <http://dx.doi.org/10.1287/opre.12.4.568>

Erdogan, G. (2017). *An Open-Source Spreadsheet Solver for Vehicle Routing Problems*. *Computers & Operations Research*, 84, 62-72. <http://dx.doi.org/10.1016/j.cor.2017.02.022>

Dantzig, G. B., & Ramser, J. H. (1959). *The truck dispatching*

problem. Management science, 6(1), 80-91. <https://doi.org/10.1287/mnsc.6.1.80>

Golden, B., Assad, A., Levy, L., & Gheysens, F. (1984). The fleet size and mix vehicle routing problem. *Computers & Operations Research*, 11(1), 49-66. [https://doi.org/10.1016/0305-0548\(84\)900078](https://doi.org/10.1016/0305-0548(84)900078)

González, J. A. (2012). *Los tranvías de vapor en España. Una historia (casi) desconocida*. [ponencia] VI Congreso de Historia Ferroviaria Vitoria Gasteiz

Hatamlou, A. (2018). *Solving travelling salesman problem using black hole algorithm*. *Soft Computing*, 22(24), 8167-8175. <https://doi.org/10.1007/s00500017-2760-y>

Larrañaga, P., Lozano, J. A. y Mühlenbein, H. (2003). Algoritmos de estimación de distribuciones en problemas de optimización combinatoria. *Inteligencia Artificial Revista Iberoamericana de Inteligencia Artificial*, 7(19), 149-168. <http://aepia.org/revista>.

López, R., González, C. y Campos, O. (2020). Metodología para el plan de marketing internacional en la exportación de la empresa CubaRon, SA. *Retos de la Dirección*, 14(1), 68-88. <http://orcid.org/0000-0002-1445-4818>

Lozano, J. A., Marmolejo, J. A. y Rodríguez, R. (2020). *Designing a resilient supply chain: An approach to reduce drug shortages in epidemic outbreaks*. *EAI Endorsed Transactions on Pervasive Health*, 6(21), 1-15 <https://doi.org/10.4108/eai.13-7-2018.164260>.

Ludvigson, S. C., Ma, S., & Ng, S. (2020). *COVID19 and the macroeconomic effects of costly disasters*. *NBER Working Paper*, (26987), 1-22. <https://www.nber.org/papers/w26987>

Martínez, Y., Madera, J. y Leguen, I. (2016). Algoritmos evolutivos con estimación de distribución celulares. *Revista Cubana de Ciencias Informáticas*, 10, 159-170. <http://rcci.uci.cu>.

Martínez, Y., Madera, J., Rodríguez, A. Y. y Barigye, S. (2019). *Cellular Estimation Gaussian Algorithm for Continuous Domain*. *Journal of Intelligent & Fuzzy Systems*, 36(5), 4957-4967. <https://doi.org/10.3233/JIFS-179042>

Martínez, Y., Rodríguez, A., Madera, J., Moya, A., Morgado, B. y Mayedo, M. B. (2019). *CUMDANCauchy-C1: a cellular EDA designed to solve the energy resource management problem under uncertainty*. [ponencia] *Paper presented at the Proceedings of the Genetic and Evolutionary Computation Conference Companion*.

Martínez-López, Y., Oquendo, H., Caballero, Y., Guerra-Rodríguez, L. E., Junco-Villegas, R., Benítez, I. & Madera, J. (2020). Aplicación de la investigación de operaciones a la distribución de recursos relacionados con la COVID-19. *Retos de la Dirección*, 14(2), 86-105. <http://retos.reduc.edu.cu/>

Mladenović, N., Hansen, P. (1997). Variable neighborhood search. *Computers & operations research*, 24(11), 1097-1100. [https://doi.org/10.1016-S0305-0548\(97\)000312](https://doi.org/10.1016/S0305-0548(97)000312)

Mostapha, S. (2015). *Solving Vehicle Routing Problem using Simulated Annealing*. <http://www.yarpiz.com>.

Taillard, É. D. (1999). *A heuristic column generation method for the heterogeneous fleet VRP*. *RAIRO-Operations Research -Recherche Opérationnelle*, 33(1), 1-14. http://www.numdam.org/item?id=RO_1999_33_1_1_0

Toth, P. & Vigo, D. (2002). *The Vehicle Routing Problem*. *Monographs on Discrete Mathematics and Applications*. SIAM

Van, P. J. y Aarts, E. H. (1987). *Simulated annealing*. In *Simulated annealing: Theory and applications*. Springer.

Yepes, V. y Medina, J. R. (junio, 2002). Optimización de rutas con flotas heterogéneas y múltiples usos de vehículos *VRPHEMSTW*. [ponencia] V Congreso de Métodos Numéricos en Ingeniería. Sociedad Española de Métodos, Madrid, España.



Capítulo 6. Aplicaciones en la Educación Superior

Capítulo 6. Aplicaciones en la Educación Superior

Yumilka Fernández Hernández

Hilda López León

Olga Lidia Pérez González

Resumen

En este capítulo se aborda la utilización de la Inteligencia Artificial en función de la atención personalizada de la experiencia educativa con el objetivo de describir contribuciones orientadas al desarrollo de metodologías y/o algoritmos para la predicción de la deserción escolar en instituciones universitarias, y al desarrollo de los sistemas expertos para la preparación de los estudiantes universitarios en la asignatura Álgebra Lineal.

Abstract

This chapter addresses the use of Artificial Intelligence in terms of personalized attention to the educational experience with the objective of describing contributions oriented to the development of methodologies and/or algorithms for the prediction of school dropout in university institutions, and to the development of expert systems for the preparation of university students in the subject Linear Algebra.

Introducción

La Agenda 2030 de la UNESCO (Naciones Unidas 2018) establece metas claras para abordar los retos educativos del siglo XXI, centradas en su Objetivo 4. Este objetivo promueve el desarrollo sostenible, la mejora de la calidad educativa, y asegura una educación inclusiva y equitativa. Además, busca fomentar el aprendizaje continuo para todos, tal como señalan Grau (2022) y Fernández *et al.* (2023).

En 2019, durante la Conferencia Internacional sobre Inteligencia Artificial celebrada en Beijing (China), se aprobó el Consenso de Beijing, primer documento que propone consejos y recomendaciones sobre cómo sacar mayor provecho de las tecnologías de IA con miras a la consecución de la Agenda 2030 de Educación. Este cónclave proporcionó un marco común para orientar la aplicación de la IA en la gestión e implementación de los procesos educativos (UNESCO, 2019).

Más tarde, en el Informe Horizon 2021 (Pelletier et al, 2021), se reflejan algunas de esas demandas, que marcan tendencias clave en la Educación Superior. En el citado informe se analizan las tecnologías y prácticas que tendrán un impacto significativo en el futuro de la enseñanza y el aprendizaje, entre las que destaca a la Inteligencia Artificial y su importancia en la personalización de la experiencia educativa, algo que ya venía adelantándose desde diferentes perspectivas científicas (Luckin *et al.*, 2016 y Dorca, 2015).

Lo anterior evidencia el reconocimiento de la importancia de esta tecnología, a través de la cual se desarrollan algoritmos y métodos que permiten a un sistema informático resolver problemas de forma inteligente, tomando en consideración el contexto y la información disponible para alcanzar el objetivo deseado; estos tienen la capacidad de aprender, razonar o de planificar, diagnosticar, entre otras (Pulido, 2022).

En ese marco, se hace un llamado a trabajar más rápido en la introducción de los resultados científicos e investigativos relacionados con la educación (Silva, 2019), y se exige a las universidades la responsabilidad de “proyectar, con visión de futuro la formación del profesional universitario desde la gestión pertinente, con calidad, eficiencia y eficacia, sustentado en el uso de las tecnologías de avanzadas para mejorar sus procesos sustantivos, lograr su pertinencia y sostenibilidad” (Triana *et al.*, 2021, p. 2).

En relación a este tema, según Triana, García y Alarcón (2022), se plantea como un desafío actual en las instituciones educativas la necesidad de reconfigurar la gestión con el objetivo de proporcionar una atención personalizada a los estudiantes desde diversas perspectivas. Esto incluye aspectos como el rendimiento académico, la adaptación de las trayectorias educativas a las particularidades y comportamientos individuales de los estudiantes para su formación integral, entre otras.

Como respuesta a estos desafíos, se presentan dos propuestas educativas donde la Inteligencia Artificial mejora la personalización de la experiencia educativa. Los autores han contribuido significativamente en dos áreas principales. La primera es el desarrollo de metodologías y algoritmos para predecir la deserción escolar. La segunda es la creación de sistemas expertos para mejorar la preparación de los estudiantes universitarios en Álgebra Lineal.

Es objetivo de este capítulo describir ambas contribuciones y los resultados relevantes que se obtienen con su implementación.

6.1 Metodologías y/o algoritmos orientados a la predicción de la deserción escolar

La deserción escolar es un flagelo que incide en todas las carreras universitarias; se considera que aproximadamente entre el 25,9 % y el 30 % de la matrícula inicial abandona sus estudios, esas cifras suelen aumentar en las carreras de ingeniería (López, Madera y Jauriga, 2021).

Las contribuciones de la Inteligencia Artificial a la problemática de la deserción escolar se extienden a múltiples áreas, incluyendo el manejo del desbalance en las clases dentro de las bases de conocimientos. Tal enfoque se apoya en la necesidad de diversas Formas de Representación del Conocimiento (FRC) para el correcto funcionamiento de los sistemas informáticos.

Esas representaciones se organizan en una base de conocimiento estructurada en atributos, clases e instancias, proporcionando una fuente detallada de información sobre el objeto de estudio. Coca y Llivina-Lavigne (2021) subrayan la importancia de estas bases de conocimiento, que se caracterizan por su capacidad para adaptarse a las necesidades específicas de las aplicaciones de la Inteligencia Artificial en el ámbito educativo, facilitando así la identificación de patrones relacionados con la deserción escolar, a partir del reconocimiento de los rasgos que definen a cada uno de los conceptos fundamentales involucrados.

Las instancias representan los objetos, elementos o individuos de un área específica del conocimiento, por ejemplo: estudiante 1, estudiante 2; los atributos encarnan otros conceptos característicos, como la edad, el sexo, etc.; la clase clasifica específicamente el elemento que se está estudiando, por ejemplo, estudiantes con signos de alarma o estudiantes sin signos de alarmas en su actividad académica.

También se usan clasificadores de caja negra para diferenciar estudiantes con y sin problemas (López, Madera y Jauriga, 2021) y la implementación de un modelo basado en conjuntos aproximados probabilísticos para manejar datos desequilibrados (aproximación inferior). Se destacan igualmente el uso de umbrales para determinar la similitud entre objetos (Ramentol *et al.*, 2019) y la selección de atributos causales mediante la aplicación de algoritmos especializados, concretamente algoritmos de selección de atributos combinados (López, *et al.*, 2021).

Los datos extraídos para la investigación provienen del Sistema de Gestión de la Nueva Universidad (SIGENU). Este sistema incluye información personal de los estudiantes, así como datos demográficos, sociales y económicos. Sin embargo, algunas instancias no contaban con todos los datos necesarios, lo que llevó a la discretización de las bases de

conocimiento. Este proceso mejoró la calidad de los resultados mediante la complementación o eliminación de datos incompletos.

Los resultados comprenden dos categorías de datos: bases de conocimientos balanceadas y no balanceadas. Para las bases no balanceadas, el primer paso fue equilibrarlas. Este balanceo es esencial para entrenar eficientemente el algoritmo clasificador durante el proceso de clasificación.

Las contribuciones fueron las siguientes:

6.1.1.1 Metodología CRISP-DM para la predicción de la deserción escolar

La metodología CRISP-DM se ha consolidado como la principal guía para proyectos de Minería de Datos (MD) Cuenta con un Grupo de Interés Especial (*Special Interest Group* - SIG), integrado por unos 200 miembros, entre ellos, proveedores de minería de datos, consultores y usuarios finales.

Esta es una metodología de libre distribución, es equitativa y se utilizó para contextualizarla en la proyección de la deserción escolar en el contexto del primer año de la carrera Ingeniería Informática de la Universidad de Camagüey *Ignacio Agramonte Loynaz*, UC, a través de tres fases: Fase 1 Construcción de la Base de Conocimiento, Fase 2 Modelado y Fase 3 Validación

Se creó una base de conocimiento a partir de los datos del SIGENU, que incluían listados de estudiantes y registros de calificaciones de las cohortes 2003-2012, con un total de 530 instancias y 36 atributos. Tras descartar tuplas y atributos irrelevantes, y realizar una meticulosa selección, limpieza, construcción e integración de datos, se conformó una base de conocimiento más refinada. Esta contenía 18 atributos y 292 instancias, de las cuales 217 fueron clasificadas *sin problemas* y 75 *con problemas*.

Se determinó al calcular su Ratio de Desequilibrio (*Imbalance Ratio*): $IR = 3.9$, que se trabajaría con una base de datos

desbalanceadas, razón por la cual se seleccionaron técnicas de balanceo de datos como *Synthetic Minority Over-sampling Technique* (SMOTE) y *CostSensitiveClassifier* (Clasificador Sensible al Peso -o al Costo).

Se aplicaron algoritmos de selección de atributos para de esta manera determinar las causas de la deserción escolar. Los atributos de selección aplicados fueron: *ChiSquaredAttributeEv*, *GainRatioAttributeEval*, *OneRAttributeEval*, *ReliefAttributeEval*, *SymmetricalUncertAttributEval*. Además, se implementó como método de búsqueda, el algoritmo Ranker.

Se seleccionaron 10 algoritmos de clasificación entre los disponibles por la herramienta de Minería de Datos WEKA, debido a que estos algoritmos, son todos del tipo “caja blanca”, donde se obtiene un modelo de salida comprensible para el usuario, porque o se obtienen reglas de clasificación del tipo “Si – Entonces” o árboles de decisión, entre ellos: *JRip*, *NNge*, *OneR*, *PART*, *Ridor*, *ADTree*, *J48*, Árbol Aleatorio de Decisión *RamdomTree*, *REPTree*, *SimpleCartç*.

La experimentación reveló los algoritmos más efectivos. La regla de clasificación *Jrip* lideró la lista, seguida por el árbol de decisión *ADTree*. En tercer lugar, hubo un empate entre los árboles de decisión *J48* y *SimpleCartç*).

6.1.2. Clasificadores de caja negra para la predicción de la deserción escolar

En la minería de datos, la clasificación es una tarea común. Cada instancia se asigna a una clase, definida por un atributo llamado etiqueta de clase. Los demás atributos se utilizan para predecir la clase de nuevas instancias.

La etiqueta de clase es un valor discreto y es conocido para cada objeto, el objetivo fue construir los modelos de clasificación que asignan la etiqueta de clase correcta a objetos no vistos antes. La exactitud del resultado de la clasificación

se midió por la tasa de errores de registros incorrectamente clasificados. Las aplicaciones incluyen el análisis de la clasificación de la aprobación del crédito, identificación de diagnósticos médicos, ubicación de las tiendas, campañas y detección de fraudes. (López León, H. & Madera, J. & Jauriga, Y., 2021)

El principal objetivo de la clasificación es desarrollar modelos o funciones que distingan entre diferentes clases o conceptos. Estos modelos son cruciales para predecir las clases de objetos con etiquetas desconocidas. Pueden adoptar diversas formas, incluyendo árboles de decisión, redes neuronales y máquinas de soporte vectorial, aplicables tanto a algoritmos de caja blanca como de caja negra (López, Madera y Jauriga, 2021).

Los algoritmos de caja negra no proveen una explicación de cómo funcionan, ocultan detalles, muestran un resultado sin explicar cómo se genera y permiten el ajuste del parámetro opcional, lo que reduce el mínimo esfuerzo necesario para utilizarlos. Las redes neuronales y las máquinas de soporte vectorial se incluyen entre las técnicas que generan modelos de caja negra.

Para evaluar esas técnicas se utilizó una base de conocimiento con 18 atributos y 292 instancias, de ellas 217 clasificados como estudiantes sin problemas y los restantes 75 con problemas. Al ser un conjunto de datos desbalanceados se le aplicaron técnicas de balanceo con algoritmos como: *SMOTE*, *Borderline_SMOTE*, *Safe_Level_SMOTE*, *SMOTE_Tomek-Links*, *SMOTE_ENN* y *Spider2*.

La selección de atributos como su nombre lo indica selecciona un subconjunto de los originales, de manera que no se afecten significativamente los porcentajes de clasificación, los cuales son, los más influyentes en la variable de salida. Tiene como ventaja que mejora el desempeño predictivo y construye modelos más eficientes.

El proceso de selección involucra varias etapas. Primeramente, está la generación de candidatos y luego su

evaluación. Después se busca un criterio de parada que puede darse por la estrategia de búsqueda, el número de iteraciones realizadas o el número de atributos seleccionados. Luego se validan los resultados comparando el error en la clasificación con y sin la selección de atributos.

Por su parte, los algoritmos de selección de atributos se distinguen por su forma de evaluar los atributos y se pueden clasificar en filtros (*filters*), que seleccionan y evalúan los atributos en forma independiente del algoritmo de aprendizaje, envoltorios (*wrappers*) los cuales usan el desempeño de algún clasificador para determinar lo deseable de un subconjunto, y los híbridos que usan una combinación de los dos criterios de evaluación en diferentes etapas del proceso de búsqueda.

De acuerdo con López, Madera y Jauriga (2021), se utilizaron varios algoritmos tanto para la selección de atributos como para la clasificación de datos. Los primeros incluyeron ABB-IEP, LVW, SA-IEP, SBS-IEP y SFS-IEP. Para la clasificación de datos, se emplearon diversos algoritmos de redes neuronales, como EvRBFN, Incr_RBFN, iRProp+, LVQ_C, MLP_BP, MLP_CG, RBFN y SOON. Además, se utilizaron algoritmos de máquinas vectoriales, como C_SVM y Nu_SVM.

Con la experimentación se demostró que existe una mejoría en los resultados usando para la clasificación algoritmos de caja negra.

6.1.3 Aplicación de algoritmos de selección de atributos combinados con algoritmos de clasificación para el pronóstico de la deserción escolar

Se trabajó en la construcción de dos bases de conocimientos balanceadas, debido a que el índice $IR = 0.9$, además, como resultado del proceso de estudio estadístico realizado para la construcción de la base de conocimiento y la discretización de los datos se determinó que en el conjunto de datos existían

atributos que están correlacionados y otros que al calcular su varianza el resultado obtenido es cercano a cero y estos creaban ruidos en los resultados de la clasificación.

Una de las bases de conocimientos tiene 20 atributos y 252 instancias, de ellas 120 estudiantes clasificados como “sin problemas” y los restantes 132 con problemas; la otra base de conocimientos tiene 16 atributos con igual cantidad de instancias, a las que se le incluyó el rasgo decisor (López León, H. & Madera, J. & Jauriga, Y., 2021).

Los algoritmos clasificadores utilizados fueron: *k-Vecinos más Cercanos*, donde *k* es un número variable (*kNN*), Bosque Aleatorio (RF), Máquina de Potenciación del Gradiente (*GBM*), Máquinas de Soporte Vectorial (SVM), Redes Neuronales *Earth*, *Naïve Bayes*, *Rpart*, J48. Estos son algoritmos de caja negra, los cuales ofrecen ventajas significativas en términos de poder y flexibilidad para modelar problemas complejos, y al mismo tiempo presentan retos en términos de interpretabilidad. Se aplicaron también algoritmos de caja blanca, para determinar un mejor predictor, ahora en una base balanceada.

Se emplearon cinco métodos de selección de rasgos tipo filtros, se implementaron en R, específicamente en la librería *FSelector*: Chi-Cuadrado (*Chi.Squared*), *oneR*, *Relief*, Incertidumbre Simétrica (*SymmetricalUncertainty*) La Incertidumbre Simétrica es una medida que combina la ganancia de información (*Information Gain*) y la entropía para evaluar la relevancia de una característica. Corrige el sesgo de la ganancia de información hacia las características con más valores, proporcionando una medida normalizada de la dependencia mutua entre la característica y la clase. Valores más altos indican mayor relevancia y la Tasa de Ganancia o Ratio de Ganancia (*GainRatio*).

Como resultado del estudio experimental se determinó que el modelo creado por la base de conocimiento, con el uso de Incertidumbre Simétrica y la Máquina de Potenciación

del Gradiente, es el que estadísticamente es más verídico, corroborado sobre un conjunto de prueba; también se comprobó que el algoritmo GBM, tanto con el algoritmo de selección de atributos *GainRatio* como con *Symmetrical Uncertainty* proporcionó los mismos rasgos: municipio, institución, lugar de nacimiento, tipo de estudiante, prueba de ingreso de Historia e Índice; como el que mejor diagnostica a los estudiantes, en cuanto a su pase de año limpio o no, o causar baja. El modelo obtenido tuvo una precisión de 79,01%.

6.1.4 Minería de Datos para predecir la deserción escolar

Se construyó una base de conocimientos, con 21 atributos y 788 ejemplos, de ellos 477 clasificados como estudiantes sin problemas y los restantes 311 como con problemas, donde se incluyó el rasgo decisor y mantiene el IR = 1.5. Se demostró que era una base desbalanceada se aplicaron cuatro métodos de balanceo de datos (López, *et al.*, 2021).

Se aplicaron varios métodos para balancear los datos, incluyendo *Down* (Submuestreo), SMOTE, ROSE (*Random Over-Sampling Examples* o Ejemplos de Sobremuestreo Aleatorio, en español) y *Upsampling* (Sobremuestreo). Mientras que *Down* elimina instancias para equilibrar la clase, los otros métodos utilizan técnicas de sobremuestreo.

Como técnicas de selección de atributos se utilizaron algoritmos de selección de característica: *Weight Symmetrical Uncertainty*, *Weight OneR*, *Correlational Features Selection* y *Recursive Feature Elimination*, y algoritmos clasificadores: *Naïve Bayes*, *NNET*, *kNN*, *MSV*, *Rpart*, *GBM*, *Random Forest*, *XGBoost* y *Treebag* (López, *et al.*, 2021).

Los mejores resultados de precisión se obtuvieron con la utilización de algoritmos de clasificación, la selección de atributos y el balanceo de datos, evaluaron con mejor desempeño los clasificadores *Random Forest*, *XGBoost*, *Naïve Bayes* y *Treebag*.

Se demostró que el clasificador *Treebag* es un predictor eficaz evaluado sobre el conjunto de prueba *BOEMDRFE_Downsampling*, sin descartar el clasificador Random Forest a partir de que ambos clasificadores obtienen valores similares, además, en cada predicción se pudo determinar las variables con mayor importancia predictivas en cada uno de los modelos evaluados.

Se concluyó que los factores: índice académico, centro de procedencia, escalafón, y los resultados de los exámenes nacionales para el ingreso a la Universidad (Matemática y Español), tienen mayor peso en la decisión del estudiante para el abandono escolar.

6.2 Sistemas expertos para el Álgebra Lineal

Un sistema experto (SE) es un programa que posee conocimientos para ejecutar una tarea un tanto complicada, la cual comúnmente es realizada por un experto humano, su potencial está dado por el conocimiento que logra sintetizar, sus algoritmos de búsqueda y métodos específicos de razonamiento; se caracterizan por incorporar en forma operativa el conocimiento. En otras palabras, puede definirse como "...aquel programa de ordenador que contiene la erudición de un especialista humano versado en un determinado campo de aplicación... un modelo computarizado de las capacidades de razonamiento y habilidades en resolución de problemas del especialista humano" (Pino *et al.*, 2001, p. 10).

En el ámbito educativo, los Sistemas Expertos (SE) pueden sintetizar y utilizar el conocimiento de docentes experimentados en diversas materias. Esto permite no solo resolver problemas de manera efectiva, sino también explicar y justificar las soluciones propuestas.

Los SE están estructurados en tres componentes principales: la base de conocimientos, la máquina de inferencia y el mecanismo de explicación. La base de conocimientos alberga información especializada, extraída del experto en el dominio, la máquina de inferencia aplica métodos para

resolver problemas utilizando esta información y el mecanismo de explicación detalla al usuario el proceso de razonamiento seguido durante la inferencia.

Dada la naturaleza abstracta del Álgebra Lineal y las dificultades que enfrentan los estudiantes en su comprensión conceptual y procedimental, se percibió la necesidad de una herramienta de autoevaluación. Esta herramienta también serviría como un recurso de consulta para aclarar dudas y apoyar el estudio independiente. Para abordar estas necesidades, se desarrolló un Sistema Experto en Álgebra Lineal (SEAL), que ofrece apoyo integral a los estudiantes en esta asignatura (Caballero *et al.*, 2011; Casas *et al.*, 2013 y Coello *et al.*, 2015).

Ese proyecto primeramente fue realizado en su primera versión solo con el tema de Espacios Vectoriales (SEALv1.0) y dada su gran aceptación se decidió generalizarlo a otros dos temas más del Álgebra Lineal –Aplicaciones Lineales y Diagonalización (SEALv2.0)– por ser los temas de mayor dificultad entre los estudiantes.

El SEAL posee como una de sus principales funcionalidades posibilitar al estudiante obtener la vía de solución a un problema de Álgebra Lineal, seleccionando sus datos e incógnitas; ese caso de uso se denomina “Obtener Vía de Solución” y es el módulo inteligente de la aplicación, que se caracteriza por su capacidad de razonamiento a partir de una base de conocimiento y en la justificación de los resultados de la clasificación.

Las características de cada versión del SEAL son las siguientes:

SEAL v1.0:

Creado en 2010, este sistema presenta una base de conocimientos enfocada en Espacios Vectoriales, un tema crucial en Álgebra Lineal. Incluye 22 rasgos, 6 clases y 66 problemas, ejemplos o instancias que se han resuelto con éxito. El algoritmo para la máquina de inferencia fue el k-vecino más cercano (k-NN). Su estrategia de clasificación es que un nuevo caso se va a clasificar en la clase más frecuente a la

que pertenecen sus k vecinos más cercanos. Utiliza funciones de distancia o similitud para generar predicciones a partir de ejemplos almacenados.

SEAL v2.0:

Desarrollado en 2011, este sistema surgió por la necesidad de incorporar nuevos temas, específicamente Aplicaciones Lineales y Diagonalización, debido a su complejidad y a los desafíos que representan para los estudiantes. Esta segunda versión implementa un Sistema Experto basado en Redes Bayesianas. Utiliza el clasificador *Naive Bayes* (Bayes Ingenuo) que se fundamenta principalmente en asignar la clasificación más probable al nuevo caso, aprendiendo de los parámetros (probabilidades) de la red.

SEAL v3.0:

Este sistema, creado en 2012, tuvo como objetivo utilizar el mismo código fuente para desarrollar una aplicación web basada en la arquitectura modelo-vista-controlador. Esta aplicación se caracteriza por una base de conocimientos centralizada a la que todos los estudiantes pueden acceder. La versión actualizada incluye un mecanismo de inferencia avanzado basado en redes neuronales artificiales del tipo MLP (Perceptrón Multicapa, *Multilayer Perceptron*, en inglés). Además, presenta mejoras significativas en la interfaz gráfica, optimizando la usabilidad y la experiencia intuitiva del usuario.

Además, se desarrolló una metodología, conocida como MEBCALSE (Metodología para la Construcción de Bases de Conocimientos de Álgebra Lineal para Sistemas Expertos), presentada en Casas *et al.* (2013).

6.2.1 Metodología MEBCALSE

La metodología MEBCALSE permite a los docentes, incluso a aquellos sin especialización en informática, representar de manera flexible el conocimiento de Álgebra Lineal en una base de datos. Esta metodología se estructura en varias etapas

clave, cada una definida por acciones específicas, descritas a continuación:

Primera acción: Extracción del conocimiento

Primero, se selecciona el tema específico dentro del Álgebra Lineal y se determina cómo adquirir el conocimiento necesario. Esto implica realizar consultas a expertos, revisar libros, analizar materiales en formatos digitales o impresos y explorar bases de datos existentes

Segunda acción: Definición de los rasgos

En la Inteligencia Artificial, los rasgos son características que definen un objeto o ejemplo. Estos pueden tener distintos valores dentro de un dominio específico. Desde el punto de MEBCALSE se precisa que los rasgos son los datos e incógnitas de un problema o ejemplo, pero también se consideran rasgos a las incógnitas del problema. Los rasgos se representan como atributos numéricos con un dominio de 0 o 1. Un rasgo toma el valor de 1 si está presente en el problema, ya sea como dato conocido o como incógnita a resolver, y 0 si no está presente.

Tercera acción: Definición de las clases

Se describen los problemas tipo utilizando la teoría de Álgebra Lineal que esté aplicándose. Las clases se consideran rasgos continuos. Aunque todos los ejemplos se categorizan dentro de cinco tipos de problemas de Álgebra Lineal, se añade una sexta clase para agrupar ejemplos atípicos o absurdos que los estudiantes podrían introducir en el Sistema Experto.

De ahí que el dominio del rasgo clase sea el siguiente: Dominio: 1,...,6 (1-6).

Para la determinación de esos problemas tipos, se analizó que los principales problemas del Álgebra Lineal pueden formularse utilizando la teoría de Espacios Vectoriales, la cual se basa en la idea de la combinación lineal de vectores, mediante la resolución de un sistema de ecuaciones en la forma

matricial $AX=Y$, donde se denotan los vectores columna de la matriz A como a_i , $i=1\dots n$. Estos se enunciaron de la siguiente forma:

- 1) Problema directo: Dados los vectores de la Combinación Lineal a_i y los coeficientes de dicha combinación lineal x_i (elementos del vector X), calcular el vector resultante de dicha Combinación Lineal Y (de componentes y_i).
- 2) Problema inverso: Dados los vectores de la Combinación Lineal a_i y el vector resultante de dicha Combinación Lineal Y (de componentes y_i). Calcular los coeficientes de dicha Combinación Lineal x_i (elementos del vector X).
- 3) Problema de dependencia de los datos: Dados los vectores a_i y el vector Y (de componentes y_i). Plantear la interrogante de la dependencia lineal de estos vectores (existencia de la x_i , únicos o infinitos) o la independencia lineal de los mismos (no existencia de la x_i).
- 4) Problema de factibilidad: Dados los vectores a_i y el vector Y (de componentes y_i). Cuestionar la posibilidad de generar el vector Y , mediante la combinación lineal de los vectores a_i , siendo los coeficientes x_i (coeficientes de la combinación lineal) únicos, infinitos o si no existen (no es posible generar el vector Y combinando linealmente los vectores a_i).
- 5) Problema de unicidad: Dados los vectores a_i y el vector Y (de componentes y_i). Cuestionar la posibilidad de generar el vector Y , mediante la combinación lineal de los vectores a_i , siendo los coeficientes x_i (coeficientes de la Combinación Lineal) únicos.

Cuarta acción: Creación de las instancias

En esta acción, el docente debe refinar la información recopilada. Esto incluye descartar ejemplos repetidos, equilibrar la cantidad de ejemplos para cada clase y modificar

los problemas que presentan datos incompletos. Para crear instancias u objetos en la base, el docente asigna valores tanto a los rasgos conocidos (datos) como a los desconocidos (incógnitas). Luego, establece una clasificación adecuada para cada instancia.

Cada instancia resultante estará compuesta por un vector de datos y una clasificación. Es importante mantener un balance en el número de instancias por cada clase.

Quinta acción: Revisión de instancias por los expertos

Una vez conformada la Base de Conocimientos, esta debe ser revisada por un grupo de expertos. Su utilidad está dada por la veracidad del contenido. Las opiniones y sugerencias de estos expertos pueden requerir volver a acciones anteriores del proceso. Este paso garantiza un refinamiento adecuado del conocimiento experto antes de su incorporación final en la base.

Sexta acción: Validación Experimental

Varias herramientas experimentales están disponibles para validar la base de conocimientos confeccionada y detectar posibles errores. Entre ellas, WEKA, se destaca como un software integral para el aprendizaje automático y la minería de datos. Las herramientas de experimentación para bases de conocimientos contienen un grupo de algoritmos programados para hacer la inferencia de la solución con un subconjunto de la misma base.

Los resultados pueden variar dependiendo del algoritmo utilizado; lo importante es asegurar una clasificación correcta de la mayoría de las instancias en varios clasificadores. Después de este análisis, el diseñador de la base de conocimientos no solo puede refinarla y comprobar su eficacia, sino también identificar el algoritmo que más se ajusta a ella.

6.2.2 Selección del clasificador óptimo para el diseño e implementación de la máquina de inferencia del SEAL v1.0

Luego de ser especificada la estructura de la base de conocimientos, se hizo necesario un mecanismo para inferir a partir de esta, la clasificación de un nuevo problema.

Para seleccionar el clasificador adecuado, la base de conocimientos se sometió a pruebas utilizando la herramienta WEKA Knn-Workshop y SmartMLP. Estos programas ofrecen soluciones para aprendizaje automático y minería de datos, incluyendo herramientas para preprocesamiento de datos, clasificación, regresión, agrupación (clusterización), asociación de reglas y visualización.

Teniendo en cuenta que los rasgos usados son relevantes y la base de conocimientos no era tan extensa se determinó no realizar selección de atributos que consiste en aplicar un algoritmo que disminuya la dimensionalidad de rasgos y que no cambie la pertenencia a la clase.

Partiendo de la necesidad de dividir el conjunto de datos, cada clasificador fue evaluado por validación cruzada (*Cross-validation*), de modo que, el conjunto de datos se dividió en k particiones mutuamente exclusivas conteniendo todas aproximadamente el mismo número de ejemplos.

En cada ronda de evaluación, se reservó un subconjunto para pruebas y se entrenó el sistema con los $k-1$ restantes. El WEKA utiliza un caso particular de este método de evaluación, la “validación cruzada dejando uno fuera” (*leaving-one-out-crossvalidation*), donde k es igual al número de ejemplos del conjunto de datos. Así, el clasificador se entrena con todas las muestras excepto una, que se excluye para la prueba.

Adicionalmente, se realizaron pruebas utilizando el método de división porcentual (*Percentage Split*), que difiere del anterior. Con esta técnica, el clasificador se evalúa según la calidad

de predicción en un porcentaje específico de los datos, que varía según el valor introducido en el campo del porcentaje.

Para identificar diferencias significativas entre los algoritmos, se emplearon pruebas no paramétricas. Se utilizaron la prueba de Wilcoxon para dos muestras relacionadas y la prueba de Friedman para k muestras relacionadas. Los resultados demostraron diferencias notables, resaltando la superioridad del clasificador kNN en precisión frente a los demás.

El informe generado por WEKA para el clasificador kNN resumió los resultados de la siguiente manera, empleando la técnica de Validación Cruzada Estratificada (*Stratified cross-validation*):

Sumario

Métrica	Valor
Clasificaciones Correctas	65 (98.4848%)
Clasificaciones Incorrectas	1 (1.5152%)
Estadística Kappa	0.9816
Error Absoluto Medio	0.0578
Error Cuadrático Medio	0.1308
Error Absoluto Relativo	20.9914%
Error Cuadrático Relativo	35.2192%
Total de Instancias	66

Matriz de Confusión

Clase Real	a	b	c	d	e	f
Clasificado como	a = 1	b = 2	c = 3	d = 4	e = 5	f = 6
a = 1	8	0	0	0	0	0
b = 2	0	13	0	0	0	0
c = 3	0	0	15	0	0	0
d = 4	0	0	0	7	0	0
e = 5	0	0	0	0	11	0
f = 6	0	0	1	0	0	11

El clasificador k-NN logró una precisión notable, clasificando correctamente el 98,48% de las instancias, es decir, 65 de 66 casos. Solo hubo una clasificación errónea. En la matriz de confusión, se observó que una instancia de la clase 6 fue clasificada erróneamente como clase 3.

- Selección del k

Para encontrar el valor óptimo de k, se experimentó con varios valores en WEKA, obteniendo los mejores resultados con k=1. Sin embargo, para validar este hallazgo, se sometió la base de conocimientos a pruebas adicionales utilizando el software kNN Workshop. Para el entrenamiento se tomaron las instancias eliminando una de cada clase, las cuales luego serían usadas para el conjunto de prueba. Este proceso confirmó que el valor óptimo de k para esta base de conocimientos es 1, lo cual llevó a su selección definitiva.

6.2.3 Generalidades sobre el módulo “Obtener Vía de Solución”

El análisis de requisitos funcionales reveló que aquellos que fueron identificados están asociados a la orientación del estudiante hacia la búsqueda de soluciones para un problema específico, sobre la base de la clasificación de este en uno de los cinco problemas tipos del Álgebra Lineal mencionados previamente. Además, es importante no solo orientar al estudiante, sino también proporcionar una justificación detallada de los resultados obtenidos en la clasificación.

Así, el módulo “Obtener Vía de Solución” se compone de varios requisitos: orientar al estudiante en la solución de un problema, clasificar los datos proporcionados por el estudiante en un tipo de problema específico, mostrar y justificar los resultados de la clasificación, guardar estos resultados en la base de datos, rechazar los resultados si son inadecuados y asignar un entrenador si es necesario.

Este módulo permite al estudiante evaluar los efectos potenciales de nuevos datos y comprender su relación, de modo que el Sistema Experto en Álgebra Lineal sea un sistema computacional con características especiales, que integra en forma operativa el conocimiento de los profesores más experimentados en la asignatura Álgebra Lineal. Esta integración le permite resolver problemas de manera inteligente y justificar sus soluciones.

Lo que distingue a este sistema de las actividades típicas en plataformas interactivas es su capacidad para ofrecer soluciones basadas en una combinación de conocimientos empíricos y razonamiento lógico. Estos elementos básicos están contenidos en dos componentes separados, aunque relacionados: una base de conocimientos y una máquina de deducción, o de inferencia.

La base de conocimientos suministra ejemplos relevantes de la materia en estudio, y la máquina de deducción aporta

habilidades de razonamiento. Estas capacidades permiten al Sistema Experto derivar conclusiones lógicas sobre la clasificación de los problemas.

Las interfaces de usuario, como en otras aplicaciones, habilitan a los estudiantes para formular consultas, proporcionar información e interactuar con el sistema de diversas maneras.

6.2.4 Procedimiento seguido para la elaboración de la base de casos de Espacios Vectoriales del SEAL v1.0

La base de conocimiento estuvo formada por rasgos cuyo dominio fue numérico, solo admisible 0 o 1; para el diseño de la propuesta, se decidió que la entrada de los datos en la interfaz de entrada, sería a través de selección, o sea, el estudiante seleccionará en pantalla los datos e incógnitas del ejercicio en cuestión, por tanto, internamente en la base de conocimiento se asume como 1 si lo selecciona y 0 si no.

Inicialmente se obtuvo como resultado 25 rasgos, además, se determinó por los expertos que los principales problemas del Álgebra Lineal pueden enunciarse en términos de la teoría de espacios vectoriales, por lo tanto, cualquier problema de este contribuye a uno de los cinco problemas tipos descritos anteriormente, y que se consideraron como referencia para la definición de las cinco clases.

Resulta difícil considerar en una base de casos todas las selecciones que puede realizar un estudiante, algunas, en un número elevado de casos, carentes de coherencia y sentido lógico y totalmente alejadas de los conceptos de Álgebra Lineal, es por ese motivo que se decidió incluir una sexta clase que agrupa estos últimos casos, luego se comenzaron a adicionar casos a la base para lo cual fue necesario realizar varias consultas a expertos. Los valores del atributo tipo clase se consideraron discretos 1, 2, 3, 4, 5 y 6.

Se consideró una base de casos de prueba con 22 rasgos,

un rasgo de clase con un dominio de 1-6 y 66 instancias o problemas, ejemplos que fueron ejecutados exitosamente sobre la base de conocimiento para validarla.

Finalmente, un caso estuvo compuesto por un vector de datos y una clasificación. Los elementos que forman el vector van a ser los datos seleccionados por el estudiante, como la información y las incógnitas que tiene del problema. La clasificación va a ser uno de los cinco problemas tipos descritos anteriormente, o la clase de los problemas absurdos.

Los atributos del 1 al 12 que representan los posibles datos de un problema planteado por el estudiante y del 13 al 22 las incógnitas:

1. Vectores de R^n
2. Sistema de matrices
3. Sistema de vectores
4. Escalares de la CL
5. Vector resultante/ matriz resultante
6. Base del EV
7. Vector genérico
8. EV (condiciones)
9. Sistema de vectores canónicos/ Matriz idéntica
10. Sistema generador
11. LI
12. LD

-
13. Vector resultante
 14. Escalares de la CL
 15. Coordenadas
 16. LD
 17. LI
 18. Expresar un vector como CL de los demás
 19. Base
 20. Vectores resultantes (vector genérico).
Subespacio generado

21. Número máximo de vectores LI/Dimensión

22. Polinomio

Las clases son los problemas tipos 1-5

1. Directo
2. Inverso
3. Dependencia de los datos
4. Factibilidad
5. Unicidad

Para cada uno de estos temas del Álgebra se realizó una base de conocimiento independiente pues, aunque se haya determinado que los cinco problemas tipos responden a diversos temas, los rasgos varían, de ahí las propuestas del SEAL v2.0 y el SEAL v3.0.

6.2.5 Diseño de la interfaz de usuario

Mediante la interfaz el estudiante plantea los problemas al Sistema Experto y este le ofrece las explicaciones necesarias. Al seleccionar la opción “Obtener vía de solución” se le mostrará una nueva ventana para que seleccione los datos de su problema y las incógnitas, los cuales serán identificados por el sistema como atributos de la base de conocimiento. El sistema derivará una clasificación y mostrará al estudiante un reporte que contendrá la explicación de cómo la obtuvo y por cual vía puede resolver su ejercicio. Cada clase se vincula con varios entrenadores, asignados aleatoriamente, evitando obtener el mismo material para clasificaciones iguales.

El entrenador servirá para que el estudiante interiorice los conceptos y definiciones de la teoría de Espacios Vectoriales implícitas en su problema. Los mismos estarán relacionados con uno de los cinco problemas de la Combinación Lineal. De esta forma el usuario obtendrá una vía de solución a su ejercicio y así podrá resolverlo por diferentes caminos.

- Módulo de aprendizaje

Una vez clasificado, el nuevo caso se guardará en una tabla temporal. Luego un experto deberá comprobar si la clasificación es acertada, de ser así, incorporará el objeto a la base de conocimiento, si no lo elimina. Este mecanismo impide que se altere el contenido de la base de conocimiento ante una clasificación errónea y permite ampliarla según el criterio de los expertos.

Conclusiones

La capacidad de predecir la deserción escolar y proporcionar medios efectivos de autoevaluación para los estudiantes son elementos cruciales para la gestión académica en las universidades. Estos aspectos son especialmente significativos desde perspectivas metodológica, pedagógica y organizacional, ya que influyen directamente en la formación integral de los estudiantes. En este contexto, las metodologías y herramientas de la Inteligencia Artificial emergen como recursos valiosos. Ofrecen rigor y flexibilidad en la representación y análisis de datos educativos, facilitando interpretaciones precisas y detalladas del desempeño y las necesidades de los estudiantes. Al aprovechar estas capacidades, las instituciones pueden mejorar la sostenibilidad de sus estrategias formativas, basándose en resultados académicos sólidos y en una atención personalizada y ajustada a cada estudiante.

Trabajos futuros

Se trabaja en una investigación enfocada en mejorar la predicción académica para estudiantes universitarios de nuevo ingreso. Esto se logrará mediante el desarrollo y la experimentación con algoritmos avanzados. Estos algoritmos se diseñan para clasificar y predecir el rendimiento académico de los estudiantes, utilizando técnicas de conjuntos aproximados, conjun-

tos borrosos y reglas de clasificación. El objetivo es alcanzar una clasificación y predicción precisas, rigurosas y versátiles, lo cual permitirá una atención más personalizada en el proceso educativo, ayudando a identificar y atender las necesidades específicas de apoyo de cada estudiante.

Referencias del capítulo

Caballero, Y., Pérez, O., Docampo, L., Casas, L., Yordi, I. y Coello, L. (2011). *Sistema experto para el álgebra lineal* [ponencia]. XII Congreso de la Sociedad Cubana de Matemática y Computación COMPUMAT2011, Villa Clara, Cuba.

Casas, L., Pérez, O., Docampo, L., Caballero, Y., Coello, L., Yordi, I. y Martín, A. (2013). Sistema inteligente para el Álgebra Lineal. *Acta Latinoamericana de Matemática Educativa*, 26, 20-32. <https://core.ac.uk/download/pdf/20482649.pdf>

Coello, L., Casas, L., Pérez, O. y Caballero, Y. (2015). Redes neuronales artificiales en la producción de tecnología. *Revista Academia y Virtualidad*, 8(1), 12-20. <https://revistas.unimilitar.edu.co/index.php/ravi/article/view/442/238>

Fernández, Y. B., Pérez, O. L. & Caballero, Y. (2023). Diagnóstico académico en el colectivo de profesores universitario con Inteligencia Artificial. *PARADIGMA*, 44(3), 395-417. <https://doi.org/10.37618/PARADIGMA.1011-2251.2023.p395-417.id1455>

Grau, F. (2022). La educación inclusiva: conceptualización y retos para los docentes. *Revista Española de Educación Comparada*, 27, 123-148. <https://library.educase.edu/media/files/library/2021/4/2021hrteachinglearning.pdf?la=en&hash=C9DEC12398593F297CC634409DFF4B8C5A60B36E>

López, H., Madera, J. & Jauriga, Y. (2021). Minería de datos para predecir la deserción escolar en carreras de ingenierías en primer año. <https://www.researchgate.net/>

[publication/356584991_Mineria_de_datos_para_predecir_la_desercion_escolar_en_carreras_de_ingenierias_en_primer_ano](#) ISBN 978-959-16-4647-7

Naciones Unidas. (2018). *La Agenda 2030 y los Objetivos de Desarrollo Sostenible: una oportunidad para América Latina y el Caribe* (LC/G.2681-P/Rev.3). <https://repositorio.cepal.org/server/api/core/bitstreams/cb30a4de-7d87-4e79-8e7a-ad5279038718/content>.

Pelletier, K., Malcolm, D., Christopher, Mark McCormack., Reeves, J., Nichole, A., Bozkurt, A., Crawford, S., Czerniewicz, L., Gibson, R., Linder, K., Mason, J. and Mondelli, Victoria. (2021). *EDUCAUSE Horizon Report, Teaching and Learning Edition*. Boulder, CO. <https://www.learntechlib.org/d/219489>

Pino, R., Gómez, A. y de Abajo, N. (2001). *Introducción a la inteligencia artificial; sistemas expertos, redes neuronales artificiales y computación evolutiva*. Ed. Servicio de Publicaciones. <https://books.google.co.cr/>



Capítulo 7. Aplicaciones en los sistemas de generación, transmisión, distribución y uso final de la energía

Capítulo 7. Aplicaciones en los sistemas de generación, transmisión, distribución y uso final de la energía

Davel Eduardo Borges Vaconcellos

Luis Benigno Corrales Barrios

Santiago Lajes Choy

Eduardo Sierra Gil

Resumen

En este capítulo se aborda el uso de sistemas expertos para la toma de decisiones. Tales sistemas se basan en técnicas de inteligencia artificial —como lógica difusa, algoritmos genéticos y redes neuronales artificiales— para el mejoramiento del desempeño energético y la calidad de la energía en sistemas eléctricos de potencia e industriales y la gestión de mantenimiento predictivo y proactivo en la industria de la energía eléctrica. Todo ello a partir de la digitalización de procesos energéticos, técnicas de preprocesamiento inteligente de datos del sector energético, algoritmos de clasificación, predicción y teoría de conjuntos difusos para seleccionar los atributos y objetos más relevantes de un conjunto de datos que se utilizará para el entrenamiento de sistemas expertos; el entrenamiento, a través de bases de datos y el almacenamiento del conocimiento adquirido. Estos temas se presentan mediante ejemplos desarrollados para la solución de diferentes problemas del campo de la Ingeniería Eléctrica.

Abstract

This chapter addresses the use of expert systems for decision making, based on artificial intelligence techniques such as fuzzy logic, genetic algorithms and artificial neural networks, for the improvement of energy performance and

power quality in power and industrial electrical systems and the management of predictive and proactive maintenance in the electric power industry, based on the digitization of energy processes, intelligent preprocessing techniques for energy sector data, classification algorithms, prediction and fuzzy set theory to select the most relevant attributes and objects from a data set to be used for training expert systems; training, through databases and storage of the acquired knowledge. These topics are presented through examples developed for the solution of different problems in the field of Electrical Engineering.

Introducción

Los modelos basados en inteligencia artificial (IA) se han aplicado ampliamente para la previsión del consumo de energía en las últimas décadas. Estos se agrupan en modelos basados en redes neuronales artificiales y admiten modelos basados en máquinas de regresión vectorial. Adicionalmente, la evaluación del desempeño de los modelos en diferentes horizontes de pronóstico es un tema crítico que no ha sido resuelto y como la Inteligencia Artificial (IA) podría ayudar a lograr los objetivos futuros de las Fuentes de Energía Renovable, se han implementado métodos de IA inspirados en la estadística y la biología en varios estudios para lograr objetivos comunes y futuros de las fuentes de energía renovable (Liu & Chen, 2018).

En el contexto de que los países promueven vigorosamente Internet de la energía, una clave para el desarrollo futuro del sistema energético es la integración de la producción de energía, el uso y ahorro de energía y la integración profunda de los eslabones en una cadena de energía y las tecnologías de la información que contribuyen en gran medida al sistema energético de manera eficiente (Yang *et al.*, 2020).

Ciertamente, la tecnología de IA puede resolver de manera efectiva muchos problemas en la integración de sistemas de energía —incluidas las características no lineales, de

alta incertidumbre y multivariables— y mejorar la eficiencia operativa del sistema de energía, la seguridad, la confiabilidad y la capacidad inteligente. Además, la Internet de la energía se reconoce como un paradigma nuevo y avanzado de redes inteligentes, donde se utilizan dispositivos de recolección de energía, dispositivos de almacenamiento de energía distribuida y varios tipos de nodos de energía, interconectados mediante la aplicación de tecnología avanzada de electrónica de potencia, tecnología de la información y tecnología de gestión inteligente (Yue & Han, 2019).

Esta Internet de la energía emergente trae una mejora notable para la sociedad desde varios aspectos con sus ventajas en alta eficiencia, fuerte flexibilidad, gran escalabilidad y confiabilidad mejorada (Feng & Liao, 2020).

Sin embargo, a su vez genera nuevos desafíos en el diseño de la arquitectura, la operación de control y la gestión de la energía. Como resultado, aún se necesita más investigación sobre cómo enfrentar estos desafíos en internet de la energía mediante la aplicación de técnicas avanzadas, como sistemas multiagente, control basado en inteligencia artificial, computación y gestión de *Big Data* en la nube, etc.

7.1 Algoritmos genéticos

Los algoritmos genéticos (AG) fueron introducidos en los años sesenta y actualmente forman parte de la familia de procedimientos de optimización basados en la inteligencia artificial (Golberg, 1989). Es una técnica de búsqueda basada en la teoría de la evolución de Darwin, que funciona análoga al proceso de cruzamiento biológico de cromosomas que comparten información genética para crear un nuevo individuo. Al final, llegan a obtener un individuo de alta adaptación, lo que en el argot matemático no significa una solución óptima, pero sí la “mejor solución” (Golberg, 1989).

Se diferencian de los algoritmos de optimización tradicionales en cuanto a que:

Son algoritmos estocásticos. Dos ejecuciones distintas pueden dar dos soluciones distintas.

Son algoritmos de búsqueda múltiple, de los cuales es posible obtener varias soluciones.

La convergencia es poco sensible a la población inicial. No existen restricciones en el espacio de búsqueda.

Son algoritmos intrínsecamente paralelos, capaces de operar simultáneamente con varias soluciones, lo que asegura por lo general un funcionamiento más rápido en las corridas.

Resultan menos afectados por los máximos locales (falsas soluciones).

Sin embargo, se le debe prestar especial atención a la adecuada selección de los parámetros del AG para evitar la no convergencia o la convergencia en tiempos superiores al deseado.

La tipología de situaciones a resolver con esta técnica de optimización de búsqueda de valores extremos, incluye problemas complejos donde resulta imposible, o al menos muy difícil, encontrar un modelo matemático que lo represente, con espacios de búsqueda discretos y restricciones no lineales de las variables que lo definen.

Muchos problemas de la vida práctica muestran comportamientos complejos y las técnicas tradicionales de optimización de la matemática, incluso algunas heurísticas, no satisfacen su solución, siendo recomendable el empleo de técnicas basadas en la inteligencia artificial.

Algunas de las condiciones de los fenómenos que muestran este comportamiento complejo son:

Capacidad de realizar transiciones entre diferentes formas de comportamiento.

Son sistemas en estado de no equilibrio.

Son sistemas abiertos.

Son sistemas no lineales y retroalimentados.

Posibilidad de pérdida de estabilidad.

Capacidad de autoorganización.

En el caso de la ingeniería eléctrica, se han utilizado con éxito en la solución de problemas como: el balance de fases en circuitos de distribución primaria (Pérez, 2010), el planeamiento de redes de distribución (Pérez, 2011), la reconfiguración de sistemas de distribución (Carreno *et al.*, 2008) y la ubicación de dispositivos FACTS (Zhang & Milanovic, 2006), así como la compensación de potencia reactiva en los sistemas de distribución primaria (Junior *et al.*, 2008) y de forma simultánea al problema de la reconfiguración (Díaz *et al.*, 2010).

7.1.1 Compensación de potencia reactiva en sistemas desbalanceados utilizando algoritmos genéticos

La compensación de potencia reactiva es un fenómeno indispensable para la mejora de la eficiencia energética de los sistemas eléctricos (León-Martínez *et al.*, 2009). Esta se realiza usualmente en dos subsistemas bien definidos: el subsistema de distribución primaria, perteneciente al proveedor eléctrico, a través del cual se hace llegar la energía disponible a los clientes, y el subsistema de suministro eléctrico o también denominado de distribución secundaria (de baja tensión por lo general), propiedad en muchos casos del cliente, el eslabón final donde se satisfacen las necesidades energéticas de los receptores.

Los métodos de selección de tales dispositivos han sido basados en técnicas de optimización, clasificadas en: analíticas, de programación matemática, heurísticas y de inteligencia artificial, encontrándose dentro de estas últimas los algoritmos genéticos (Ng *et al.*, 2000).

Las técnicas tradicionales (analíticas y de programación matemática) son determinísticas y requieren una formulación matemática exhaustiva del modelo de compensación, a los

efectos del empleo de la herramienta específica de optimización.

Por ejemplo, si se pretende utilizar la técnica simplex, es necesario que el modelo de optimización esté representado por una función objetivo lineal, con restricciones de dominio lineal; en tanto si se requiere utilizar la técnica de programación cuadrática secuencial, la función objetivo deberá representarse a través de una ecuación cuadrática (González, 2002). La premisa de adaptar el problema al modelo de solución es una limitante significativa para el empleo de estas técnicas. Un grupo importante de técnicas heurísticas también poseen este tipo de limitante.

Formulación del problema de adaptabilidad

El problema en cuestión, en este caso, es: determinar la cantidad de unidades de compensación de cada capacidad disponible, necesarias para compensar de forma fija simétrica y asimétrica o escalonada simétrica la potencia reactiva en el sistema de suministro eléctrico, para que se obtenga el máximo de Valor Actual Neto (VAN).

Como consideraciones del problema aproximado (condiciones iniciales) fueron planteadas las siguientes:

Serán considerados en el análisis todos los efectos de la compensación de reactivo, que puedan ser expresados desde el punto de vista económico (inversiones en los bancos, gastos de explotación, beneficios por la mejora del factor de potencia y pérdidas de energía, así como por la liberación de carga en caso necesario).

Los análisis de flujo de cargas serán considerados de forma trifásica, teniendo en cuenta la existencia de desbalances en el sistema.

El modelo de gráficos de cargas deberá permitir la consideración de más de un comportamiento patrón, en los periodos de tiempo necesarios para su caracterización más exacta.

Para la compensación se utilizarán bancos de capacitores fijos (simétricos y asimétricos) o bancos controlados escalonados simétricos.

La compensación fija se aplicará en todos los posibles nodos del sistema, siempre de forma exclusiva o derivados de aquellos nodos con compensación escalonada simétrica.

La compensación escalonada simétrica no deberá utilizarse en nodos derivados con igual tipo de compensación.

Los efectos sobre la tensión del sistema (baja o alta) deberán ser considerados en el análisis, de modo tal que no se excedan los límites de desviaciones establecidos en las normas.

Modelo propuesto

Dados:

Datos del sistema eléctrico.

Datos del comportamiento de las cargas.

Datos de facturación del servicio eléctrico en el último año de explotación de la instalación.

Datos de los bancos de capacitores disponibles.

Datos de análisis económico.

Se necesita minimizar la función de adaptabilidad

Fitness:

$$Fitness = Kt - (Bfp - Bc - Cp - Ca) \cdot \sum_{j=1}^n \frac{1}{(1+\tau D)^j} + 10^7 \cdot Penal \quad (1)$$

Donde:

Kt: Costos de inversión de los bancos de capacitores (se incluyen los costos asociados a las unidades de compensación en relación con el nivel de tensión, los costos auxiliares en armarios y accesorios, costos de los dispositivos de control y costos de montaje).

Bc: Beneficios por la reducción de pérdidas de energía activa.

Bfp: Beneficios por la mejora del factor de potencia que incluye las pérdidas de energía eléctrica.

Cp: Gastos de explotación (pérdidas de los bancos).

Ca: Gastos de explotación (amortización de los bancos).

TD: Tasa de descuento.

n: Años de explotación.

Penal: Función de penalización.

La función de penalización incorpora las restricciones al espacio de búsqueda del algoritmo y es la suma de magnitudes constantes determinadas a partir de los aspectos siguientes:

Disponibilidad de unidades capacitadoras.

Capacidad máxima de compensación (limitación del espacio de búsqueda).

Cumplimiento de las normativas de desviaciones de tensión (máxima y mínima).

Como se puede notar, la expresión (1) sin incluir las penalizaciones representa el Valor Actual Neto (invertido en signo) obtenido a partir de los efectos económicos de la compensación.

El pseudocódigo específico del algoritmo diseñado para implementar el modelo de compensación es el siguiente:

[Datos] Lectura de datos necesarios.

[Inicio] Generar una población aleatoria de n

Individuos (*PopulationSize*, n) con cromosomas, expresados en el vector fila [x], codificados en forma binaria (*PopulationType*, 'bitstring') de longitud nvars de acuerdo al problema.

[Fitness] Evaluar la función Fitness (-VAN+Penal) para cada cromosoma x de la población. En este caso, la calidad de la evaluación estará dada por la minimización de esta.

[Nueva población] Crear una nueva población, repitiendo los siguientes pasos hasta que se cumpla la condición de parada, en este caso por el número de m generaciones preestablecidas en el algoritmo (*Generations*, m).

[Selección] Seleccionar dos cromosomas padres de una población de acuerdo a su Fitness.

[Crossover] Con un porcentaje de cruzamiento

(*CrossoverFraction*) crear una nueva descendencia.

[Mutación] Con un porcentaje de mutación de acuerdo con la función seleccionada (*MutationFcns*) modificar la nueva población.

[Aceptar] Colocar los nuevos descendientes en la nueva población.

[Remplazar] Usar la nueva población generada para una futura corrida del algoritmo.

[Test] Si la condición de salida se satisface, parar y retornar la mejor solución (cromosoma x) para la población actual.

[Loop] Regresar al paso 3.

Método de compensación de la potencia reactiva

A partir de la formulación del problema de optimización y el modelo matemático obtenido se puede establecer el método de compensación, que se compone de los aspectos siguientes:

Obtención de los datos necesarios para la aplicación, provenientes de los registros documentales de las empresas o trabajos previos de obtención de estos.

Establecimiento del modelo de optimización deseado para determinar la cantidad de unidades de cada tipo necesarias para compensar bajo las condiciones de compensación previamente especificadas:

- Con bancos fijos simétricos.
- Con bancos fijos simétricos y asimétricos.
- Con bancos controlados simétricos.
- Con bancos fijos simétricos de conjunto con bancos controlados simétricos.
- Con bancos fijos simétricos y asimétricos de conjunto con bancos controlados simétricos.

Ajuste de los parámetros de corrida del algoritmo genético de optimización, para asegurar la convergencia. En este caso, se sugiere iniciar con los parámetros especificados a continuación y bajo el criterio de corridas sucesivas del algoritmo, comprobar la convergencia de la solución ante la variación de los parámetros. Los principales parámetros de corrida obtenidos en la experimentación para la solución del

problema de la compensación a partir del algoritmo genético simple son:

- Tamaño de la población: 100 individuos.
- Tipo de variables de exploración: cadena de bit.
- Criterio de parada: Por el valor de la tolerancia de la función de adaptabilidad y por máximo número de generaciones (50) excedida.
- Probabilidad de cruzamiento: 0.8
- Función de mutación: tipo Gaussiana.

A partir de la cantidad de unidades de compensación de cada tipo y bajo las condiciones de compensación especificadas, seleccionar los parámetros técnicos de los dispositivos para cada variante de alternativa de solución (al menos 10) que maximice la función objetivo establecida en el modelo con una desviación máxima de un 5 %. Fundamentar en cada caso la factibilidad a partir del análisis económico, empleando como sugerencia los métodos del Valor Actual Neto (VAN) y el Plazo de Retorno de la Inversión (PRI), así como los indicadores económicos que lo avalan (costos de inversión, gastos de explotación, beneficios por la mejora del factor de potencia y pérdidas, entre otros).

Efectuar la toma de decisiones por parte del especialista acerca de qué variante asumir, tomando en consideración factores principalmente relacionados con la disponibilidad real del proveedor de bancos de capacitores y la inversión inicial a realizar, establecida en la futura relación contractual.

Proceder a la adquisición y montaje de la variante de bancos de capacitores asumida.

Fase experimental

Para validar el método de compensación, se realizaron dos tipos de ensayos:

Pruebas a escala del modelo y sus valores extremos de solución. La finalidad de estas pruebas es comprobar el funcionamiento correcto del algoritmo bajo condiciones iniciales de frontera, donde son evidentes o han sido estudiadas y se conocen las respuestas probables de solución.

Pruebas de comprobación en relación con otros modelos.

En este caso, se comprueban los resultados del algoritmo, en relación con determinadas condiciones iniciales asumidas por otros modelos y tomadas en consideración desde el propio algoritmo propuesto, a fin de comprobar la superioridad de este.

En ambos ensayos se utilizaron dos instalaciones de suministro eléctrico reales, pertenecientes a hoteles del sector consumidor terciario de la República de Cuba, cuyos esquemas topológicos se muestran en las Figuras 1 y 2.

Para las pruebas a escala del modelo y sus valores extremos

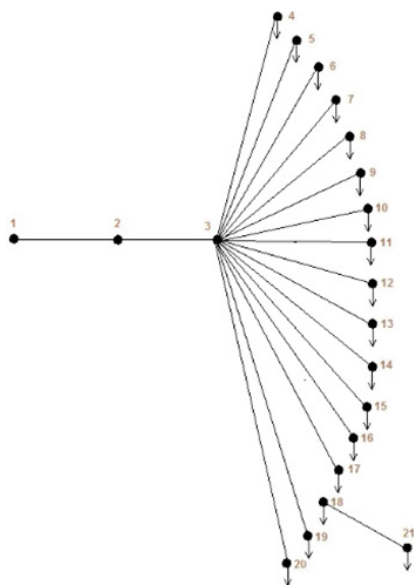


Figura 1. Esquema topológico del sistema eléctrico (caso de estudio 1).

de solución se implementaron 12 casos experimentales a razón de comprobar su efectividad, ante variaciones de las condiciones iniciales. Para asegurar una respuesta correcta (adecuada convergencia) se realizaron al menos 10 corridas de cada uno de los 12 casos.

Los casos experimentales de pruebas a escala con valores

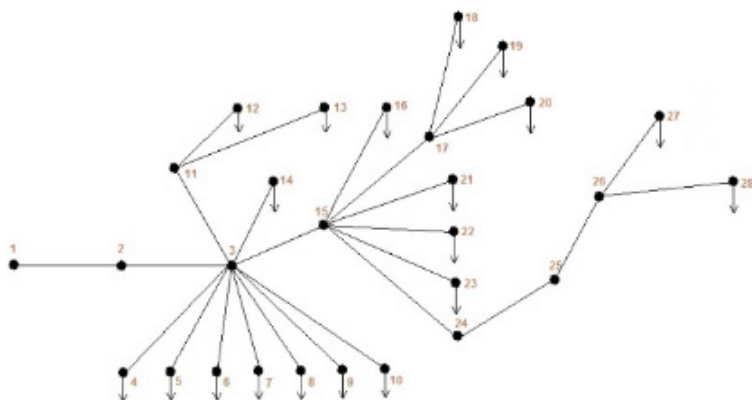


Figura 2. Esquema topológico del sistema eléctrico (caso de estudio 2).

extremos del modelo y sus resultados fueron:

Considerando cero los gastos de explotación y los beneficios (sólo se consideran los costos de inversión): Como se esperaba no se ubica compensación en ninguno de los nodos del sistema.

Considerando cero los costos de inversión y por ende amortización: La compensación se ubica indistintamente en los nodos de mayores pérdidas de energía, asegurando un factor de potencia entre 0,96 y 1.

Considerando cero los gastos de explotación (amortización y pérdidas de las unidades): La compensación se ubica indistintamente en los nodos de mayores pérdidas de energía, asegurando un factor de potencia entre 0,96 y 1, pero como regla ligeramente más bajo que en el caso anterior.

Considerando cero los gastos de pérdidas de las unidades: La compensación se ubica indistintamente en los nodos de mayores pérdidas de energía, con un resultado similar al caso anterior.

Considerando cero los gastos de amortización de las unidades: La compensación se ubica indistintamente los nodos de mayores pérdidas de energía, con un resultado similar al caso anterior.

Considerando cero los gastos de inversión y explotación: La compensación se ubica prácticamente en todos los nodos, asegurando un factor de potencia muy cercano a 1.

Considerando solo los beneficios por la mejora del factor de potencia (sin pérdidas): La compensación se ubica a nivel global, asegurando un factor de potencia indistintamente entre 0,96 y 1.

Considerando solo los beneficios por la reducción de pérdidas: La compensación se ubica prácticamente en todos los nodos, asegurando un factor de potencia muy cercano a 1.

Considerando cero los gastos de explotación (amortización y pérdidas de las unidades), así como la reducción de pérdidas: La compensación se ubica a nivel global, asegurando un factor de potencia muy cercano a 0,96.

Considerando cero los beneficios por la reducción de pérdidas: La compensación se ubica a nivel global, asegurando un factor de potencia cercano a 0,96, similar al caso anterior.

Considerando cero los beneficios por la mejora del factor de potencia: No se ubica compensación en ninguno de los nodos del sistema.

Considerando cero los beneficios por la mejora del factor de potencia y los gastos de explotación: No se ubica compensación en ninguno de los nodos del sistema.

En la experimentación solo se utilizaron unidades compensadoras trifásicas, para lograr un adecuado nivel de comparación en relación con los modelos existentes.

Este algoritmo muestra un comportamiento estocástico. Las soluciones en diferentes corridas son por naturaleza diferentes, aunque convergen a un valor estable de la función objetivo VAN con un porcentaje muy bajo de desviación (inferior al 5 %).

Los resultados obtenidos en las pruebas experimentales a escala con el modelo para valores extremos coinciden con los encontrados en la bibliografía sobre el tema, lo que corrobora la validez del modelo y el algoritmo implementado.

Para las pruebas de comprobación en relación con otros modelos se simularon sus condiciones iniciales y se compararon los resultados con respecto al modelo propuesto.

Los parámetros de calidad empleados en la comparación fueron el valor de la función de adaptabilidad (VAN), el tiempo de corrida del algoritmo y los resultados de convergencia.

Para el caso de la ubicación de bancos fijos (simétricos y/o asimétricos), existe un predominio de compensación con bancos asimétricos, dado el desbalance de los circuitos del sistema. Se pudo comprobar que la compensación se ubica en aquellos nodos de mayores niveles de pérdidas de energía y en correspondencia con el grado de asimetría de la potencia reactiva.

El otro caso evaluado fue la ubicación global de un banco controlado, a partir de valores estándares de unidades capacitadoras, que aseguran un factor de potencia luego de la compensación de máxima bonificación (0,96), según la reglamentación vigente en la República de Cuba.

Para ambos casos se observa que se obtienen diferentes respuestas por corridas con un máximo y aproximado VAN, lo que refuerza el comportamiento estocástico del algoritmo, a la vez que posibilita al especialista la toma final de decisiones.

Esta técnica de inteligencia artificial fue aplicada también con buenos resultados a la compensación de potencia reactiva en sistemas de distribución primaria de energía. En este caso, un aspecto básico para el método presentado, descansa en la obtención del modelo de la red, lo que se traduce en la formación de la matriz admitancia de la misma.

Esta matriz se compone adicionando todas las submatrices admitancia de los diferentes elementos que integran el sistema

eléctrico: transformadores, alimentadores y otros elementos pasivos como bancos de condensadores, etc.

Como consideraciones del problema aproximado (condiciones iniciales) se plantean las siguientes:

Serán considerados en el análisis todos los efectos de la compensación de reactivo que puedan ser expresados desde el punto de vista económico.

Los análisis de carga serán considerados en cada una de las fases, ya que se tiene en cuenta la existencia de desbalances en el sistema.

Para la compensación se utilizarán bancos de condensadores fijos simétricos y/o asimétricos.

La compensación se aplicará en todos los posibles nodos del sistema.

Se limitará la cantidad de unidades condensadoras disponibles para la compensación.

Se limitará la capacidad de compensación de potencia reactiva para los estados de máxima demanda.

Los efectos sobre la tensión del sistema (baja o alta) deberán ser considerados en el análisis, de modo tal que no se excedan los límites de desviaciones establecidos en las normas

Como variables de decisión del problema se consideran las cantidades de unidades de compensación de cada capacidad monofásica disponible a ubicar en cada nodo. Por lo tanto, los cromosomas estarán representados por la matriz $[x]$, compuesta por los elementos (cadena de genes) x_{ij} , donde i es el número de nodos por fases y j la cantidad de capacidades disponibles. Debido a lo anterior, los elementos x_{ij} de la matriz $[x]$ deben ser positivos y enteros ($x \in \mathbb{N}$). Esta condición deberá tenerse en cuenta a la hora de generar cada individuo de la población. Para esto, como alternativa cada elemento x_{ji} a su vez, es codificado en forma binaria por una cadena de n bit.

Como función de adaptación (Fitness) en primera instancia se toma el -VAN, añadiendo una función de penalización (Penal) incrementada por una cantidad suficientemente grande (10^7), tal y como se presentó en la ecuación (1).

La función de penalización se compone de la suma de penalizaciones especificadas cuando se incumplen las condiciones siguientes:

El número de unidades de compensación de un tipo excede de la cantidad de unidades disponibles de ese tipo. Esta restricción permite tener en cuenta la cantidad disponible de unidades compensadoras de cada tipo.

El máximo valor de potencia reactiva de compensación a ubicar en una fase de forma global es mayor que el máximo valor de potencia reactiva de fase global del sistema en cualquier período de medición.

El máximo valor de potencia reactiva de compensación a ubicar en una fase para cada nodo es mayor que el máximo valor de potencia reactiva de fase global del sistema en cualquier período de medición.

Las tres restricciones anteriores se consideran excluyentes del algoritmo, por lo que los individuos (soluciones) formados bajo esas condiciones tendrán una función de adaptación igual al valor de las penalizaciones acumuladas.

El mayor valor de la tensión efectiva en cualquier nodo para cualquier estado de carga es mayor que la máxima tensión permisible.

El menor valor de la tensión efectiva en cualquier nodo para cualquier estado de carga es menor que la mínima tensión permisible.

Las últimas dos condiciones permiten amonestar a aquellos individuos (soluciones) no deseados, luego de la evaluación de los efectos de la compensación.

Se asumen como valores límites de las desviaciones de tensión -10% a +10% de la tensión del sistema en el nodo

de entrega del suministro. De esta forma, la función Fitness para los individuos (soluciones) en este caso quedará definida según la ecuación (2).

$$Fitness = \left\{ Kt - (Bc + Cs - Cp - Ca) \cdot \sum_{j=1}^n \frac{1}{(1+TD)^j} + 10^7 \cdot (Penal_4 + Penal_5) \right\} \quad (2)$$

Donde:

Cs: Beneficios por la liberación de la capacidad de carga.

Para las condiciones especificadas en cada caso.

Todas las condiciones fijadas por la función Penal permiten restringir el espacio de búsqueda, lo que favorece el tiempo de corrida del algoritmo.

Finalmente, el modelo de optimización, quedaría expresado de la siguiente forma:

Datos del sistema:

Diagrama monolineal de conexiones de los nodos y fases (incluye sistema de potencia).

Datos de los elementos del sistema (transformadores y conductores).

Para transformadores: nodos de ubicación, topología de conexión de los devanados, capacidad nominal, resistencia y reactancia en % y valores de las derivaciones (tap).

Para conductores: nodos de ubicación, fases, secciones por fase, longitud, tensión permisible y tipo de conductor.

Tensiones en los nodos.

Datos del comportamiento de las cargas:

Matriz potencia activa y reactiva de fase compuesta por los elementos P_{ik} y Q_{ik} donde i es el número de nodos y k . la cantidad de estados de carga de comportamiento regular.

Intervalos de medición Δt_k .

Datos de los bancos de condensadores disponibles:

Unidades fijas de compensación disponibles. Cantidades disponibles de las unidades fijas.

Costo de inversión de las unidades fijas disponibles. Costos de inversión de los accesorios.

Costo de inversión de la instalación.

Datos de análisis económico:

Años de explotación (n). Tasas de descuento (TD).

A partir de la formulación del problema de optimización y el modelo matemático obtenido se puede establecer el método de compensación, que se compone de los aspectos siguientes:

Obtención de los datos necesarios para la aplicación.

Establecimiento del modelo de optimización deseado para determinar la cantidad de unidades de cada tipo necesarias para compensar bajo las condiciones previamente especificadas:

Con bancos fijos simétricos.

Con bancos fijos simétricos y asimétricos.

Ajuste de los parámetros de corrida del algoritmo genético de optimización, para asegurar la convergencia. En este caso se sugiere iniciar con los parámetros de referencia sugeridos por la herramienta computacional y bajo el criterio de corridas sucesivas del algoritmo, comprobar la convergencia de la solución ante la variación de los parámetros. Es recomendable realizar corridas iniciales bajo condiciones de frontera donde sean conocidos los resultados.

A partir de la cantidad de unidades de compensación de cada tipo y bajo las condiciones de compensación especificadas, seleccionar los parámetros técnicos de los dispositivos para cada alternativa de solución que maximice la función objetivo del modelo con una desviación máxima preestablecida. Fundamentar en cada caso la factibilidad a partir del análisis económico, donde se emplea como sugerencia los métodos del Valor Actual Neto (VAN), así como los indicadores económicos que lo avalan (costos de inversión, gastos de explotación y beneficios por la liberación de carga).

Efectuar la toma de decisiones por parte del especialista

acerca de cuál alternativa considerar, para lo cual se toman en consideración factores principalmente relacionados con la disponibilidad real del proveedor de bancos de condensadores y la menor inversión inicial a realizar. Un criterio recomendado para la toma de decisión a partir de las alternativas, puede ser el Plazo de Recuperación de la Inversión (PRI).

Proceder a la adquisición y montaje de la variante de bancos de condensadores seleccionada.

Fase experimental

Se diseñó una herramienta computacional y se aplicó el método en la compensación de potencia reactiva de un circuito de distribución primaria de la ciudad de Las Tunas, Cuba.

Un resumen de los principales resultados obtenidos permite considerar que:

El ahorro de energía activa promedio como resultado de la reducción de las pérdidas es mayor 2,2 veces cuando se ubican bancos simétricos y asimétricos en relación a la ubicación de bancos simétricos. Esto se debe a que la compensación por fases permite una mayor reducción de pérdidas al compensar el desbalance.

Se logra una reducción del desbalance promedio de las tensiones, al mitigarse la sobrecompensación.

La magnitud de compensación de reactivo, así como la cantidad de nodos a ubicar es mayor cuando se ubican bancos simétricos y asimétricos en relación con la ubicación de bancos simétricos.

En consecuencia, el VAN se incrementa 2,1 veces con la combinación de bancos simétricos y asimétricos, ajustados a las demandas de potencia reactiva por fase en relación con las asimetrías del sistema y la potencia reactiva equivalente.

Los límites de tensión en todos los casos se mantienen en los valores de tolerancia establecidos del 10 %, siendo de 5,8 %.

7.2 Redes Neuronales

7.2.1 Clasificación de fallas con redes neuronales para grupos electrógenos

Con el incremento del grado de dependencia de la sociedad moderna de los sistemas y procesos tecnológicos complejos, su disponibilidad y correcto funcionamiento se han convertido en una cuestión estratégica, donde las tareas de diagnóstico y clasificación de fallos juegan un rol muy importante con el fin de garantizar y mantener en operación continua y confiable al proceso, los fallos pueden provocar desde una reducción del desempeño hasta un daño que provoque paradas en la producción. La generación distribuida de energía eléctrica a través de los grupos electrógenos instalados, no está ajena a sufrir fallas. Este ejemplo muestra un sistema para el diagnóstico y clasificación de fallos para la unidad de motores Diesel (MTU) del Emplazamiento de Grupo Eléctrico Camagüey 1.

Para la implementación del sistema de diagnóstico se utilizó una red neuronal *Multilayer-Perceptron* (MLP). Las redes MLP entrenadas por el algoritmo *BackPropagation*, son consideradas quizás las más generales de las Redes Neuronales Artificiales. Son usadas en problemas de predicción, clasificación, reconocimiento de patrones, estimación de parámetros y resolución de señales.

Este algoritmo de entrenamiento de las Redes Neuronales, tiene la ventaja de no necesitar un conocimiento previo de la forma de la señal analítica tratada, o sea, puede ser usado para modelar el sistema, que es de gran utilidad en los casos en que el modelo matemático que describe el sistema es desconocido. Además, el trabajo con este tipo de redes puede utilizarse con información heterogénea, o sea, los vectores de entrada pueden contener variables de diferente naturaleza, lo que en principio permite su uso en problemas de naturaleza muy diferentes.

Diseño del sistema de clasificación basado en las redes neuronales artificiales

El proceso del diseño consiste en los siguientes pasos:

1. Construcción de los conjuntos de entrenamiento y prueba.
2. Selección de la arquitectura de la Red Neuronal Artificial.
3. Entrenamiento.
4. Evaluación de la red entrenada usando patrones de prueba.

Construcción de los conjuntos de entrenamiento y prueba
La red primero debe ser entrenada con un conjunto de entrenamiento el cual incluye ejemplos de datos en buen funcionamiento y otros con fallo. Al finalizar el entrenamiento, la red deberá estar lista para reconocer los ejemplos aprendidos, y clasificar otros nuevos basándose en las generalizaciones hechas a partir del aprendizaje (Caballero, 2007 y Caballero *et al.*, 2007). Para usarla se emplea como una función, la cual se evalúa y da un resultado. En la evaluación, la red recibe un vector de entrada de componentes reales que identifica a un patrón determinado, y luego de “analizarlo”, devuelve la clase o patrón al cual debe pertenecer dicho vector.

El entrenamiento es el proceso durante el cual la red MLP aprende los ejemplos que se le enseñen y clasifica a cada ejemplo del conjunto de entrenamiento y, en dependencia del error que comenta, se rectificará ella misma, para cuando vuelva a evaluar a ese ejemplo, intentar hacerlo mejor. Luego de clasificar a todos los ejemplos, se considera que la red ha efectuado un paso del algoritmo de entrenamiento. Este proceso se ejecutará hasta que se cumpla una condición de parada.

Selección de la arquitectura de la Red Neuronal Artificial

Sin lugar a dudas, la topología de red neuronal mejor conocida y más utilizada en la actualidad para la solución de problemas de clasificación es la MultiLayerPerceptron (MLP).

La literatura, para soluciones del tipo de problema, hace referencia a la utilización de redes multicapa con conexiones hacia delante (feed-forward), aprendizaje supervisado y corrección de error aplicando el algoritmo de propagación de los errores hacia atrás (backpropagation).

Partiendo de lo anterior definimos una red MLP de 2 capas, una capa oculta con 30 neuronas y 12 en la salida, la primera capa con función de transferencia tansig y la segunda purelin.

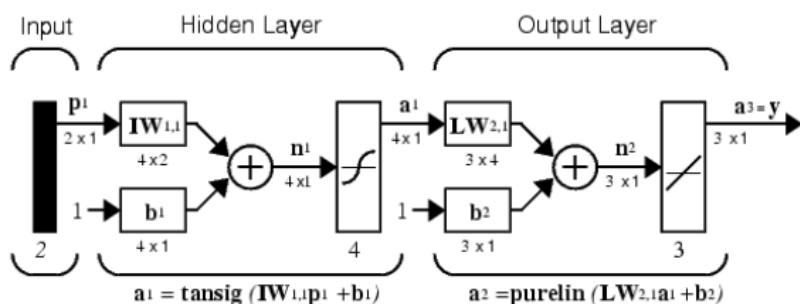


Figura 3. Ejemplo de topología de red Feedforward

El método a seguir para desarrollar la red neuronal multicapa, fue el siguiente:

Para la primera red, que funciona como generador de residuos, como punto de partida, se comenzaron las pruebas a partir de la topología propuesta: 30 neuronas en la capa oculta.

Inicialmente se hicieron las pruebas con 5 modelos, cambiando las funciones de entrenamiento:

Modelo 1: Algoritmo de entrenamiento (traingd), arquitectura 12:30:12 (12 neuronas en la capa de entrada, 30 en la capa oculta y 12 en la capa de salida).

Modelo 2: Algoritmo de entrenamiento trainbfg, arquitectura 12:30:12.

Modelo 3: Algoritmo de entrenamiento trainlm, arquitectura 12:30:12.

Modelo 4: Algoritmo de entrenamiento trainlm, arquitectura 12:40:12.

Modelo 5: Algoritmo de entrenamiento `rainlm`, arquitectura 4:16:4.

Los resultados del entrenamiento con conjuntos de prueba y validación, para la primera red, que funciona como generador de residuos, y con un Algoritmo de entrenamiento `rainlm`, arquitectura 4:16:4. se muestran en la figura 4.

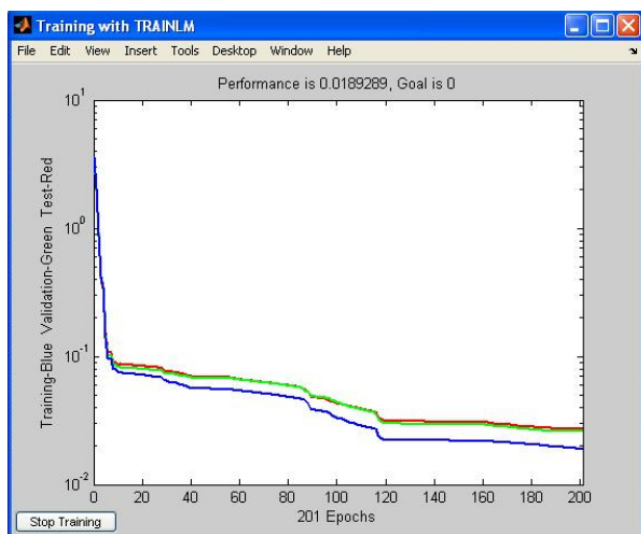


Figura 4. Resultado del entrenamiento del modelo 5

La simulación de la red con un conjunto de prueba, mostró resultados satisfactorios, comprobándose que esta red sí clasifica correctamente un conjunto de datos nuevos para ella, con un índice de 99,9% de certeza.

Para la segunda red, que funciona como un clasificador de residuos, se tiene la experiencia de la anterior, por lo que se selecciona una red con 16 neuronas en la capa oculta y con la función de entrenamiento `trainlm`.

En la figura 5, se muestra el resultado del entrenamiento realizado a la red # 2.

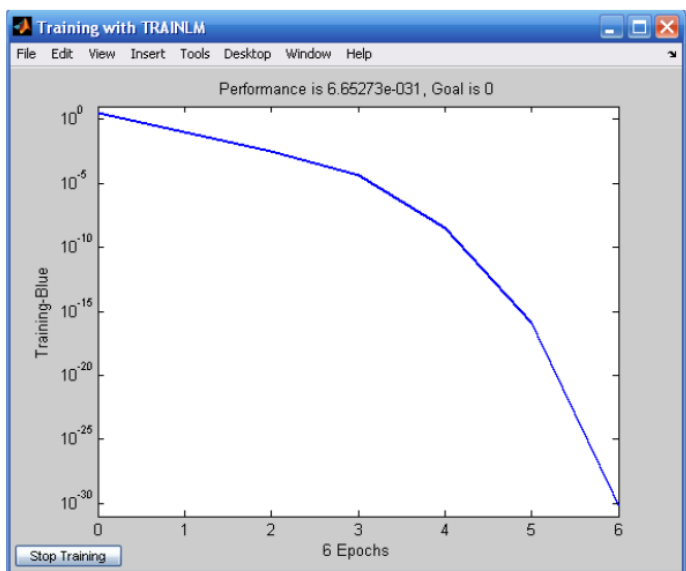


Figura 5. Gráfica que muestra el entrenamiento realizado a la red # 2.

Evaluación de la red entrenada usando patrones de prueba

Para evaluar la capacidad de generalización de las 2 redes se utilizó un conjunto de patrones de prueba (150 patrones) y se calculó el porcentaje de clasificaciones incorrectas ante estos nuevos vectores de entrada. Para la selección de la red a usar se tuvo en cuenta los resultados de:

1. Convergencia de la red en el entrenamiento (figura 6).
2. Capacidad de generalización (figura 7).

En la figura 6, se muestra un resumen de los resultados obtenidos durante el entrenamiento y con el conjunto de prueba en la figura 7.

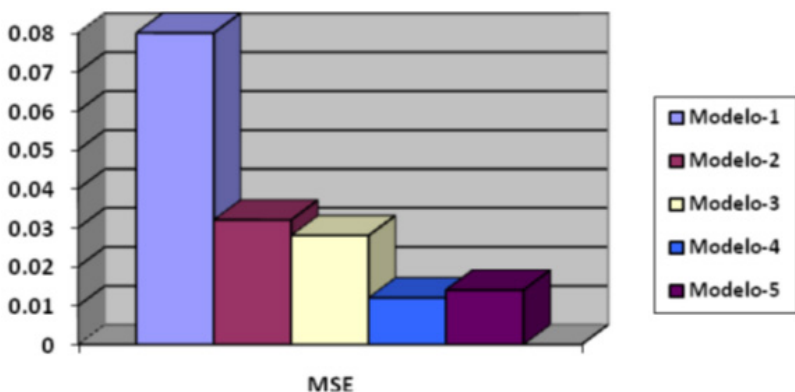


Figura 6. Resultados del entrenamiento por la función MSE.

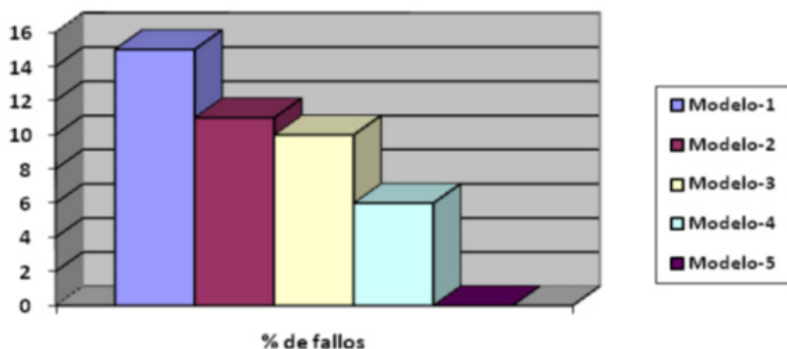


Figura 7. Resultados del entrenamiento por el % de clasificaciones incorrectas del conjunto de pruebas

7.3 Lógica difusa

7.3.1 Aplicación de la lógica difusa al mantenimiento por diagnóstico de redes eléctricas de distribución

Los métodos de diagnóstico de líneas aéreas de distribución están en pleno desarrollo en la actualidad y se basan normalmente en análisis estadístico de interrupciones

y la inspección de las líneas con el objetivo de optimizar, en un momento determinado, la relación entre el costo de dar mantenimiento y el costo de las fallas, si el mantenimiento no se ejecutara.

Como se conoce un conjunto difuso se diferencia de un conjunto ordinario por el hecho de que la pertenencia al conjunto puede tomar valores menores que uno, o sea, existen elementos que no pertenecen del todo al conjunto (Thaker & Nagori, 2018).

Por otro lado, un enunciado lógico puede ser representado mediante la pertenencia o la membresía en el conjunto correspondiente (Feng *et al.*, 2019; Thaker & Nagori, 2018). Por ejemplo, la afirmación: “La línea está en mal estado técnico” se puede sustituir por: La línea pertenece al conjunto “Líneas en mal estado técnico”; de este modo la sentencia “Esta línea está en muy mal estado técnico” se puede sustituir por: La membresía de esta línea en el conjunto “Líneas en mal estado técnico” es igual a 1. Como se puede apreciar, la aplicación de la lógica difusa permite representar el nivel de deterioro de un elemento mediante los números del intervalo $[0,1]$, lo cual a su vez facilita la aplicación de la matemática y la computación a las apreciaciones vagas e imprecisas de los resultados de observación.

Fuzzificación de las entradas

Las variables que se utilizan como entrada al sistema de inferencia representan

- 1) El estado técnico de la línea en función del Indicador global del nivel de deterioro
- 2) La calidad del tiempo medio entre fallas de la línea
- 3) El período de vida característica.
- 4) La calidad del tiempo medio de reparación de la línea

Para la “fuzzificación” de la primera de las variables, “Estado Técnico”, se representa por tres funciones de membresía trapezoidales, que se muestran en la figura 8.

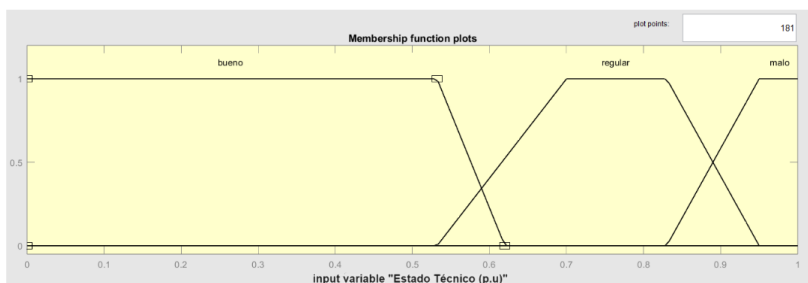


Figura 8. Representación gráfica de la fuzzificación de la variable Estado Técnico, funciones de pertenencia de la variable “Estado Técnico” con $\gamma_b = \alpha_e = 0.532$, $\delta_e = 0.7$, $\gamma_e = 0.8$, $\chi_e = 0.9$, $\alpha_m = 0.8$ y $\delta_m = 0.9$ (Estos valores se obtienen como resultado del proceso de agrupación difusa (Fuzzy Clustering) de la información estadística.

La variable de entrada “Tiempo medio entre fallas” se representa por tres funciones de membresía triangulares, que se muestran en la figura 9.

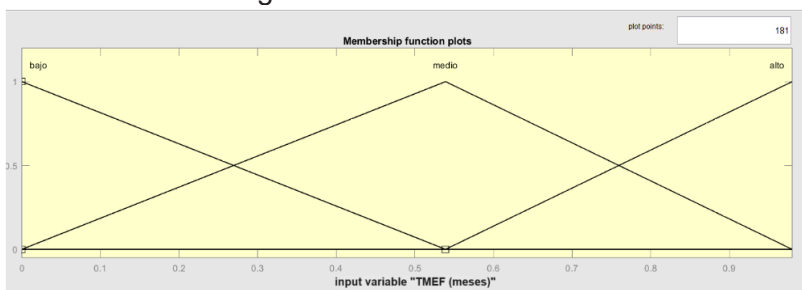


Figura 9. Representación gráfica de la fuzzificación de la variable “Tiempo medio entre fallas”

Aunque las empresas distribuidoras controlan la tasa de fallas, como un indicador de confiabilidad no se considera un valor máximo para esta; sin embargo, este límite se puede obtener a partir de otros indicadores de confiabilidad que sí están normados, como el SAIFI en Estados Unidos, el FMIK o el FTIK en Chile, el FES en Colombia, etc. A partir del análisis

de las normativas de varios países para los indicadores antes mencionados.

El período de vida característica, que como ya se mencionó, responde a la curva de vida característica de cualquier elemento, conocida como “curva de bañera”; se representa por dos funciones trapezoidales y un “singleton”, como se representa en la figura 10.

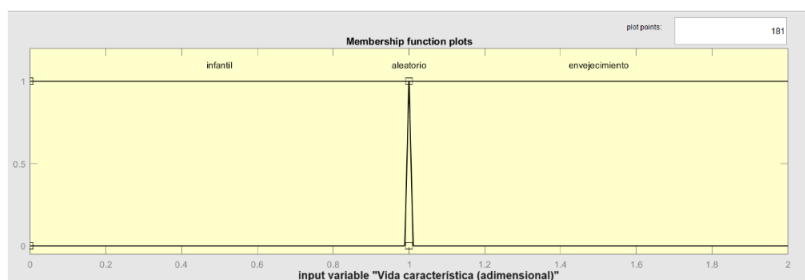


Figura 10. Representación gráfica de la fuzzificación de la variable Período de vida característica.

Finalmente, la variable tiempo medio de reparación “TMR” se representa por dos funciones trapezoidales y una triangular como se muestra en la figura 11.

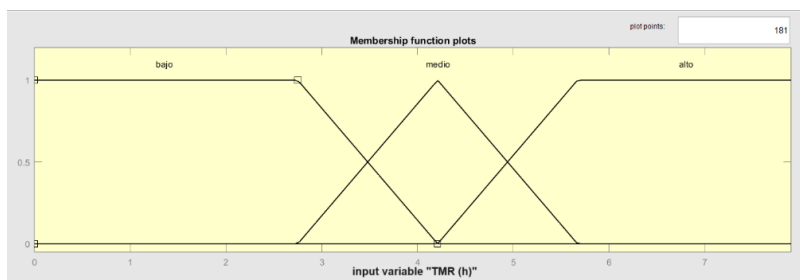


Figura 11. Representación gráfica de la fuzzificación de la variable “Tiempo medio reparación” (Utilizando la información estadística y gráficos de control)

Defuzzificación

Al observar un proceso de inferencia cuando el consecuente es un conjunto de singletons, y la defuzzificación se realiza con el método del “centro de gravedad”, resulta desde el punto de vista matemático una regla.

IF ((x_1 es A_1) AND (x_2 es B_1) AND ... AND (x_n es w_1))
THEN (u es S_1)

Donde S_1 es singletons, puede describirse así:

IF ((x_1 es A_1) AND (x_2 es B_1) AND ... AND (x_n es w_1))
THEN ($u = C_1$)

Donde C_1 es el valor central del singleton S_1 y se puede interpretar como la función lineal

$y = A_n x_n + A_{n-1} x_{n-1} + \dots + A_1 x + A_0$
de orden cero.

Prosiguiendo la generalización, Sugeno ha propuesto representar el consecuente como una función de orden “n” de las variables del antecedente; el caso particular de función de primer orden se presenta a continuación:

IF((x_1 es A_1) AND...AND(x_n es A_n)) THEN ($u = a_1 * x_1 + a_2 * x_2 + \dots + a_n * x_n + C$),

Donde $a_1, a_2, \dots, a_n, C$ son constantes

De este modo, el sistema Takagi-Sugeno se diferencia del de Mamdani por el hecho de sustituir el proceso de “fuzzificación-defuzzificación” por el simple cálculo del consecuente como una función de orden “n” de las variables del antecedente; aquí, en la región de solapamiento de membresías del antecedente, el valor del consecuente se obtiene mediante la interpolación entre dos funciones correspondientes de este último (Ontiveros-Robles *et al.*, 2019).

La función de membresía del consecuente se representa por Singletons que expresan los diferentes períodos de mantenimiento y cuyo valor depende de la importancia crítica del tramo o circuito, según la siguiente expresión:

$$S = \frac{\text{Plazo de mantenimiento (meses)}}{1 + (ICR_{j \text{ normalizada}} \cdot (1 - e^{-\alpha \cdot I_{eq}}))}$$

Donde

$ICR_{j \text{ normalizada}}$: Importancia crítica normalizada del tramo o

del circuito j (p.u).

$e^{-\alpha I_{eq}}$: Componente de la función de confiabilidad que depende del estado de carga de la línea.

$I_{eq} = I_l / K_d$: Corriente equivalente que circularía por la línea en régimen de máxima demanda si la carga estuviera concentrada en el extremo de la misma.

I_l : Corriente que circula por la primera sección de línea del tramo o circuito en régimen de máxima demanda.

K_d : Coeficiente de distribución de la carga.

α : Coeficiente de atenuación (0.001-0.0001)

Un ejemplo de función del consecuente se puede apreciar en la figura 12.

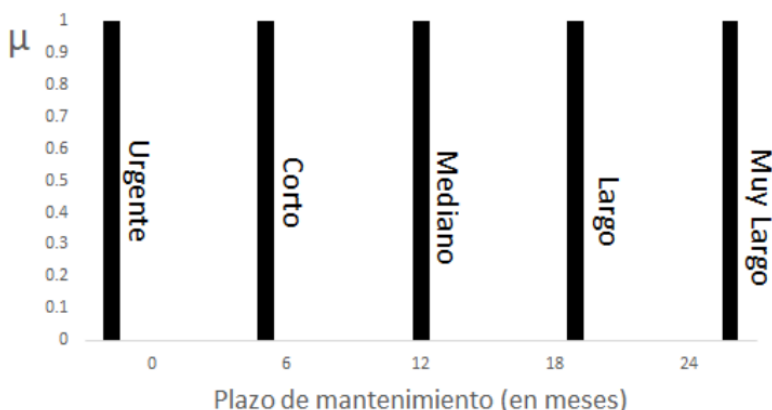


Figura 12. Funciones membresía del consecuente (Singletons) para una importancia crítica igual a 0

Como resultado de la “desfuzzificación”, utilizando el método de centro de gravedad, se obtienen las superficies de comportamiento del consecuente en función de las variables del antecedente. Un ejemplo ilustrativo aparece en las figuras 14, 15 y 16, este ejemplo representa el comportamiento de las variables para el circuito las 500, este circuito suministra un área rural, tiene un nivel de tensión de 13.2 kV, una longitud de 61.16 kmy una densidad de carga de 30.28 kVA/km .

Como se puede observar la variable más influyente en el comportamiento del plazo de mantenimiento es el “Estado técnico” de la línea, obtenido a partir del indicador del nivel de deterioro que refleja los resultados del celaje, considerando que las cuatro variables de entrada tienen el mismo peso.

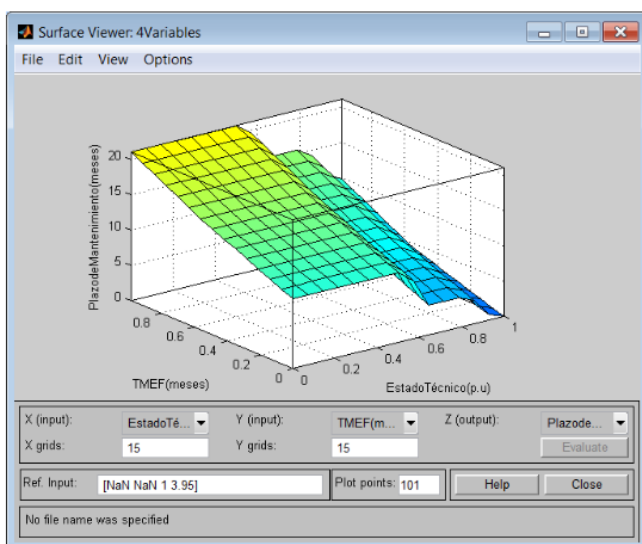


Figura 13. Superficie de comportamiento de la variable del consecuente, en función de la variación de las variables del antecedente “Estado Técnico” y TMEF, en el sistema de inferencia de Sugeno para el circuito Las 500

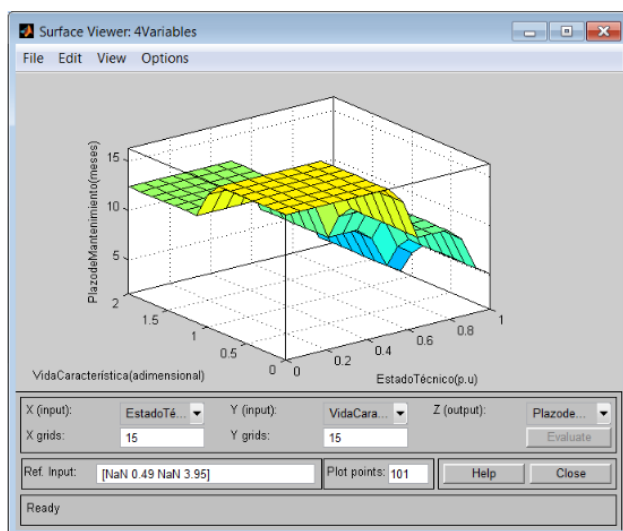


Figura 14. Superficie de comportamiento de la variable del consecuente, en función de la variación de las variables del antecedente “Estado técnico” y “Vida Característica”, en el sistema de inferencia de Sugeno para el circuito Las 500

Conclusiones

En los procesos de generación, transmisión, distribución y uso final de la energía se presentan fenómenos complejos difíciles de abordar, en muchos casos, con métodos y herramientas convencionales, debido a elementos como la no linealidad, la incertidumbre que se introduce en la obtención de los valores de las variables involucradas, etc.

Las técnicas de IA facilitan el desarrollo de sistemas expertos para resolver estos problemas de una forma eficiente, tienen además la ventaja de ser ajustables, a través del aprendizaje, según la experiencia acumulada por el personal, las condiciones iniciales y características particulares del entorno donde se aplican.

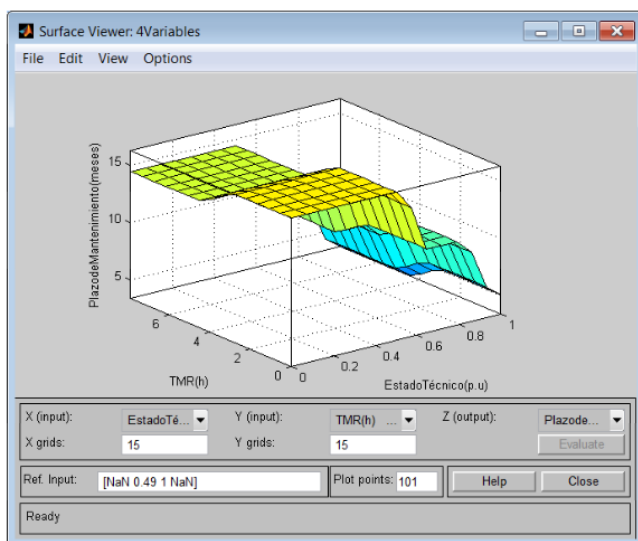


Figura 15. Superficie de comportamiento de la variable del consecuente, en función de la variación de las variables del antecedente “Estado técnico” y “TMR”, en el sistema de inferencia de Sugeno para el circuito Las 500

Trabajos futuros

En esta temática se continúa trabajando para abordar problemas de suma importancia que afectan hoy al sector electroenergético, cómo es el caso del pronóstico de demanda y de generación, en un contexto donde la intermitencia de las fuentes renovables de energía impacta en la continuidad del suministro, o el dimensionamiento y ubicación óptima de fuentes renovables de energía y sistemas de almacenamiento distribuidos para su integración; de esta forma garantizan la flexibilidad en los sistemas eléctricos de potencia.

Referencias del capítulo

Borges Vasconcellos, D., Pérez Abril, I., León Martínez, V. (2012). Compensación de potencia reactiva en sistemas

desbalanceados utilizando algoritmos genéticos Ingeniare. *Revista chilena de ingeniería* 20, 284–292.

Borges Vasconcellos, D.E., Concepción Céspedes, Y., (2017). Compensación de potencia reactiva en sistemas de distribución primaria de energía, aplicando algoritmos genéticos *Ingeniería Energética* 38, 25–34.

Caballero, Y., Arco, L., Bello, R., Salgado, Y., Márquez, Y., León, P., Álvarez, D. (2007). Nuevas medidas de la Teoría de los Conjuntos Aproximados para la evaluación de sistemas de información en Bioinformática. Presented at the presentado en II Congreso Internacional de Bioinformática y Neuroinformática Informática.

Cagman, N., Enginoglu, S., Citak, F. (2011). Fuzzy soft set theory and its applications *Iranian journal of fuzzy systems* 8, 137–147.

Carreno, E.M., Romero, R., Padilha-Feltrin, A. (2008). An Efficient Codification to Solve *Distribution Network Reconfiguration for Loss Reduction Problem IEEE Transactions on Power Systems* 23, 1542–1551.

Corrales Barrios, L., Ramírez Vázquez, A. (2013). Clasificación de fallas con redes neuronales para grupos electrógenos. *Ingeniería Energética* 34, 137–150.

Díaz R, H., Harnisch V, I., Sanhueza H, R., Olivares R, R. (2010). Ingeniare. Reconfiguración y ubicación de condensadores en sistemas de distribución: metodología de solución simultánea usando algoritmos genéticos *Revista chilena de ingeniería* 18, 144–153.

Feng, C., Liao, X. (2020). Constructing phase boundary in AgNbO₃ antiferroelectrics: pathway simultaneously achieving high energy density and efficiency. *Journal of Cleaner Production* 258, 120630.

Golberg, D.E. (1989). Genetic algorithms in search, optimization, and machine learning Addison wesley 1989, 36

González, J. (2002). *Compensación de potencia reactiva*

en sistemas contaminados con armónicos. Tesis de doctorado. Universidad Central Marta Abreu de Las Villas, Cuba.

Junior, I.S., Junior, S.C., Oliveira, E., Costa, J., Pereira, J., Garcia, P. (2008). Identification and antifungal susceptibility of fungi isolated from dermatomycoses. *IEEE Trans. Power Syst* 23, 1619–1626.

León-Martínez, V., Montañana-Romeu, J., Roger-Folch, J., Navarro, A.C. (2009). Why reactive compensators do not improve the efficiency correctly in unbalanced circuits. Presented at the *Proceedings of the XIX IMEKO World Congress Fundamental and Applied Metrology*, Lisbon, Portugal, Citeseer, pp. 6–11.

Liu, L., Chen, S. (2018). The Application of Artificial Intelligence Technology in Energy Internet, in: 2018 2nd IEEE Conference on Energy Internet and Energy System Integration (EI2). Presented at the 2018 2nd *IEEE Conference on Energy Internet and Energy System Integration (EI2)*, pp. 1–5.

Mota Caballero, Y. (2007). Cuba. Universidad de Camagüey, Facultad de Informática.

Ng, H.N., Salama, M.M.A., Chikhani, A.Y. (2000). *IEEE Transactions on Power Delivery* 15, 387–392.

Pérez, I. (2010). *Revista IEEE América Latina* 8.

Pérez, R.N.C. (2011). *Ingeniería Energética* 32, 77-a.

Sierra, E., Lajes, S., Alexandrovna, M., Barrios, F. (2007). *Ingeniería Energética* 28, 45–52.

Sierra, E., Lajes, S., Filiberto, Y., Barrios, F. (2013). *Dyna* 80, 31–39.

Takagi, T., Sugeno, M. (1985). *IEEE Transactions on Systems, Man, and Cybernetics* SMC-15, 116–132.

Yang, X., Gao, T., Ma, S., Gong, X., Huang, Z., Hao, D. (2020). Application and Future of Modern Energy Internet Based on Artificial Intelligence, in: 2020 35th Youth Academic Annual Conference of Chinese Association of Automation (YAC). Presented at the 2020 35th Youth Academic Annual Conference of Chinese Association of Automation (YAC), pp. 889–894.

Yue, D., Han, Q.-L. (2019). *IEEE Transactions on Systems, Man, and Cybernetics: Systems* 49, 1551–1553.

Zhang, Y., Milanovic, J. (2006). *Optimal placement of FACTS devices for voltage sag mitigation based on genetic algorithms*. Presented at the 12th international conference on Harmonics and Quality of Power.

Zimmermann, H.-J. (2011). Fuzzy set theory—and its applications. *Springer Science & Business Media*.

Feng, F., Fujita, H., Ali, M.I., Yager, R.R., Liu, X. (2019). Another View on Generalized Intuitionistic Fuzzy Soft Sets and Related Multiattribute Decision Making Methods. *IEEE Transactions on Fuzzy Systems* 27, 474–488. <https://doi.org/10.1109/TFUZZ.2018.2860967>

Ontiveros-Robles, E., Melin, P., Castillo, O. (2019). Relevance of Polynomial Order in Takagi-Sugeno Fuzzy Inference Systems Applied in Diagnosis Problems, in: 2019 IEEE International Conference on Fuzzy Systems (FUZZ-IEEE). Presented at the 2019 *IEEE International Conference on Fuzzy Systems (FUZZ-IEEE)*, pp. 1–6. <https://doi.org/10.1109/FUZZ-IEEE.2019.8859028>

Sierra, E., Lajes, S., Alexandrovna, M., Barrios, F. (2007). Perfeccionamiento del diagnóstico de las líneas aéreas de distribución utilizando la lógica difusa. *Ingeniería Energética* 28, 45–52.

Sierra, E., Lajes, S., Filiberto, Y., Barrios, F. (2013). Modelo difuso para la determinación del período de mantenimiento de redes eléctricas, a partir de los datos del celaje. *Dyna* 80, 31–39.

Thaker, S., Nagori, V. (2018). Analysis of Fuzzification Process in Fuzzy Expert System. *Procedia Computer Science, International Conference on Computational Intelligence and Data Science* 132, 1308–1316. <https://doi.org/10.1016/j.procs.2018.05.047>

Índice de acrónimos y términos especializados

Agricultura de precisión: La agricultura de precisión puede utilizar tecnologías de la IA para mejorar la gestión de las actividades agrícolas, pero no es exclusivamente una forma de esta; es decir que, puede aprovechar la IA como una herramienta para optimizar las operaciones agrícolas, pero no se limita únicamente a ella. No son conceptos con significados idénticos, aunque sí están estrechamente relacionados. La IA puede desempeñar un papel importante en la agricultura de precisión al procesar grandes cantidades de datos, identificar patrones y tomar decisiones automatizadas.

Algoritmo “Optimización por Enjambre de Partículas” o *Particle Swarm Optimization* (PSO): Se trata de una técnica basada en población, que tiene sus orígenes en el comportamiento social y cognitivo de ciertas especies (de ahí la alusión al enjambre). Utiliza múltiples partículas que forman un “enjambre”, en el que cada una de ellas se refiere a una solución candidata y “vuela” en un área, buscando la mejor solución. Es una herramienta útil para la optimización de diseños.

Algoritmo de ahorros de Clarke y Wright: Es una heurística diseñada para resolver el problema del enrutamiento de vehículos (*Vehicle Routing Problem*, VRP). Este algoritmo es particularmente conocido por su simplicidad y eficacia para generar soluciones de alta calidad en un tiempo computacionalmente razonable, especialmente en comparación con los enfoques de optimización exacta que pueden ser impracticables para problemas de gran tamaño.

Algoritmo k-NN: El k- vecinos más cercanos, k-NN por sus siglas en inglés, es un algoritmo de aprendizaje automático (AA), que pertenece a los algoritmos de aprendizaje supervisado simples. Puede ser utilizado para resolver problemas de clasificación- predicción. Posee gran aplicabilidad en la Ingeniería Civil debido a la inmensa cantidad de datos que puede analizar, facilitando así que las clasificaciones o predicciones sean realizadas de forma más eficiente. Es un algoritmo versátil y poderoso que se destaca por su simplicidad y efectividad en una amplia gama de problemas de aprendizaje automático.

Algoritmos de “caja blanca”: Son algoritmos de aprendizaje automático y análisis de datos en los que es posible comprender y explicar el proceso de toma de decisiones del modelo de manera transparente y clara, por lo que puede rastrearse cómo llegaron a sus conclusiones.

Algoritmos genéticos: Es una herramienta basada en la idea del proceso de selección natural desarrollado por Darwin. Así, a partir de una población inicial, dada por un conjunto de soluciones, se imitan diferentes momentos de ese proceso: herencia, mutación, cruzamiento, selección, etc., tras lo cual, los individuos resultantes serán las posibles soluciones al problema, que serán mejores cada vez, tal y como sucede en la evolución de la naturaleza. Estas soluciones se obtienen de manera rápida y eficiente. Poseen gran aplicabilidad en la realización de predicciones y en la optimización.

Analítica de Macrodatos (AM), en inglés *Big Data* (BD): Se refiere a la capacidad de analizar y obtener información significativa a partir de grandes conjuntos de datos, lo que puede ser especialmente útil en un amplio rango de aplicaciones en campos como el *marketing*, la atención médica, la logística, la investigación científica y muchas otras áreas. Aunque es

frecuente el uso de este anglicismo en textos técnicos en español, puede usarse y es reconocido en esta lengua.

Aprendizaje Automático (AA): Se refiere a una rama de la inteligencia artificial centrada en el desarrollo de algoritmos y modelos que permiten a las computadoras aprender y mejorar su rendimiento en tareas específicas a través de la experiencia y el análisis de datos, sin necesidad de una programación explícita. *Machine Learning* (ML), en inglés. Aunque este anglicismo es ampliamente utilizado y reconocido en el campo de la IA, en textos en español, su uso en español es igualmente común y comprensible en la literatura científico-tecnológica, por eso se ha preferido.

Aprendizaje Profundo: *Deep Learning*, en inglés, la sustitución de este anglicismo por su versión en español, se está haciendo cada vez más común en textos en esta lengua. Es un subcampo dentro del AA, que se centra en redes neuronales artificiales y puede aprender automáticamente las características relevantes a partir de los datos.

ARIMA (Autoregressive Integrated Moving Average): Es un modelo estadístico utilizado para analizar y predecir series temporales. Este modelo combina componentes de autorregresión (AR), integración (I) y medias móviles (MA) para capturar patrones complejos en los datos temporales. El modelo ARIMA es ampliamente utilizado en la econometría, las finanzas, la meteorología y otras áreas donde se requiere analizar y predecir datos secuenciales.

Base de casos de Espacios Vectoriales: En el contexto de sistemas de razonamiento basado en casos (*Case-Based Reasoning* - CBR) aplicados a la matemática o más específicamente al Álgebra Lineal, sería una colección

estructurada de casos relacionados con problemas de Espacios Vectoriales que han sido previamente resueltos. Los Espacios Vectoriales son conceptos fundamentales en Álgebra Lineal y tienen múltiples aplicaciones en ciencia, ingeniería, informática y otras. Un “caso” en este contexto suele referirse a un problema específico junto con su solución y una explicación detallada de cómo se llegó a esta.

Bosque Aleatorio (RF): El acrónimo de este algoritmo procede de *Random Forest*, su denominación en inglés. Es un método de aprendizaje en conjunto (*ensemble learning*) utilizado tanto para tareas de clasificación como de regresión, bioinformática, reconocimiento de imágenes, análisis de riesgos financieros, entre otros. Posee reconocida flexibilidad y facilidad de uso, y una herramienta poderosa en el ámbito del aprendizaje automático.

Capacitated Vehicle Routing Problem: Problema de Enrutamiento de “Vehículos con Capacidad”, en español. El uso en inglés de ese término (CVRP) es habitual en textos en español debido a la falta de una versión corta y estandarizada en esta lengua, por lo que se mantiene el anglicismo. En el contexto de la informática y la optimización matemática es una colección de algoritmos, modelos y herramientas diseñadas para resolver el problema de enrutamiento de vehículos con capacidad (CVRP). Las bibliotecas CVRP son utilizadas en la investigación operativa y en industrias como logística, distribución y transporte, donde la optimización de rutas es crucial para reducir costos y mejorar la eficiencia operativa.

Chi-Cuadrado: Es una prueba que mide la dependencia entre las variables categóricas. Se utiliza para evaluar si la presencia (o ausencia) de una característica y la clase objetivo están relacionadas. Un valor alto del Chi-Cuadrado indica que la hipótesis de independencia entre la característica y la clase puede ser rechazada, sugiriendo que la característica es importante para la clasificación.

Clusterización: Es un término que técnicamente no existe en la lengua española; sin embargo, en el argot de las ciencias de la computación, se utiliza comúnmente en el contexto de la minería de datos y el análisis de datos para describir el proceso de formar grupos coherentes o clústeres a partir de un conjunto de datos.

Coefficiente de Determinación R^2 : El coeficiente de determinación, denotado como R^2 o R cuadrado, es una medida estadística utilizada en el análisis de regresión para evaluar la bondad de ajuste de un modelo. Este coeficiente proporciona una indicación de la cantidad de variación en la variable dependiente que puede ser explicada por la(s) variable(s) independiente(s) en el modelo. En otras palabras, mide qué tan bien los valores observados pueden ser predichos por el modelo, en comparación con la media de los valores observados. Es una herramienta útil que, no obstante, debe ser utilizado junto con otras métricas y pruebas estadísticas para obtener una evaluación completa del modelo.

Computación en la Nube: En inglés *Cloud Computing* es un modelo que permite el acceso remoto a recursos de computación, como servidores, almacenamiento, bases de datos, redes, software, aplicaciones, entre otros, a través de Internet.

Correlational Features Selection (CFS): Puede traducirse como “Selección de Características Correlacionales”, aunque su uso en español no es común en los textos científico-técnicos. Es un método de selección que evalúa subconjuntos de características en función de su correlación con la clase objetivo y la correlación entre ellas. Su objetivo principal es encontrar un subconjunto de características que estén altamente correlacionadas con la clase y poco correlacionadas entre sí. Esto ayuda a reducir la redundancia y mejora la capacidad predictiva del conjunto seleccionado.

CRISP-DM: Acrónimo de *Cross-Industry Standard Process for Data Mining* o “Proceso Estándar de la Industria Cruzada para la Minería de Datos”, en español. Es una metodología, ampliamente utilizada en el campo de la minería de datos y el análisis predictivo, que se desarrolló a finales de la década de 1990 y sigue siendo relevante hoy debido a su enfoque estructurado y flexible para resolver problemas en esa área del conocimiento. Aunque puede traducirse al español, como se ha indicado, no es común su uso en esta lengua.

Earth: Algoritmo también conocido en algunos contextos como “MARS”, por su denominación original en inglés *Multivariate Adaptive Regression Splines*, es una técnica de modelado estadístico que se utiliza para la regresión y, en ciertos casos, para la clasificación. Fue desarrollado originalmente por Jerome Friedman en 1991; es eficaz para modelar relaciones no lineales y complejas entre variables mediante el uso de *splines* (o *esplines*) adaptativos multivariantes. El nombre al uso de este algoritmo es en su original inglés, no suele llamársele Tierra en los textos técnicos.

Envenenar datos: En el contexto de la IA se refiere a la manipulación intencionada de los datos utilizados para entrenar un modelo de AA con el fin de alterar su comportamiento o rendimiento de manera perjudicial, provocando que el modelo aprenda a tomar decisiones incorrectas en procesos tan importantes como la detección de *malware*, la clasificación de correo electrónico no deseado como *spam*, o la seguridad cibernética, son algunos ejemplos. La diferencia clave entre los tipos de envenenamiento radica en el momento en que se aplican las manipulaciones maliciosas.

Error Porcentual Absoluto Medio (MAPE): *Mean Absolute Percentage Error* (MAPE), por sus siglas en inglés. Aunque es común su uso en español en textos técnicos, especialmente

en aquellos relacionados con estadística, análisis de datos, pronósticos y aprendizaje automático, el acrónimo en inglés es el que está al uso y es el que se reconoce internacionalmente. Esta es una métrica utilizada para evaluar la precisión de modelos de predicción o pronóstico. Proporciona una medida del error como un porcentaje, lo que facilita la interpretación de la precisión del modelo independientemente de la escala de los datos. El MAPE es particularmente popular en el campo de la predicción de series temporales y en situaciones donde es importante entender el tamaño del error en términos relativos. Esta métrica expresa el error como un porcentaje del valor real, lo que permite una interpretación intuitiva: un MAPE de 5 significa que, en promedio, las predicciones del modelo se desvían un 5 % de los valores reales. Es muy útil para comparar la precisión entre diferentes modelos o conjuntos de datos con diferentes escalas.

Espacios Vectoriales: En el aprendizaje automático, los datos se representan a menudo como vectores en un espacio vectorial, lo que permite aplicar técnicas de álgebra lineal para entrenar modelos y realizar operaciones matemáticas en los datos. Esta es una de las razones por las que la teoría de espacios vectoriales es fundamental en el aprendizaje automático y en la inteligencia artificial.

Estimation of Distribution Algorithm (EDA): Es una clase de algoritmo evolutivo que se basa en la creación de modelos estadísticos para representar la población de soluciones. En lugar de utilizar operadores de cruzamiento o mutación como en los algoritmos genéticos, los EDAs actualizan estos modelos estadísticos en cada iteración para guiar la búsqueda hacia regiones del espacio de soluciones que tienen mayor probabilidad de contener óptimos.

Función de transferencia “*purelin*”: El término *purelin* proviene de *pure linear*, o “puramente lineal”, en español. Es

una denominación específica utilizada en algunas bibliotecas y entornos de programación para referirse a la función lineal o identidad como función de activación en redes neuronales. Esta función se caracteriza por no aplicar ninguna transformación no lineal a su entrada; es decir, la salida es igual a la entrada.

Función de transferencia “*tansig*”: El término “tansig” proviene de la contracción de “tangente sigmoide” y se refiere a la función de activación tangente hiperbólica. Es una de las funciones de activación no lineales utilizadas en redes neuronales, incluidos los perceptrones multicapa (MLP). Esta función es particularmente popular en las capas ocultas de las redes neuronales debido a su naturaleza no lineal, lo que permite a la red aprender y modelar relaciones complejas entre las entradas y las salidas.

Gobierno Cognitivo: Este es un concepto de dirección que aprovecha las capacidades de la IA y el análisis de datos para mejorar la gestión gubernamental, la toma de decisiones y la prestación de servicios públicos, lo cual debe conducir a un gobierno más eficiente, transparente y orientado a datos.

Grupo de Interés Especial (*Special Interest Group* - SIG): Esta expresión no es exclusiva de la Minería de Datos. Se refiere, en este caso, a un grupo de profesionales e interesados que se han unido para contribuir, desarrollar, promover y apoyar la metodología CRISP-DM.

IA invisible: La inteligencia artificial invisible se refiere a sistemas y aplicaciones de IA que funcionan en segundo plano de manera automática y sin intervención humana aparente. Los usuarios no necesitan saber que están utilizando IA, simplemente experimentan sus beneficios de manera natural. En lugar de tener que interactuar con una pantalla o un dispositivo, la tecnología se adaptará al entorno y proporcionará

información relevante de manera intuitiva. Por ejemplo, los filtros de correo electrónico, el autocompletado de búsquedas en Google, la corrección automática de texto y gramática en procesadores como Microsoft Word, la detección de fraude en tarjetas de crédito, o de actividad inusual en cuentas de correo, etc.

IA responsable: Es un enfoque para desarrollar, evaluar e implementar sistemas de IA de manera ética, transparente y justa. La esencia de este concepto es que los sistemas de IA deben diseñarse y operarse de modo que se tengan en cuenta aspectos como privacidad, seguridad, equidad, inclusión y rendición de cuentas, de modo que respeten los derechos humanos, no contribuyan a las diferentes formas de discriminación, y tengan un impacto positivo sobre la sociedad y el medioambiente.

IMBNPBASIR SEL-CLASS: Es un algoritmo de clasificación que permite la adquisición del conocimiento e incluye la construcción de los prototipos que mejor representan las clases del problema por ser una forma interpretable de formalización del conocimiento para los especialistas.

Industria 4.0: Término también conocido como la Cuarta Revolución Industrial. Se refiere a la transformación de la industria y la fabricación mediante la integración de tecnologías digitales avanzadas. Esta revolución se basa en la interconexión de sistemas, la automatización, el análisis de datos y la IA para mejorar la eficiencia, la productividad y la toma de decisiones en la industria. La expresión “I4.0”, por su parte, se utiliza para denotar de forma simplificada la Industria “4.0”; es una abreviatura o forma corta de referirse a ella.

Inteligencia Geoespacial: Es la capacidad de recopilar, analizar y utilizar datos geográficos y espaciales para obtener

información valiosa y tomar decisiones informadas, con base en la combinación de información geográfica, que incluye datos de mapas y ubicaciones geográficas, con análisis avanzados y herramientas tecnológicas para resolver problemas específicos. El Sistema de Posicionamiento Global (GPS) es una de sus aplicaciones más comunes.

Internet de las cosas (IoT): Es la conexión de dispositivos y máquinas a través de Internet para recopilar y compartir datos en tiempo real. En los textos técnicos, sobre todo para audiencias especializadas, aunque la traducción del término en español es de uso común, no se reconoce el acrónimo en esta lengua; por lo que se usa IoT, por sus siglas en inglés.

J48: Es una implementación en Java del algoritmo C4.5, que a su vez es una extensión del algoritmo ID3 (*Iterative Dichotomiser 3*, o *Dicotomizador Iterativo*) para la generación de árboles de decisión. C4.5. Fue desarrollado por Ross Quinlan en los años 90, como una herramienta de aprendizaje automático para clasificación supervisada. Es conocido, además, por su inclusión en la popular suite de software de aprendizaje automático *Weka*.

JRip, NNge, OneR, PART, Ridor, ADTree, J48, Árbol Aleatorio de Decisión RandomTree, REPTree, SimpleCartç: Son algoritmos de aprendizaje automático implementados en WEKA. Con sus propias características y enfoque en diferentes tipos de problemas de clasificación. WEKA los implementa de modo que los usuarios pueden elegir el que mejor se adapte a sus necesidades específicas de análisis de datos.

k-Vecinos más Cercanos (kNN): El acrónimo de este algoritmo procede de *k-Nearest Neighbors*, su denominación original en inglés; es un método simple pero poderoso utilizado para clasificación y regresión en aprendizaje automático.

Aunque k-NN es conocido principalmente por su uso en ese tipo de aprendizaje, también puede integrarse en un sistema experto, especialmente en contextos donde la similitud entre casos es importante.

Lógica Difusa: *Fuzzy Logic*, en inglés, es un enfoque que se aparta de la lógica booleana tradicional, en la que cada proposición tiene un valor de verdad de verdadero o falso. La lógica difusa permite grados de verdad que se expresan en valores entre 0 y 1. Es más flexible y se asemeja más a la forma en que los humanos interpretamos y trabajamos con la información en situaciones reales y ambiguas. Introducida por el Dr. Lotfi Zadeh en la década de 1960 como una forma de modelar la incertidumbre y la imprecisión de la vida real, es ampliamente utilizada en sistemas de control, sistemas expertos y otras aplicaciones de IA.

Máquina de Potenciación del Gradiente (GBM): Acrónimo del nombre original en inglés *Gradient Boosting Machine*, es un algoritmo de aprendizaje en conjunto que utiliza el método de potenciación del gradiente para construir modelos predictivos precisos. Se basa en la idea de combinar secuencialmente múltiples modelos débiles, de manera que cada modelo subsiguiente intenta corregir los errores del modelo anterior. No se refiere a una máquina física, sino a un enfoque de modelado de datos que utiliza el algoritmo de impulso de gradiente.

Máquinas de Soporte Vectorial (SVM): En inglés *Support Vector Machines*, de donde parte el acrónimo que se usa. Estas son un conjunto de algoritmos de aprendizaje supervisado utilizados para clasificación y regresión. Fueron desarrolladas originalmente en la década de 1990 por Vladimir Vapnik y sus colegas y se han convertido en la actualidad en herramientas de gran precisión en el campo del aprendizaje automático debido a su capacidad para manejar datos de alta dimensión

y su flexibilidad para modelar diferentes tipos de relaciones complejas entre los datos.

Metaheurísticas: Estas son estrategias de alto nivel que guían otros algoritmos heurísticos para explorar el espacio de soluciones de problemas de optimización, con el objetivo de encontrar soluciones óptimas o aproximadas en un tiempo razonable. Son especialmente útiles para problemas complejos y de gran escala donde los métodos tradicionales, como la programación lineal o los algoritmos exactos, pueden no ser eficientes.

Método de Validación Cruzada de K-Particiones (*K-Folds Cross-Validation*): Es una técnica estadística utilizada para evaluar la capacidad de generalización de un modelo en el campo del aprendizaje automático y la estadística. Este método es especialmente útil para estimar cómo un modelo predictivo se desempeñará en la práctica cuando se le presenten datos nuevos e independientes. Es común encontrar el término tanto en su versión original en inglés como en su traducción al español en textos técnicos, es perfectamente comprensible en español. Este método es fundamental para evitar el sobreajuste, garantizar que el modelo sea robusto y que tenga un buen desempeño en datos no vistos. Al utilizar toda la muestra de datos tanto para entrenamiento como para prueba, la validación cruzada de K-particiones ofrece una estimación más fiable de la capacidad predictiva del modelo en comparación con una simple división entre datos de entrenamiento y prueba.

Método *OneR* (*One Rule*): Es un algoritmo simple de clasificación que genera una única regla para cada característica para predecir la clase objetivo. La característica que produce la regla con el menor error de clasificación se considera la más importante. Aunque es simple, *OneR* puede ser muy efectivo.

Métodos de selección de rasgos tipo filtros: Son técnicas utilizadas en el preprocesamiento de datos para el aprendizaje automático. Estos métodos evalúan la importancia de las características de los datos basándose en medidas estadísticas, sin involucrar un modelo predictivo.

Minería de Datos (MD), *Data Mining*, en inglés (DM): Es un proceso de exploración y análisis de grandes conjuntos de datos para descubrir patrones, tendencias y relaciones ocultas en ellos. Se utiliza para extraer información valiosa y conocimiento útil a partir de grandes cantidades de datos, lo que puede ser beneficioso en la toma de decisiones y la resolución de problemas en diversas áreas.

MSV: “Máquinas de Soporte Vectorial”, SVM por su denominación original en inglés: *Support Vector Machines*. Estas son modelos supervisados potentes para clasificación y regresión, conocidos por su capacidad para manejar espacios de alta dimensión y encontrar márgenes óptimos de separación entre clases.

***Naïve Bayes*:** Es un algoritmo de aprendizaje supervisado basado en el teorema de Bayes. La “ingenuidad” (*naïve*) se refiere a la suposición simplificadora de que las características (o atributos) que describen a una instancia son condicionalmente independientes entre sí, dada la clase de la instancia. A pesar de su aparente simplicidad, el *Naïve Bayes* suele ser efectivo para datos como el filtrado de *spam* y la clasificación de documentos. Es ampliamente utilizado en tareas de clasificación, en general. Aunque no es incorrecto su uso en español, es más común encontrarlo en inglés, sobre todo en textos técnicos.

NNET (*Feed-Forward Neural Networks*): Se centra en redes neuronales artificiales, específicamente en las redes

de tipo *feedforward* o redes neuronales de capa adelante, donde las conexiones entre las unidades no forman ciclos. Este término a menudo se asocia con paquetes o funciones en software estadístico que implementan este tipo de redes neuronales para clasificación o regresión.

NP-Completo: Término referido a una clasificación en la teoría de la complejidad computacional que describe un conjunto de problemas de decisión, cuya importancia radica en su capacidad para modelar y abordar una amplia gama de problemas prácticos en ciencias de la computación, optimización, investigación de operaciones, y más allá, proporcionando un marco común para entender la dificultad computacional y las limitaciones de lo que se puede calcular eficientemente.

Objetos Conectados Inteligentes: Son dispositivos físicos que están equipados con sensores, capacidad de procesamiento y conectividad a redes, como Internet, que les permite recopilar datos, comunicarse y tomar decisiones de manera autónoma o en colaboración con otros dispositivos o sistemas.

PMD: No es un acrónimo que tenga un solo significado bien definido; en su comprensión es importante tener en cuenta el contexto. El sentido en que se está utilizando en este texto puede interpretarse como “promedio del módulo de las diferencias entre los valores experimentales y los que se obtienen por la predicción” y se refiere a una métrica de error utilizada para evaluar modelos de predicción o pronóstico. Aunque este acrónimo no es estándar en la literatura de inteligencia artificial o estadística, en este caso, se trata de una medida similar al Error Absoluto Medio (MAE).

Problem with Time Windows (FSMVRPTW): Problema de Enrutamiento de Vehículos con Tamaño y Mezcla de Flota y Ventanas de Tiempo, o Problema de Enrutamiento de Vehículos de Flota Mixta y Tamaño con Ventanas de Tiempo: No es común su uso traducido al español, por lo que se mantiene el acrónimo en inglés. En esencia, se trata de un problema de optimización complejo que generalmente implica determinar la flota óptima de vehículos de diferentes capacidades para cumplir con un conjunto de entregas o servicios, teniendo en cuenta restricciones de tamaño y mezcla de flota, así como ventanas de tiempo durante las cuales los servicios deben realizarse. Combina elementos del clásico problema de ruteo de vehículos (VRP) con la complejidad adicional de manejar flotas mixtas y restricciones de tiempo.

Problemas de Fleet Size and Mix (FSM): Se refieren a una categoría específica dentro de los problemas de ruteo de vehículos (*Vehicle Routing Problems*, VRP); buscan determinar no solo cómo asignar y enrutar una flota de vehículos para satisfacer una serie de demandas de entrega o servicio, sino también cuál es la composición óptima de esa flota en términos de tamaño y tipos (o mezcla) de vehículos. El objetivo es minimizar el costo total, que puede incluir costos de operación de los vehículos, costos de mantenimiento, costos de establecimiento de rutas, entre otros, mientras se satisfacen todas las demandas y restricciones operativas.

Q_{exp} / Q_m : Los “valores experimentales”, a menudo denotados como Q_{exp} , se refieren a los datos reales obtenidos a través de la observación, experimentación o medición. Estos valores sirven como el “estándar de oro” o criterio de verdad para comparar y evaluar las predicciones hechas por un modelo. La relación entre los valores experimentales (Q_{exp}) y las predicciones (Q_m) es crítica para evaluar la precisión y efectividad de un modelo. Analizar la discrepancia entre estos valores permite entender la efectividad del modelo para replicar o predecir fenómenos reales.

QSAR: Siglas en inglés de Relaciones Cuantitativas Estructura-Actividad (*Quantitative Structure-Activity Relationship*). Se trata de un enfoque en la química y la farmacología que busca establecer una correlación cuantitativa entre la estructura química de una molécula y su actividad biológica o farmacológica. Los avances en la informática, en particular la IA, han mejorado la precisión y la aplicabilidad de los modelos QSAR. Esta es en la actualidad una herramienta importante en la investigación farmacológica y química que permite prever las reacciones de determinados fármacos. El acrónimo al uso es el del inglés.

Raíz del Error Cuadrático Medio (RMSE): *Root Mean Square Error*, en inglés. Su uso español en textos técnicos es común y ampliamente reconocido, especialmente en áreas relacionadas con estadística, aprendizaje automático, análisis de datos y otras disciplinas que involucran modelado predictivo y evaluación de algoritmos, sin embargo, el acrónimo que se reconoce es el del inglés. El RMSE es una de las métricas más populares para cuantificar el error entre los valores predichos por un modelo y los valores reales observados. Su uso es frecuente debido a que proporciona una medida clara y directa del tamaño promedio de los errores, escalada en las mismas unidades de la variable de interés, lo cual facilita su interpretación. Además, al elevar al cuadrado los errores antes de promediarlos, el RMSE da un peso mayor a los errores grandes, lo que puede ser deseable en muchos contextos donde estos son más problemáticos que los pequeños.

Rastreo de Datos: *Data Tracking*, en inglés. Es un término que se refiere al proceso de seguir y registrar datos a medida que se generan, se mueven o se utilizan en una red o sistema de información.

Recursive Feature Elimination (RFE): “Eliminación Recursiva de Características”, en español, aunque no suele

utilizarse en esta lengua en la literatura especializada. Es una técnica de selección de características, que elimina iterativamente los rasgos menos importantes de un conjunto de dato, utilizando para esa eliminación, el reentrenamiento del modelo y la repetición del proceso, hasta alcanzar el número deseado de características. Este método es efectivo para identificar un subconjunto óptimo o reducido de características, las cuales contribuyen más significativamente a la precisión del modelo.

Redes Neuronales Artificiales (RNA): Una red neuronal es un método de la inteligencia artificial que enseña a las computadoras a procesar datos de una manera que está inspirada en la forma en que lo hace el cerebro humano. Se trata de un tipo de proceso de aprendizaje automático llamado aprendizaje profundo, que utiliza los nodos o las neuronas interconectados en una estructura de capas que se parece al cerebro humano. Crea un sistema adaptable que las computadoras utilizan para aprender de sus errores y mejorar continuamente. Son capaces de aprender de datos, reconocer patrones, procesar secuencias, hacer generalizaciones, así como, organizar y representar datos de manera automática. De esta forma, las redes neuronales artificiales intentan resolver problemas complicados, como la realización de resúmenes de documentos, el reconocimiento de rostros o voces, con mayor precisión.

Redes neuronales del tipo Perceptrón Multicapa (MLP): El acrónimo se da por su nombre en inglés "*Multilayer Perceptron*". Estas son un tipo de red neuronal artificial, estructurada en múltiples capas, lo que las hace parte de la categoría de redes neuronales de alimentación directa (*feedforward neural networks*). Son capaces de modelar relaciones complejas y no lineales entre las entradas y las salidas, lo que las hace muy útiles para una amplia gama de

tareas de aprendizaje automático, incluyendo clasificación, regresión, y reconocimiento de patrones.

Relief: Es el nombre propio de un algoritmo en el campo del aprendizaje automático, que evalúa la importancia de las características, analizando cómo estas se distinguen entre instancias que están cercanas unas a otras. Para cada instancia, busca las instancias más cercanas de la misma clase (instancias cercanas) y de clases diferentes (instancias lejanas) y actualiza un puntaje de importancia para cada característica basándose en cuánto contribuye a diferenciar entre instancias cercanas y lejanas.

Rpart (Recursive Partitioning and Regression Trees: “Partición Recursiva y Árboles de Regresión”, en español. Se refiere a un paquete de algoritmos en el lenguaje de programación R, utilizado para la creación de árboles de decisión para clasificación y regresión. Al ser el nombre de un paquete de software, *Rpart* se mantiene igual en todos los idiomas para evitar confusiones, por lo que, generalmente, en textos técnicos se mantiene el anglicismo.

Scheduling: Anglicismo que literalmente se traduce como “programación, planificación o asignación de recursos”. Aunque el significado de planificación eventualmente puede superponerse con el de *scheduling*, el primero tiene un alcance más amplio y puede aplicarse en diferentes contextos y disciplinas; mientras que *scheduling*, se refiere específicamente a la asignación temporal de recursos o tareas a lo largo del tiempo, con el objetivo de optimizar algún proceso. Es común su uso en idioma inglés en la informática.

Simulated Annealing: “Recocido Simulado”, en español. Es una técnica de optimización inspirada en el proceso físico de calentamiento y enfriamiento controlado de materiales

para aumentar su tamaño de grano y reducir sus defectos, proceso que lleva a un estado de menor energía y, por lo tanto, a una mayor estabilidad del material. En el contexto de la optimización, el Recocido Simulado es utilizado para encontrar una solución aproximada a problemas de optimización global, especialmente aquellos que presentan un espacio de búsqueda complejo con múltiples mínimos locales. La metaheurística SA es particularmente útil para problemas donde los métodos de búsqueda convencionales pueden quedar atrapados en óptimos locales, impidiendo encontrar el óptimo global o una solución satisfactoria cercana a este.

Situación de “vigilancia”: Tiene lugar en las ciudades inteligentes, y puede incluir la monitorización de datos de tráfico, sensores de seguridad, cámaras de videovigilancia, análisis de datos de redes sociales y más, con el objetivo de mejorar la gestión urbana y la calidad de vida de los ciudadanos. Es muy importante equilibrar los beneficios de la vigilancia con la protección de la privacidad y la seguridad de los datos de los ciudadanos.

Smart Buildings: Son edificaciones inteligentes que utilizan tecnología avanzada para automatizar y optimizar sistemas como la iluminación, la climatización y la seguridad, con el fin de mejorar la eficiencia energética y la comodidad de los ocupantes.

Smart Grids: Son sistemas eléctricos modernos que integran tecnologías de comunicación y control para gestionar de manera eficiente la generación, transmisión y distribución de energía eléctrica, permitiendo la incorporación de fuentes de energía renovable y la respuesta a la demanda en tiempo real.

SMOTE, Borderline_SMOTE, Safe_Level_SMOTE, SMOTE_Tomek-Links, SMOTE_ENN y Spider2: Estos algoritmos tienen como objetivo mejorar la calidad del sobremuestreo en conjuntos de datos desbalanceados, abordando diferentes aspectos del desafío que representa la superposición de clases y las áreas de incertidumbre entre clases. Estas técnicas son particularmente útiles en problemas de clasificación donde es esencial manejar adecuadamente la clase minoritaria.

The Opal TM: Opal es un sensor portátil para investigación diseñado para ofrecer el máximo control. Opal es la tecnología básica de todas las soluciones. Aunque puede personalizar un sistema de cualquier tamaño, una de las soluciones estándar entre las que elegir es Mobility Lab (Laboratorio de movilidad): para medir la marcha espacial/temporal y el equilibrio.

Travelling Salesman Problem (TSP): “Problema del Viajante de Comercio” o “Problema del Agente Viajero”, en español, es un problema clásico en la teoría de la optimización y la investigación de operaciones. Su objetivo es encontrar la ruta más corta posible que debe seguirse para visitar una serie de lugares y regresar al punto de partida, logrando en este recorrido una distancia total mínima. El término es ampliamente reconocido en español, aunque en los textos técnicos suele mantenerse el acrónimo en inglés.

Treebag: El término proviene de la combinación de *Tree* (árbol, en inglés) y *Bagging*, refiriéndose específicamente al ensamblaje de árboles de decisión mediante la técnica de remuestreo *Bootstrap*. No es un término comúnmente usado en español; en el lenguaje técnico es ampliamente aceptado el anglicismo, especialmente en el campo del aprendizaje automático y la minería de datos. Esta técnica se basa en generar múltiples muestras de un conjunto de datos, con

reemplazo, para estimar la distribución de una estadística o parámetro.

Variable Neighborhood Search: Búsqueda de Vecindario Variable, en español. Es una metaheurística para la resolución de problemas de optimización que se basa en la idea de cambiar sistemáticamente el vecindario dentro del cual se busca una solución mejor. Este enfoque permite explorar eficazmente el espacio de búsqueda y ayuda a evitar quedar atrapado en óptimos locales, mejorando así las posibilidades de encontrar una solución global óptima o cercana al óptimo.

Vehicle Routing Problem (VRP): Problema de Ruteo de Vehículos, en español, es un problema clásico de optimización combinatoria. Su objetivo principal es determinar las rutas óptimas para una flota de vehículos que deben realizar entregas o recogidas en varios puntos, con el fin de minimizar el costo total, que puede ser expresado en términos de distancia total recorrida, tiempo total de viaje, costo de combustible, o una combinación de estos y otros factores.

Ventanas horarias: *Time windows* , en inglés, son restricciones de tiempo aplicadas a las visitas en problemas de enrutamiento, como el Problema de Enrutamiento de Vehículos con Ventanas de Tiempo (*Vehicle Routing Problem with Time Windows*, VRPTW). En este contexto, una ventana horaria define el rango de tiempo durante el cual se debe realizar una entrega, una recogida, o cualquier otro tipo de servicio en una ubicación específica. El uso de estas ventanas es común en la logística de distribución en contextos donde el tiempo de servicio es un factor crítico.

Weight OneR: Este es un término de difícil y muy poco usada traducción al español. Su significado podría ser “OneR Ponderado” o “Ponderación de OneR”, lo que indica que

incorpora algún mecanismo de ponderación para evaluar o mejorar la selección de características o la construcción de reglas de decisión. Este algoritmo es conocido por su simplicidad, a partir del uso de una única característica para hacer predicciones.

Weight Symmetrical Uncertainty: Puede traducirse al español como “Incertidumbre Simétrica Ponderada” pero – aunque la traducción directa proporciona una idea bastante clara de su significado– es común que se mantenga el anglicismo en la literatura científica y técnica. Se refiere a una medida derivada de la teoría de la información que evalúa la cantidad de información mutua entre dos variables, normalizando este valor para corregir el sesgo hacia variables con un mayor número de estados o categorías. Tal “ponderación” se vincula con ajustes adicionales o consideraciones hechas para evaluar la importancia relativa de diferentes características o variables en un conjunto de datos, basándose en esta medida de incertidumbre simétrica. Este algoritmo es especialmente útil en escenarios donde se desea reducir la dimensionalidad de los datos antes de aplicar modelos predictivos, mejorando así la eficiencia del modelado y potencialmente la precisión de las predicciones.

WEKA: Acrónimo de “*Waikato Environment for Knowledge Analysis*” (Entorno Waikato para el Análisis de Conocimiento). Se trata de un software de aprendizaje automático y minería de datos desarrollado por la Universidad de Waikato en Nueva Zelanda; Es de fácil uso y posee una amplia gama de herramientas para tareas de preprocesamiento de datos, clasificación, regresión, *clustering* (“*agrupamiento*”) asociación de reglas, y visualización. Su reconocimiento internacional está dado, fundamentalmente, por su acrónimo del inglés. El término en español, aunque posee una traducción comprensible, no se usa.

XGBoost: Acrónimo de *eXtreme Gradient Boosting* que literalmente significa “Potenciación del Gradiente Extremo”; traducción, cuyo uso en la literatura especializada es muy limitado. Este es un algoritmo de aprendizaje supervisado que se utiliza para problemas de clasificación, regresión y *ranking*. Es una implementación eficiente y escalable de árboles de decisión potenciados por gradientes. *XGBoost* posee gran rendimiento y velocidad en comparación con otros algoritmos de potenciación del gradiente.

K-Folds Cross-Validation: El método Validación Cruzada de K-Particiones o Validación Cruzada de *K-Fold* es una técnica utilizada para evaluar la capacidad de generalización de un modelo en el campo del aprendizaje automático y la estadística. Es común encontrarlo tanto en su versión original en inglés como en su traducción al español en textos técnicos, es perfectamente comprensible en español. Resulta especialmente útil para estimar cómo un modelo predictivo se desempeñará en la práctica cuando se le presenten datos nuevos e independientes. Es fundamental para evitar el sobreajuste, garantizar que el modelo sea robusto y que tenga un buen desempeño en datos no vistos. Al utilizar toda la muestra de datos tanto para entrenamiento como para prueba, la validación cruzada de K-particiones ofrece una estimación más fiable de la capacidad predictiva del modelo en comparación con una simple división entre datos de entrenamiento y prueba.

Ediciones Universidad de Camagüey
Marzo 2024