

Comparative Analysis of Time Series Forecasting Methods for Predicting Spare Parts Demand in Supply Chains

Ítalo P. Caliari¹, Ravi B. D. Figueiredo¹, Diogenes W. F Silva¹, Edierley B. Messias¹,
Mario Haddad Neto¹

¹Fundação Paulo Feitoza - FPFtech, Av. Danilo de Matos Areosa, 1170, Distrito Industrial, Manaus, Amazonas, 69075-351, Brazil.

{italo.caliari, ravi.figueiredo, diogenes.silva, edierley.messias,
mario.haddad}@fpf.br

Abstract. *Accurate demand forecasting prevents stockouts and excess inventory, but sparse data for many components hampers model building. To address this, interchangeable components were grouped, enabling analysis of aggregated warranty, repair, and Consumption Index (CI) data for 3,000 groups. Classical time series techniques, including the Simple Moving Average (SMA) and Exponential Moving Average (EMA) were evaluated along with different SARIMAX implementations. Performance analysis was conducted using the Root Mean Square Error (RMSE) and Sufficiency metrics. The results of this study highlight the effectiveness of simple approaches compared to more complex methods, such as neural networks and XGBoost, which did not outperform baseline algorithms in terms of RMSE and Sufficiency. However, simpler methods have limitations in capturing complex dynamics.*

1. Introduction

Companies in the electronics and microcomputers sector manufacture and distribute a diverse range of products, including computers, peripherals, and mobile devices, all while navigating a highly competitive and dynamic environment. A significant challenge these companies face is managing the demand for components in technical service centers, particularly in fulfilling Service Level Agreements (SLAs) for repairing products that develop defects during the warranty period. Meeting these SLAs is crucial, as they establish clear expectations regarding response times, resolution deadlines, and the availability of necessary parts.

The logistics of ensuring that service centers have the right materials in the appropriate quantities and at the right time are complex. Accurate demand forecasting is essential for optimizing inventory management and ensuring that service centers are adequately stocked. Inadequate forecasting can lead to stockouts, resulting in service delays and customer dissatisfaction. This dissatisfaction can directly impact a company's reputation and lead to lost customers. Conversely, excess inventory incurs unnecessary costs and resource waste, negatively affecting profitability. Companies that struggle to balance supply and demand face high storage costs and risks of product obsolescence.

Demand for repair parts often exhibits seasonal fluctuations influenced by factors such as weather changes, new product launches, and shifts in consumer behavior. These variations complicate demand forecasting, necessitating robust predictive models that can effectively capture these dynamics.

Despite the challenges, there is limited research specifically addressing the demand for components in technical assistance companies, although related fields have been explored extensively [Fanoodi et al. 2019, Ameer and Fatahi Valilai 2025, Ji and Wang 2017]. For instance, Mahin *et al.* (2025) presented an innovative approach to sales prediction through advanced machine learning methods, highlighting the effectiveness of the Voting Regressor in improving forecasting accuracy and optimizing supply chain operations [Rahman Mahin et al. 2025]. Similarly, Espinel *et al.* (2025) explored the integration of econometric modeling and deep learning strategies to improve forecasting accuracy within the pharmaceutical supply chain, ultimately benefiting patient care and operational efficiency [Sierra Espinel and Suarez Barón 2025]. Andrianakis (2024) *et al.* developed and implemented PredictionSCMS, an integrated supply chain management system designed to optimize product flow, prevent sales losses due to stock shortages, and enhance resource efficiency through inventory management and supply chain automation [Andrianakis et al. 2024].

This paper offers a comprehensive analysis of various predictive models for forecasting the replacement of parts in technical assistance for electronic components, drawing insights from supply chain literature across different sectors. The research systematically compares the performance of diverse approaches across various forecast horizons and component life-cycle stages. It evaluates a broad spectrum of algorithms, ranging from traditional methods such as Simple Moving Average (SMA) and Exponential Moving Average (EMA) to advanced techniques like auto optimizable SARIMAX algorithms, neural networks, Convolutional Neural Networks (CNN), Long Short-Term Memory (LSTM), Prophet [SJ and B. 2017], and XGBoost. The primary objective of the study is to identify the most accurate and efficient models for demand forecasting within the context of technical assistance in the electronics industry, thereby optimizing inventory management and enhancing customer service. The findings demonstrate how advanced forecasting techniques, real-time data utilization, and proactive risk management can significantly enhance supply chain operations, agility, and competitiveness. Furthermore, the paper underscores the critical importance of collaboration across the supply chain, emphasizing that information sharing among stakeholders can lead to improved decision-making and increased resilience. Overall, this multifaceted approach not only streamlines operations but also bolsters the agility and competitiveness of the supply chain, providing valuable insights into the field of supply chain management.

2. Methodology

2.1. Database

One of the primary challenges faced in this study was the limited availability of data for many individual components, which significantly hindered the effective development and evaluation of predictive models. In numerous instances, historical demand records were inadequate for conducting a robust analysis. To address this limitation, the study adopted the concept of component subgroups or interchangeable components. Interchangeability refers to the characteristic of components and sub-assemblies used in modern manufactured products, where each unit produced is designed to be identical to every other unit. The fundamental principle is that all units must conform to design specifications within a narrowly defined allowable range of variation. These specifications

typically encompass all characteristics of the component that are relevant to the manufacturing process and the customer's requirements, ensuring consistency and reliability in production [Swamidass 2000].

This approach facilitated the development of predictive models using aggregated data from subgroups with similar usage characteristics. A dedicated engineering technical team was responsible for identifying and grouping interchangeable components, ensuring that the subgroups accurately reflected the relationships among the parts. By combining information from various components, it was possible to create a comprehensive database comprising the replacements usage history of technical assistance for over 3,000 component subgroups. This strategy not only alleviated issues related to data scarcity but also provided valuable insights into consumption patterns across related component groups, allowing for more precise adjustments in parts supply planning strategies.

For each interchangeable subgroup, monthly historical data was available regarding the number of components under warranty (*warranty*), the number of items that were considered damaged and were replaced in technical assistance (*replacements*), and the Replacements Index (RI) provided by $100 * replacements / warranty$.

2.2. Pre-processing

Before training the predictive models, time series data normalization is performed. This step is crucial because if the values in the dataset are on different scales, it can result in a model with suboptimal performance, as certain features may disproportionately influence the training process [Kherdekar and Naik 2024]. In this study, SARIMAX, Prophet, XGBoost, and Neural Networks were trained both with and without normalization, utilizing two normalization techniques: Min-Max Scaling and Standard Scaling. This approach allows for a comprehensive evaluation of the impact of normalization on model performance.

2.3. Prediction models

This subsection describes the forecasting methods used to model and predict replacements at technical assistance centers.

2.3.1. Zeroes forecast

For zero forecasts, all predictions are set to zero. This type of prediction serves as a crucial baseline for comparison with other methods, allowing for the evaluation of the impact of not taking any action in anticipating technical assistance needs.

2.3.2. Naive forecast

For naive forecasts, all predictions are set to the value of the last observation in the training set [Hyndman and Athanasopoulos 2021]. The naive model was developed using historical data from the Replacements Index (RI). The predicted RI values were then multiplied by the warranty data to estimate the number of replacements at technical assistance.

2.3.3. Random Walk

The random walk model is a statistical forecasting approach that posits that future values in a time series are equal to the most recent observed value plus a random error term. This assumption suggests that future movements are inherently unpredictable and have an equal probability of increasing or decreasing. Consequently, the random walk model serves as the foundation for naive forecasts [Hyndman and Athanasopoulos 2021]. It can be expressed as:

$$\hat{y}_{T+1} = y_T + \epsilon_T, \quad (1)$$

where $\hat{y}_{T+1}$ is the forecast, y_t is the last observed value, and ϵ_T is random noise.

The random walk model was developed using historical data from the Replacements Index (RI). The predicted RI values were then multiplied by the warranty data to estimate the number of replacements at technical assistance. The Random Walk algorithms implemented in this paper were based on the default configurations provided by the Python library `StatsForecast 2.0.3`.

2.3.4. Simple Moving Average (SMA)

The Simple Moving Average (SMA) is a statistical method employed in time series analysis and forecasting. It calculates the average of a selected set of data points over a specified period. In this approach, each data point in the time series is given equal weight, meaning that no weighting factors are applied to any of the observations [Hansun 2013]. It can be expressed as:

$$SMA_t = \frac{Y_{t-1} + Y_{t-2} + \cdots + Y_{t-n}}{n}, \quad (2)$$

where SMA_t is a simple Moving Average forecast for observation t , Y_t denotes the data point at observation t , and n indicates the number of observations within the moving window.

The SMA model was developed using historical data from the Replacements Index (RI). The predicted RI values were then multiplied by the warranty data to estimate the number of replacements at technical assistance. The SMA algorithms implemented in this paper were based on the default configurations provided by the Python library `Pandas 3.0.0`.

2.3.5. The Exponential Moving Average (EMA)

The Exponential Moving Average (EMA) is a variant of the moving average that gives greater weight to more recent data points, making it more responsive to new information compared to the Simple Moving Average (SMA). Unlike traditional moving averages, the EMA is mathematically defined as a limit of an exponential average with

a specific weighting factor [Klinker 2011]. The weights assigned to older data points decrease exponentially, ensuring that they never reach zero. It can be expressed as:

$$EMA_t = \frac{(1 - \alpha)Y_{t-1} + (1 - \alpha)^2Y_{t-2} + \dots + (1 - \alpha)^nY_{t-n}}{(1 - \alpha) + (1 - \alpha)^2 + \dots + (1 - \alpha)^n}, \quad (3)$$

where EMA_t denotes the Exponential Moving Average forecast for observation t , Y_t denotes the data point at observation t , n indicates the number of observations within the moving window, and α represents the degree of weighting decrease, a constant smoothing factor between 0 and 1.

The EMA model was developed using historical data from the Replacements Index (RI). The predicted RI values were then multiplied by the warranty data to estimate the number of replacements at technical assistance. The EMA algorithms implemented in this paper were based on the default configurations provided by the Python library `Pandas 3.0.0`.

2.3.6. Seasonal AutoRegressive Integrated Moving Average with Exogenous Variables - SARIMAX

Another well-established statistical algorithm available is the SARIMAX model [Box et al. 2015]. SARIMAX, which stands for Seasonal AutoRegressive Integrated Moving Average with Exogenous Variables, is a statistical model used for time series forecasting that integrates several key components: seasonality (S), which captures seasonal patterns in the data; auto-regression (AR), which utilizes past values to predict future outcomes; differencing (I), which eliminates trends to achieve stationarity; moving averages (MA), which leverage past forecast errors to enhance predictions; and external variables (X) that can help explain the primary behavior of the time series. The model is typically denoted as SARIMAX(p, d, q)(P, D, Q, M), where p, d, and q represent the orders of the respective components [Box et al. 2015]. The parameter p denotes the order of the autoregressive component, specifying the number of past values utilized in the model. d represents the degree of differencing applied to achieve stationarity in the time series. The parameter q indicates the order of the moving average component. D denotes the number of seasonal differences required to achieve stationarity in the seasonal component of the time series, similar to d , which helps eliminate seasonal trends. Q refers to the number of seasonal moving average terms, representing the number of lagged forecast errors from the seasonal period included in the model. Finally, M indicates the number of periods within each season.

The SARIMAX model was developed using historical data from the monthly *replacements* as the primary forecasting variable, while monthly *warranty* data was incorporated as an exogenous variable. The SARIMAX algorithms discussed in this paper were developed using the default configurations from the Python libraries `pmdarima 2.1.1` and `StatsForecast 2.0.3`, with a focus on yearly seasonality. The parameters p, d, q, P, D, Q, and M for the SARIMAX model were automatically selected using the `autoarima` functions available in both libraries.

2.3.7. Other studied approaches

In addition to traditional approaches, this study also evaluated various modern techniques for time series forecasting, including neural networks, CNN (Convolutional Neural Networks), LSTM (Long Short-Term Memory), as well as machine learning methods such as Prophet and XGBoost. These methodologies were considered aiming to identify the most effective approach for the forecasting task at hand.

2.4. Metrics

The evaluation of predictive models' quality is essential to ensure accuracy and effectiveness in predictions, especially in contexts where data-driven decisions have a significant impact. The Root Mean Square Error (RMSE) and Sufficiency metrics were used to measure the performance of the developed predictive models.

2.4.1. Root Mean Square Error (RMSE)

RMSE was chosen due to its capacity to penalize larger errors more severely. This characteristic is particularly important in the current problem domain, where significant deviations can result in substantial disruptions to inventory planning. Consequently, RMSE was favored over other metrics such as MAE (Mean Absolute Error), MASE (Mean Absolute Scaled Error), and MAPE (Mean Absolute Percentage Error). RMSE quantifies the difference between the values predicted by the model and the actual observed values, providing a clear indication of how well the model is performing in forecasting the data. One of the advantages of RMSE is that its value is expressed in the same units as the variable of interest, which facilitates the interpretation of results and allows for a more intuitive assessment of the model's performance. A lower RMSE indicates that the model performs better, as the predicted values are closer to the actual values [Liemohn et al. 2021]. The RMSE is calculated as follows:

$$RMSE = \sqrt{\frac{\sum_{i=1}^n (y_i - \tilde{y}_i)^2}{n}}, \quad (4)$$

where y_i is the real value of the observation i , $\tilde{y}_i$ is the predicted value of the observation i , and n is the total number of observations.

2.4.2. Sufficiency

The Sufficiency Metric is an important tool for evaluating the quality of predictive models, especially in contexts where meeting Service Level Agreements (SLAs) is critical. It is defined as the proportion of predictions that meet or exceed the actual demand value. In other words, if a predictive model forecasts a demand of 100 units and the actual demand is 90 units, that prediction is considered sufficient. Sufficiency can be expressed as the ratio between the number of sufficient predictions and the total number of predictions, multiplied by 100.

This metric helps assess the model's effectiveness in predicting values that meet demand, making it crucial for operations that rely on service levels. Models with high

sufficiency are preferable, as underestimating demand can lead to financial losses or customer dissatisfaction. It is important to use sufficiency alongside other performance metrics, such as RMSE, for a more comprehensive evaluation. Sufficiency is calculated as follows:

$$S = \frac{N_s}{N_t} \times 100, \quad (5)$$

where N_s is the number of forecasts that meet or exceed the actual demand value, and N_t is the total number of predictions.

2.5. Experimental Setup

The different algorithms were evaluated using roughly 3,000 subgroup-specific forecasting models, allowing a comprehensive assessment of each approach's effectiveness across a large and diverse set of real world demand series. This large-scale evaluation provides insight into both predictive accuracy and the robustness of the algorithms. Models were trained on data up to each reported date and evaluated on their six-month-ahead forecasts. Evaluation windows span July 2021 through October 2024, ensuring the results reflect performance across multiple market conditions and time periods and therefore provide a reliable measure of each method's temporal stability and generalizability. Computation was performed in a Databricks environment using PySpark, leveraging distributed cluster resources for scalable model training and evaluation.

3. Results and Discussion

The results presented in this section reflect the overall performance of the forecasting models derived from approximately 3,000 analyzed subgroups, each with its own model. This extensive array of models enables a thorough evaluation of the effectiveness of the employed approaches, offering valuable insights into the efficacy of each method in addressing the demand forecasting problem. Additionally, it highlights the predictive capabilities and robustness of the various algorithms used. Results for the neural networks, LSTM, CNN, Prophet, and XGBoost experiments have been omitted from this section, as these models either failed to converge or did not outperform the baseline methods. These results indicate that the insufficient availability of data may have been a limiting factor for these algorithms. Instead of omitting any mention of these algorithms in this study, it is important to acknowledge that, although these approaches are popular and promising, they still struggle to model smaller datasets and in some cases, more traditional methods remain preferable.

Figure 1 illustrates the RMSE results for the algorithms evaluated in forecasting the demand for electronic components over a six-month period. The robustness of these approaches was assessed by simulating the models' effectiveness across different time frames, spanning from July 2021 to October 2024. This extended evaluation is crucial for validating the reliability of the algorithms, ensuring that the results remain applicable and relevant over time.

Figure 2 presents box-plots displaying the aggregated RMSE values from the experiments for each evaluated method, highlighting the mean ($\times$), median ($—$), and the 25% ($Q1$) and 75% ($Q3$) quartiles. Additionally, the box-plots mark outliers—values that

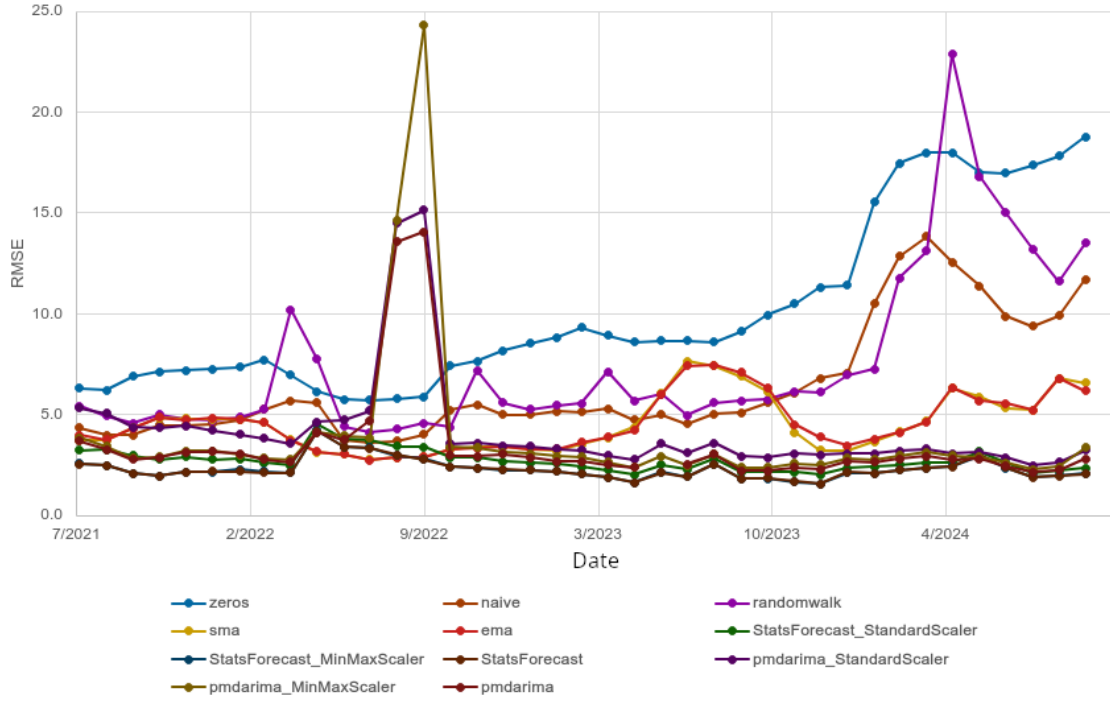


Figure 1. RMSE results for the six-month forecast for models with different pre-processing techniques.

lie significantly outside the data distribution, as determined by the inter-quartile range ($IQR = Q3 - Q1$). The lower limit is defined as $Q1$ minus 1.5 times the IQR , while the upper limit is $Q3$ plus 1.5 times the IQR . Values that fall outside these limits are considered outliers and are represented in the box-plot as individual points ($\bullet$). This graphical representation allows for a clear visual analysis of the dispersion and central tendency of the results, facilitating comparisons between the different forecasting methods.

Based on Figures 1 and 2, it can be observed that simpler methods such as zero forecasting, naive forecasting, and random walk exhibit relatively high RMSE values, mostly ranging between 4.5 and 12. This indicates that these methods, which do not account for historical patterns or data complexities, perform poorly in demand forecasting, reflecting their limitations in capturing the real dynamics of the problem. This outcome is anticipated, as these algorithms are commonly employed as performance baselines that more advanced methods should aim to exceed.

The SMA and EMA algorithms demonstrated relatively lower RMSE values compared to the reference methods, indicating that, despite their simplicity and lack of optimization, these approaches still contribute to the planning process. The RMSE results for these techniques were observed to range predominantly between 3.0 and 5.5, demonstrating their effectiveness in smoothing short-term fluctuations and identifying general trends. However, the nature of these approaches may make them less sensitive to abrupt demand changes, which limits their ability to capture relevant trends and seasonality in the business context. This limitation can be a disadvantage in scenarios where demand is volatile or seasonal, as the inability to quickly react to these variations may result in inaccurate forecasts and, consequently, inventory management challenges. Therefore, al-

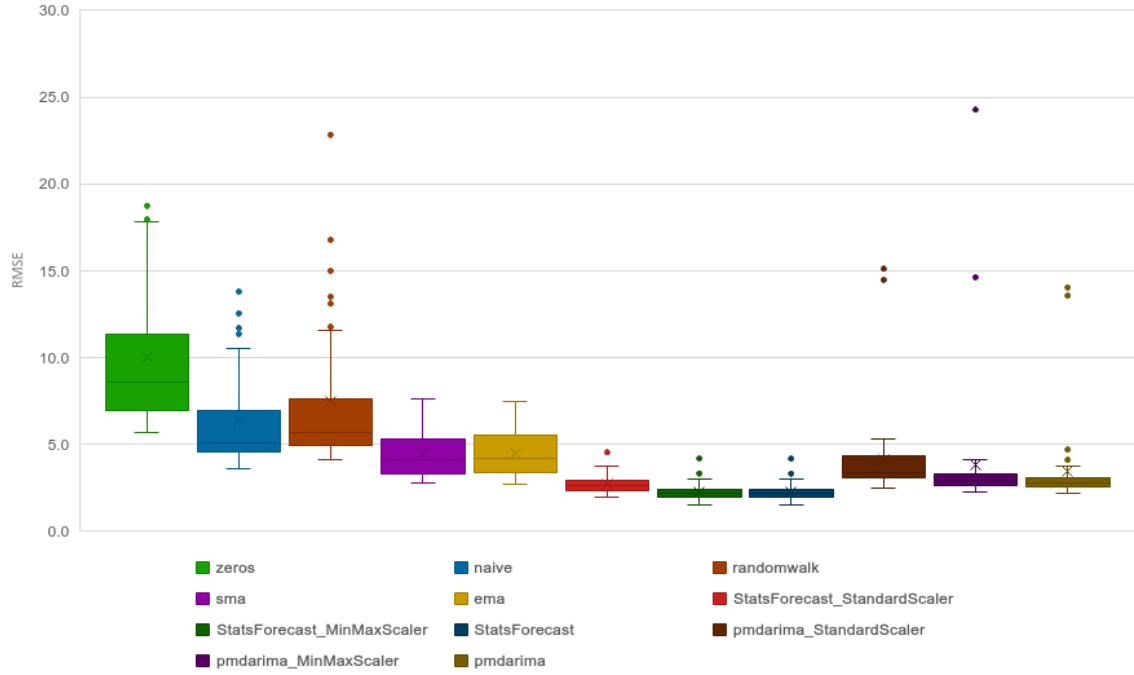


Figure 2. Box-plot for RMSE results for the six-month forecast for models with different pre-processing techniques.

though SMA and EMA are useful tools for demand analysis, it is essential to consider their limitations when applying them in dynamic and constantly changing environments, suggesting the need to complement them with more complex *models* depending on the application’s criticality.

The results obtained for the SARIMAX algorithms from the `StatsForecast` and `pmdarima` libraries demonstrate remarkable performance in demand forecasting. The `StatsForecast` models predominantly exhibited RMSE values ranging between 2.0 and 3.5, indicating superior performance compared to simpler methods such as SMA and EMA. On the other hand, the `pmdarima` results proved equally promising, with RMSE values mostly varying between 2.5 and 4.0. The comparison between `StatsForecast` and `pmdarima` results suggests that both are more effective in demand forecasting when compared to SMA and EMA methods. In summary, the results from `StatsForecast` and `pmdarima` auto-optimized SARIMAX algorithms demonstrate greater effectiveness in demand prediction.

Figures 3 and 4 present the Sufficiency results for the evaluated algorithms in forecasting the demand for electronic components. The results presented in Figures 3 and 4 reveal that simpler methods, such as the naive forecast and random walk, exhibit relatively low Sufficiency indices, ranging between 85.0% and 87.5%. This performance is particularly concerning when we consider that approximately 74.0% of the evaluated cases show no part replacements, as indicated by the zero-forecast results.

The SMA and EMA algorithms demonstrated slightly higher Sufficiency values compared to the naive and random walk methods. These methods predominantly guaranteed between 86.0% and 88.0% sufficiency for SLA compliance, demonstrating their limited yet slightly superior effectiveness over baseline methods. This limitation is antic-

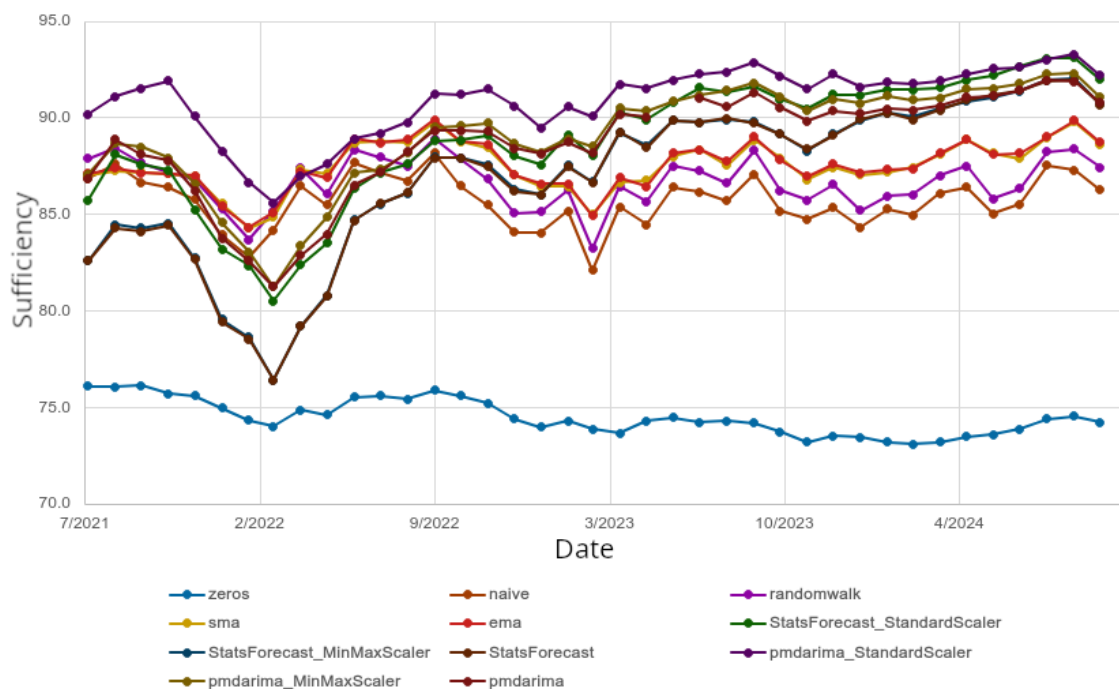


Figure 3. Sufficiency results for the six-month forecast for models with different pre-processing techniques.

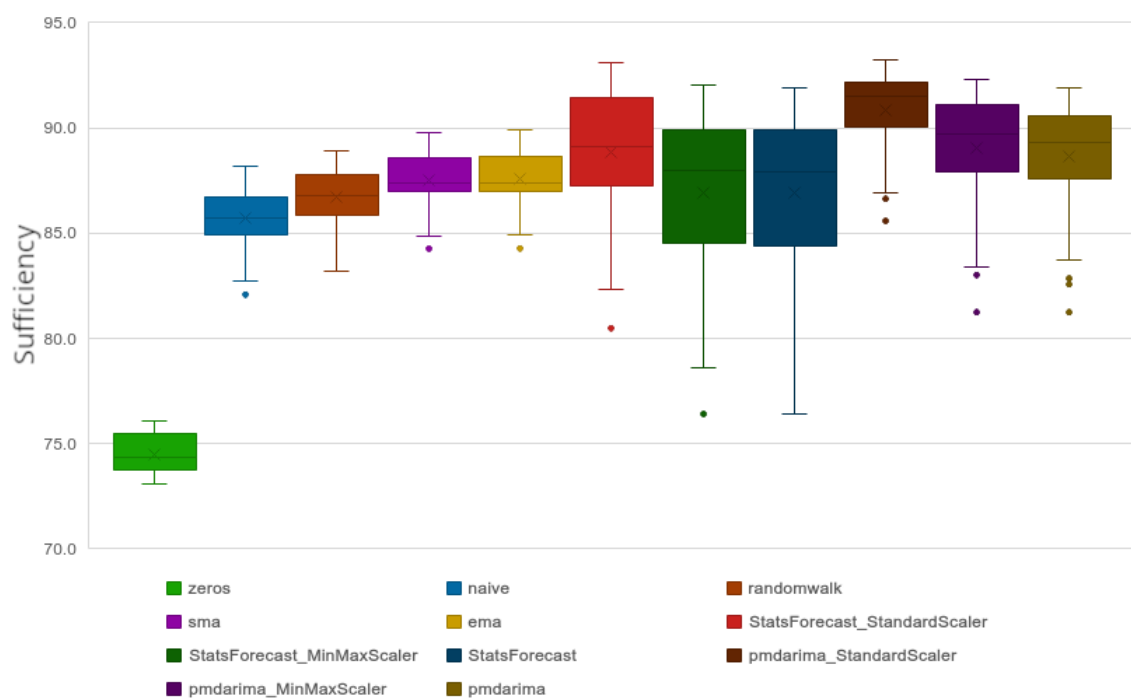


Figure 4. Sufficiency box-plot for the six-month forecast for models with different pre-processing techniques.

ipated, as these algorithms do not account for seasonality effects or cyclical events.

In contrast, the `StatsForecast` models, and particularly `pmdarima`, demonstrate considerably superior performance. The results indicate that when using pre-processing techniques such as `Standard Scaling` and `Min-Max Scaling`, the sufficiency can be further increased, reaching values up to 93.0% in some cases. This suggests that data normalization can improve the models' ability to learn patterns and trends, resulting in higher SLA compliance.

As a final evaluation of the results, `pmdarima` combined with `Standard Scaling` pre-processing technique achieves the highest sufficiency rates, predominantly ranging between 90.0% – 93.0%. Concurrently, this experiment demonstrates well-optimized and controlled RMSE values.

Given that full SLA compliance offers organizations substantial benefits, such as increased customer satisfaction, enhanced market reputation, reduced operational costs, and improved competitive positioning, this metric is prioritized, especially in the context of maintaining effective RMSE control. Consequently, the `pmdarima` method with `Standard Scaling` pre-processing was selected as the demand forecasting solution for technical support operations.

4. Conclusion

This study evaluated a spectrum of forecasting methods for electronic component demand and demonstrated that method selection materially affects forecast accuracy and operational outcomes. Simple benchmarks (naive, random walk, SMA, EMA) are useful reference points but generally fail to capture the complexity of demand signals. In contrast, SARIMAX-based approaches implemented via `StatsForecast` and `pmdarima` consistently delivered lower RMSE and higher sufficiency rates, indicating superior ability to model trends, seasonality, and exogenous effects. More complex machine-learning and deep-learning models (neural networks, LSTM, CNN, Prophet, XGBoost), although promising, suffered from scarce per-subgroup data availability, which reduced their practical effectiveness in this setting.

Data pre-processing (standard scaling, min-max scaling) improved model performance, underscoring the importance of normalization in forecasting pipelines. Structuring the dataset into meaningful subgroups was critical: subgrouping mitigated temporal data sparsity, revealed otherwise-hidden patterns, and relied on domain engineering to ensure consistent component classification. Operationally, the `pmdarima` SARIMAX approach combined with appropriate pre-processing emerged as the most effective solution for technical-support demand forecasting.

Finally, integrating predictive models into an automated system linked with internal and inventory-management platforms modernized material planning, improved SLA compliance, and enhanced customer satisfaction, operational efficiency, and cross-functional coordination. These results recommend prioritizing robust, statistically grounded models with careful pre-processing and subgrouping when implementing demand-forecasting systems in component-driven operations.

References

- Ameer, T. and Fatahi Valilai, O. (2025). Predictive exploratory data analysis of shopfloor cnc machine operation through a machine learning model. *Journal of Open Innovation: Technology, Market, and Complexity*, 11(2):100559.
- Andrianakis, I., Gkatas, V., Eleftheriadis, N., Ellinidis, A., and Avramidou, E. (2024). Predictionscms: The implementation of an ai-powered supply chain management system. *International Journal of Industrial and Manufacturing Engineering*, 18(3):65 – 71.
- Box, G. E. P., Jenkins, G. M., Reinsel, G. C., and Ljung, G. M. (2015). *Time Series Analysis: Forecasting and Control*. 5 edition.
- Fanoodi, B., Malmir, B., and Jahantigh, F. F. (2019). Reducing demand uncertainty in the platelet supply chain through artificial neural networks and arima models. *Computers in Biology and Medicine*, 113:103415.
- Hansun, S. (2013). A new approach of moving average method in time series analysis. In *2013 Conference on New Media Studies (CoNMedia)*, pages 1–4.
- Hyndman, R. and Athanasopoulos, G. (2021). *Forecasting: principles and practice*,. OTexts: Melbourne, Australia.
- Ji, W. and Wang, L. (2017). Big data analytics based fault prediction for shop floor scheduling. *Journal of Manufacturing Systems*, 43:187–194.
- Kherdekar, V. A. and Naik, S. (2024). Impact of feature normalization techniques for recognition of speech for mathematical expression. In Senjyu, T., So-In, C., and Joshi, A., editors, *Smart Trends in Computing and Communications*, pages 109–117, Singapore. Springer Nature Singapore.
- Klinker, F. (2011). Exponential moving average versus moving exponential average. *Math. Semesterber.*, 58(1):97–107.
- Liemohn, M. W., Shane, A. D., Azari, A. R., Petersen, A. K., Swiger, B. M., and Mukhopadhyay, A. (2021). Rmse is not enough: Guidelines to robust data-model comparisons for magnetospheric physics. *Journal of Atmospheric and Solar-Terrestrial Physics*, 218:105624.
- Rahman Mahin, M. P., Shahriar, M., Das, R. R., Roy, A., and Reza, A. W. (2025). Enhancing sustainable supply chain forecasting using machine learning for sales prediction. *Procedia Computer Science*, 252:470–479. 4th International Conference on Evolutionary Computing and Mobile Sustainable Networks.
- Sierra Espinel, A. I. and Suarez Barón, M. J. (2025). Applying deep learning and forecasting techniques to the pharmaceutical supply chain. *Procedia Computer Science*, 253:2791–2800. 6th International Conference on Industry 4.0 and Smart Manufacturing.
- SJ, T. and B., L. (2017). Forecasting at scale. *PeerJ Preprints*, 5:e3190v2.
- Swamidass, P. M. (2000). *INTERCHANGEABILITY*, pages 293–294. Springer US, New York, NY.