

Comparación de métodos estadísticos y aprendizaje automático para la mejora del proceso de Slotting

Comparison of statistical methods and machine learning for Slotting process improvement

Agusto Guillermo Sánchez Ferrer

<https://orcid.org/0000-0003-3643-3046>
agusto.sanchez@unmsm.edu.pe

Jose Herrera Quispe

<https://orcid.org/0000-0002-8207-9714>
jherreraqu@unmsm.edu.pe

Universidad Nacional Mayor de San Marcos, Lima, Perú

RECIBIDO: 25/03/2024 - ACEPTADO: 15/06/2024 - PUBLICADO: 31/07/2024

RESUMEN

Los centros de distribución o almacenes de suministros enfrentan desafíos significativos en el proceso de selección de pedidos. Un problema particular es la congestión de operadores en la zona de selección debido a la distribución inexacta de productos en ubicaciones limitadas. Para abordar este problema, ha surgido el proceso de "Slotting", que tiene como objetivo asignar eficientemente productos a ubicaciones específicas utilizando métodos analíticos, mejorando la preparación de pedidos al mismo tiempo que se reducen los costos operativos. Sin embargo, las variaciones estacionales y las fluctuaciones en las tendencias contribuyen a las oscilaciones en la demanda de productos. Por lo tanto, se sugiere integrar métodos estadísticos de pronóstico y aprendizaje automático para monitorear y anticipar cambios de productos en zonas de selección. En este contexto, esta investigación presenta dos enfoques. El primero es proponer un nuevo modelo de proceso de "Slotting" basado en la revisión de la literatura, y el segundo es una comparación entre métodos como el suavizado exponencial doble y triple; el modelo estacional autorregresivo integrado de media móvil; redes neuronales recurrentes Long Short-Term Memory y Gated Recurrent Unit, para pronosticar la demanda de productos. Para validar los métodos, se realizó una prueba con 50 productos y una base de datos histórica de 7 años, desde enero de 2016 hasta mayo de 2023. Esta investigación sugiere un enfoque híbrido en el que los métodos utilizados pueden ser una estrategia práctica para lograr la mejor precisión en el pronóstico, optimizando así el proceso de "Slotting".

Palabras clave: Slotting, Pronóstico, Asignación de Producto, Redes Neuronales Recurrentes, SARIMA, Suavizado Exponencial.

ABSTRACT

Distribution centers or supply warehouses face significant challenges in the order-picking process. A particular issue is the congestion of operators in the picking zone due to the inaccurate distribution of products in limited locations. To address this issue, "Slotting" process has emerged, aiming to efficiently allocate products to specific locations using analytical methods, enhancing order preparation while concurrently reducing operational costs. However, seasonal variations and fluctuations in trends contribute to oscillations in product demand. Therefore, it is suggested to integrate statistical forecasting methods and machine learning to monitor and anticipate changes of products in picking zones. In this context, this research presents two approaches, the first is to propose a new "Slotting" process model based on the literature review and the second is a comparison between methods such as double and triple exponential smoothing; Seasonal Autoregressive Integrated Moving Average; Long Short-Term Memory and Gated Recurrent Unit were proposed to forecast product demand. To validate the methods, a test was carried out with 50 products and a 7-year historical database from January 2016 to May 2023. This research suggests a hybrid approach in which the methods utilized can be a practical strategy to achieve the best forecasting accuracy, optimizing the Slotting process.

Keywords: Slotting, Forecasting, Product Allocation, Recurrent Neural Networks, SARIMA, Exponential Smoothing.

I. INTRODUCCIÓN

En los últimos años, las empresas de cadena de suministro que gestionan centros de distribución han enfrentado desafíos logísticos en el proceso de selección de pedidos para sus clientes y sucursales. La selección de pedidos es el proceso más crítico, representando entre el 50% y el 65% del costo operativo total (Jaramillo, Camilo, Cuellar Molina, & Cogollo Flórez, 2020). Uno de los desafíos es el Problema de Asignación de Ubicación de Stock (SLAP), que implica asignar productos a ubicaciones dentro del almacén y es particularmente desafiante cuando se manejan numerosos productos dentro de espacios limitados (Viveros, González, Mena, Kristjanpoller, & Robledo, 2021). En respuesta al desafío SLAP, han surgido prácticas innovadoras, todas orientadas a reducir los costos de selección de pedidos y mejorar la productividad operativa (Fontana, Leyva López, Virgínio Cavalcante, & Solano Noriega, 2020). El concepto de "Slotting" se enfoca principalmente en optimizar la disposición de productos dentro de las zonas de (Jaramillo, Camilo, Cuellar Molina, & Cogollo Flórez, 2020). Para lograr esto, se toman decisiones estratégicas sobre la categorización de productos y se determina el número necesario de ubicaciones para satisfacer las demandas operativas (Viveros, González, Mena, Kristjanpoller, & Robledo, 2021).

Varios estudios ofrecen recomendaciones valiosas para implementar el "Slotting", enfatizando la utilización eficiente del espacio en almacenes a través de diversas estrategias y métodos analíticos para la categorización de productos (Jaramillo, Camilo, Cuellar Molina, & Cogollo Flórez, 2020). Sin embargo, los cambios estacionales y las variaciones en las tendencias de la demanda de productos requieren herramientas analíticas sólidas para garantizar cálculos unitarios precisos (Papcun, y otros, 2019). La estadística y el aprendizaje automático

han ganado prominencia en el mundo empresarial debido a su capacidad para construir modelos predictivos precisos que anticipan eventos futuros en varios sectores empresariales; técnicas como el Suavizamiento Exponencial, el Modelo Autorregresivo Integrado de Media Móvil Estacional y las Redes Neuronales Recurrentes tienen el potencial de revolucionar la toma de decisiones y mejorar los resultados comerciales al descubrir patrones y variaciones en series temporales. Sin embargo, la efectividad de las predicciones depende de la metodología apropiada para comprender el negocio y del análisis de la calidad de los datos disponibles (Chicco, Warrens, & Jurman, 2021).

Este artículo proporciona una comparación exhaustiva entre métodos estadísticos y técnicas de aprendizaje automático, logrando pronósticos precisos para la demanda de productos. Además, se introduce un modelo optimizado de "Slotting" que facilitó la extracción y procesamiento de datos para tomar decisiones de asignación de productos en la zona de selección. Se llevó a cabo un escenario experimental utilizando un conjunto de datos de 50 productos durante 72 períodos mensuales. Los modelos seleccionados incluyen suavizado exponencial doble y triple, el Modelo Autorregresivo Integrado de Media Móvil Estacional (SARIMA) y Redes Neuronales Recurrentes como la Unidad Recurrente Gateada (GRU) y la Memoria a Corto y Largo Plazo (LSTM). Estos modelos fueron evaluados utilizando la métrica de error porcentual absoluto medio (MAPE).

Para lograr esto, la sección 2 proporciona una revisión de literatura sobre el "Slotting" aplicado en almacenes, análisis estadístico de series temporales y aprendizaje automático. La sección 3 define los tipos de Slotting, métodos estadísticos y de aprendizaje automático. La sección 4 describe la metodología utilizada para partes, procesos y

características. La sección 5 presenta la ejecución de un caso de prueba basado en 50 productos y un historial de 8 años desde enero de 2016 hasta mayo de 2023. La sección 6 presenta los resultados del caso de prueba. La sección 7 discute los resultados. Finalmente, la sección 8 presenta las conclusiones y hallazgos.

II. REVISION LITERARIA

En esta sección se describen estudios de Slotting aplicados a Almacenes, análisis estadístico de series temporales y aprendizaje automático.

A. Slotting aplicado a los Almacenes

Los problemas de almacenamiento en almacenes logísticos para productos emergen como un tema central que abarca múltiples áreas de investigación, referido como SLAP (Problema de Asignación de Ubicación de Stock). Para resolver esto, se implementan métodos para mejorar el almacenamiento, conocidos como "Slotting". Los primeros métodos analíticos utilizados en la investigación de zonas de selección son COI y ABC (Viveros, González, Mena, Kristjanpoller, & Robledo, 2021) (Fontana, Leyva López, Virginio Cavalcante, & Solano Noriega, 2020). En el estudio de (Rios, Morillo-Torres, Olmedo, Coronado-Hernandez, & Gatica, 2022), se aplica un modelo de Slotting matemático para mejorar la eficiencia en los tiempos de selección y asignar estratégicamente productos de alta rotación e inventario de clase ABC. De manera similar, los investigadores en (Aase & Petersen, 2022) implementan un Slotting basado en clases que reasigna productos dividiéndolos en cuatro clases (A, B, C, D) y analiza la frecuencia de productos en pedidos. La investigación de (Cardona & Gue, 2020) propone un método de Slotting diferente en áreas de almacenamiento de pallets que proporciona resultados positivos en ahorro de espacio y costos operativos. Además, el estudio de (Viveros, González, Mena, Kristjanpoller, & Robledo, 2021) aplica un modelo de Slotting similar llamado SLAP-VH, enfocado en el almacenamiento aplicado de productos en pallets de varios niveles utilizando almacenamiento basado en clases. Tiene beneficios como mejoras en las operaciones de selección y tiempos de viaje reducidos para movimientos de grúas. Existe la posibilidad de combinar Slotting con otros métodos de mejores prácticas. En el estudio de (Espinoza Sanchez, Sanchez Santos, Marcelo Lastra, Quiroz Flores, & Alvarez Merino, 2021), los autores combinan Slotting con métodos y técnicas de "Lean Warehousing" y "Wave Picking". Los resultados mostraron cambios positivos en el indicador de "A Tiempo, Completo" (OTIF), que mide

la eficiencia y precisión de la entrega de productos a los clientes. Finalmente, el estudio de (Lavender, Sun, Xu, & Sommerich, 2021) enfatiza estudios que ponen en práctica Slotting sin considerar la ergonomía de los operadores durante los procesos de reabastecimiento y selección en las zonas de selección. Los resultados muestran los criterios de peso esperados para el reabastecimiento y la velocidad de selección de productos.

B. Análisis estadístico de series temporales y aprendizaje automático

Dentro del proceso de análisis de series temporales en aprendizaje automático, se emplean diversos métodos para pronosticar valores futuros basados en datos históricos. Una técnica ampliamente utilizada es el suavizado exponencial, que analiza el comportamiento de los datos a lo largo del tiempo. En la investigación realizada por (Sidqi & Sumitra, 2019), se llevó a cabo un análisis comparativo entre los métodos de suavizado exponencial simple y doble para la predicción de ventas. Notablemente, los resultados indicaron que el suavizado exponencial simple demostró un rendimiento superior, alcanzando una métrica MAPE del 20%. De manera similar, en el estudio de (Silitonga, Himawan, & Damanik, 2020), emplear el suavizado exponencial doble predijo el número de nuevas admisiones de estudiantes en un programa universitario. Los resultados obtenidos mostraron un Error Porcentual Medio del 11.72%. Por otro lado, los modelos ARIMA y SARIMA han mostrado valiosas contribuciones en el análisis y pronóstico. Investigaciones recientes, como la de (Falatouri, Darbanian, Brandtner, & Udokwu, 2020) han corroborado que el modelo SARIMA ofrece un rendimiento de pronóstico superior, especialmente para productos con comportamiento estacional, cuando se incorporan variables externas utilizando SARIMAX. Este hallazgo subraya la eficacia de SARIMA en la gestión de la cadena de suministro minorista. En cuanto al modelo ARIMA, el estudio de (Fatima Mohamad, y otros, 2023) comparó estos dos modelos, evaluando sus métricas AIC y BIC más bajas, y los resultados indicaron la superioridad de SARIMA. De manera similar, en el estudio de (Yamak, Yujian, & Gadosey, 2019) se realizó una comparación entre los modelos ARIMA, GRU y LSTM para pronosticar el precio de Bitcoin. Los hallazgos indicaron que el modelo ARIMA logró una precisión superior con un MAPE del 2.76% junto con un tiempo de implementación de 0.0007 segundos, lo que lo hace óptimo para un procesamiento eficiente en tiempo. Sin embargo, en el estudio de (Airlangga, Rachmat, & Lapihu, 2019), los modelos que emplean suavizado exponencial

simple, doble y triple, así como el modelo de redes neuronales Backpropagation, demostraron aceptabilidad según sus métricas de error bajos, pero Backpropagation resultó superior a los demás, logrando un MAPE del 1.5% y un tiempo de procesamiento promedio de 20.2329 segundos. Además, en el estudio de (Can Ozdemir, Buluş, & Zor, 2022), se evaluaron las redes neuronales LSTM y GRU, siendo el método GRU la opción óptima con un valor MAPE del 6.986%. Sin embargo, la diferencia con LSTM fue del -0.074%, lo que hace que ambos modelos sean aceptables. Sin embargo, el tiempo de procesamiento computacional de GRU es un 33% más corto que el de LSTM.

III. FUNDAMENTOS TEORICOS

A. Slotting

El objetivo principal del Slotting es optimizar la distribución de productos a sus respectivas ubicaciones dentro de las zonas de picking, con el fin de reducir los costos operativos y los tiempos de preparación de pedidos (Jaramillo, Camilo, Cuellar Molina, & Cogollo Flórez, 2020). El Slotting ha sido objeto de investigación reciente sobre los tipos y criterios de los métodos de almacenamiento de productos aplicados en las zonas de picking. A continuación, se presenta una descripción detallada de estos aspectos:

1. **Almacenamiento aleatorio:** Se refiere a asignar los productos a cualquier ubicación en un almacén que sea fácil de gestionar sin necesidad de un método analítico (Aase & Petersen, 2022).
2. **Almacenamiento dedicado:** Implica asignar productos a ubicaciones cercanas a la zona de picking basándose en el método Cube-per-Order-Index (COI) (Jaramillo, Camilo, Cuellar Molina, & Cogollo Flórez, 2020). COI optimiza el espacio de almacenamiento considerando el volumen o la frecuencia de los productos, con el objetivo de asignar productos de alta rotación a ubicaciones más accesibles o cercanas en las zonas de picking. Mientras que los productos de baja rotación se almacenan en ubicaciones centrales o finales (Bahrami, Piri, & Aghezzaf, 2019). COI tiene como objetivo minimizar el tiempo de viaje y reducir los movimientos físicos requeridos para acceder a los productos (Lorenc & Lerher, 2020).
3. **Almacenamiento por clase:** Implica categorizar productos en varias clases, asignadas a

ubicaciones específicas basadas en un análisis matemático conocido como el principio de Pareto (Bahrami, Piri, & Aghezzaf, 2019). El criterio comúnmente consiste en tres categorías principales denominadas ABC, donde cada letra (A, B y C) representa diferentes porcentajes de la variable de demanda, frecuencia, rotación, volumen o valor de los productos (Aase & Petersen, 2022). En esta clasificación, la categoría 'A' abarca aproximadamente el 80%, la categoría 'B' contiene alrededor del 15% y la categoría 'C' incluye el 5% restante (Rios, Morillo-Torres, Olmedo, Coronado-Hernandez, & Gatica, 2022)

B. Análisis estadístico de series temporales

Una serie temporal es una secuencia de datos registrados cronológicamente, que representa observaciones correlacionadas a lo largo del tiempo. Estos datos históricos se utilizan para capturar tendencias y patrones cíclicos, y para construir modelos de pronóstico apropiados. Las series temporales son prevalentes en diversas industrias, empleadas para predecir eventos futuros como el pronóstico del tiempo, el rendimiento de las acciones, la inflación, la propagación de enfermedades, el mantenimiento predictivo, la gestión de inventarios y otras aplicaciones (Blázquez-García, Conde, Mori, & Lozano, 2021) (Kumar Dubey, Kumar, García-Díaz, Kumar Sharma, & Kanhaiya, 2021)

1. Suavizado exponencial

El suavizado exponencial enfatiza los datos recientes y suaviza gradualmente los datos más antiguos, lo que lo hace útil para identificar cambios en las tendencias en series temporales y para predecir patrones no lineales (Sidqi & Sumitra, 2019). Los tipos de suavizado exponencial son los siguientes:

a. Método de Holt o Suavizamiento Doble Exponencial:

Es una técnica utilizada en el análisis de datos de series temporales para pronosticar datos que tienen una tendencia o componente estacional. Este método es notable por su eficacia cuando se enfrenta a conjuntos de datos con tendencias complejas que resultan desafiantes para que los métodos convencionales capturen eficientemente, especialmente aquellos que muestran patrones discernibles de crecimiento o declive con el tiempo (Silitonga, Himawan, & Damanik, 2020). La aplicación de esta técnica implica la incorporación de dos coeficientes de suavizado: alfa (α) y beta (β).

La fórmula general para el Suavizado Exponencial Doble es:

$$\hat{y}_{t+1} = \alpha \cdot y_t + (1 - \alpha) \cdot (\hat{y}_t + T_t) \quad (1)$$

$$T_{t+1} = \beta \cdot (\hat{y}_{t+1} - \hat{y}_t) + (1 - \beta) \cdot T_t \quad (2)$$

Donde:

- $\hat{y}_{t+1}$ es el valor de pronóstico para el siguiente período.
- $\hat{y}_t$ es el valor real en el período t.
- $\hat{y}_t$ es el valor pronostico en el período t.
- T_{t+1} es la tendencia estimada para el próximo período.
- α es el parámetro de suavizado para el componente de ponderación ($0 \leq \alpha \leq 1$).
- β es el parámetro de suavizado para el componente de tendencia ($0 \leq \beta \leq 1$).

b. Holt-Winters o Suavizado Triple Exponencial: Está diseñado para capturar y predecir patrones de datos con tendencia y estacionalidad. Es eficaz para analizar conjuntos de datos con una tendencia discernible y patrones estacionales repetidos. El proceso de triple suavizado actualiza el nivel, la tendencia y los componentes estacionales dentro de cada ciclo, proporcionando un enfoque integral y preciso para el pronóstico de series temporales. La implementación de esta técnica implica una cuidadosa consideración de tres coeficientes de suavizado: alfa (α), beta (β) y gamma (γ) (Emmanuel K., Wagala, & Muriithi, 2020). La fórmula general para el suavizado triple exponencial es:

$$\hat{y}_{t+m} = (l_t + m \cdot b_t) \cdot s_{t-m+1} \quad (3)$$

$$l_t = \alpha \cdot \frac{y_t}{s_{t-m}} + (1 - \alpha) \cdot (l_{t-1} + b_{t-1}) \quad (4)$$

$$b_t = \beta \cdot (l_t - l_{t-1}) + (1 - \beta) \cdot b_{t-1} \quad (5)$$

$$s_t = \gamma \cdot \frac{y_t}{l_t} + (1 - \gamma) \cdot s_{t-m} \quad (6)$$

Donde:

- $\hat{y}_{t+m}$ es el valor previsto para m períodos venideros.
- l_t componente de nivel en el periodo t.
- b_t componente de tendencia en t.
- s_t componente estacional en t para una temporada de duración m.
- y_t es el valor real en el período t.
- α , β y γ son parámetros suaves ($0 \leq \alpha, \beta, \gamma \leq 1$)
- m el número de estaciones.

2. Seasonal Autoregressive Integrated Moving Average:

SARIMA extiende el método ARIMA, incorporando una variable temporal estacional denotada como "S". Este método combina técnicas autorregresivas (AR), integradas (I) y de promedio móvil (MA) para pronosticar datos de series temporales con patrones estacionarios y no estacionarios (Fatima Mohamad, y otros, 2023). El componente AR captura relaciones entre observaciones pasadas y presentes, el componente MA modela relaciones entre errores actuales y pasados, y el componente I implica diferenciaciones para lograr estacionariedad en la serie temporal. SARIMA consta de cuatro pasos: identificar variables y evaluar la estacionariedad de la serie, determinar la combinación óptima de autorregresión y promedio móvil, seleccionar el modelo más eficiente, verificar la precisión del modelo y explorar mejoras, y finalmente, usar el modelo para predecir datos futuros dentro de un intervalo de confianza (Fatima Mohamad, y otros, 2023). En cuanto a la fórmula general para un modelo ARIMA (p, d, q), representa el orden de los componentes autorregresivos, de diferenciación y de promedio móvil, respectivamente. El

$$\text{Arima } y_t = C + \varphi_1 \cdot y_{t-1} + \varphi_2 \cdot y_{t-2} + \dots + \varphi_p \cdot y_{t-p} + \theta_1 \cdot \varepsilon_{t-1} + \theta_2 \cdot \varepsilon_{t-2} + \dots + \theta_q \cdot \varepsilon_{t-q} + \varepsilon_t \quad (7)$$

$$\text{Sarima } y_t = C + \varphi_1 \cdot y_{t-1} + \varphi_2 \cdot y_{t-2} + \dots + \varphi_p \cdot y_{t-p} + \theta_1 \cdot \varepsilon_{t-1} + \theta_2 \cdot \varepsilon_{t-2} + \dots + \theta_q \cdot \varepsilon_{t-q} + \varphi_p \cdot y_{t-p} + \varphi_{2p} \cdot y_{t-2s} + \dots + \varphi_{np} \cdot y_{t-ns} + \theta_{1p} \cdot \varepsilon_{t-s} + \theta_{2p} \cdot \varepsilon_{t-2s} + \dots + \theta_{np} \cdot \varepsilon_{t-ns} + \varepsilon_t \quad (8)$$

mismo principio se extiende a SARIMA con órdenes estacionales adicionales.

Donde:

- y_t valor observado en el momento t .
- C es constante.
- $\phi_1, \phi_2, \dots, \phi_p$ parámetros autorregresivos.
- e_t error o residuo en t .
- $\theta_1, \theta_2, \dots, \theta_q$ parámetros de media móvil.
- φ_P Coeficientes autorregresivos estacionales.
- θ_P Coeficientes de media móvil estacional.

El algoritmo Autoarima encuentra los parámetros que producen el mejor ajuste del modelo ARIMA o SARIMA para un conjunto de datos de series temporales dado. Para hacer esto, evalúa múltiples combinaciones de estos parámetros y selecciona el modelo con los valores más bajos de AIC (Criterio de Información de Akaike) y BIC (Criterio de Información Bayesiano), lo que significa un mejor equilibrio entre el ajuste del modelo y la complejidad.

C. Machine Learning (ML)

Esto se refiere a una parte de la inteligencia artificial que dota a una computadora con la capacidad de comprender datos a través de código en varios algoritmos de aprendizaje. Las Redes Neuronales Recurrentes (RNN), están especialmente diseñadas para manejar datos secuenciales o temporales tienen conexiones de retroalimentación que les permiten modelar dependencias temporales en los datos. Estas RNN operan mediante un proceso de entrenamiento y evaluación, que culmina en la generación de un modelo capaz de predecir información futura o analizar patrones en secuencias de datos (Nabipour, Nayyeri, Jabani, S., & Mosavi, 2020).

Long Short-Term Memory (LSTM): Es una red neuronal recurrente que puede capturar dependencias a largo plazo en datos secuenciales y es altamente efectiva en la predicción de series temporales. El núcleo puede mantener y actualizar un estado de celda a lo largo del tiempo, lo que le permite almacenar información en secuencias largas y selectivamente olvidar o actualizar información (Nabipour, Nayyeri, Jabani, S., & Mosavi, 2020). Una celda LSTM tiene tres puertas primarias: la puerta de entrada, la puerta

de olvido y la puerta de salida. Las operaciones de la celda LSTM se definen:

$$\text{Puerta de Entrada } (i_t) = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (9)$$

$$\text{Puerta de Olvidar } (f_t) = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (10)$$

$$\begin{aligned} \text{Estado de Celda Candidata } (\tilde{C}_t) \\ = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c) \end{aligned} \quad (11)$$

$$\text{Estado de la Celda } (C_t) = f_t \cdot C_{t-1} + i_t \cdot \tilde{C}_t \quad (12)$$

$$\text{Puerta de Salida } (o_t) = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (13)$$

$$\text{Estado Oculto } (h_t) = o_t \cdot \tanh(C_t) \quad (14)$$

Donde:

- i_t si la puerta de entrada se activa en el momento t .
- f_t si la puerta de olvido está activada en t .
- C_t es el estado de la celda en el tiempo t .
- $\tilde{C}_t$ es el estado candidato de la celda en t .
- h_t es el estado oculto (salida) en t .
- o_t es la activación de la puerta de salida en t .
- x_t es la entrada en el tiempo t .
- W_i, W_f, W_c y W_o son matrices de pesos b_i, b_f, b_c y b_o son vectores de sesgo.
- σ es la función de activación sigmoide.
- $\tanh$ es la función de activación tangente hiperbólica.

Las LSTMs permiten que la red aprenda cuándo almacenar información en el estado de la celda (C_t), cuándo olvidar información (f_t) y cuándo emitir información (h_t). Esta gestión dinámica de la memoria hace que las LSTMs sean particularmente efectivas para la modelización de datos de pronóstico de series temporales (Gao, y otros, 2020).

Gated Recurrent Unit (GRU): Representa una variante especializada de las redes neuronales recurrentes diseñada para modelar y capturar dependencias dentro de los datos de series temporales, ganando una amplia popularidad en el ámbito del análisis de series temporales. La característica distintiva de las GRUs radica en su capacidad para capturar y propagar eficientemente información a través de varios pasos de tiempo mientras mitigan

el problema del gradiente desvaneciente, un desafío común en las redes neuronales recurrentes tradicionales (Yamak, Yujian, & Gadosey, 2019). Las operaciones de la celda GRU están definidas:

Donde:

$$z_t = \sigma.(W_z.[h_{t-1}, x_t]) \quad (15)$$

$$r_t = \sigma.(W_r.[h_{t-1}, x_t]) \quad (16)$$

$$\tilde{h}_t = \tanh.(W_h.[r_t.h_{t-1}, x_t]) \quad (17)$$

$$h_t = (1 - z_t).h_{t-1} + z_t.\tilde{h}_t \quad (18)$$

- z_t La activación de la puerta de actualización en el tiempo t.
- r_t La activación de la puerta de reinicio en el tiempo t.
- x_t La entrada en el tiempo t.
- $\tilde{h}_t$ El estado oculto candidato en el tiempo t.
- h_t estado oculto actualizado (salida) en el tiempo t.
- W_z, W_r, W_h son matrices de peso.

Las GRUs proponen la puerta de reinicio r_t y la puerta de actualización z_t que controlan la información a través de la celda. La puerta de actualización decide cuántos datos se retienen en el estado oculto anterior h_{t-1} y la puerta de reinicio determina cuántos datos se olvidan en el estado anterior. La puerta de reinicio y la entrada x_t controlan el estado oculto candidato $\tilde{h}_t$, que se actualiza al estado oculto. Las GRUs ofrecen un equilibrio entre capturar dependencias a largo plazo en secuencias y eficiencia computacional (Gao, y otros, 2020).

IV. PROPUESTA DE SLOTTING

En esta sección, ofrecemos una exposición detallada de nuestro modelo dirigido a mejorar el proceso de asignación de ubicaciones. Comenzamos delineando las diversas fases del modelo, proporcionando una elaboración extensa del modelo conjunto integral. Además, se muestra el diagrama de sus actividades en la Figura 1, y las funcionalidades se describen en la Tabla 1.

A. Entrada

La sección de "Entrada" juega un papel fundamental como fuente de datos históricos esenciales para el funcionamiento del centro de distribución o almacén. En esta fase, se almacena información detallada sobre los productos, su historial de pedidos y datos de ubicación de productos dentro del almacén. Generalmente, establecer una conexión con el sistema histórico del almacén es necesario para extraer los datos requeridos para el procesamiento posterior. Además, es esencial evaluar cuidadosamente el alcance y la calidad de los datos históricos. Delinear adecuadamente el rango de datos históricos es crucial para garantizar que la información obtenida sea representativa y útil para análisis posteriores.

B. Procesamiento

1. Procesamiento de datos

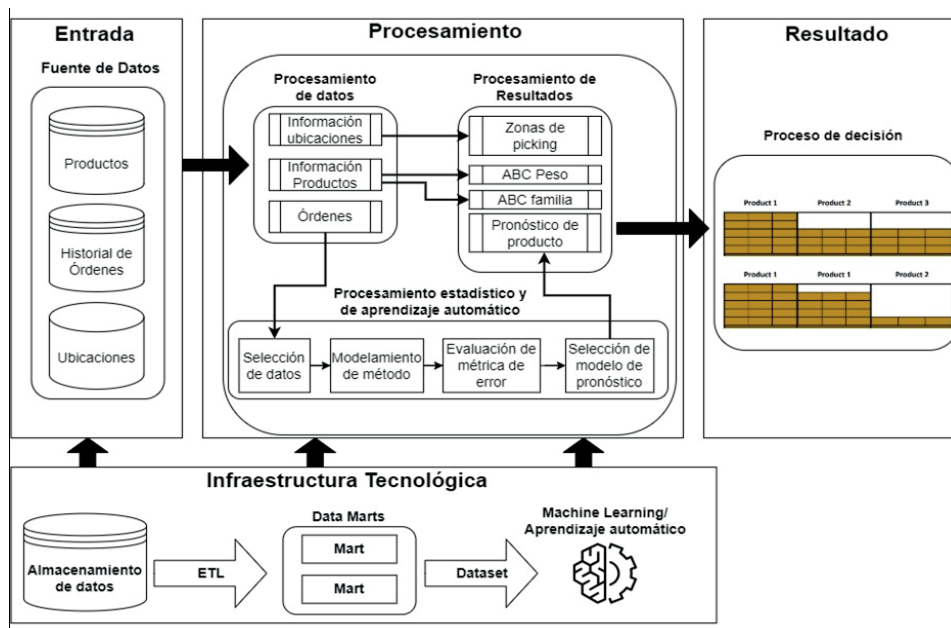
La etapa de procesamiento de datos abarca tres componentes clave: En primer lugar, el componente "Información de Ubicaciones", que almacena datos vitales sobre las ubicaciones del centro de distribución, incluidas dimensiones, tipos de instalaciones y clasificación. En segundo lugar, "Información de Productos" recopila detalles como dimensiones de caja, peso, categoría y cualquier característica especial relevante para el almacenamiento. Finalmente, "Historial de Pedidos" es crucial ya que resume la información sobre el historial de pedidos recopilada durante el mes, calculando totales mensuales por producto a partir de unidades diarias obtenidas. En casos donde falte información para alguno de los meses, se aplican métodos de relleno para completar los registros faltantes.

2. Procesamiento estadístico y de aprendizaje automático

Primero, en la fase de "Selección de Datos", se eligen secuencialmente los registros de productos ya resumidos por mes para su procesamiento. Segundo, durante la etapa de "Modelado de Métodos", se aplican métodos analíticos seleccionados utilizando los datos previamente seleccionados como conjunto de datos para el análisis. Esto permite una comprensión más profunda de la información recopilada y la generación de pronósticos. Tercero, en el paso de "Evaluación de la Métrica de Error", se evalúan los resultados de pronóstico producidos por los métodos analíticos utilizando la métrica MAPE. Esta evaluación proporciona una medida cuantitativa de la precisión del pronóstico, mejorando significativamente

Figura 1

Metodología para mejorar el proceso de Slotting con un enfoque estadístico y de aprendizaje automático.



Fuente: elaboración propia.

la confiabilidad del modelo analítico. Cuarto, en la actividad de "Selección del Modelo de Pronóstico", se examinan los resultados de la métrica de error MAPE, donde se prefieren los métodos que producen valores de error más bajos. En los casos en que múltiples métodos exhiben valores de error por debajo del 10% para diferentes productos, se considera la posibilidad de combinarlos.

3. Procesamiento de resultados

En la fase de "Zonas de Picking", se recopilan datos para áreas de almacenamiento específicas, determinadas por las características del producto y las políticas del centro de distribución. El "Peso ABC" proporciona parámetros esenciales para determinar el nivel de altura de almacenamiento adecuado para prevenir accidentes que involucren al personal de picking, y la "Familia ABC" permite agrupar los tipos de productos en la misma zona según sus características. Finalmente, el "Pronóstico de Productos" evalúa la demanda anticipada de productos para la próxima temporada, ayudando a decidir si asignar productos a otras ubicaciones capaces de cumplir con los requisitos de demanda de unidades.

4. Resultado

En esta sección, el gerente del almacén de productos inicia el proceso de toma de decisiones,

que implica analizar el curso de acción óptimo para cada producto. Esto puede incluir reubicar el producto a un área diferente dentro del almacén o expandir el número de ubicaciones de almacenamiento. Las decisiones están guiadas por información derivada del informe de procesamiento de resultados, considerando factores como las zonas de picking, el peso ABC, la familia ABC y el pronóstico de productos. Este análisis exhaustivo capacita al gerente para tomar decisiones informadas que mejoren la eficiencia del almacenamiento de productos y aborden las demandas en evolución dentro del almacén.

5. Infraestructura tecnológica

Esta sección proporciona soporte tecnológico a las tres áreas mencionadas anteriormente, que son las siguientes: el almacenamiento histórico del centro de distribución, la utilización de la herramienta ETL (Extract, Transform, Load) para migrar datos esenciales y su posterior almacenamiento en un Datamart. Después de seleccionar los datos en el proceso de aprendizaje automático, la ejecución se realiza dentro del entorno de Google Colab. Finalmente, este enfoque integrado garantiza un flujo de información sin problemas y facilita los procedimientos de aprendizaje automático para obtener información mejorada impulsada por datos.

Tabla 1*Descripción de las partes del modelo de Slotting.*

Partes	Procesos	Características
Entrada	Fuente de datos	Base de datos histórica de pedidos de productos.
		Base de datos de dimensiones de productos: alto, ancho, largo y peso.
Procesamiento	Procesamiento de datos	Base de datos de dimensiones de ubicación: alto, ancho, largo, vano y pasillo.
		Depuración del historial de pedidos.
		Información del Producto.
		Información sobre las ubicaciones.
	Procesamiento estadístico y de aprendizaje automático	Selección de datos: selección de un conjunto de datos para su análisis o procesamiento.
		Modelado de métodos: ejecución del modelo con el conjunto de datos.
		Evaluación de métricas de error: evaluación del valor previsto y real.
Resultado	Proceso de decisión	Selección del modelo de pronóstico: seleccione el mejor modelo según la métrica de error.
		Zona de picking: capacidad de unidades que se puede almacenar el producto en las ubicaciones.
		Peso ABC: clasificación del peso del producto en la zona de nivel y altura del local.
		Familia ABC: clasificación del tipo de producto en el pasillo de la zona de picking.
	Procesamiento de resultados	Previsión de producto: previsión de las unidades de producto para los siguientes periodos.
		Toma de decisiones del usuario para el almacenamiento de productos teniendo en cuenta el resultado del procesamiento.

Fuente: elaboración propia.

V. MATERIAL Y METODOS

A. Caso de estudio

Evaluamos la estrategia de inventario de un centro de distribución mayorista farmacéutico con el objetivo de pronosticar el inventario de productos. La empresa ha experimentado un crecimiento en su línea de productos a lo largo de los años y cambios en la demanda, lo que ha resultado ser un desafío en la actualización de las cantidades en sus ubicaciones. Por lo tanto, llevaremos a cabo una evaluación de 50 productos utilizando un conjunto de datos históricos que abarca los últimos ocho años, desde enero de 2016 hasta mayo de 2023.

B. Las herramientas utilizadas

Un equipo de computadora sirvió como entorno de prueba con bases de datos de SQL Server y SQL Data Studio instaladas. Se estableció un proceso ETL para la migración de datos. Posteriormente, se empleó la herramienta Google Colab para implementar los modelos de aprendizaje automático.

C. Los métodos utilizados

Esta investigación empleará varios modelos de pronóstico, incluyendo suavizado exponencial doble y triple. Además, utilizaremos la función Autoarima de la biblioteca 'pmdarima'. Esta función ajusta automáticamente los parámetros adecuados para los

modelos ARIMA y SARIMA. Además, exploraremos las capacidades de las redes neuronales utilizando LSTM y GRU como enfoques alternativos de modelado predictivo.

D. Métrica

Los modelos de pronóstico evaluados utilizaron la métrica de Error Porcentual Absoluto Medio (MAPE, por sus siglas en inglés). En esta evaluación, nuestro enfoque se guió por los criterios propuestos en el estudio de referencia (Sidqi & Sumitra, 2019), que describe pautas para evaluar los intervalos de porcentaje de MAPE. Estos intervalos se presentan en la Tabla 2. Para este estudio, elegimos centrarnos en el rango de MAPE del 0% al 20% como criterio principal. Cualquier valor de MAPE que exceda el 20% se considera un rango inaceptable para pronósticos confiables y se excluye de consideración adicional. La fórmula general para MAPE es:

Donde:

$$MAPE = \frac{100\%}{N} \sum_{i=0}^{N-1} \frac{|y_i - \hat{y}_i|}{y_i}$$

- N: Representa el número total de observaciones o puntos de datos en el conjunto de datos.

- y_i : es el valor real u observado en la posición i en el conjunto de datos.
- $\hat{y}_i$: Representa la predicción o estimación correspondiente al valor real y_i en la posición

Tabla 2*Rango de puntuación MAPE.*

Rango (%)	Capacidad de predicción
<10	Excelente capacidad de previsión
10 - 20	Buena capacidad de previsión
20 - 50	Capacidad de previsión razonable
>50	Mala capacidad de previsión

Fuente: elaboración propia.

E. Procedimiento

Para llevar a cabo el procedimiento experimental, se recopilaron datos del centro de distribución mayorista en el caso de estudio. La información relevante se extrajo de su sistema de gestión de bases de datos a través de un proceso ETL, utilizando la herramienta Visual Studio para transferirla a una base de datos en un entorno de prueba. Posteriormente, se aplicaron transformaciones, se gestionaron datos faltantes e información relacionada con pedidos, productos y ubicaciones. Estos datos se agruparon en períodos mensuales correspondientes a los últimos seis años para facilitar su utilización en la ejecución de modelos de aprendizaje automático y métodos estadísticos. Los modelos y métodos se evaluaron utilizando métricas de error para seleccionar el enfoque analítico más apropiado. Una vez obtenidos los resultados de pronóstico, se realizaron análisis con la ponderación ABC, la familia ABC y las zonas de picking para evaluar las opciones disponibles para cada producto. Finalmente, se entregó un informe que contenía recomendaciones sobre la zona de picking más adecuada para el almacenamiento de productos al usuario responsable

VI. RESULTADOS

En la prueba del modelo, los resultados revelaron variaciones significativas en las métricas de error de MAPE entre los modelos para el conjunto específico de 50 productos, como se muestra en la Tabla 3. Para los valores de MAPE por debajo del 10%, el método de Doble Suavizado toma la delantera con 25 productos, seguido por el Triple Suavizado con 22, SARIMA con 19, LSTM con 17 y GRU con 12. En el rango del 10% al 20%, el Triple Suavizado lidera con 26 productos, seguido por SARIMA con

25, Doble Suavizado con 23, GRU con 22 y LSTM con 20. Dentro de la categoría de MAPE que abarca el 20% al 50%, GRU lidera con 15 productos, LSTM con 10, SARIMA con 3, tanto Doble como Triple Suavizado con 2 productos cada uno. Para valores de MAPE que exceden el 50%, LSTM y SARIMA emergen como los modelos más comunes con 3 cada uno y GRU con 3 productos. En resumen, el Doble Suavizado sobresale en el pronóstico de la demanda de productos, especialmente dentro del rango de MAPE por debajo del 10%.

Tabla 3*Número de productos en el intervalo MAPE.*

Rango (%)	GRU	LSTM	(S) ARIMA	SMOOTH. DOBLE	SMOOTH. TRIPLE
< 10	12	17	19	25	22
10 - 20	22	20	25	23	26
20 - 50	15	10	3	2	2
> 50	1	3	3		

Fuente: elaboración propia.

A continuación, se presentaron los resultados obtenidos para los valores mínimos, máximos y promedio de MAPE en la Tabla 4. En cuanto al porcentaje mínimo de MAPE, el modelo de suavizado exponencial doble se desempeñó admirablemente con un valor del 2.38%, seguido de cerca por el suavizado exponencial triple con el 2.39% y luego (S) ARIMA con el 2.44%. En cuanto al valor máximo de MAPE, el modelo (S)ARIMA exhibió el porcentaje más alto con un 187.79%, seguido por el modelo LSTM con un 139.86% y GRU con un 51.24%. Finalmente, al analizar el valor promedio de MAPE, el enfoque de suavizado exponencial doble registró una cifra del 10.75%, seguido por el suavizado exponencial triple, GRU, (S)ARIMA y LSTM con valores de 11.61%, 17.78%, 17.96% y 19.35%, respectivamente. Estos datos resumidos proporcionan una visión general de las tendencias generales en el rendimiento de los modelos evaluados.

Tabla 4*Rango de puntuación MAPE.*

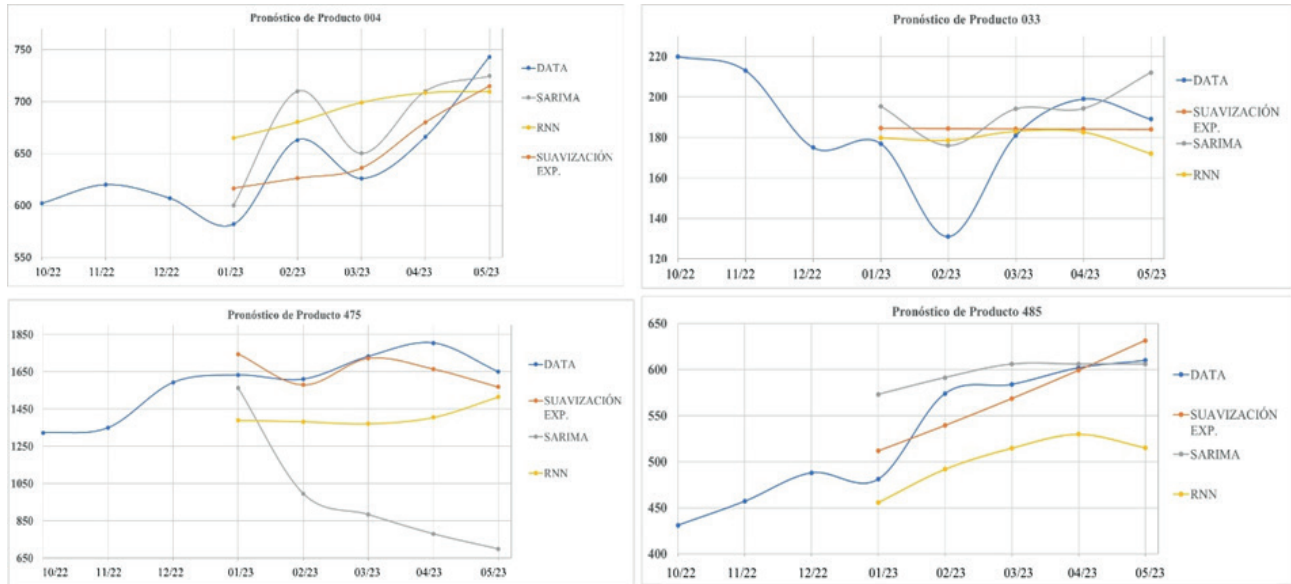
Método	Máximo	Mínimo	Promedio
GRU	51.24%	3.25%	17.78%
LSTM	139.86%	2.71%	19.35%
(S)ARIMA	187.79%	2.44%	17.96%
DOBLE EXP	29.97%	2.38%	10.75%
TRIPLE EXP	48.84%	2.39%	11.61%

Fuente: elaboración propia.

Finalmente, los modelos se consolidaron en sus respectivas categorías, como se muestra en la Ta-

Figura 2

Se pronosticó los periodos de enero a mayo de 2023. Se seleccionó el modelo de suavizado exponencial más adecuado. Para (S)ARIMA, se eligieron los valores sugeridos para p , d y q de autoarima. Posteriormente, se compararon las redes neuronales GRU y LSTM, seleccionando la más adecuada para nuestra investigación.



Fuente: elaboración propia.

Tabla 5. Después de la observación, el rango MAPE se muestra por debajo del 10%, el suavizado exponencial, (S)ARIMA y RNN tienen 28, 19 y 18 productos respectivamente. En el rango MAPE entre 10% y 20%, (S)ARIMA lidera con 25 productos, seguido de Exponential Smoothing y RNN, con 22 y 21 productos respectivamente. En el MAPE entre el rango 20% y 50%, RNN tiene el mayor número de productos con 11, seguido de (S)ARIMA con 3. Además, en MAPE por encima del 50%, (S)ARIMA tiene 3 productos. En resumen, considerando la cantidad de productos, el Suavizado Exponencial tiende a sobresalir en la mayoría de los rangos de MAPE entre 0% y 20%. Los pronósticos de productos a partir de los métodos se pueden ver en la Figura 2.

Tabla 5

Resultados de la métrica de error MAPE por categorías.

Rango (%)	RNN	(S)ARIMA	Suavización Exponencial
< 10	18	19	28
10 - 20	21	25	22
20 - 50	11	3	
> 50		3	

Fuente: elaboración propia.

VII. DISCUSION DE RESULTADOS

Nuestros resultados de métricas de error MAPE en pronósticos de demanda se alinean con estudios similares realizados por (Sidqi & Sumitra, 2019), quienes obtuvieron resultados comparables al evaluar el modelo de suavizado exponencial doble en pronósticos de ventas de productos, logrando un valor del 24%, que difiere de nuestro resultado del 10.75%. Además, los autores (Silitonga, Himawan, & Damanik, 2020) exploraron enfoques de suavizado exponencial doble y reportaron un valor de error alto del 27.72%, con una variación del 16.97% en nuestro modelo de investigación, lo que hace que nuestros resultados sean más significativos. Sin embargo, en ambos estudios, las métricas de error se encuentran dentro del rango del 20% al 50%, lo que puede considerarse razonable para modelos de pronóstico. En el estudio de (Airlangga, Rachmat, & Lapihu, 2019), utilizaron modelos de suavizado exponencial triple, revelando un valor del 3% en comparación con nuestro valor mínimo de 2.39%. También introdujeron un modelo de red neuronal llamado "Backpropagation", que produjo un valor del 1.5%, demostrando ser muy preciso en comparación con nuestros resultados para los

modelos LSTM y GRU, que fueron del 2.71% y 3.25%, respectivamente. Estos resultados son similares a los reportados en el estudio de (Yamak, Yujian, & Gadosey, 2019), que enfatizó la precisión del modelo LSTM en un 6.8% y del modelo GRU en un 3.97%. Aunque los resultados se consideran excelentes para el pronóstico, se destaca que el modelo de retropropagación es mucho mejor que los modelos LSTM y GRU. Además, en comparación con la investigación realizada por (Kumar Dubey, Kumar, García-Díaz, Kumar Sharma, & Kanhaiya, 2021), que analiza el consumo de energía, es destacable que la función SARIMA logró un MAPE del 5.23%, contrastando con nuestro resultado del 2.71%. Esta disparidad del 2.52% entre nuestros estudios subraya los diversos desafíos y dinámicas inherentes a sus respectivos campos de estudio. Por último, para futuras investigaciones, recomendamos emplear intervalos de MAPE más refinados, ya que los intervalos utilizados actualmente son bastante amplios para el pronóstico de la demanda. A continuación, la Tabla 6 presenta un conjunto más detallado de intervalos preparados por nosotros.

Tabla 6

Nuevo rango de puntuación MAPE.

Rango (%)	Capacidad de predicción
<10	Excelente capacidad de pronóstico
10-20	Buena capacidad de pronóstico.
20-30	Capacidad de pronóstico aceptable.
30-40	Poca capacidad de pronóstico.
>40	Capacidad de pronóstico inexacta

Fuente: elaboración propia.

VIII. CONCLUSIONES

El análisis comparativo de métodos de pronóstico para un conjunto de 50 productos ha proporcionado información valiosa para optimizar el proceso de asignación de ubicaciones. Podemos concluir lo siguiente:

- Los métodos de suavizado exponencial doble y triple exhibieron un rendimiento robusto en el pronóstico para el 100% de todos los productos con un MAPE de menos del 20%. Esta eficacia les permite capturar eficazmente las tendencias y variaciones estacionales, estableciéndose como opciones confiables para una amplia gama de productos.
- El (S)ARIMA demostró un rendimiento efectivo para el 88% con un MAPE de menos del 20%. Este método es particularmente adecuado para casos con patrones de datos relativamente

sencillos, ofreciendo un enfoque simplificado para obtener pronósticos precisos.

- Las redes neuronales recurrentes (RNN), incluidos los modelos LSTM y GRU, demostraron su eficacia para el 78% con un MAPE de menos del 20%. Si bien las RNN son excelentes para capturar dependencias complejas en datos secuenciales, su rendimiento puede variar según el tamaño y la complejidad del conjunto de datos.

En conclusión, la elección del método de pronóstico debe alinearse con las características específicas de los productos o series temporales analizadas. Los métodos de suavizado exponencial doble y triple son las opciones más sólidas para un conjunto diverso de productos, luego (S)ARIMA ofrece automatización para patrones más simples y, finalmente, las RNN destacan en capturar patrones de datos intrincados. Un enfoque híbrido presentado en nuestra metodología que adapta los métodos a categorías de productos puede ser una estrategia práctica para lograr la mejor precisión en el pronóstico y optimizar el proceso de asignación de ubicaciones. El monitoreo y la evaluación continuos serán cruciales para adaptarse a los cambios en los patrones de datos y el comportamiento del producto con el tiempo. Este enfoque, centrado en mejorar el diseño de productos en las zonas de picking, promete traer mejoras significativas a la gestión de inventario al permitir la evaluación de cambios en las zonas de picking y la adición de nuevas ubicaciones para adaptarse a las variaciones en la demanda con el tiempo. A su vez, garantiza la eficiencia en la gestión de inventarios en entornos logísticos y una reducción en los costos asociados con la reducción de riesgos.

REFERENCIAS BIBLIOGRÁFICAS

- [1] Aase, G., & Petersen, C. (2022). A Decision Support System for Re-Slotting a Case Pick Distribution Center. *Open Journal of Business and Management*, 10(4), 1923-1935.
- [2] Airlangga, G., Rachmat, A., & Lapihu, D. (2019). Comparison of exponential smoothing and neural network method to forecast rice production in Indonesia. *Telecommunication Computing Electronics and Control*, 17(3), 1367-1375.
- [3] Bahrami, B., Piri, H., & Aghezzaf, E.-H. (2019). Class-based storage location assignment: an overview of the literature. In 16th International

- Conference on Informatics in Control, Automation and Robotics, 390-397.
- [4] Blázquez-García, A., Conde, A., Mori, U., & Lozano, J. (2021). A review on outlier/anomaly detection in time series data. *ACM Computing Surveys*, 54(3), 1-33.
 - [5] Can Ozdemir, A., Buluş, K., & Zor, K. (2022). Medium- to long-term nickel price forecasting using LSTM and GRU networks. *Resources Policy*, 78, 102906.
 - [6] Cardona, L., & Gue, K. (2020). Layouts of unit-load warehouses with multiple slot heights. *Transportation Science*, 54(5), 1332-1350.
 - [7] Chicco, D., Warrens, M., & Jurman, G. (2021). The coefficient of determination R-squared is more informative than SMAPE, MAE, MAPE, MSE and RMSE in regression analysis evaluation. *PeerJ Computer Science*, 7, 623.
 - [8] Espinoza Sanchez, N., Sanchez Santos, P., Marcelo Lastra, G., Quiroz Flores, J., & Alvarez Merino, J. (2021). Implementation of Lean and Logistics Principles to Reduce Non-conformities of a Warehouse in the Metalworking Industry. 10th International Conference on Industrial Technology and Management (ICITM), 89-93.
 - [9] Falatouri, T., Darbanian, F., Brandtner, P., & Udokwu, C. (2020). Predictive Analytics for Demand Forecasting – A Comparison of SARIMA and LSTM in Retail SCM. *Procedia Computer Science*, 993-1003.
 - [10] Fatina Mohamad, A., Mat Jasin, A., Asmat, A., Canda, R., Ismail, J., & Md Soom, A. (2023). Sales Analytics Dashboard with ARIMA and SARIMA Time Series Model. In 2023 IEEE 13th Symposium on Computer Applications & Industrial Electronics (ISCAIE), 106-112.
 - [11] Fontana, M., Leyva López, J., Virgínio Cavalcante, C., & Solano Noriega, J. (2020). Multi-criteria assignment model to solve the storage location assignment problem. *Investigacion Operacional*, 41(7), 1019-1029.
 - [12] Gao, S., Huang, Y., Zhang, S., Han, J., Wang, G., Zhang, M., & Lin, Q. (2020). Short-term runoff prediction with GRU and LSTM networks without requiring time step optimization during sample generation. *Journal of Hydrology*, 589, 125188.
 - [13] Jaramillo, D., Camilo, J., Cuellar Molina, M., & Cogollo Flórez, J. (2020). Slotting y picking: una revisión de metodologías y tendencias. *Revista chilena de ingeniería*, 514-527.
 - [14] Kumar Dubey, A., Kumar, A., García-Díaz, V., Kumar Sharma, A., & Kanhaiya, K. (2021). Study and analysis of SARIMA and LSTM in forecasting time series data. *Sustainable Energy Technologies and Assessments*, 47, 101474.
 - [15] Lavender, S., Sun, C., Xu, Y., & Sommerich, C. (2021). Ergonomic considerations when slotting piece-pick operations in distribution centers. *Applied Ergonomics*, 97(103554).
 - [16] Lorenc, A., & Lerher, T. (2020). PickupSimulo–Prototype of Intelligent Software to Support Warehouse Managers Decisions for Product Allocation Problem. *Applied Sciences*, 10(23), 8683.
 - [17] Nabipour, M., Nayyeri, P., Jabani, H., S., S., & Mosavi, A. (2020). Predicting Stock Market Trends Using Machine Learning and Deep Learning Algorithms Via Continuous and Binary Data; a Comparative Analysis. *IEEE Access*, 150199 - 150212.
 - [18] Papcun, P., Cabadaj, J., Kajati, E., Romero, D., Landryova, J. L., Vascak, J., & Zolotova, I. (2019). Augmented reality for humans-robots interaction in dynamic slotting “chaotic storage” smart warehouses. In *Advances in Production Management Systems. Advances in Production Management Systems. Production Management for the Factory of the Future*, 556, 633-641.
 - [19] Rios, J., Morillo-Torres, D., Olmedo, A., Coronado-Hernandez, J., & Gatica, G. (2022). A faster optimal model slotting in rack positions with mono SKU pallet. *Procedia Computer Science*(203), 588-593.
 - [20] Sidqi, F., & Sumitra, I. (2019). Forecasting product selling using single exponential smoothing and double exponential smoothing methods. *IOP Conference Series: Materials Science and Engineering*, 662(3), 032031.
 - [21] Silitonga, P., Himawan, H., & Damanik, R. (2020). Forecasting acceptance of new students using double exponential smoothing method. *Journal of Critical Review*, 7(1), 300-305.
 - [22] Viveros, P., González, K., Mena, R., Kristjanpoller, F., & Robledo, J. (2021). Slotting Optimization Model for a Warehouse with Divisible First-Level Accommodation Locations. *Applied Sciences*, 11(936), 936.

- [23] Yamak, P., Yujian, L., & Gadosey, P. (2019). A comparison between arima, lstm, and gru for time series forecasting. In Proceedings of the 2019 2nd international conference on algorithms, computing and artificial intelligence, 49-55.

Financiamiento:

Propia.

Conflictos de interés:

Los autores declaran no tener conflictos de interés.

Contribuciones de autoría:

Agusto Guillermo Sánchez Ferrer: Conceptualización, visualización, curación de datos, software, roles/ escritura: borrador original, Metodología, redacción: revisión y edición, Investigación.

José Herrera Quispe: Análisis formal, Administración de proyectos, Supervisión, Validación, Recursos.