

Técnicas de inteligencia artificial usadas en el control, seguridad y optimización de procesos químicos: una revisión sistemática de la literatura.

Esteban Alejandro Gomez Jurado^{*a}, Nilson de Jesús Marriaga Cabrales ^{a, b}
(co-orientador), José Sebastián López Vélez ^{a, c} (orientador)

^aEscuela de Ingeniería Química, Facultad de Ingeniería, Universidad del Valle, Cali 760042, Colombia.

^bGrupo de investigación en procesos avanzados de oxidación para tratamientos biológicos y químicos (GAOX), Universidad del Valle.

^cGrupo de investigación en sistemas complejos adaptativos (GISCA), Instituto de pensamiento complejo.

ABSTRACT

La inteligencia artificial (IA) se ha consolidado como una tecnología clave en la transformación digital de la industria química. Este trabajo presenta una revisión sistemática de la literatura sobre aplicaciones de IA en plantas químicas, organizada en tres categorías funcionales: control de procesos, seguridad de procesos y optimización. La metodología implementada incluye técnicas bibliométricas como herramientas de análisis semántico y de acoplamiento bibliográfico, lo que permitió identificar modelos representativos y tendencias tecnológicas relevantes en el periodo 2014 – 2024. Entre los hallazgos más destacados se encuentra la tendencia de arquitecturas basadas en redes neuronales recurrentes (RNN), especialmente LSTM y GRU, las cuales han demostrado una alta eficacia en escenarios dinámicos y multivariantes. Además, se observa una creciente integración de enfoques híbridos, como los modelos que incluyen enfoques no supervisados o células del tipo *Physics-informed*, mejorando su robustez e interpretabilidad. Finalmente, se identifican desafíos clave para la adopción industrial, como la escasez de datos etiquetados y la necesidad de modelos explicables, señalando que las tendencias emergentes en modelos semi-supervisados, auto-adaptativos y con arquitectura híbrida constituyen rutas prometedoras hacia sistemas más autónomos, eficientes y seguros.

Palabras clave: inteligencia artificial (IA); control de procesos; monitoreo de procesos; optimización de procesos, detección de anomalías; seguridad de procesos

1. Introducción

La cuarta revolución industrial, en especial la industria 4.0, ha traído consigo una integración sin precedentes entre tecnologías digitales, físicas y biológicas, transformando profundamente los sistemas de producción, gestión y operación. La inteligencia artificial (IA) se posiciona como un eje central de innovación, permitiendo el desarrollo de modelos capaces de aprender, predecir y tomar decisiones de manera autónoma en distintos escenarios industriales. Esto es posible gracias a contribuciones en ámbitos como el internet de las cosas (IoT), IA, modelos de *Machine Learning* (ML), entre otros (Thon et al, 2021). En el ámbito de los procesos químicos la IA ha demostrado un impacto significativo en todas las escalas del proceso en

* Correspondencia: gomez.esteban@correounivalle.edu.co

plantas a nivel industrial, así como también en etapas preliminares de diseño o en fases de operación en el control de procesos (Bunian, S. et al., 2024).

El campo de los procesos químicos enfrenta crecientes desafíos relacionados con la complejidad, no linealidad y dinamismo de los procesos industriales modernos. La IA ha emergido como una herramienta transformadora, con capacidad para superar las limitaciones de los enfoques tradicionales. El grupo de Hajjar et al. (2016) señala que las técnicas de IA han demostrado una notable capacidad para abordar problemas complejos en modelado, control, optimización y diagnóstico de fallos, gracias a su flexibilidad, adaptabilidad y capacidad para operar sobre sistemas con información incompleta o altamente no lineal. Autores como Lee et al. (2018) destacan que el deep learning, ha generado avances significativos en tareas previamente difíciles, como la predicción secuencial, la toma de decisiones bajo incertidumbre y el control adaptativo de procesos dinámicos.

La IA está transformando profundamente la manera en que se enfrentan los desafíos técnicos en los procesos químicos. El uso de modelos de aprendizaje profundo y aprendizaje automático ha ganado protagonismo debido a su capacidad de autoaprendizaje, adaptación dinámica a distintas condiciones operativas y eficiencia en la gestión de datos complejos (Venkat, V., 2018). Dada la creciente complejidad de los sistemas químicos modernos y la presión por mejorar la eficiencia, la seguridad y la sostenibilidad, resulta esencial comprender cómo estas tecnologías están siendo adoptadas en la industria (Zhao, J. et al. 2022).

Los beneficios derivados de la implementación de modelos IA en la industria de los procesos químicos no se limitan únicamente a mejoras operativas, sino que también se extienden a aspectos clave como la seguridad de procesos. Entre los impactos más destacados se encuentra la reducción de costos asociados a insumos y servicios industriales, así como la capacidad de anticipar eventos no deseados dentro de las plantas (Zeinab, H. et al., 2018). Por ejemplo, la localización de posibles fugas químicas se convierte en una tarea crítica, dada la urgencia de una pronta respuesta con el fin de mitigar los daños secundarios y terciarios. En esta línea, en el trabajo de Shin et al. (2019) han subrayado la importancia de diseñar sistemas de respuesta de emergencia integrados que se apoyen en técnicas de IA para identificar rápidamente la fuente de la fuga, utilizando datos de sensores fijos distribuidos en la planta. Esta capacidad predictiva y localizada no solo mejora la respuesta operativa, sino que también reduce el margen de error en decisiones críticas, constituyéndose como un apoyo directo a la seguridad industrial y de procesos.

En este trabajo se llevó a cabo una revisión sistemática de la literatura sobre las aplicaciones de la IA en los procesos químicos, con el fin de establecer las técnicas IA que otorgan mayores beneficios. Para ello, se desarrolló una sección metodológica en la que se describe el proceso seguido para la selección objetiva de los artículos analizados. Esta etapa incluyó la construcción de la ecuación de búsqueda, que permitió obtener una muestra inicial de publicaciones; posteriormente, se aplicaron criterios de exclusión para reducir dicha muestra. Finalmente, se empleó la herramienta VOSviewer para visualizar y validar la correlación entre los artículos seleccionados.

En la sección *resultados* se presenta una breve explicación de los modelos propuestos en los estudios revisados, enfocados en tres grandes áreas de aplicación: control de procesos, seguridad de procesos y optimización. En esta parte del análisis se describe tanto el modelo utilizado como el entorno en el que fue implementado.

En la sección de discusión, se agrupan los indicadores de rendimiento identificados según su categoría y se proporciona información que permite establecer comparaciones entre los distintos modelos presentados. Adicionalmente, se incluyen representaciones gráficas con porcentajes de recurrencia, lo que facilita la identificación de tendencias relevantes en la literatura. Todo ello con el objetivo de aportar una visión crítica sobre los alcances y limitaciones actuales de la aplicación de la IA en plantas químicas.

Finalmente, las conclusiones derivadas de esta revisión presentan los modelos con mayor recurrencia y efectividad aplicados en procesos químicos, además de observaciones pertinentes en base al análisis de los documentos estudiados.

2. Metodología

Para garantizar una revisión sistemática, el presente estudio adoptó una metodología estructurada que permitió de forma objetiva: reducir, clasificar y filtrar los artículos más relevantes en el uso de IA en plantas químicas. Este proceso metodológico define el alcance temático y temporal del estudio, a la vez que establece las bases para la selección de artículos de calidad científica, soportada tanto en criterios bibliométricos como en herramientas de análisis visual.

El procedimiento se inició con la formulación de una ecuación de búsqueda empleada en la base de datos Scopus para delimitar las publicaciones a un conjunto manejable y pertinente de estudios. La construcción de la ecuación se realizó utilizando operadores booleanos y términos clave como “*Artificial Intelligence*”, “*Machine Learning*”, “*Chemical Engineering*” y “*Process Control*”. Esta búsqueda inicial se enfocó en una recuperación amplia, seguida de un proceso iterativo de refinamiento.

Para optimizar esta estrategia, se utilizó el software VOSviewer, mediante el cual se generaron mapas de densidad de palabras clave que permitieron visualizar la frecuencia y relación entre términos utilizados en los artículos indexados. En la Figura 1 se muestra este mapa de red para las palabras clave encontradas en los artículos tipo review que abordan estos temas. Con base en este análisis semántico, se refinaron las ecuaciones de búsqueda incluyendo sinónimos técnicos y términos aplicados que aumentaron la precisión de la búsqueda.

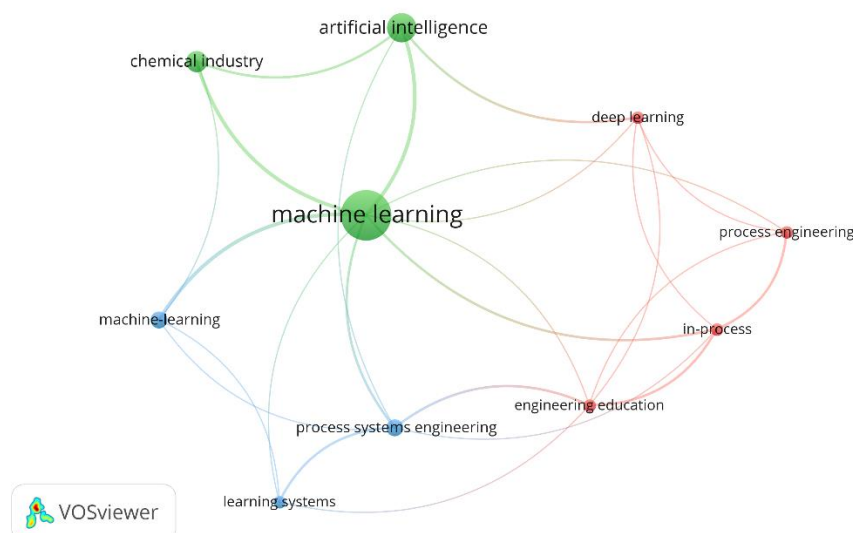


Figura 1. Mapa de red para palabras clave de artículos review.

Este proceso de ajuste se reflejó en una reducción significativa en la cantidad de resultados obtenidos, esto se evidencia en el cambio del volumen de publicaciones presentados en la Figura 2 en comparación con el volumen de artículos obtenidos mediante la ecuación de búsqueda refinada. La aplicación de filtros como el idioma (inglés), tipo de documento (artículos científicos), área de estudio (ingeniería química), documentos open access (excluidos), estado (publicación finalizada) y periodo (2014–2024) contribuyó a consolidar un volumen inicial de artículos mediante la ecuación de búsqueda: *("IA" OR "ARTIFICIAL INTELLIGENCE" OR "MACHINE LEARNING" OR "DEEP LEARNING" OR "MACHINE-LEARNING") AND ("CHEMICAL ENGINEERING" OR "PROCESS ENGINEERING" OR "CHEMICAL PLANT" OR "CHEMICAL INDUSTRY")*.

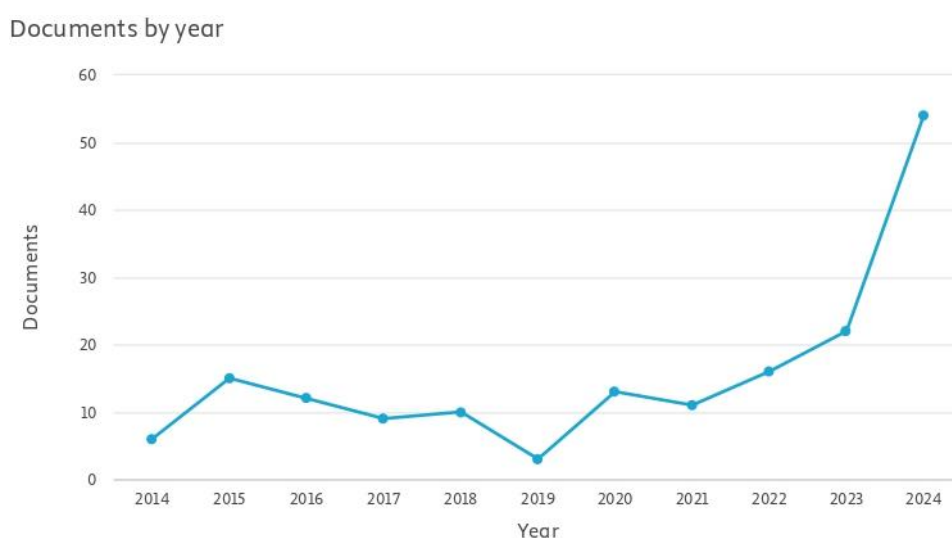


Figura 2. Tendencia de artículos por año para la ecuación: *("IA" OR "ARTIFICIAL INTELLIGENCE") AND ("CHEMICAL ENGINEERING" OR "PROCESS CONTROL")*. Elsevier, (s.f.-a).

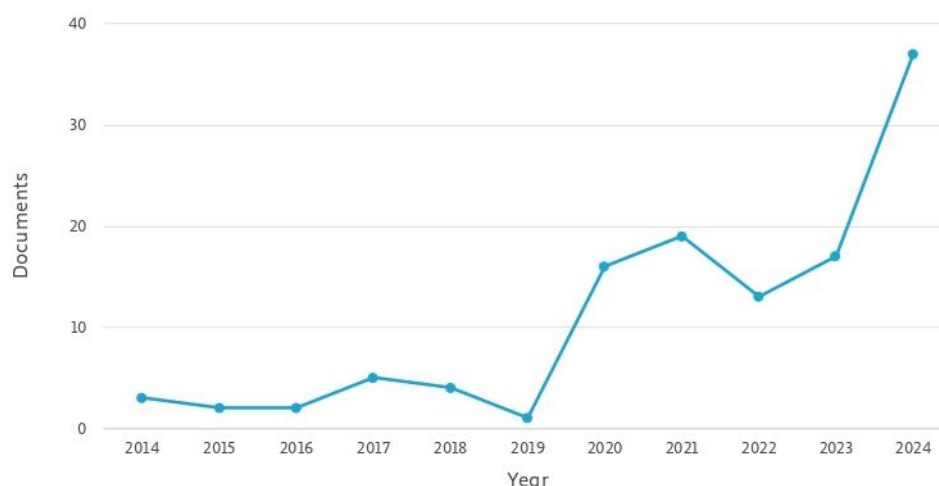


Figura 3. Tendencia de artículos por año para la ecuación: ("IA" OR "ARTIFICIAL INTELLIGENCE" OR "MACHINE LEARNING" OR "DEEP LEARNING" OR "MACHINE-LEARNING") AND ("CHEMICAL ENGINEERING" OR "PROCESS ENGINEERING" OR "CHEMICAL PLANT" OR "CHEMICAL INDUSTRY"). Elsevier, (s.f.-b).

Posteriormente, se incorporaron criterios bibliométricos para evaluar la calidad e impacto de las publicaciones. Para ello, se extrajeron indicadores como el cuartil de la revista (Q), el índice SJR y el *CiteScore*, que fueron integrados en una métrica compuesta denominada *metric score*, este indicador se presenta en la siguiente ecuación.

$$metric\ score = SJR * \frac{CiteScore}{cuartil} \quad (1)$$

Esta métrica permitió valorar las publicaciones de forma cuantitativa, priorizando aquellas con mayor visibilidad científica. Los valores bibliométricos asociados a cada revista son presentados en la Tabla 1.

Tabla 1.

Revistas asociadas a los 124 artículos con sus respectivos indicadores bibliométricos.

JOURNAL	Q	SJR	Cite score	Articulos	Metric Score	JOURNAL	Q	SJR	Cite score	Articulos	Metric Score
ACS Applied Polymer Materials	1.5	0.982	7.2	1	4.71	Education for Chemical Engineers	1	0.933	8.8	7	8.21
ACS Sustainable Chemistry and Engineering	1	1.664	13.8	4	22.96	Fluid Phase Equilibria	2	0.569	5.3	1	1.51
Advanced Powder Technology	1	0.813	9.5	1	7.72	Gas Science and Engineering	1	1.161	11.2	1	13
AIChE Journal	1.5	0.734	7.1	5	3.47	Industrial and Engineering Chemistry Research	1	0.811	7.4	5	6
Asia-Pacific Journal of Chemical Engineering	3	0.34	3.5	1	0.4	International Journal of Chemical Reactor Engineering	3	0.307	2.7	1	0.28
Biofuels, Bioproducts and Biorefining	2	0.748	7.8	1	2.92	International Journal of Multiphase Flow	1	0.939	7.3	1	6.85
Biomedical Engineering - Applications, Basis and Communications	4	0.195	1.5	1	0.07	Journal of Advanced Manufacturing and Processing	2	0.645	4.5	1	1.45
Bioresource Technology	1	2.576	20.8	2	53.58	Journal of Chemical and Engineering Data	2	0.567	5.2	1	1.47
Canadian Journal of Chemical Engineering	2	0.402	3.6	5	0.72	Journal of Chemical Information and Modeling	1	1.396	9.8	5	13.68
Chemical Engineering and Processing - Process Intensification	1	0.739	7.8	1	5.76	Journal of Chemical Technology and Biotechnology	1.5	0.624	7	6	2.91
Chemical Engineering Communications	2	0.473	5.5	1	1.3	Journal of Food Process Engineering	2	0.588	5.7	2	1.68
Chemical Engineering Journal	1	2.852	21.7	3	61.89	Journal of Loss Prevention in the Process Industries	1	0.689	7.2	2	4.96
Chemical Engineering Research and Design	2	0.661	6.1	4	2.02	Journal of Membrane Science	1	1.848	17.1	1	31.6
Chemical Engineering Science	1	0.817	7.5	5	6.13	Journal of the Taiwan Institute of Chemical Engineers	1	0.849	9.1	1	7.73
Chemical Engineering Transactions	3	0.258	1.4	9	0.12	Korean Journal of Chemical Engineering	2	0.541	4.6	3	1.24
Chemical Processing	4	0.1	0	2	0	Lab on a Chip	1	1.246	11.1	1	13.83
Chemie-Ingenieur-Technik	3	0.449	3.4	2	0.51	Petroleum Science and Technology	3	0.325	2.9	2	0.31
Chemometrics and Intelligent Laboratory Systems	1.5	0.667	7.5	1	3.34	Powder Technology	1	0.97	9.9	3	9.6
Chinese Journal of Chemical Engineering	1.5	0.697	6.6	2	3.07	Process Integration and Optimization for Sustainability	3	0.445	4.3	1	0.64
Colloid and Polymer Science	2	0.424	4.6	1	0.98	Process Safety and Environmental Protection	1	1.293	11.4	4	14.74
Computers and Chemical Engineering	1	0.903	8.7	18	7.86	Separation and Purification Technology	1	1.533	14	2	21.46
Drying Technology	1	0.69	7.4	1	5.11	Theoretical Foundations of Chemical Engineering	4	0.226	1.2	2	0.07

Con base en el *metric score*, se calculó un índice adicional denominado *impact score*, que pondera el número de citas recibidas por año de publicación ajustado por la calidad de la revista. Esta clasificación permitió priorizar los artículos más influyentes, facilitando la transición hacia un conjunto de publicaciones más reducido y focalizado.

$$Impact\ score = \frac{Citas}{tiempo\ de\ publicación} * metric\ score \quad (2)$$

En la etapa final de esta fase metodológica, se aplicó nuevamente el software VOSviewer, esta vez para realizar un análisis de acoplamiento bibliográfico entre los artículos seleccionados. Esta técnica identifica la fuerza de relación temática entre documentos a partir de las referencias bibliográficas compartidas. El resultado fue la identificación de un grupo cohesionado de 15 artículos interrelacionados, que presentan un núcleo temático y representan el estado del arte en la aplicación de la IA en procesos químicos. Este acople se presenta en la Figura 4, correspondiente al mapa de densidad por acoplamiento bibliográfico).

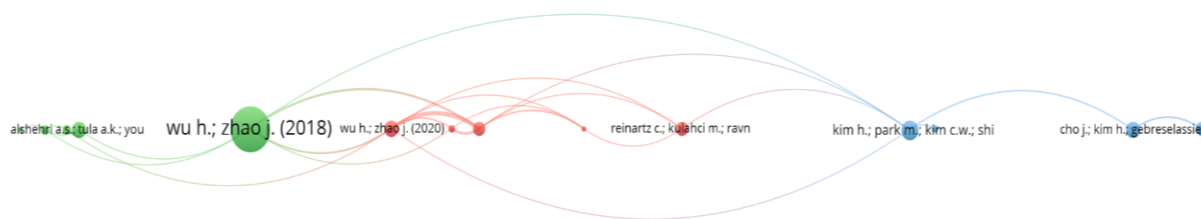


Figura 4. Visualización de redes del mapa de emparejamiento bibliográfico.

Adicionalmente, se presenta en la Figura 5 un mapa de co-ocurrencia de palabras clave con el fin de reforzar la validez de la selección final.

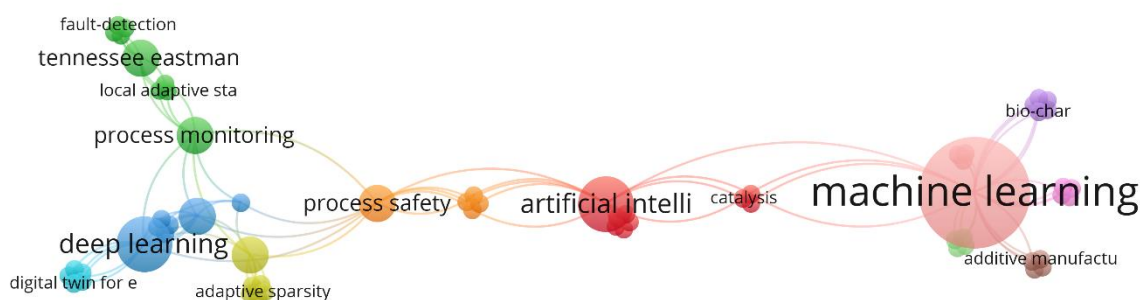


Figura 5. Mapa de densidad basado en la Co-ocurrencia de palabras clave.

Finalmente, este enfoque metodológico permitió no solo reducir el volumen de información inicial, sino también asegurar que los artículos seleccionados cumplieran con criterios de actualidad, calidad editorial, impacto científico y coherencia temática.

3. Resultados

En los últimos años, la literatura ha documentado una amplia gama de aplicaciones de la IA en la ingeniería de procesos, donde se han desarrollado modelos con propósitos diversos. Por ejemplo, el trabajo de Zhao et al. (2022) proporciona una visión general del papel que tienen las herramientas inteligentes para la detección y diagnóstico de fallas (FDD) en procesos químicos, enfatizando el papel crucial que tienen este tipo de herramientas, mediante ejemplos de accidentes reales y ampliamente documentados (como el caso de la planta de Texas City en 2005) se ilustra cómo la ausencia de sistemas inteligentes de monitoreo y diagnóstico oportuno puede derivar en consecuencias catastróficas. En contraste, los autores plantean que los modelos de FDD inteligentes basados en datos pueden operar con gran precisión incluso en entornos complejos, siendo capaces de detectar patrones anómalos en etapas tempranas y proporcionando diagnósticos interpretables para la toma de decisiones operativas humanas.

Con base a los artículos resultantes, este trabajo reconoce tres tópicos o categorías principales de aplicación en las que se concentran los esfuerzos actuales de investigación y desarrollo: el control y monitoreo de procesos, la seguridad de procesos y la optimización de métodos

tradicionales para temas ya estudiados como el diseño de reactores o la determinación de propiedades de compuestos. A continuación, se presenta una breve contextualización de cada una de estas áreas, sirviendo como preámbulo a la compilación y análisis de los modelos más representativos dentro de cada categoría.

3.1. *Control de procesos*

El control de procesos es una de las áreas más estudiadas en la aplicación de la IA dentro de los procesos químicos. Aquí, los modelos de IA se utilizan principalmente para mejorar la estabilidad operativa, anticipar comportamientos anómalos y reducir la intervención humana. En particular, los modelos de diagnóstico y detección de fallas han demostrado ser herramientas eficaces para mitigar riesgos, aumentar el tiempo de actividad de las plantas y permitir una respuesta más ágil ante desviaciones operativas (Zhao, J. et al. 2022).

Un punto de inflexión en el desarrollo y evaluación de herramientas para el campo de FDD ha sido el uso del proceso *Tennessee Eastman* (TEP) como benchmark de referencia. Este sistema, que emula un proceso químico industrial con múltiples variables interconectadas, se ha convertido en una plataforma estandarizada para probar el desempeño de nuevos modelos FDD, especialmente en entornos controlados y simulados. Su popularidad radica en que permite replicar condiciones dinámicas reales, evaluar distintos modos de operación, y generar fallas intencionales bajo criterios definidos, proporcionando así un marco reproducible y comparativo. Esta característica ha facilitado el desarrollo de modelos más robustos frente a datos ruidosos o escenarios multimodales, un reto común en entornos industriales reales. La Figura 6 presenta el diagrama de instrumentación y proceso (P&ID) del sistema, con el fin de contextualizar su estructura antes de abordar los modelos evaluados en esta plataforma (Ravn, O. et al. 2018).

En el contexto de este trabajo, el uso del TEP como herramienta de evaluación resulta fundamental para validar modelos inteligentes que apuntan a una mayor autonomía, adaptabilidad y precisión diagnóstica. Gracias a su estructura simulada pero representativa de procesos industriales reales, el TEP permite medir el desempeño de distintas técnicas de detección y diagnóstico de fallas bajo condiciones bien definidas. En particular, indicadores como la tasa de detección de fallas (FDR) y la tasa de falsas alarmas (FAR) se han establecido como métricas estándar para evaluar la eficacia de estos modelos. Así, los resultados obtenidos a partir de este benchmark ofrecen una base objetiva para comparar enfoques y valorar su potencial de aplicación en la industria química (Downs and Vogel, 1993).

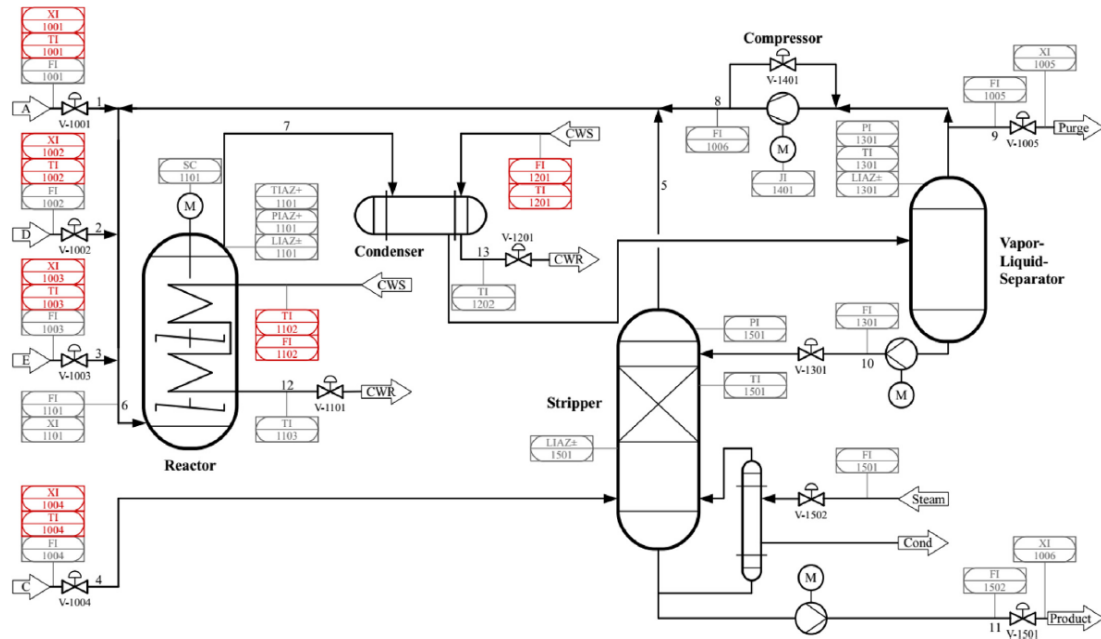


Figura 6. P&ID del TEP (Bathelt et al., 2015).

En las secciones siguientes se presenta una recopilación de modelos destacados que han sido propuestos y evaluados directamente sobre el proceso Tennessee Eastman y reportan resultados en función de los indicadores mencionados.

3.1.1. Modelo VAE-LSTM con interferencia bayesiana

En el trabajo de Zhao et al. (2022), el grupo de investigación propone un modelo innovador de detección de fallos para procesos químicos, basado en un *autoencoder* (AE) con arquitectura *Long short-term memory* (LSTM) e inferencia bayesiana. La Figura 7 presenta el esquema básico de una célula LSTM, mientras que la Figura 8 muestra una representación esquemática de la arquitectura AE, destacando las secciones *encoder* y *decoder* como parte de su estructura clásica.

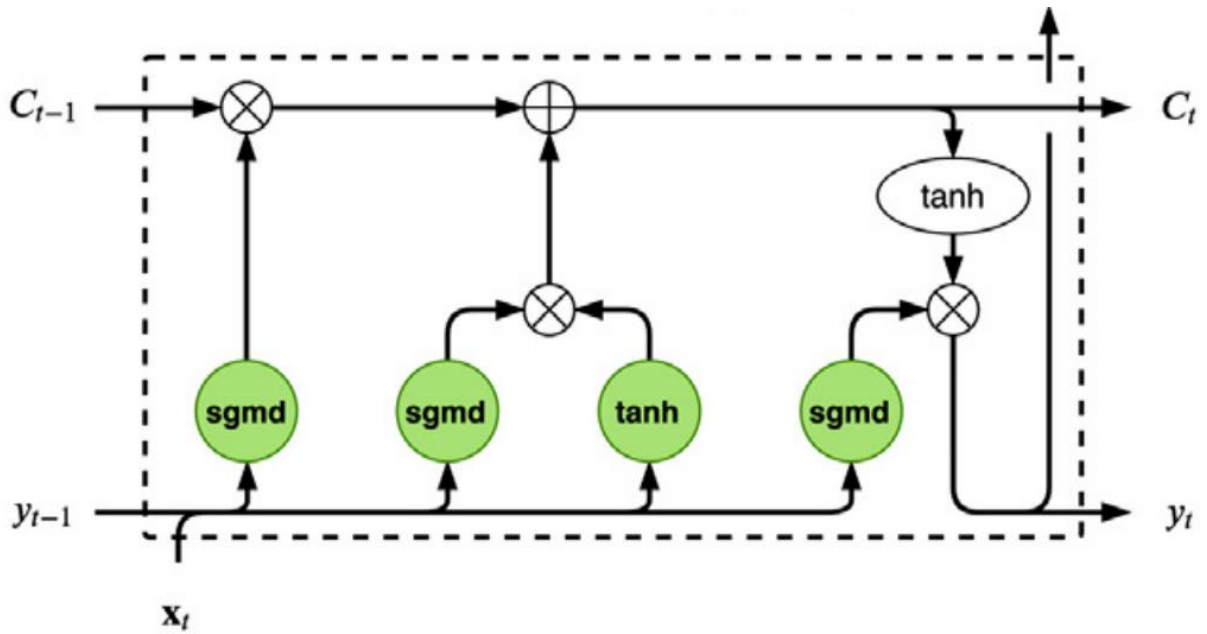


Figura 7. Diagrama de una célula de arquitectura LSTM (Viana et al., 2021).

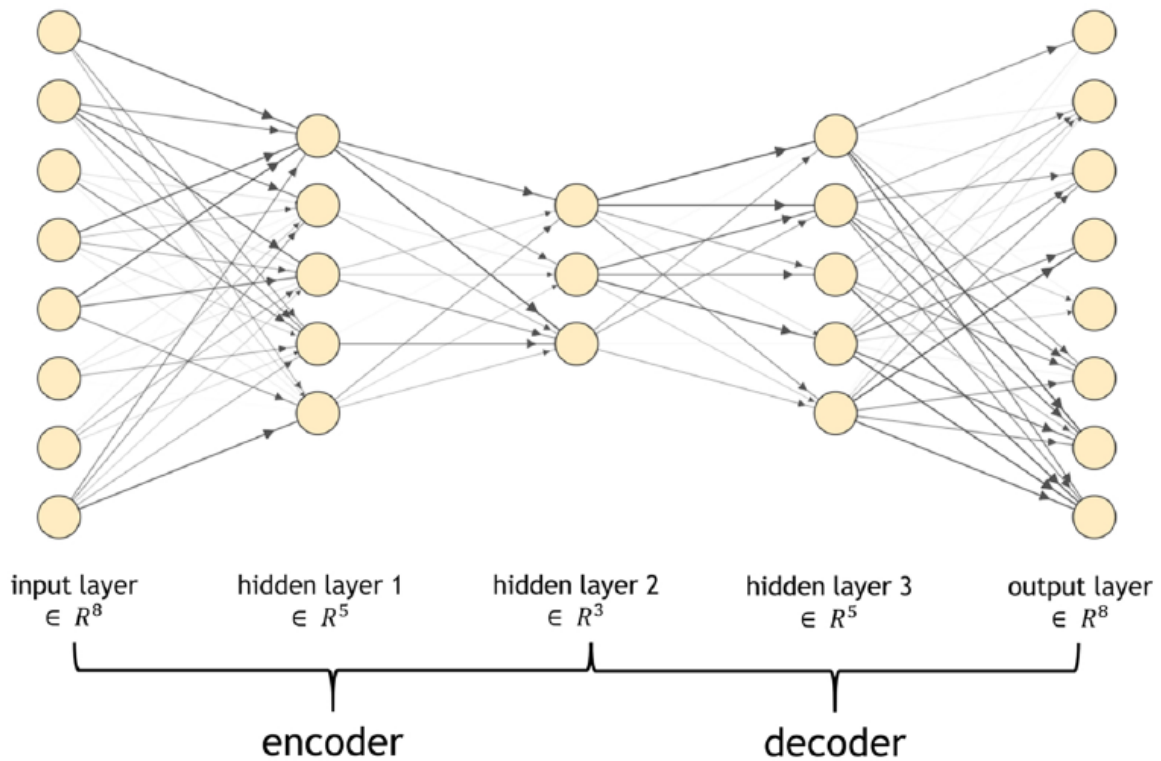


Figura 8. Diagrama representativo de arquitectura AE (Zhao et al., 2022).

El modelo híbrido VAE-LSTM con interferencia bayesiana se compone de un conjunto de arquitecturas de redes neuronales profundas, como los AE, los cuales son redes neuronales que buscan reconstruir sus entradas tras comprimirlas en un espacio latente de menor dimensión. Esto obliga al modelo a aprender una representación condensada (las características más importantes) de los datos de entrada. Estas arquitecturas son comúnmente entrenadas con datos normales, de modo que el AE genera reconstrucciones precisas para dichos datos y presenta

errores notables al intentar reconstruir entradas anómalas, lo cual se evidencia en un aumento del error de reconstrucción. El modelo incorpora una red con arquitectura LSTM, una variante especializada de las redes neuronales recurrentes (RNN) diseñada para capturar dependencias temporales a largo plazo en secuencias multivariantes. Esta capacidad resulta fundamental en contextos industriales, donde los datos de sensores evolucionan en el tiempo y presentan dinámicas complejas. La inclusión de LSTM dentro del AE permite considerar explícitamente la evolución temporal de las variables, lo que incrementa la sensibilidad del modelo ante desviaciones sutiles o graduales. Este modelo, alcanzó un FDR de 95.5% en el benchmark TEP.

Este modelo aprende una representación probabilística de los datos normales, incluyendo sus dependencias temporales, y detecta fallos mediante el análisis del error de reconstrucción o la divergencia respecto a la distribución latente aprendida. Su implementación ha demostrado una excelente capacidad para capturar patrones anómalos en procesos industriales complejos, incluso con condiciones dinámicas y datos ruidosos.

3.1.2. *Deep convolutional neural network (DCNN)*

En el estudio de Zhao et al. (2018), se propone un método de FDD en procesos químicos basado en una red neuronal convolucional profunda (DCNN), con la que alcanzaron un FDR promedio de 88.2% en el benchmark del TEP. Este modelo representa una evolución respecto a los enfoques tradicionales, al explotar la capacidad de las DCNN para aprender representaciones jerárquicas y robustas directamente desde datos sin requerir ingeniería manual de características.

Las redes neuronales convolucionales (CNN) son una clase de modelos de aprendizaje profundo especialmente diseñados para extraer patrones locales de los datos mediante filtros convolucionales aplicados en ventanas deslizantes. Una CNN típica se compone de capas convolucionales (que detectan características locales), capas de *pooling* (que reducen la dimensionalidad y ayudan a evitar sobreajuste), y capas completamente conectadas (FC) que realizan la clasificación final. Cuando estas redes se apilan en múltiples capas, se convierten en DCNN, capaces de extraer representaciones complejas de alto nivel a partir de datos de entrada crudos. La Figura 9 presenta una representación esquemática de una capa del algoritmo CNN.

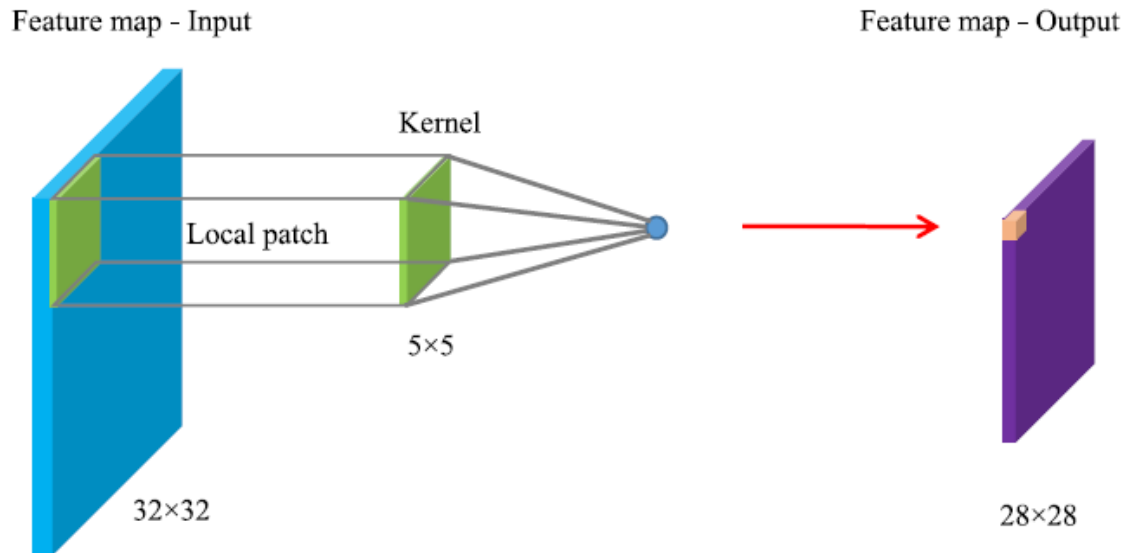


Figura 9. Diagrama representativo de arquitectura CNN (Zhao et al., 2018).

En este trabajo, los datos del proceso (originalmente secuencias temporales multivariadas) son reorganizados como matrices bidimensionales (espacio-temporales), similares a imágenes, donde las filas representan ventanas temporales y las columnas corresponden a variables de proceso. Esta transformación permite que la DCNN detecte patrones tanto en el dominio espacial (entre variables) como en el temporal (evolución de las variables), aspecto crítico en procesos químicos dinámicos.

Durante el entrenamiento, la DCNN aprende a clasificar estas matrices en distintas categorías de fallos o estado normal, utilizando un enfoque supervisado con la función *softmax* en la salida. El modelo entrenado puede luego ser desplegado en línea, donde evalúa matrices de datos recientes para identificar desviaciones anómalas en tiempo real. Entre las ventajas del modelo DCNN destacan su capacidad de generalización, baja tasa de falsos positivos y su eficiencia computacional frente a otros métodos basados en redes profundas no supervisadas.

3.1.3. *Enhanced Temporal algorithm-coupled optimized adaptive sparse PCA (ETA-OASPCA)*

En el estudio de Dong et al. (2023), se presenta un modelo híbrido denominado *Enhanced Temporal Algorithm-coupled Optimized Adaptive Sparse PCA* (ETA-OASPCA), con el objetivo de mejorar la capacidad del análisis de componentes principales (PCA) para diagnosticar fallos en procesos químicos dinámicos y multivariados. A continuación, se presenta en la Figura 10 el diagrama representativo del modelo propuesto por Dong et al (2023).

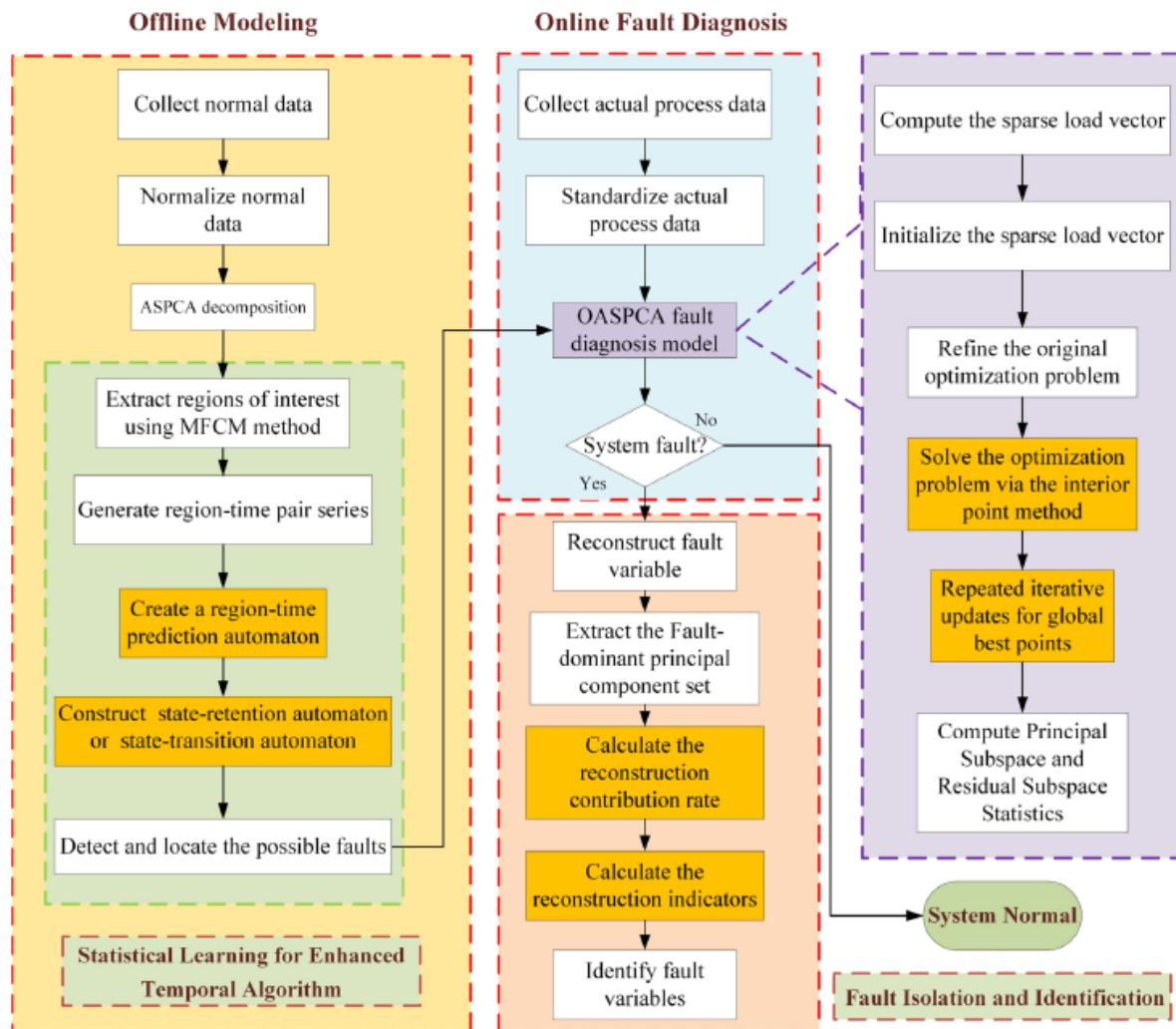


Figura 10. Diagrama representativo del modelo ETA-OASPCA (Dong et al., 2023).

Este modelo se compone de dos módulos principales: el primero es el *Optimized Adaptive Sparse PCA* (OASPCA), una extensión del PCA que introduce penalizaciones dinámicas para lograr esparsidad adaptativa en los vectores de carga, permitiendo identificar de manera más clara qué variables tienen mayor influencia en los componentes principales. La optimización se realiza mediante un método de punto interior iterativo, que asegura la convergencia a soluciones globales, evitando mínimos locales. Además, OASPCA divide el espacio de datos en subespacios principal y residual, sobre los cuales se calculan las estadísticas CT^2 y CSPE para detectar desviaciones anómalas. El segundo componente del modelo es el *Enhanced Temporal Algorithm* (ETA), basado en lógica temporal extendida mediante operadores como *Hold* y *Transition*, que permiten modelar y detectar cambios de estado en datos multivariantes y no estacionarios. A través del uso de *Maximum Fuzzy C-Means* (MFCM), se identifican regiones de interés temporal, que luego son integradas en un autómata de predicción de tiempo-región capaz de anticipar comportamientos anómalos. Finalmente, una vez detectada una anomalía, se aplica un método de aislamiento e identificación de fallos basado en reconstrucción por contribución (RBC), que localiza las variables responsables del fallo. Este modelo logró un desempeño sobresaliente en el benchmark TEP, alcanzando un FDR de 99.4%, un FAR de 0.37% y una latencia de diagnóstico de 4.11 minutos.

3.1.4. Local adaptive standardization and variational auto-encoder bidirectional long short-term memory (LAS-VB)

El grupo de Zhao et al. (2020), introduce un método de aprendizaje profundo auto-adaptativo para el monitoreo de procesos con múltiples modos de operación, denominado *Local adaptive standardization and variational auto-encoder bidirectional long short-term memory* (LAS-VB), el cual integra una estrategia de estandarización local adaptativa (LAS) con un modelo de AE variacional con capas LSTM bidireccionales (VAE-BiLSTM). De igual forma, como para los modelos anteriores en la Figura 11 se presenta el diagrama de flujo del modelo propuesto.

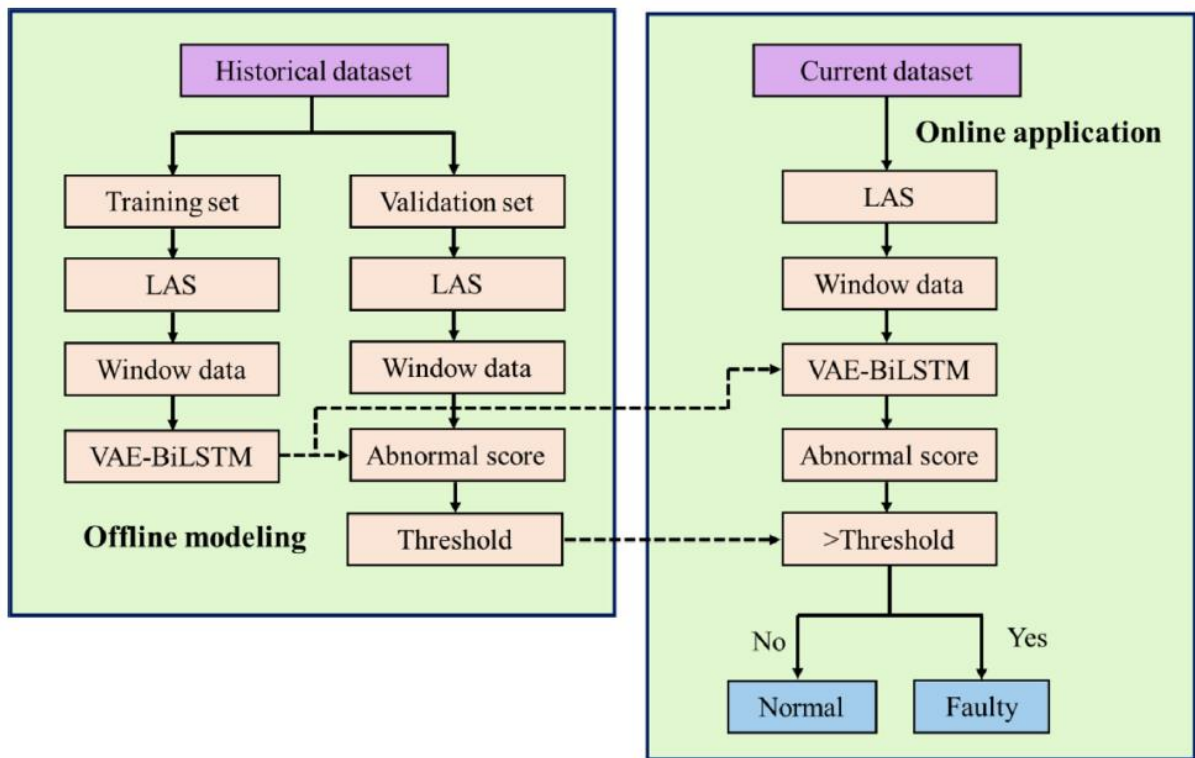


Figura 11. Diagrama de flujo del modelo LAS-VB (Zhao et al., 2020).

Este modelo fue desarrollado con el propósito de superar las limitaciones de los métodos tradicionales al enfrentar nuevos modos de operación no observados en los datos históricos, un desafío recurrente en entornos industriales complejos y que es repetitivo en las investigaciones de esta área de estudio. El componente LAS estandariza dinámicamente los datos en una ventana deslizante local, utilizando la media y desviación estándar calculadas dentro de esa ventana, sin depender de datos previos o clasificaciones por modo. Esto permite detectar tendencias inestables en las señales del proceso, característica fundamental para identificar fallos iniciales.

Posteriormente, los datos estandarizados son introducidos al modelo VAE-BiLSTM (VB), una arquitectura que combina la codificación probabilística del *variational autoencoder* (VAE) con la capacidad del bilateral LSTM (BiLSTM) de capturar dependencias temporales bidireccionales en series multivariadas. Esta integración permite que el modelo detecte desviaciones sutiles tanto hacia adelante como hacia atrás en el tiempo, mejorando la sensibilidad y adaptabilidad ante variaciones dinámicas. En la práctica, el sistema calcula un puntaje de anomalía basado en la pérdida de reconstrucción del VAE y lo compara contra un umbral obtenido por estimación de densidad en el conjunto de validación. Si el puntaje excede este límite, se considera que hay una condición anómala. En pruebas realizadas con el TEP, el modelo LAS-VB alcanzó FDR superiores al 90% en todos los modos evaluados, incluyendo modos no presentes durante el entrenamiento, y presentó una FAR con valores bajos como 0.3% en condiciones normales.

3.1.5. Three step high fidelity positive unlabeled (THPU)

En el trabajo de Zhao y Zheng (2022), se propone nuevamente un método de aprendizaje profundo semi-supervisado para la detección de fallos en procesos químicos, denominado *Three-step High-fidelity Positive-Unlabeled* (THPU). A continuación, se presenta el diagrama de flujo del modelo propuesto en la Figura 12.

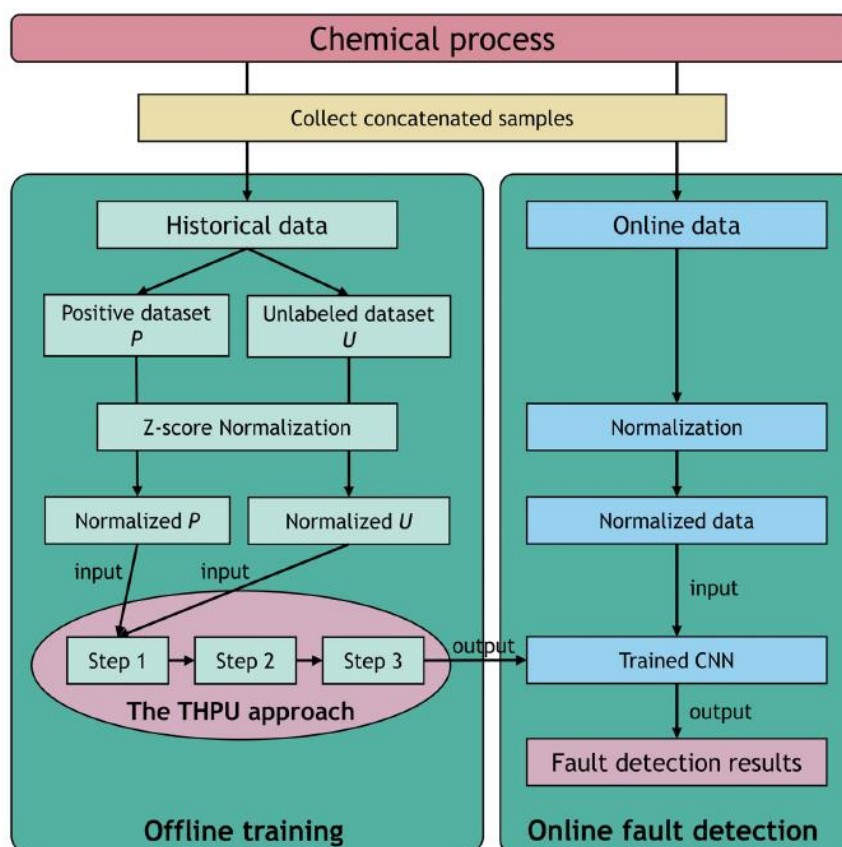


Figura 12. Diagrama representativo del modelo ETA-OASPCA (Dong et al., 2023).

Este modelo aborda uno de los principales desafíos en aplicaciones industriales reales, la escasez de datos etiquetados. La propuesta se basa en aprendizaje positivo-no etiquetado (PU), un paradigma que permite entrenar clasificadores binarios a partir de un conjunto pequeño de datos etiquetados como positivos (normalidad) y una gran cantidad de datos no etiquetados. El modelo THPU opera en tres etapas: en la primera, se emplea el algoritmo *Self-Training Nearest Neighbors* (STNN), una mejora de los algoritmos de autoetiquetado, para construir un conjunto confiable de muestras positivas (RP1) desde los datos no etiquetados, utilizando distancia euclidiana y validación mediante un *Stacked Autoencoder* (SAE) que filtra muestras con altos errores de reconstrucción. En la segunda etapa, el algoritmo *Positive and Negative Data Recognition* (PNDR) identifica tanto muestras confiables positivas adicionales como negativas (RP2 y RN), también utilizando SAE para estimar la confiabilidad de las muestras según su error de reconstrucción. Finalmente, en la tercera etapa, se entrena una CNN sobre los conjuntos RP2 y RN como clasificador binario para la detección en línea de fallos. Esta red CNN incorpora múltiples capas convolucionales, de agrupamiento y completamente conectadas para aprender representaciones espaciales y temporales de los datos del proceso. Aplicado al benchmark TEP, el método propuesto alcanzó un FDR promedio de 92.31%, un FAR de 2.81% y una latencia promedio de 11 minutos.

3.1.6. *Machine learning (ML) aided Model predictive control (MPC) with multi-objective optimization (MOO) and multi-criteria decision making (MCDM) methodology (ML-aided MPC-MOO-MCDM)*

En el estudio de Wu et al. (2023), se propone una metodología integral de control predictivo basada en modelos asistidos por aprendizaje automático, denominada *Machine learning (ML)-aided MPC-MOO-MCDM*. Este enfoque busca mejorar el desempeño de procesos químicos complejos al incorporar simultáneamente múltiples objetivos de control y toma de decisiones en la estructura tradicional del *Model Predictive Control* (MPC). El modelo se compone de tres etapas principales: (1) Predicción mediante modelos de aprendizaje automático, (2) Optimización multiobjetivo (MOO) y (3) Toma de decisiones multicriterio (MCDM).

Inicialmente, en la primera etapa, el modelo predictivo de MPC es reemplazado por un modelo LSTM entrenado a partir de datos simulados del proceso. Esta red neuronal recurrente permite capturar dinámicas no lineales y temporales del sistema sin necesidad de derivar modelos fisicoquímicos explícitos, esto es particularmente útil en sistemas donde la modelación de primer principio es difícil o costosa. En la segunda etapa, se formula el problema de control como una optimización multiobjetivo, donde se consideran simultáneamente funciones como la minimización del error de seguimiento, el consumo de reactivo junto con el consumo energético y la estabilidad del sistema (medida por una función de *Lyapunov*). Para resolver este problema se utiliza el algoritmo *Non-dominated Sorting Genetic Algorithm II* (NSGA-II), que permite encontrar un conjunto de soluciones de Pareto óptimas que representan distintos compromisos entre los objetivos.

Finalmente, la tercera etapa introduce mecanismos de toma de decisiones multicriterio para seleccionar una solución óptima desde el frente de Pareto. Se aplican los métodos *CRITIC* (que

asigna pesos a los objetivos en función de su varianza y correlación) y *PROBID* (que selecciona la mejor solución combinando la distancia a soluciones ideales y promedio). Además, se emplea la técnica *Weighted Sum Technique* (WST) para integrar dinámicamente los pesos calculados a lo largo de las iteraciones del MPC, permitiendo que el sistema priorice diferentes objetivos según las condiciones del proceso. Aplicado a un reactor de mezcla completa (CSTR), el modelo se representa en un diagrama de bloques en la Figura 13.

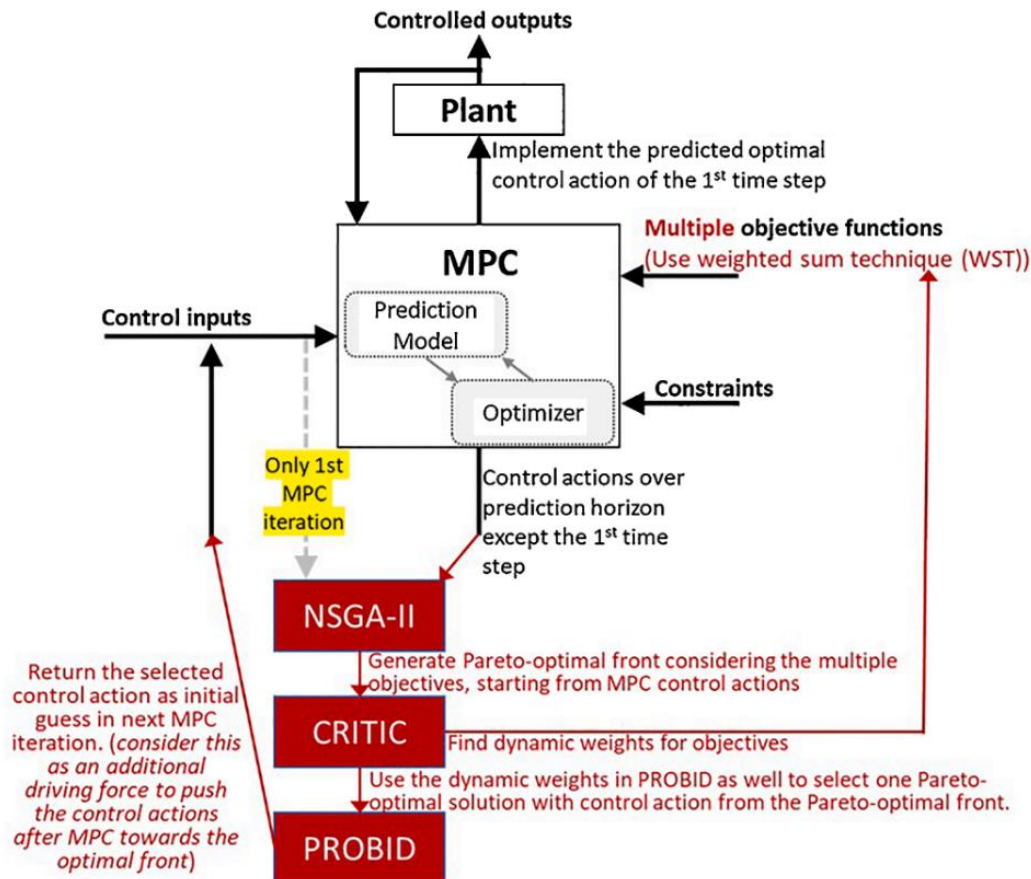


Figura 13. Diagrama de bloques del modelo ML-aided MPC-MOO-MCDM (Wu et al., 2023).

Este modelo, logró estabilizar el sistema en 12 iteraciones de MPC (equivalente a 0.12 h) bajo diversas condiciones iniciales, demostrando una mejora significativa en la rapidez de convergencia, el consumo energético y la estabilidad del sistema, sin comprometer la estabilidad del lazo cerrado. De forma más específica, al considerar simultáneamente los tres objetivos (control del sistema, optimización de insumos y estabilidad operativa) se identifican configuraciones que maximizan individualmente cada uno de ellos, así como una configuración que considera los tres objetivos al mismo tiempo. En el escenario que prioriza el control del sistema, se observa una mejora del 68 % en el tiempo requerido para alcanzar el estado controlado, en comparación con el escenario multiobjetivo. Para el ahorro de insumos se observa una reducción de 14% en ΔCA_0 y 10% en ΔQ en el escenario donde se prioriza el control del sistema junto con el ahorro de insumos. Para la estabilidad del sistema se observa una reducción de ~99.87% en energía del sistema en el escenario donde se prioriza este objetivo junto con el control del sistema. Para los tres objetivos se observa un ~10% ahorro en insumos y un 68% más estable que el control predictivo convencional. La metodología fue evaluada

tanto con modelos de primer principio como con el modelo LSTM, demostrando su flexibilidad y robustez para su aplicación en entornos industriales con múltiples objetivos.

3.2. Seguridad de procesos

La detección y localización automatizada de puntos de fuga en plantas químicas es una tarea de gran importancia dentro de los sistemas de seguridad de procesos. Una identificación precisa del origen de la fuga permite activar protocolos de respuesta más eficaces y reducir la propagación del evento. En este contexto, los trabajos de Shin et al. (2018, 2019) constituyen una referencia al desarrollar e implementar modelos de clasificación basados en aprendizaje supervisado, entrenados con datos obtenidos mediante simulaciones dinámicas de fluidos computacional (CFD). Estos estudios proponen un enfoque práctico basado en sensores fijos ubicados en el perímetro de la planta (*fence monitoring*), lo que elimina la necesidad de sensores móviles o mallas densas. A partir de este esquema, se evalúa el desempeño de distintos algoritmos de clasificación DNN, LSTM, FNN y *Random Forest* (RF) para la localización del punto de fuga entre múltiples posiciones posibles.

A continuación, se describen las características técnicas de cada modelo, así como los resultados obtenidos en términos de precisión de clasificación (*accuracy*), definida como la proporción de predicciones correctas respecto al total de casos evaluados. Es decir, la capacidad del modelo para identificar correctamente la ubicación del punto de fuga entre el conjunto de sensores. La métrica se calcula comparando la salida del modelo con la etiqueta real del escenario de fuga, ambas definidas durante la generación de datos mediante simulaciones CFD.

Adicionalmente, los modelos DNN, LSTM, FNN comparten una raíz común en las redes neuronales artificiales (ANN), de las cuales derivan variantes como las *feedforward neural network* (FNN) y las redes neuronales profundas (DNN). Estas arquitecturas constituyen la base sobre la que se construyen modelos más avanzados como las redes recurrentes LSTM. Las ANN constituyen una clase fundamental de modelos dentro del aprendizaje automático supervisado, estas arquitecturas están inspiradas en la estructura neuronal biológica y se componen de capas de nodos (neuronas artificiales) interconectadas, que transforman entradas mediante funciones de activación no lineales. Una de las variantes más comunes de las ANN es la FNN, en la cual la información fluye unidireccionalmente desde la capa de entrada hacia la de salida, sin ciclos o retroalimentación. Las FNN son ampliamente utilizadas para tareas de clasificación y regresión sobre datos tabulares o estructurados, pero presentan limitaciones al procesar información secuencial o con dependencia temporal.

Asimismo, cuando estas arquitecturas se componen por múltiples capas ocultas intermedias, se denominan DNN. Las DNN tienen la capacidad de aprender representaciones jerárquicas y no lineales de alta complejidad, lo que las hace aptas para problemas con gran dimensionalidad y estructuras de datos complejas. En aplicaciones como la detección de fallos o la identificación de patrones espaciales en sensores, las DNN permiten mejorar el rendimiento frente a modelos superficiales. Sin embargo, al carecer de mecanismos para modelar relaciones temporales, su

rendimiento puede verse superado por arquitecturas recurrentes como LSTM en tareas que implican series de tiempo.

Finalmente, en la Figura 14 se presenta el diagrama esquemático de una DNN usado en el trabajo de Shin et al. (2018). Este esquema representa la mayor complejidad alcanzada a través de los estudios de Shin et al. en 2018 y 2019 sin tener en cuenta los modelos que no emplean redes neuronales o son capaces de procesar información secuencial o con dependencia temporal como los modelos LSTM.

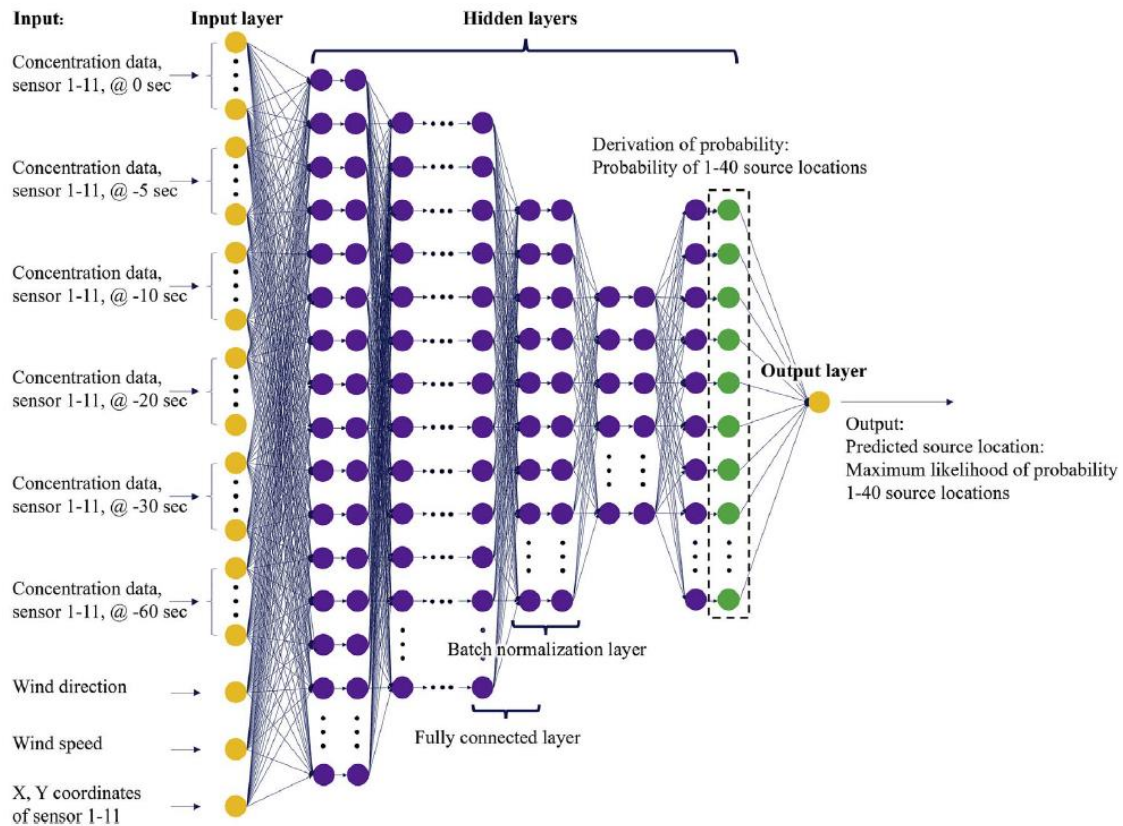


Figura 14. Diagrama de flujo del RF (Shin et al., 2018).

3.2.1. Deep Neural Network (DNN) & Random Forest (RF)

En el trabajo de Shin et al. (2018), se emplea un modelo basado en una Deep Neural Network (DNN) con 25 capas ocultas para abordar el problema de localización de fugas presentado previamente. El modelo se entrena con datos generados por simulaciones de CFD de 640 escenarios, integrando información de velocidad/dirección del viento y concentración química en los sensores. Cada conjunto de entrada representa un estado del sistema y es clasificado en una de las 40 ubicaciones potenciales de fuga. La arquitectura DNN busca aprender representaciones jerárquicas profundas, lo que permite una identificación de patrones de fuga. Sin embargo, al no incorporar explícitamente relaciones temporales ni mecanismos de memoria, su desempeño es limitado frente a arquitecturas recurrentes, alcanzando una precisión del 75.43%.

En el mismo estudio, como método alternativo basado en aprendizaje supervisado, aplican un clasificador RF para la misma tarea de localización de fugas. A continuación, se presenta en la Figura 15 el diagrama de flujo del modelo RF empleado en este estudio.

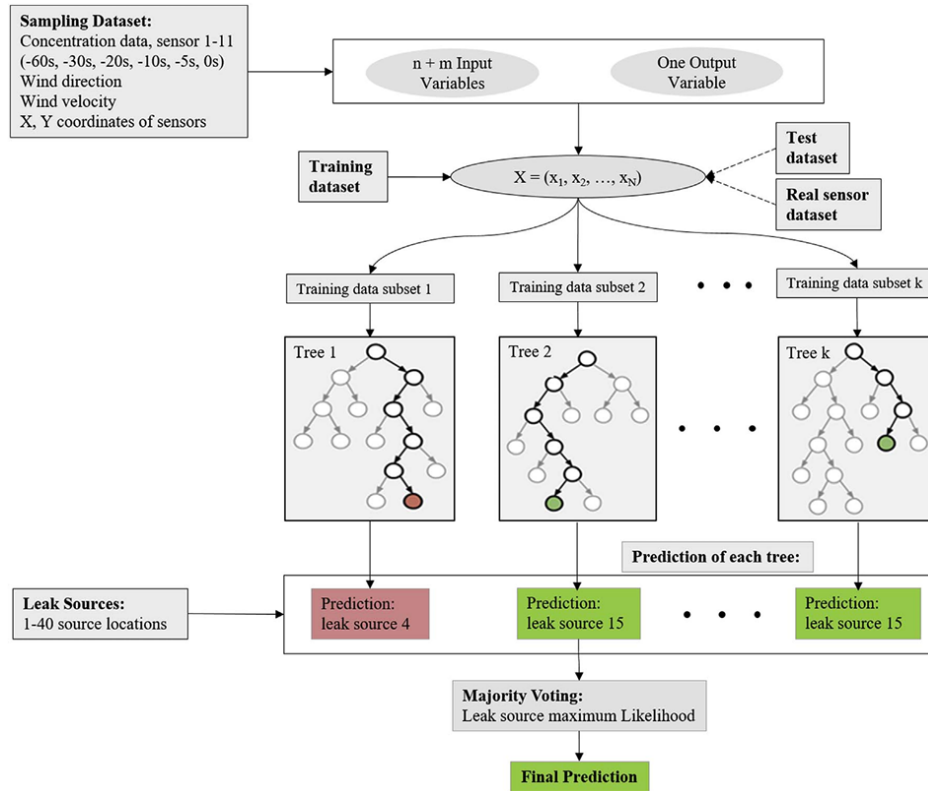


Figura 15. Diagrama de flujo del RF (Shin et al., 2018).

RF es un algoritmo no paramétrico que construye múltiples árboles de decisión sobre subconjuntos aleatorios de datos y vota por la clase más frecuente. En este caso, el modelo se entrena con 13 variables por muestra (concentración, velocidad y dirección del viento) y utiliza 100 árboles de decisión para clasificar cada estado en una de las 40 clases posibles. A pesar de su simplicidad estructural y de no requerir entrenamiento profundo, el modelo RF muestra una notable precisión del 86.33%, destacándose por su robustez frente al sobreajuste y su bajo costo computacional en comparación con las redes neuronales profundas. Este rendimiento lo posiciona como una alternativa eficiente y confiable en entornos con limitaciones de recursos o tiempo de cómputo.

3.2.2. Recurrent Neural Network with Long Short-Term Memory (LSTM) & Feed-forward neural network (FNN).

En el año 2019 Shin et al., continua su estudio de detección de fugas implementando un modelo basado en FNN como grupo de control para comparar su rendimiento con el modelo LSTM. Esta red recibe la misma cantidad de datos (780 variables por muestra: 13 variables $\times$ 60 s), pero los procesa sin considerar la secuencia temporal, ya que carece de retroalimentación. Aunque FNN logra identificar patrones espaciales en los datos simulados, su limitación para

captar relaciones temporales entre observaciones afecta su precisión. A pesar de configurarse con la misma profundidad que el modelo LSTM (6 capas ocultas), el FNN alcanza una precisión promedio 20.45% inferior a la del modelo LSTM en condiciones equivalentes, lo que resalta la ventaja del aprendizaje secuencial en este tipo de problemas.

En el mismo estudio, se entrena el modelo LSTM con datos de sensores virtuales generados por simulaciones de 640 escenarios de fuga en una planta química real, considerando 40 tanques y 16 direcciones de viento distintas. La arquitectura LSTM, diseñada para capturar patrones en series temporales, se implementa en forma multicapa (hasta 3 capas LSTM apiladas) para procesar secuencias de datos de 60 segundos, junto con concentraciones medidas por 11 sensores. Debido a su capacidad para retener información relevante y olvidar datos que pueden considerarse ruido, el modelo alcanza un desempeño superior, con una precisión del 97.08% en la localización del punto de fuga, destacándose por su capacidad para modelar dinámicas complejas de dispersión sin necesidad de sensores móviles.

3.3. Optimización

La optimización con apoyo de técnicas IA en procesos químicos abarca múltiples enfoques, dependiendo de la variable objetivo: desde el rendimiento de productos, eficiencia energética, hasta la reducción del tiempo de simulación o la precisión en predicciones complejas. A pesar de esta diversidad, los modelos analizados en esta sección comparten un objetivo común: mejorar la precisión y eficiencia en el modelado de sistemas complejos, evaluados mediante métricas cuantitativas como el coeficiente de determinación (R^2) y el error cuadrático medio (RMSE).

Los estudios seleccionados implementan algoritmos de aprendizaje automático incluyendo ANN, RF, *Gradient Boosting Regression* (GBR), *Gated Recurrent Unit* (GRU), redes bayesianas y AE variacionales aplicados a la predicción de propiedades fisicoquímicas, desempeño de materiales, procesos dinámicos y simulaciones de alta fidelidad. A continuación, se describen los enfoques implementados, los principios que los sustentan y los resultados alcanzados.

3.3.1. Artificial neural network (ANN)

En el trabajo de Kim et al. (2022), se propone una metodología basada en ANN para predecir y optimizar el desempeño catalítico en la producción de hidrógeno mediante la reacción de desplazamiento de agua-gas (WGSR). En lugar de emplear modelos cinéticos tradicionales, los autores desarrollan un modelo ANN entrenado con un conjunto de 419 datos experimentales extraídos de la literatura. Las entradas del modelo incluyen 34 descriptores catalíticos y operativos, mientras que la salida es la conversión de CO (como indicador de rendimiento). La arquitectura óptima de la red, determinada por búsqueda manual, está compuesta por cuatro capas ocultas (con 39, 15, 10 y 1 neuronas, respectivamente) y fue entrenada mediante *resilient backpropagation* y función de activación del tipo sigmoide. El modelo final alcanzó un coeficiente de determinación $R^2 = 0.997$, un error cuadrático medio (MSE) de 0.003 y una raíz

del error cuadrático medio (RMSE) de 0.058, lo que demuestra su capacidad para capturar relaciones no lineales complejas entre variables catalíticas y rendimiento del proceso.

3.3.2. *Random forest (RF) & Gradient Boosting Regression (GBR)*

En el estudio de Leng et al. (2023), el modelo RF fue empleado para predecir y optimizar simultáneamente el rendimiento, el área superficial específica (SSA) y el volumen total de poro (TPV) del biocarbón, en el marco de un esquema multiobjetivo (MT). El propósito es identificar las combinaciones óptimas de composiciones de biomasa y condiciones de pirólisis (temperatura, tiempo de residencia y velocidad de calentamiento) que permitan obtener biochar con propiedades adecuadas para aplicaciones en adsorción o almacenamiento de energía.

El modelo empleado para este fin fue RF, su salida es la media de todas las predicciones individuales, lo que permite manejar relaciones no lineales y datos heterogéneos con buena capacidad de generalización. El modelo RF alcanzó un coeficiente de determinación promedio R^2 de 0.90 y un RMSE de 5.67 en la predicción simultánea de los tres objetivos, demostrando un equilibrio adecuado entre precisión y robustez computacional, especialmente en contextos donde se manejan múltiples variables objetivo.

Como modelo alternativo, se utilizó GBR para entrenar predictores separados para cada objetivo único (ST), alcanzando un rendimiento superior al del modelo RF. GBR es un algoritmo de ensamblado secuencial que construye árboles de decisión de manera aditiva, donde cada nuevo árbol corrige los errores residuales del conjunto anterior. Esta estrategia lo hace especialmente eficaz para modelar relaciones complejas entre variables, como las existentes entre los parámetros del proceso y las propiedades fisicoquímicas del biocarbón. En el enfoque ST, el modelo GBR obtuvo un R^2 de 0.91 a 0.94 y un RMSE de 4.25, superando levemente al modelo RF tanto en precisión como en estabilidad. Además, el modelo permitió identificar variables como la temperatura de pirólisis, el contenido de materia volátil y las cenizas como los factores más influyentes en el diseño de biochar con alta SSA y TPV.

3.3.3. *Group Contribution models (GC) Gaussian process method (GP) Machine learning model (ML)*

El grupo de estudio de Alshehri et al. (2021), se propone una estrategia avanzada para la predicción de propiedades termofísicas puras de compuestos químicos basada en la integración de modelos de contribución de grupos *Group Contribution* (GC) con técnicas de aprendizaje automático ML. El enfoque GC parte del principio de que la propiedad de una molécula puede estimarse como la suma de las contribuciones individuales de los grupos funcionales que la componen, ponderadas por su cantidad en la estructura (Gmehling, Tiegs, & Knipp, 1990). Esta metodología ha sido ampliamente utilizada por su simplicidad y aplicabilidad en ausencia de datos experimentales, pero presenta limitaciones importantes al tratarse de un enfoque lineal que no captura interacciones complejas entre grupos, efectos de vecindad ni contribuciones de segundo o tercer orden.

Para superar estas limitaciones, los autores desarrollan un modelo híbrido denominado GC-ML, en el cual los vectores de entrada representan la frecuencia de aparición de 424 grupos funcionales y subestructuras moleculares codificadas mediante descriptores estructurales. Esta información es procesada por un modelo de regresión basado en *Gaussian Processes* (GP), una técnica bayesiana no paramétrica que permite modelar relaciones no lineales con alta precisión, al tiempo que cuantifica la incertidumbre asociada a cada predicción. A diferencia del modelo GC-simple, que utiliza directamente coeficientes ajustados por regresión lineal ordinaria, el modelo GC-ML aprende automáticamente patrones complejos a partir de datos experimentales, sin asumir una forma funcional predeterminada para la relación entre estructura y propiedad. Este modelo fue entrenado con una base de datos de más de 24,000 compuestos y más de 50,000 puntos de datos experimentales, abarcando un conjunto de 25 propiedades puras relevantes en simulación y diseño de procesos químicos, como entalpía de formación, punto de ebullición, densidad, constante dieléctrica, entre otras. Los resultados obtenidos muestran un desempeño significativamente superior frente al enfoque tradicional, alcanzando valores de $R^2 = 0.98$ y reduciendo el RMSE, validando su capacidad predictiva y generalización. En este contexto, el rol del ML no es reemplazar el método GC, sino potenciarlo, actuando como un sistema de aprendizaje flexible que corrige y complementa las deficiencias del modelo clásico, permitiendo capturar relaciones de mayor complejidad estructural sin perder la interpretabilidad de la representación basada en grupos.

3.3.4. Gated Recurrent Unit - Physics informed (GRU-PI)

En el estudio de Aydin & Asrav (2023), se propone un modelo híbrido denominado GRU-PI, que combina una arquitectura de red neuronal recurrente de tipo GRU con una metodología de entrenamiento *physics-informed* (PI), orientado a la optimización de la predicción de concentraciones en procesos dinámicos, como un reactor semi-batch o una planta de tratamiento de aguas residuales. A continuación, se presenta el diagrama de bloques empleado por Aydin & Asrav en la Figura 16.

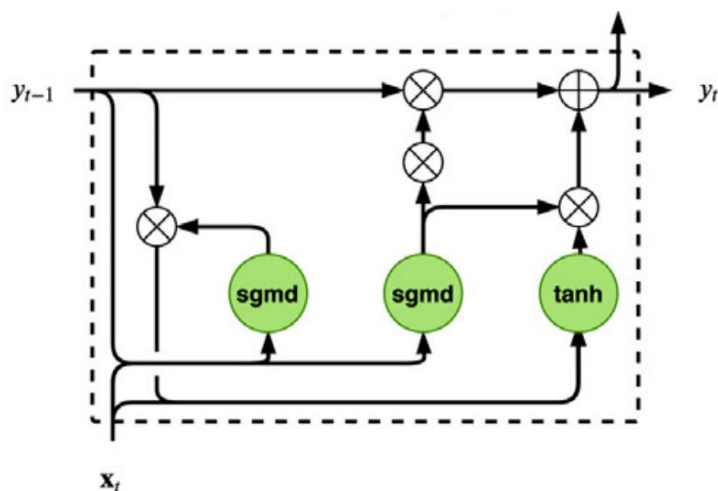


Figura 16. Diagrama de bloques de una célula del modelo GRU (Viana et al., 2021).

El objetivo principal es modelar el comportamiento dinámico de estos sistemas con alta precisión, especialmente cuando no se dispone de modelos mecánicos completos. Las GRU son una variante de las RNN tradicionales diseñadas para evitar el problema del desvanecimiento del gradiente de manera similar a como lo hacen los métodos LSTM, mediante el uso de puertas que regulan el flujo de información en el tiempo. A diferencia de LSTM, las GRU tienen una estructura más simple (con solo dos puertas: actualización y reinicio), lo que reduce la complejidad computacional sin comprometer la capacidad de capturar dependencias temporales de largo alcance.

La innovación del enfoque GRU-PI radica en la incorporación de conocimiento físico explícito durante el entrenamiento del modelo. Esto se logra mediante dos estrategias complementarias: Inicialmente, incluir una función de pérdida multiobjetivo que penaliza tanto el error de predicción como la violación de una forma discretizada de las ecuaciones diferenciales ordinarias (ODE) que gobiernan el sistema. Seguidamente, diseñar celdas recurrentes híbridas que integran directamente nodos PI (calculados, por ejemplo, mediante la regla de Euler) junto a nodos puramente *data-driven*. Esta arquitectura permite al modelo capturar tanto las dinámicas conocidas como aquellas que son difíciles de modelar explícitamente. En las pruebas realizadas, el modelo GRU-PI obtuvo un desempeño destacado con $R^2 = 0.979$, $MSE = 0.005$ y $RMSE = 0.51$, superando consistentemente a modelos sin información física. Adicionalmente, los autores muestran que este enfoque permite obtener arquitecturas más robustas y compactas mediante técnicas de optimización de hiperparámetros, como búsqueda en malla, optimización bayesiana y algoritmos genéticos, logrando resultados similares con menor número de nodos por capa.

3.3.5. *Convolutional-Variational Autoencoder-Variational Bayesian neural network (Conv-VAE-VBnn)*

En el trabajo de Shi et al. (2023), se propone un modelo híbrido de aprendizaje profundo denominado Conv-VAE-VBnn para la predicción en tiempo real de la dispersión espacial de nubes inflamables tras una liberación de gas natural en instalaciones industriales congestionadas. El objetivo es reducir drásticamente la carga computacional asociada a simulaciones de alta fidelidad, como las basadas en CFD, sin comprometer la precisión en la estimación del área afectada por la pluma gaseosa. El modelo se basa en una arquitectura *end-to-end* que combina tres componentes: un AE convolucional probabilístico (Conv-VAE) para reducir dimensionalidad y extraer características latentes del espacio de concentración, una red bayesiana variacional (VBnn) para aprender la relación no lineal entre parámetros del escenario (como tasa de fuga y velocidad del viento) y las variables latentes, finalmente una etapa de decodificación para reconstruir la concentración espacial con estimación de incertidumbre.

Complementariamente, el modelo se entrena usando una estrategia de transferencia de aprendizaje basada en fidelidad variable, en la que primero se preentrena con grandes volúmenes de datos sintéticos generados mediante el modelo de dispersión gaussiana (baja fidelidad), y luego se ajusta con un número reducido de simulaciones CFD (alta fidelidad)

usando un esquema de *fine-tuning* por capas congeladas. Esta metodología permitió una reducción del 72% en el número de simulaciones de alta fidelidad requeridas, manteniendo una precisión elevada ($R^2 = 0.96$, $RMSE = 6.52 \times 10^{-5}$) en la predicción de la concentración espacial. Los autores argumentan que esta aproximación es especialmente útil para entornos industriales a gran escala, donde los modelos convencionales carecen de capacidad de respuesta en tiempo real y el coste computacional impide su uso continuo. Además, el modelo ofrece estimaciones probabilísticas de la incertidumbre asociada a la predicción, lo cual representa una mejora significativa frente a enfoques deterministas que podrían subestimar riesgos en zonas críticas.

4. Discusión

La estructuración de esta sección corresponde a una tendencia identificada en la literatura, la cual divide las aplicaciones de IA en los procesos químicos según su propósito operacional. Esta clasificación permitió integrar los indicadores de desempeño más relevantes *FDR*, *Accuracy* y R^2 con el fin de facilitar una comparación objetiva entre enfoques y tecnologías.

Inicialmente, la categoría de control de procesos se cimenta en el uso sistemático del proceso TEP como benchmark. Esta plataforma ampliamente utilizada, proporciona condiciones de evaluación homogéneas que posibilitan una comparación directa del desempeño entre modelos, lo que le otorga un carácter normativo en la validación de metodologías inteligentes (Bathelt et al., 2015).

En este contexto, es importante destacar el trabajo de Wu et al. (2023) por reflejar el potencial de estos enfoques en escenarios reales más allá del entorno simulado. Este grupo de estudio implementa una arquitectura compleja basada en control predictivo asistido por IA en un reactor CSTR. Este modelo, que integra LSTM, optimización multiobjetivo y métodos de toma de decisiones multicriterio, representa un avance sustancial hacia el diseño de sistemas de control predictivos y adaptativos con resultados ampliamente positivos en cuanto a estabilidad, ahorro energético y eficiencia operativa.

De la misma manera, es notable el modelo propuesto por Zhao y Zheng (2022), que introduce una metodología robusta de aprendizaje semi-supervisado capaz de operar en contextos con pocos datos etiquetados, una limitación común en la industria (descrita por el mismo Zhao en 2022). Este modelo no solo presenta un rendimiento sobresaliente frente a otros métodos PU y supervisados, sino que enfatiza la necesidad de diseñar soluciones viables para entornos industriales reales, donde los datos etiquetados son costosos o difíciles de obtener. La reiteración del uso del benchmark TEP en la mayoría de los artículos de esta categoría confirma tanto su valor como estándar como la escasa disponibilidad de casos industriales con datos accesibles, situación que limita la generalización inmediata de muchos de estos modelos.

A continuación, se presenta la tabla 2, la cual compila los modelos presentados para el control de procesos jerarquizados por la magnitud de su indicador de rendimiento.

Tabla 2.

Listado de modelos presentados para el control de procesos ordenados de mayor a menor magnitud de indicador.

Modelo	Componentes / Arquitectura	FDR	FAR	Latencia	Ventaja principal
ML-aided MPC-MOO-MCD M (Wu, 2023)	LSTM predictor + NSGA-II (MOO) + CRITIC + PROBIT + WST	-	-	0.12 h (12 iter.)	Optimización dinámica de múltiples objetivos con estabilidad garantizada; control predictivo adaptable con bajo esfuerzo computacional.
ETA-OASPCA (Dong et al., 2023)	Enhanced Temporal Algorithm + Optimized Adaptive Sparse PCA	99.4%	0.37%	4.11 min	Integra patrones temporales y reducción adaptativa de dimensionalidad; bajo FAR y buen aislamiento de variables contribuyentes.
VAE-LSTM (Zhao et al., 2022)	Variational Autoencoder + LSTM bidireccional + Inferencia bayesiana	95.5%	-	-	Captura estructura latente y dinámica temporal; útil para procesos multivariados complejos; buena generalización.
THPU (Zhao & Zheng, 2022)	Self-Training Nearest Neighbors + SAE + CNN + PU learning	92.31%	2.81%	11.0 min	Funciona con muy pocas etiquetas; combina autoetiquetado, aprendizaje profundo y clasificación binaria robusta.
LAS-VB (Zhao et al., 2020)	Local Adaptive Standardization + Variational Autoencoder + BiLSTM	>90% (hasta 99.8%)	~0.3%	-	Capaz de detectar fallos en modos no vistos; robusto ante datos inestables; sensible a desviaciones suaves.
DCNN (Zhao et al., 2018)	Deep Convolutional Neural Network	88.2%	-	-	Aprende patrones espacio-temporales sin necesidad de ingeniería de características; rápido y escalable.

La categoría de seguridad presenta estudios realizados por Shin (2018, 2019) los cuales evidencian cómo la IA permite optimizar la localización de fugas mediante el uso de sensores fijos y simulaciones CFD, eliminando la necesidad de soluciones invasivas o de alto costo. La tabla 3, presenta la compilación de modelos presentados por Shin et al. (2018, 2019) ordenados de manera descendiente en base a su indicador de rendimiento.

Tabla 3.

Listado de modelos presentados para la predicción de fugas en procesos químicos ordenados de mayor a menor magnitud de indicador.

Modelo	Componentes / Arquitectura	Entrada de datos	Accuracy	Ventaja principal
LSTM (Shin, 2019)	Red Neuronal Recurrente con memoria a largo plazo	780 variables (13 × 60s) incluyendo viento y concentraciones	97.08%	Captura dependencias temporales, ideal para datos secuenciales; alta precisión
Random Forest (Shin, 2018)	Ensamble de árboles de decisión (100 árboles)	13 variables por muestra (sensores + viento)	86.33%	Robusto, interpretable y eficiente; buena precisión sin redes profundas
FNN (Shin, 2019)	Red Neuronal Alimentada Hacia Adelante (6 capas)	Mismo conjunto de entrada que LSTM	76.63%	Simplicidad y menor coste computacional, pero limitada por falta de contexto temporal
DNN (Shin, 2018)	Red Neuronal Profunda con hasta 25 capas	Mismas 13 variables que RF	75.43%	Buen aprendizaje de representaciones, pero sin manejo de series temporales

Esta categoría evidencia una marcada preferencia por los métodos con enfoques supervisados, esto se debe principalmente a la disponibilidad de datos etiquetados proporcionados por los entornos de simulación, así como también sucede en la categoría de control con el benchmark TEP. Esta tendencia se ve reflejada en la Figura 17, donde se muestra el análisis de recurrencia en el enfoque de los modelos vistos en las tres categorías de este trabajo.

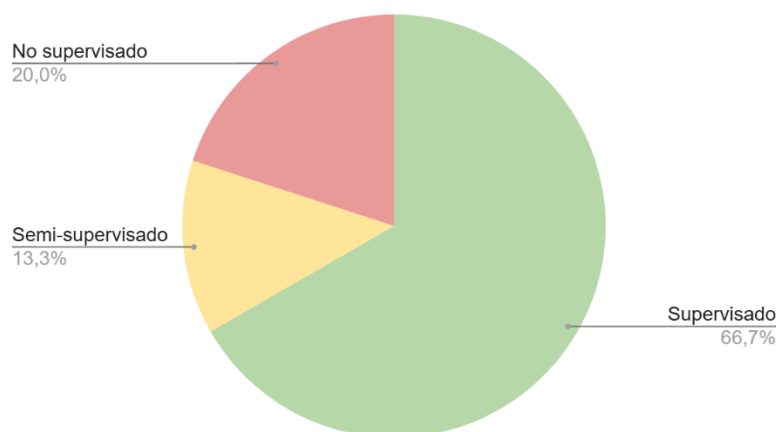


Figura 17. Distribución de enfoque de modelos en los artículos estudiados.

La categoría de optimización representa un reto metodológico mayor. A diferencia del control o la seguridad, en este grupo los modelos, aunque comparten métricas como R^2 o RMSE, están diseñados para optimizar objetivos diferentes: desde propiedades físico-químicas hasta rendimiento energético o parámetros de diseño de reactores. Esta dispersión complica la comparación directa entre estudios, aunque no impide identificar tendencias relevantes. La Tabla 4, presenta los modelos vistos en este estudio pertenecientes a esta categoría.

Tabla 4.

Listado de modelos presentados para la optimización de parámetros de operación ordenados de mayor a menor magnitud de indicador.

Modelo	Componentes / Arquitectura	Objetivo de optimización	R^2	RMSE	Ventaja principal
ANN (Kim, 2022)	Red neuronal artificial	Predicción del rendimiento catalítico (conversión de CO)	0.99	0.058	Alta precisión y robustez con múltiples variables de entrada en procesos catalíticos complejos
GRU-PI (Aydin, 2023)	GRU + celdas physics-informed	Predicción dinámica de concentración en procesos no lineales	0.98	0.51	Integra conocimiento físico en la arquitectura recurrente; mejora precisión y estabilidad
GC-ML (Alshehri, 2021)	Gaussian Process sobre vectores de GC	Estimación de propiedades moleculares puras	~ 0.98	6.96	Mejora sobre GC clásico; permite predicción con incertidumbre en diseño molecular
Conv-VAE-VBnn (Shi, 2023)	Autoencoder convolucional + red bayesiana variacional	Reducción de simulaciones CFD en dispersión de gases	0.96	6.52×10^{-5}	Reduce 72% del costo computacional; alta precisión con manejo de incertidumbre
GBR (Leng, 2023)	Gradient Boosting (single-target)	Predicción individual de propiedades del biocarbón	0.91–0.94	4.25	Mayor precisión que RF; captura relaciones complejas entre parámetros del proceso y propiedades
RF (Leng, 2023)	Random Forest (multi-target)	Predicción de SSA y TPV de biocarbón	0.90	5.67	Manejo robusto de datos heterogéneos con capacidad para modelar múltiples salidas

Leng (2023) demuestra como los modelos RF y GBR no sólo predicen con precisión las propiedades del biocarbón, sino que también orientan directamente el diseño experimental, lo cual sugiere los métodos de IA como una herramienta guía en el diseño de reactores. De la misma manera, otro trabajo de especial interés es el de Aydin & Asrav (2023), donde se propone el modelo Conv-VAE-VBnn, el cual logra combinar eficiencia computacional, alta precisión y manejo de incertidumbre. Esta arquitectura permite una reducción sustancial del número de simulaciones de alta fidelidad necesarias para mantener una calidad de resultados,

lo que la posiciona como una solución práctica para contextos industriales donde el costo computacional es un factor limitante.

Cabe señalar que, a lo largo de los artículos analizados, se observa una tendencia creciente hacia modelos más interpretables y comprensibles para el operador humano, como lo señala Zhao et al. (2022). Esta característica resulta crítica, ya que promueve la apropiación de la tecnología y reduce la probabilidad de errores derivados del desconocimiento o la desconfianza frente a herramientas automatizadas. Esta observación adquiere aún más relevancia cuando se considera la necesidad de interacción fluida entre sistemas inteligentes y operadores en entornos industriales complejos.

Adicionalmente, tal como se muestra en la Figura 18, se identifica una alta recurrencia en el uso de RNN y sus derivados como LSTM y GRU dentro de los modelos aplicados a sistemas dinámicos en procesos químicos.

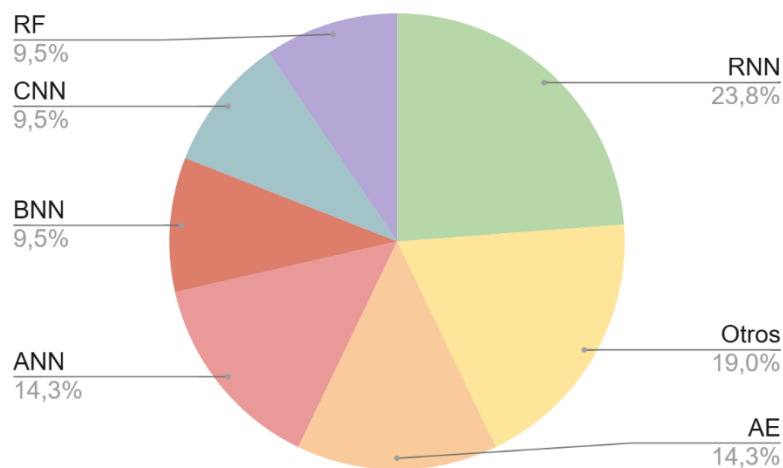


Figura 18. Distribución de frecuencia por componente dentro de los modelos presentados.

Esta tendencia se debe a que dichas arquitecturas están especialmente diseñadas para capturar dependencias temporales y patrones secuenciales en datos, lo cual resulta esencial en procesos industriales caracterizados por señales multivariadas y comportamiento dinámico. Su capacidad para mantener la memoria de estados a largo plazo les confiere una ventaja significativa frente a otras arquitecturas más tradicionales, particularmente en tareas de monitoreo, detección de fallas y predicción de comportamiento dinámico. En este contexto, trabajos recientes, como los de Aydin & Asrav (2023) y Bracconi (2022), han comenzado a explorar la integración de enfoques PI en estas arquitecturas, demostrando que la incorporación explícita de conocimiento físico en los modelos RNN permite mejorar la generalización en condiciones no vistas y también reducir el sobreajuste y aumentar la robustez frente a incertidumbres estructurales.

Aydin & Asrav (2023) proponen el uso de funciones de pérdida multiobjetivo que integran directamente las ecuaciones diferenciales del sistema, así como celdas híbridas que combinan nodos físicos y *data-driven* dentro de redes LSTM o GRU, logrando una mejora sustancial del

rendimiento predictivo en escenarios industriales simulados. De forma complementaria, Bracconi (2022) enfatiza que los modelos híbridos, entrenados con datos y reforzados con leyes físicas, pueden asistir en el diseño multiescala de reactores intensificados, ofreciendo predicciones confiables con menores requerimientos computacionales. Estas aproximaciones sugieren que el uso de variantes de RNN con capas PI no solo representa una línea de avance actual, sino que podría convertirse en una de las principales tendencias en el desarrollo de modelos predictivos robustos y explicables en la ingeniería de procesos.

En síntesis, la IA no solo ha demostrado ventajas cuantitativas en términos de precisión, velocidad y eficiencia, sino que también aporta beneficios cualitativos sustanciales, tales como la reducción de la dependencia del conocimiento experto codificado, la posibilidad de operar bajo condiciones dinámicas no vistas y la capacidad de ofrecer predicciones interpretables con manejo razonable de la incertidumbre. Estos hallazgos reflejan el notable potencial transformador de la IA en los procesos químicos, consolidando su papel como una tecnología clave en el tránsito hacia plantas más autónomas, eficientes y seguras.

5. Conclusiones

La presente revisión permitió a través de un estudio bibliométrico identificar, clasificar y analizar una serie de métodos IA aplicados a procesos químicos, centrándose en tres áreas fundamentales: control de procesos, seguridad de procesos y optimización. A través del análisis comparativo de indicadores de rendimiento como FDR, *Accuracy* y R^2 .

Uno de los hallazgos más relevantes del estudio es la identificación de tendencias en estas áreas de estudio. Es el caso del benchmark TEP, el cual representa un entorno estándar para la validación de modelos de detección y diagnóstico de fallas. Este benchmark ha facilitado el desarrollo de modelos reproducibles, robustos y comparables. No obstante, este benchmark evidencia una limitación importante: la escasez de bases de datos industriales etiquetadas y accesibles, lo cual restringe la generalización de los modelos a escenarios reales. En este contexto, destacan propuestas como las de Wu et al. (2023) y Zhao y Zheng (2022), quienes incorporan arquitecturas complejas y enfoques de aprendizaje semi-supervisado o autoaprendizaje para abordar este desafío con éxito en entornos simulados y reales.

También, se observó una preferencia hacia los modelos de enfoque supervisado, presentes a lo largo de toda la categoría de seguridad de procesos, esto debido al aprovechamiento de simulaciones CFD como fuentes de datos etiquetados, lo cual permite entrenar modelos precisos sin comprometer la operación industrial.

Las arquitecturas LSTM y GRU han demostrado ser altamente eficaces a través de las categorías abarcadas en este estudio, gracias a su capacidad de modelar secuencias temporales multivariadas. Esta capacidad se beneficia significativamente al integrarse con otros modelos de enfoque no supervisados como lo son las variantes de los AE, generando arquitecturas híbridas con una capacidad mejorada para manejar variables con retroalimentación. En

múltiples casos revisados, esta combinación ha permitido fortalecer el desempeño predictivo y la estabilidad de los modelos, posicionándolas como un estándar emergente.

Los modelos IA aplicados a la optimización evidenciaron una amplia diversidad de propósitos. Aunque esta diversificación de estudios dificulta la comparación directa, es evidente que los modelos híbridos y aquellos con arquitectura PI, representan un gran potencial de mejora a los modelos tradicionales, particularmente en tareas donde las relaciones causales son parcialmente conocidas pero complejas de modelar con ecuaciones explícitas.

De manera transversal, se identificó una tendencia creciente hacia la interpretabilidad de los modelos, orientada a facilitar su adopción por parte de operadores humanos. Esta característica, sumada a la integración de enfoques híbridos, representa un avance hacia sistemas más robustos, explicables y confiables. Estudios como los realizados por Aydin & Asray (2023) y Bracconi (2022), evidencian cómo esta combinación puede llevar a modelos capaces de servir como asistentes de diseño de reactores, trasladando así el alcance de estas herramientas a diferentes etapas de los procesos químicos.

En conjunto, los hallazgos de esta revisión confirman que la inteligencia artificial tiene el potencial de redefinir profundamente las operaciones en plantas químicas, no solo desde una perspectiva de rendimiento técnico, sino también como catalizador de una transformación digital centrada en la autonomía, la eficiencia y la resiliencia operativa. El avance hacia modelos más accesibles, interpretables y adaptativos constituye un paso necesario para su integración efectiva en la industria, abriendo la puerta a nuevos escenarios de diseño, operación y toma de decisiones inteligentes en los procesos químicos.

6. Referencias

Alshehri, A. S., Tula, A. K., You, F., & Gani, R. (2022). Next generation pure component property estimation models: With and without machine learning techniques. *AIChE Journal*, 68(6). <https://doi.org/10.1002/aic.17469>

Asrav, T., & Aydin, E. (2023). Physics-informed recurrent neural networks and hyperparameter optimization for dynamic process systems. *Computers and Chemical Engineering*, 173. <https://doi.org/10.1016/j.compchemeng.2023.108195>

Bathelt, A., Ricker, N. L., & Jelali, M. (2015). Revision of the Tennessee eastman process model. *IFAC-PapersOnLine*, 28(8), 309–314. <https://doi.org/10.1016/j.ifacol.2015.08.199>

Bi, X., Qin, R., Wu, D., Zheng, S., & Zhao, J. (2022). One step forward for smart chemical process fault detection and diagnosis. *Computers and Chemical Engineering*, 164. <https://doi.org/10.1016/j.compchemeng.2022.107884>

Bracconi, M. (2022). Intensification of catalytic reactors: A synergic effort of Multiscale Modeling, Machine Learning and Additive Manufacturing. *Chemical Engineering and Processing - Process Intensification*, 181. <https://doi.org/10.1016/j.cep.2022.109148>

Cho, J., Kim, H., Gebreselassie, A. L., & Shin, D. (2018). Deep neural network and random forest classifier for source tracking of chemical leaks using fence monitoring data. *Journal of Loss Prevention in the Process Industries*, 56, 548–558. <https://doi.org/10.1016/j.jlp.2018.01.011>

Downs, J. J., & Vogel, E. F. (1993). A PLANT-WIDE INDUSTRIAL PROCESS PROBLEM CONTROL (Vol. 17, Issue 3).

Elsevier. (s.f.-a). *Analyze search results: TITLE-ABS-KEY (("IA" OR "ARTIFICIAL INTELLIGENCE") AND ("CHEMICAL ENGINEERING" OR "PROCESS CONTROL")) AND PUBYEAR > 2013 AND PUBYEAR < 2025 AND (LIMIT-TO (SUBJAREA , "ENGI")) AND (LIMIT-TO (DOCTYPE , "ar")) AND (LIMIT-TO (LANGUAGE , "English")) AND (EXCLUDE (OA , "all"))*. Scopus. [https://www-scopus-com.bd.univalle.edu.co/term/analyzer.uri?sort=plf-f&src=s&sid=43d4bdc9b13e5be4d89e52c29f52f511&sot=a&sdt=a&cluster=scosubjabbr%2c"ENGI"%2ct%2bscosubtype%2c"ar"%2ct%2bscolang%2c"English"%2ct%2bscofreetoread%2c"all"%2cf&sl=142&s=TITLE-ABS-KEY%28%28"IA"+OR+"ARTIFICIAL+INTELLIGENCE"+%29+AND+%28"CHEMICAL+ENGINEERING"+OR+"PROCESS+CONTROL"+%29%29+AND+PUBYEAR+>+2013+AND+PUBYEAR+<+2025&origin=resultslist&count=10&analyzeResults=Analyze+results](https://www-scopus-com.bd.univalle.edu.co/term/analyzer.uri?sort=plf-f&src=s&sid=43d4bdc9b13e5be4d89e52c29f52f511&sot=a&sdt=a&cluster=scosubjabbr%2c)

Elsevier. (s.f.-b). *Analyze search results: TITLE-ABS-KEY (("IA" OR "ARTIFICIAL INTELLIGENCE" OR "MACHINE LEARNING" OR "DEEP LEARNING" OR "MACHINE-LEARNING") AND ("CHEMICAL ENGINEERING" OR "PROCESS ENGINEERING" OR "CHEMICAL PLANT" OR "CHEMICAL INDUSTRY")) AND PUBYEAR > 2013 AND PUBYEAR < 2025 AND (LIMIT-TO (SUBJAREA , "ENGI")) AND (LIMIT-TO (DOCTYPE , "ar")) AND (LIMIT-TO (PUBSTAGE , "final")) AND (LIMIT-TO (LANGUAGE , "English")) AND (EXCLUDE (OA , "all"))*. Scopus. [https://www-scopus-com.bd.univalle.edu.co/term/analyzer.uri?sort=plf-f&src=s&sid=0f72ada9d94239f96fabec06428bac43&sot=a&sdt=a&cluster=scosubjabbr%2c"ENGI"%2ct%2bscosubtype%2c"ar"%2ct%2bscopubstage%2c"final"%2ct%2bscolang%2c"English"%2ct%2bscofreetoread%2c"all"%2cf&sl=248&s=TITLE-ABS-KEY%28%28"IA"+OR+"ARTIFICIAL+INTELLIGENCE"+OR+"MACHINE+LEARNING"+OR+"DEEP+LEARNING"+OR+"MACHINE-LEARNING"%29+AND+%28"CHEMICAL+ENGINEERING"+OR+"PROCESS+ENGINEERING"+OR+"CHEMICAL+PLANT"+OR+"CHEMICAL+INDUSTRY"%29%29+AND+PUBYEAR+>+2013+AND+PUBYEAR+<+2025&origin=resultslist&count=10&analyzeResults=Analyze+results](https://www-scopus-com.bd.univalle.edu.co/term/analyzer.uri?sort=plf-f&src=s&sid=0f72ada9d94239f96fabec06428bac43&sot=a&sdt=a&cluster=scosubjabbr%2c)

Gmehling, J., Tiegs, D., & Knipp, U. (1990). A comparison of the predictive capability of different group contribution methods. *Fluid Phase Equilibria*, 54(C), 147–165. [https://doi.org/10.1016/0378-3812\(90\)85077-N](https://doi.org/10.1016/0378-3812(90)85077-N)

Hajjar, Z., Tayyebi, S., & Ahmadi, M. H. E. (2018). Application of AI in Chemical Engineering. In *Artificial Intelligence - Emerging Trends and Applications*. InTech. <https://doi.org/10.5772/intechopen.76027>

Kim, C., Won, W., & Kim, J. (2022). Early-Stage Evaluation of Catalyst Using Machine Learning Based Modeling and Simulation of Catalytic Systems: Hydrogen Production via

Water-Gas Shift over Pt Catalysts. ACS Sustainable Chemistry and Engineering, 10(44), 14417–14432. <https://doi.org/10.1021/acssuschemeng.2c03136>

Kim, H., Park, M., Kim, C. W., & Shin, D. (2019). Source localization for hazardous material release in an outdoor chemical plant via a combination of LSTM-RNN and CFD simulation. Computers and Chemical Engineering, 125, 476–489. <https://doi.org/10.1016/j.compchemeng.2019.03.012>

Lee, J. H. (2011). Model predictive control: Review of the three decades of development. In International Journal of Control, Automation and Systems (Vol. 9, Issue 3, pp. 415–424). <https://doi.org/10.1007/s12555-011-0300-6>

Lee, J. H., Shin, J., & Realff, M. J. (2018). Machine learning: Overview of the recent progresses and implications for the process systems engineering field. Computers and Chemical Engineering, 114, 111–121. <https://doi.org/10.1016/j.compchemeng.2017.10.008>

Li, H., Ai, Z., Yang, L., Zhang, W., Yang, Z., Peng, H., & Leng, L. (2023). Machine learning assisted predicting and engineering specific surface area and total pore volume of biochar. Bioresource Technology, 369. <https://doi.org/10.1016/j.biortech.2022.128417>

Reinartz, C., Kulahci, M., & Ravn, O. (2021). An extended Tennessee Eastman simulation dataset for fault-detection and decision support systems. Computers and Chemical Engineering, 149. <https://doi.org/10.1016/j.compchemeng.2021.107281>

Shi, J., Xie, W., Li, J., Zhang, X., Huang, X., Usmani, A. S., Khan, F., & Chen, G. (2023). Real-time plume tracking using transfer learning approach. Computers and Chemical Engineering, 172. <https://doi.org/10.1016/j.compchemeng.2023.108172>

Thon, C., Finke, B., Kwade, A., & Schilde, C. (2021). Artificial Intelligence in Process Engineering. Advanced Intelligent Systems, 3(6). <https://doi.org/10.1002/aisy.202000261>

Venkatasubramanian, V. (2019). The promise of artificial intelligence in chemical engineering: Is it here, finally? AIChE Journal, 65(2), 466–478. <https://doi.org/10.1002/aic.16489>

Viana, F. A. C., Nascimento, R. G., Dourado, A., & Yucesan, Y. A. (2021). Estimating model inadequacy in ordinary differential equations with physics-informed neural networks. Computers and Structures, 245. <https://doi.org/10.1016/j.compstruc.2020.106458>

Wang, Z., Tan, W. G. Y., Rangaiah, G. P., & Wu, Z. (2023). Machine learning aided model predictive control with multi-objective optimization and multi-criteria decision making. Computers and Chemical Engineering, 179. <https://doi.org/10.1016/j.compchemeng.2023.108414>

Wu, H., & Zhao, J. (2018). Deep convolutional neural network model based chemical process fault diagnosis. Computers and Chemical Engineering, 115, 185–197. <https://doi.org/10.1016/j.compchemeng.2018.04.009>

Wu, H., & Zhao, J. (2020). Self-adaptive deep learning for multimode process monitoring. Computers and Chemical Engineering, 141.

<https://doi.org/10.1016/j.compchemeng.2020.107024>

Zhang, J., Dai, Y., Feng, Z., & Dong, L. (2023). An enhanced temporal algorithm- coupled optimized adaptive sparse principal component analysis methodology for fault diagnosis of chemical processes. *Process Safety and Environmental Protection*, 174, 663–680. <https://doi.org/10.1016/j.psep.2023.04.036>

Zheng, S., & Zhao, J. (2022). High-fidelity positive-unlabeled deep learning for semi-supervised fault detection of chemical processes. *Process Safety and Environmental Protection*, 165, 191–204. <https://doi.org/10.1016/j.psep.2022.06.058>