

Comparing Machine Learning Forecasting Models Based on Accuracy and Efficiency for Predicting Demand in a Food and Beverage Company

**Ray Baltazar Alunen, Cyrene Franchesca
Molina Raven Franchesca Quesada
and Chloe Nicole Reyes**

Department of Industrial Engineering
Faculty of Engineering
University of Santo Tomas, Manila
Philippines

raybaltazar.alunen.eng@ust.edu.ph, cyrenefranchesca.molina.eng@ust.edu.ph,
ravenfranchesca.quesada.eng@ust.edu.ph, chloenicole.reyes.eng@ust.edu.ph

Engr. Delfin Jacob, MSIE, PIE
Department of Industrial Engineering
Faculty of Engineering
University of Santo Tomas, Manila
Philippines
drjacob@ust.edu.ph

Abstract

Accurate demand forecasting significantly benefits the food and beverage industry – minimizing waste, optimizing inventory, and meeting market demands. Traditional forecasting methods fail to capture non-linear relationships and external factors such as holidays, weather conditions, and macroeconomic factors. This study developed a machine learning-based forecasting framework by working on alcoholic beverages as a representative product category, given their growing per capita consumer spending and consumption. The research aimed to evaluate the performance of different ML models in comparison to one another and other traditional forecasting techniques. The models also incorporated feature selection and hyperparameter tuning to optimize predictions and were assessed through different accuracy metrics. Results demonstrated that XGBoost excelled in both accuracy and computational efficiency. Feature selection using correlation enhanced computational efficiency, but led to a slight reduction in forecast accuracy. Additionally, hyperparameter tuning methods of Random Search outperformed Grid Search in both accuracy and execution time. Overall, the study recommended adding more factors, leveraging algorithms for other applications, incorporating hyperparameter tuning, and investing in data.

Keywords

Demand forecasting, alcoholic beverages, machine learning, feature selection, hyperparameter tuning

1. Introduction

The Philippine food and beverage (F&B) industry plays a critical role in satisfying national food requirements, yet it faces numerous challenges, particularly in supply chain efficiency and demand forecasting. This sector is rapidly

evolving, increasing investments and shifting consumer preferences toward processed and ready-to-eat foods. Retail F&B sales in the Philippines were projected to grow by 10% annually starting in 2021, highlighting the urgent need for better alignment between industry operations and customer demand.

However, the industry struggles with implementing sophisticated supply chain practices, particularly in forecasting demand. The tropical climate exacerbates the challenge by accelerating the spoilage of perishable goods, emphasizing the need for accurate forecasting to ensure efficient resource allocation, inventory management, and production planning. Many establishments lack the tools and models to anticipate demand effectively, leading to inefficiencies such as overstocking, understocking, and increased operational costs.

In the alcoholic beverage segment, which forms a significant part of the F&B industry, these challenges are even more pronounced. Alcohol consumption patterns are influenced by seasonality, public holidays, and socio-economic conditions. Warmer weather and cultural events drive higher consumption, while preferences vary across demographics. With consumer spending on alcoholic beverages expected to grow substantially in the coming years, understanding and predicting demand in this sector is more critical than ever. Despite its economic significance, research on demand forecasting for alcoholic beverages, especially in the Philippine context, remains limited. Current practices often overlook crucial external factors such as regional weather patterns, macroeconomic indicators, and cultural nuances. These gaps lead to suboptimal decisions, resulting in waste and lost revenue opportunities.

This study focused on developing a machine learning-based framework that integrated external variables to improve demand forecasting accuracy. Unlike traditional methods, this approach captured complex and non-linear relationships in data, enabling more precise predictions. The framework was designed to address the specific needs of the alcoholic beverage sector while also being adaptable to other categories within the F&B industry. This research sought to revolutionize supply chain practices in the Philippine F&B industry by addressing critical gaps in forecasting accuracy and enhancing its resilience and adaptability to meet consumer demands.

1.1 Objectives

This study aimed to quantify the performance and accuracy of a proposed forecasting framework utilizing different machine learning algorithms, basic hyperparameter tuning models, and feature selection or a lack thereof. Machine learning algorithms, namely Random Forest Regression and Boosted Regression Trees (e.g., Gradient-Boosted, XGBoost, AdaBoost), were evaluated. The study also aimed to evaluate the performance of forecasting algorithms with grid search and random search hyperparameter tuning. The presence of feature selection is also assessed. The best framework is determined by comparing the computational efficiency and forecasting accuracy using a range of metrics.

2. Literature Review

Accurately forecasting demand is crucial in effectively managing a food and beverage (F&B) outlet. It allows for precise inventory ordering that minimizes pre-consumer waste and is a key input for revenue, operations, and product management (Lasek et al. 2016). However, most F&B outlets are small businesses with limited access to advanced forecasting technologies, prompting experience- and intuition-based forecasting heuristics. This results in inaccurate predictions as such techniques exclude external factors such as weather conditions, holidays, and events (Groene and Zakharov 2024; Posch et al. 2022).

In addition to manual forecasting and related heuristics, studies also discuss applying other traditional forecasting models in the F&B industry. Jiang et al. (2020) also used a multiple regression model to study the demand for vodka products and compared the results with an autoregressive integrated moving average (ARIMA) model. A Box-Jenkins method has been further applied in forecasting overall alcohol consumption and consumption by alcohol type (e.g. wine, beer, spirits) concerning socio and economic factors (Slováčková et al. 2016). Hallak et al. (2022) employed Poisson regression in predicting the demand for healthy beverages through the willingness to pay a premium; the data about motivational, demographic, and behavioral variables was also analyzed. Tirkeş et al. (2017) employed exponential smoothing and Holt-Winters in forecasting the demand for jam and sherbet products, exhibiting low reliance on historical data by assigning increasing weights to recent observations.

Forecasting in the F&B industry is further characterized by its complexity and is influenced not just by historical sales data. Demand tends to peak particularly on Fridays and Saturdays, holidays, summer months, and even specific meal

times such as lunch or dinner (Lasek et al. 2016). Meanwhile, variations in weather can support seasonality trends (Bujisic et al. 2016). Hirche et al. (2021), using seasonal ARIMA, included the effects of holidays in analyzing the sales of alcoholic beverages, in addition to considering day temperature and geographical locations. Public holidays generally boost sales, especially in locales with warmer temperatures. Regarding the macro perspective, Sadik-Zada and Niklas (2021) state that alcohol consumption volume is significantly influenced by economic expansion and contraction. This is further supported by Collins (2016), with increasing alcohol consumption being triggered by periods of recessions and unemployment. Unemployment can generally decrease food spending (Restrepo et al., 2021) as consumers would have lower disposable income; overall, demand for F&B products would decline, particularly in restaurants and premium food segments (Kumar 2021). Location-related factors, on the other hand, such as city sales and geographic distribution, can determine where most demand occurs (Nassibi et al. 2023).

Over time, machine learning (ML) models have been used; by leveraging ML, food companies can better adapt to market fluctuations and optimize their supply chains (Falatouri et al. 2022). ML models outperform traditional linear models, often affected by human error; ML techniques provide higher predictive capacity and can handle big datasets (Tsoumakas 2019). ML models recursively look for statistical regularities and adapt to evolving patterns to reduce error and increase accuracy, even without pre-defined assumptions (Chan et al. 2022). Priyadarshi et al. (2019) compared random forest regression, gradient-boosted regression, and extreme gradient-boosting (XGBoost) regression models in determining the optimal daily quantity of fresh produce. XGBoost regression best performs in terms of accuracy (Yerragudipadu et al. 2023); the model is further applied by related studies in hourly sales data of a fast-food outlet (Groene and Zakharov 2024) and daily customer sales and POS data (Tanizaki et al. 2019). ML models outperform heuristic forecasts, reducing prediction errors by 22% to 33% in root mean squared error and 19% to 31% in mean average error (Groene and Zakharov 2024).

Other ML algorithms include random forest regression, which breaks down the best prediction into sub-decisions (Gianey and Choudary 2017). The model is easily interpretable but prone to overfitting. Boosted decision trees combine multiple decision trees and iteratively focus on the learning times of incorrectly forecasted instances by increasing the weight allocation of those cases to improve predictive accuracy (Prajwala 2015; Tanizaki et al. 2019). Adaptive boosting is the first application to update the dataset incrementally and adjust coefficients, focusing on those with higher inaccuracies (Azmi and Baliga 2020).

To improve the accuracy of ML models and save computational resources, irrelevant factors are removed, and dimensions are reduced through feature selection (FS). FS allows the removal of irrelevant and noisy predictors (Zebarri et al. 2020). The most widely used FS method is the Pearson correlation coefficient, which is computationally simple and efficient (Venkatesh and Anuradha 2019). In machine learning, hyperparameters exist outside the model itself and are not directly learned from the data. Different techniques have been developed to tune hyperparameters and avoid manual selection. The most used and simplest ones to implement include Grid Search and Random Search (Mantovani 2016; Aguiar and Cano 2023). Grid search assesses all combinations of hyperparameters to determine which produces the best-performing model but is subject to the curse of dimensionality. Random search is more practical as it searches for combinations strategically and randomly (Liashchynskyi and Liashchynskyi 2021). Evaluation metrics, on the other hand, assess forecasting performance and how well the models can capture trends (Rodrigues et al. 2024):

- **Mean Squared Error (MSE).** MSE quantifies the mean squared differences between the predicted and true values, penalizing larger errors (Aci and Yergok 2023). Generally, an MSE closer to zero is favorable.
- **Mean Absolute Error (MAE).** The MAE measures the average magnitude in which the model's predictions are incorrect, measuring maximum losses; hence, the ideal value is always zero (Wiyanti et al. 2021).
- **Root Mean Squared Error (RMSE).** RMSE is the mean squared error between the forecasted and true values and is sensitive toward outliers due to squared prediction errors. A low RMSE is preferred (Yoo and Oh 2020).
- **Coefficient of Determination (R^2).** R^2 explains the proportion of the variation in the dependent variable that can be explained by the independent variable(s); a value closer to 1 indicates a more substantial relationship (Aci and Yergok 2023).
- **Execution Time.** Du et al. (2013) measured execution time in seconds to quantify how long it takes for a model to train and test a particular data set.

In the debate on which evaluation metric demonstrates superiority over the other, it is often necessary to use a combination of metrics to assess model performance comprehensively. MSE and RMSE is optimal for normally

distributed errors (Hodson 2022) and account for the variance of errors, resulting in the metrics being more sensitive to outliers (Brassington 2017). MAE is preferred for Laplacian errors and is more interpretable in quantifying average error (Robeson and Willmott 2023). On the other hand, Chicco et al. (2021) indicate that R^2 is more informative in regression analysis evaluation.

The different studies highlighted the significance of demand forecasting in the F&B industry and alcoholic beverages. Accurate demand forecasting methodologies enable organizations to align their operations with the market demand. Traditional forecasting methods have been widely used but display reliance on time-series analysis, which is often limited by their inability to account for complex and nonlinear relationships in demand patterns. Meanwhile, machine learning techniques can offer significant advantages in processing larger data and capturing complex patterns, especially when external factors are considered apart from historical patterns that further enhance forecast accuracy. Additionally, through feature selection and hyperparameter tuning, ML methods can address overfitting and interpretability to improve model reliability. Hence, this allows predictions to adapt to market fluctuations and demand shifts. Additionally, common evaluation metrics were also mentioned to assess the forecasting performance of different forecasting models. These metrics can quantify the accuracy of predictions. Overall, ML models can enhance supply chain efficiency by demonstrating reliable forecasts, addressing current limitations, and considering external variables.

3. Methods

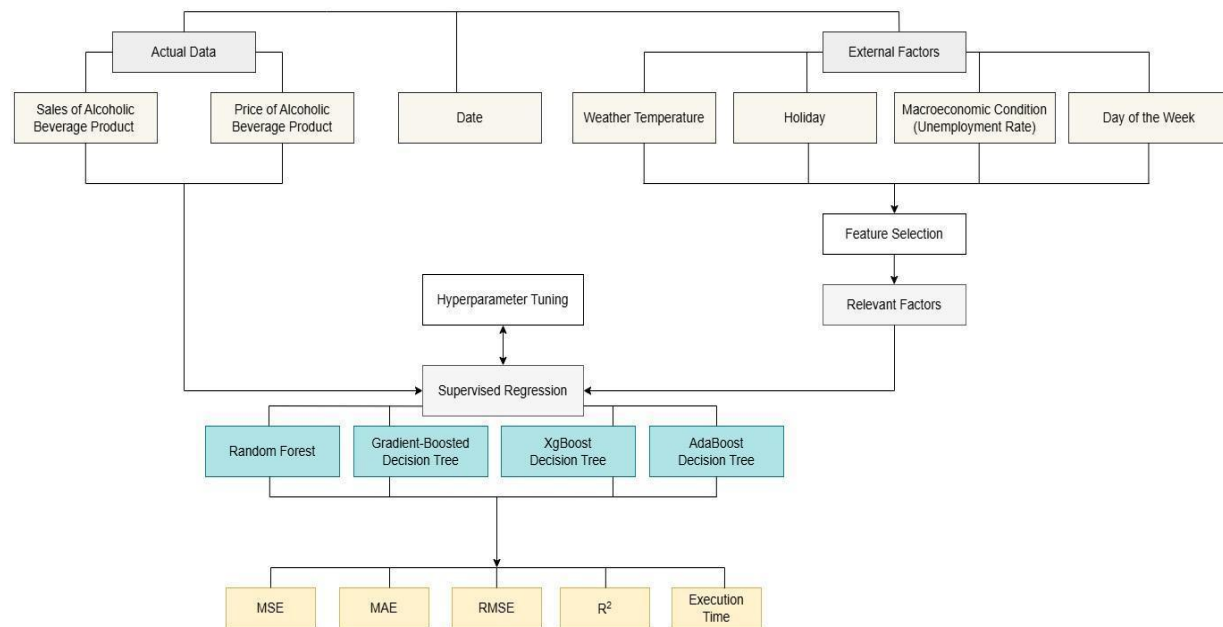


Figure 1. Conceptual Framework

Historical sales data were gathered from a food and beverage company in the Philippines. These sales data were aligned by date and combined into a single data frame with other external predictors gathered via public online data. These external factors are measured and collected as follows:

Table 1. Measures of Variables

Factor	Unit of Measurement
Sales	Quantity sold
Macroeconomic Conditions Unemployment Rate	Percentage
Day of the week	1- Monday, 2 - Tuesday, 3 - Wednesday, 4 - Thursday, 5 - Friday, 6 - Saturday, 7 - Sunday
Weather Temperature	Average temperature in degF
Holiday	0 - No holiday, 1 - with holiday

After pre-processing, the data underwent feature selection. Specifically, Pearson's Correlation Coefficient and factors with negligible correlation (less than positive and negative 0.3) were removed. Two data sets were further processed, one with feature selection having fewer dimensions and another without. The resulting datasets were used to train and test the machine learning model; 80% of the data were used for training, while 20% were used for testing. Several machine learning models were evaluated: Random Forest Regression, Gradient Boosting, Extreme Gradient Boosting, and Adaptive Boosting. Each machine learning algorithm was sourced from Scikit Learn to Python programming. Additionally, the hyperparameters for each machine learning algorithm were determined by grid search and random search in separate trials. Ten folds of cross-validation were used to ensure satisfactory model training. The resulting combinations are shown in the tree diagram below:

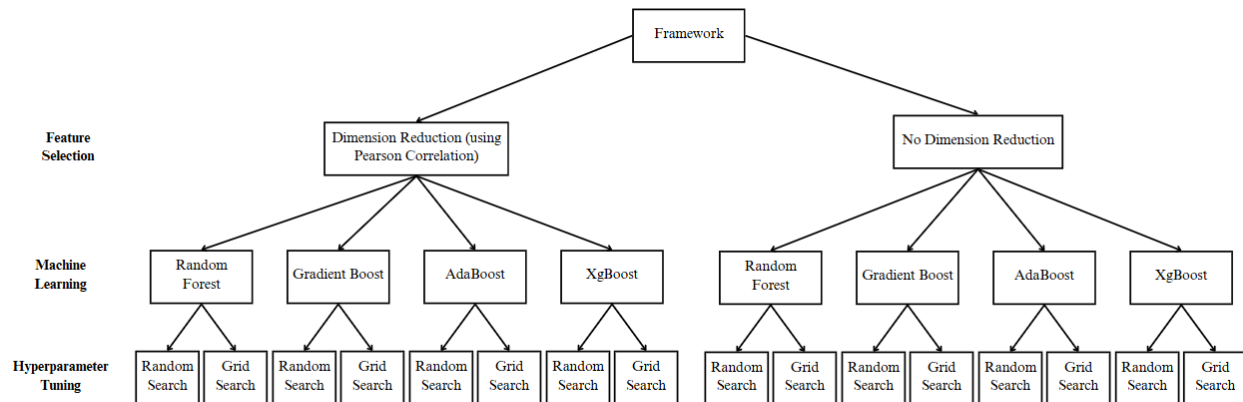


Figure 2. Tree Diagram of Framework Combinations

Each combination or trial was evaluated using metrics of accuracy and computational efficiency. These metrics were the basis of comparison for assessing the best framework.

4. Data Collection

The research began with acquiring historical sales data from a restobar in Quezon City, Philippines, covering 2021 to 2024. The dataset, specifically focusing on alcoholic beverage sales, was categorized into three primary product groups: Product 1 - Bucket (Beer), Product 2 - Cocktails, and Product 3 - Beer (Bottle). This internal data served as the foundational dataset for analyzing past demand patterns.

Supplementary data was sourced from online resources to include possible external factors influencing demand. This included macroeconomic indicators such as the unemployment rate (macroeconomic factor), weather temperature

data, and holiday dates. Subsequently, the online and company sales data were merged, ensuring all information aligned with the corresponding dates.

5. Results and Discussion

5.1 Numerical Results

Results of accuracy metrics and execution time were compared for models with and without feature selection and between hyperparameter tuning methods: Grid Search and Random Search. The values were aggregated by averaging the value across all three products tested. This approach provided a summary value, ensuring consistency.

Feature Selection

As presented in Table 2, feature selection may not always enhance model performance, significantly when the data is already well-structured and selected features are not strongly correlated with the target variable. Models such as the random forest have the necessary robustness to irrelevant or redundant features. Additionally, models such as XGBoost have an inherent ability to assign important features during training due to their tree-based splitting mechanisms and regularization capabilities, likely contributing to the minimal impact of explicit feature selection on accuracy.

Table 2. Comparison of feature selection for machine learning algorithms

Average of Values	Feature Selection							
	Random Forest		Gradient-Boosted Decision Tree		XGBoost Decision Tree		AdaBoost Decision Tree	
	Absent	Present	Absent	Present	Absent	Present	Absent	Present
ET	207.80	199.54	62.19	53.20	12.14	6.80	9.31	9.70
MAE	5.72	5.83	5.87	5.86	5.82	5.85	6.67	5.81
MSE	86.09	87.10	86.40	88.01	85.42	85.69	111.66	87.92
R ²	0.42	0.41	0.42	0.39	0.42	0.41	0.40	0.41
RMSE	8.71	8.75	8.74	9.20	8.71	8.71	10.20	8.83

Hyperparameter Tuning

Table 3 displays minimal variability between the sets of results. In contrast, more data and dimensions can exponentially increase the difference between the computational resources they require as per the curse of dimensionality. Random Search resulted in an overall higher accuracy, producing models that are equally good or better than Grid Search in less time. However, it is at the companies' discretion to consider whether the option is best based on their trade-off analyses.

Table 3. Comparison of hyperparameter tuning for machine learning algorithms

Average of Values	Hyperparameter Tuning							
	Random Forest		Gradient-Boosted Decision Tree		XGBoost Decision Tree		AdaBoost Decision Tree	
	Grid Search	Random Search	Grid Search	Random Search	Grid Search	Random Search	Grid Search	Random Search
ET	390.43	16.91	88.01	27.37	10.64	8.31	11.70	7.30
MAE	5.76	5.79	5.85	5.89	5.80	5.87	5.82	6.66
MSE	86.15	87.04	86.88	87.52	85.19	85.92	88.88	110.70
R ²	0.42	0.42	0.40	0.41	0.42	0.41	0.39	0.42
RMSE	8.72	8.74	8.47	9.47	8.69	8.73	8.91	10.12

Comparison of Machine Learning Algorithms

For product 1 (*bucket of beers*), XGBoost is best performing in terms of the goodness of the model, while AdaBoost minimizes absolute error and also has the quickest execution time. For product 2 (*cocktails*), the random forest model best describes the model's goodness and minimizes absolute error. The model with the fastest execution time is AdaBoost. Lastly, for product 3 (*beers, by bottle*), both XGBoost and AdaBoost display superior performance regarding the model's goodness, with XGBoost also characterized by the quickest execution time. Meanwhile, random forest minimizes absolute error. Table 4 displays such results.

Table 4. Comparison of machine learning algorithms by products

Comparison of Machine Learning Algorithms by Products												
Average of Values	for Product 1 (<i>Bucket of Beers</i>)				for Product 2 (<i>Cocktails</i>)				for Product 3 (<i>Beer, by bottle</i>)			
	RF	GB	XG	AB	RF	GB	XG	AB	RF	GB	XG	AB
ET	181.04	61.14	12.06	7.51	199.94	29.85	6.36	6.14	230.03	82.08	10.00	14.86
MAE	8.45	8.56	8.54	8.33	3.09	3.12	3.15	4.51	5.78	5.92	5.81	5.89
MSE	157.46	157.09	154.93	160.81	22.65	24.22	23.18	59.40	79.68	80.30	78.57	79.16
R ²	0.51	0.51	0.52	0.50	0.45	0.41	0.43	0.41	0.29	0.29	0.30	0.30
RMSE	12.56	12.53	12.45	12.68	4.69	5.42	4.81	6.96	8.93	8.96	8.86	8.90

Let: RF (*Random Forest*), GB (*Gradient-Boosted*), XG (*XGBoosted*), AB (*AdaBoost*)

Overall, no machine learning outperformed another significantly. The ML models with the highest accuracies mainly varied from XgBoost and Random Forest. Regarding execution time, XgBoost and AdaBoost significantly surpassed the other ML models, proving their computational efficiency. Considering the importance of balancing these two considerations, time and accuracy, XgBoost is the best machine learning algorithm.

To further establish the validity of ML algorithms as a forecasting tool or their prospective to be a better predictor, the different models were compared against traditional forecasting models: (1) Multiple Linear Regression, and (2) Exponential Smoothing; results are displayed in Table 5. All the investigated machine learning models had a lower error value in all metrics (MAE, MSE) and a higher accuracy in terms of r^2 . However, the multi-linear regression has a relatively adequate r^2 , indicating it can capture the variability of the dataset to an extent. However, it failed to minimize the absolute error crucial in forecasting demand. Conversely, the exponential smoothing presented a relatively low error but a very low r^2 . This implied an oversimplified model that may work like this dataset but cannot

be certainly generalizable or applied to other models. In summary, machine learning algorithms outdid traditional forecasting models in maximizing accuracy, capturing variability, and minimizing errors.

Table 5. Overall comparison of machine learning algorithms and against traditional methods

Overall Comparison of Machine Learning Algorithms					Against Traditional Methods	
Average of Values	Random Forest	Gradient-Boost	XGBoost	AdaBoost	Multi-Linear Regression	Exponential Smoothing
ET	203.67	57.69	9.47	9.50	-	-
MAE	5.77	5.87	5.83	6.24	10.44	6.83
MSE	86.60	87.20	85.56	99.79	71.54	89.64
R ²	0.42	0.40	0.42	0.40	0.39	0.07
RMSE	9.31	9.34	9.25	9.99	8.46	9.47

5.2 Graphical Results

The figures below summarize algorithm performance, compare it with traditional methods, show the effects of feature selection and hyperparameter tuning, and provide insights for selecting optimal machine learning strategies for accurate and efficient modeling.

Comparison of Machine Learning Algorithms

The figure below shows the visual comparative analysis of all the machine learning algorithms regarding execution time, MAE, MSE, and R². Considering the very minimal discrepancies, Figure 3 reveals that XgBoost and Random Forest exhibited the lowest errors and highest variability, while XgBoost and AdaBoost demonstrated the lowest execution times. Figure 3 suggests that XgBoost is the best option for high accuracy and fast execution time.

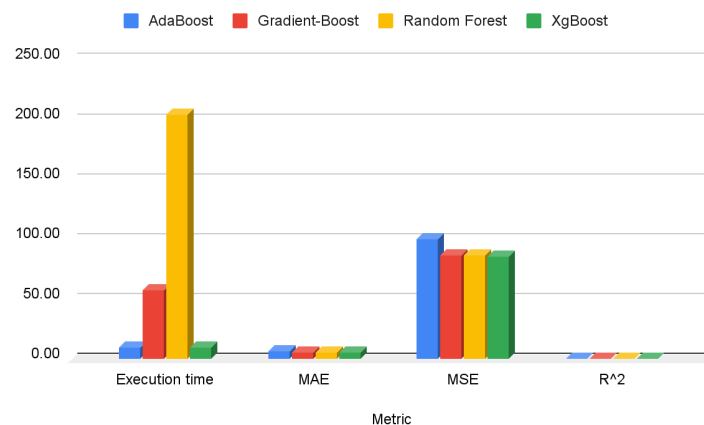


Figure 3. Comparison of Machine Learning Algorithms

Comparison of ML Algorithms Against Traditional Methods

As depicted in Figure 4, the bar graph compares various machine learning algorithms and traditional forecasting methods across key metrics. The results indicate that machine learning models, particularly XgBoost and Random Forest, generally exhibited lower error values (MAE and MSE) and higher predictive power (R-squared) than traditional methods. Although multiple linear regression demonstrated a relatively high R-squared, indicating a good fit, its error values were notably higher. Similarly, exponential smoothing showed low error but a significantly lower R-squared, suggesting a model that may not generalize well beyond the specific dataset.

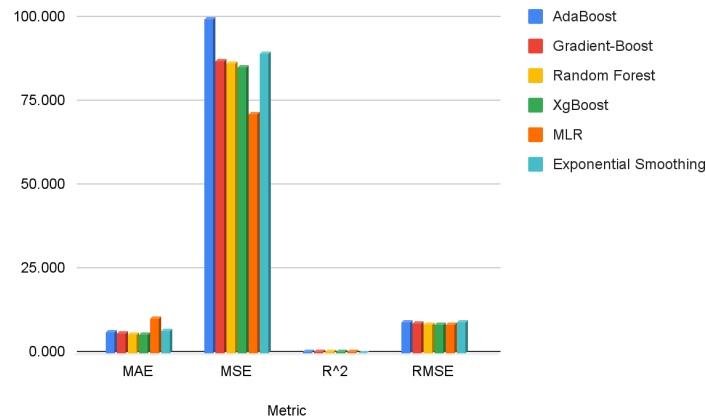


Figure 4. Comparison of ML Algorithms Against Traditional Methods

Comparison of Feature Selection for all MLs

Figure 5 illustrates the overall impact of feature selection on the performance of the machine learning models. The results indicate that feature selection led to improvements in accuracy, as supported by lower Mean Absolute Error (MAE) and Mean Squared Error (MSE) despite the increase in execution time. Furthermore, feature selection did not significantly impact the model's ability to capture variability, as indicated by equal R-squared values. Feature selection can enhance model performance and efficiency.

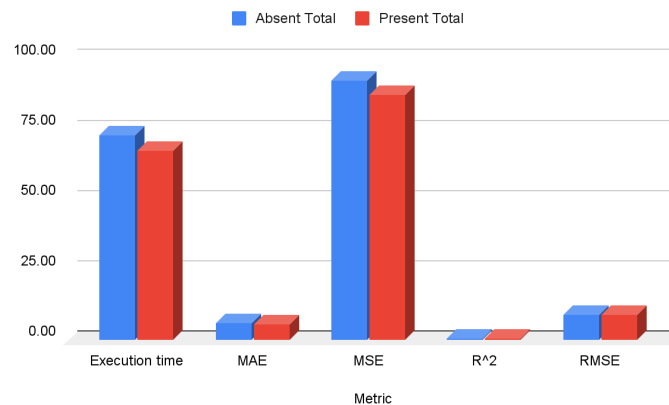


Figure 5. Comparison of Feature Selection for all MLs

Comparison of Hyperparameter Tuning for all MLs

Figure 6, visualizes the performance comparison of Grid Search and Random Search, two hyperparameter tuning methods applied to the machine learning algorithms. The results indicate a significant difference in computational cost, with Grid Search requiring significantly higher execution time than Random Search. However, Grid Search demonstrated overall accuracy, supported by lower Mean Absolute Error (MAE) and Mean Squared Error (MSE). These findings suggest that while Grid Search is computationally more expensive, it can potentially lead to more optimal hyperparameter configurations and improved model performance.

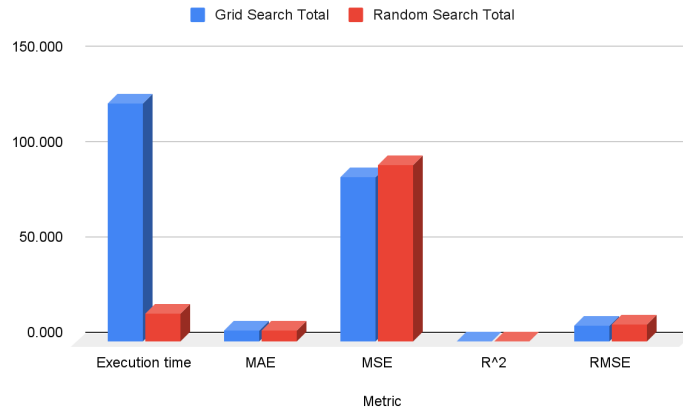


Figure 6. Comparison of Feature Selection for all MLs

5.3 Proposed Improvements

The developed model is a preliminary work that can be further expanded and adapted for general use to produce more accurate forecasts and broaden its applicability in the dynamic market. Different machine learning algorithms can be compared to determine the primary algorithm that shows superior accuracy and computational efficiency; the study, however, recommends XGBoost to ensure a balance between accuracy and execution speed. Random search hyperparameter tuning is recommended even without sufficient knowledge of the ML algorithms. Additionally, as machine learning requires big data, data spanning at least two (2) years is suggested to train and test the model in developing accurate predictions.

Future studies can suggest using other external factors, such as price and consumer behavior to better capture demand variability. In expanding the model's applicability, more products can be incorporated to assess how well it generalizes across different product categories with distinct demand patterns and seasonality. In addition to this, expanding the sample size – a larger dataset spanning multiple years and/or diverse regions – would provide a more robust mode, reducing overfitting and improving the model's generalizability.

6. Conclusion

This study successfully developed and evaluated a demand forecasting framework for food and beverage products using machine learning, meeting all research objectives and contributing uniquely to demand forecasting. To enhance prediction accuracy, the framework incorporated non-linear factors such as macroeconomic indicators, weather conditions, holidays, and day-of-the-week effects. The research assessed the performance of multiple machine learning algorithms, including Random Forest, XGBoost, AdaBoost, and Gradient Boosting Machines, focusing on forecast accuracy and computational efficiency. Additionally, it explored the impact of hyperparameter tuning methods and feature selection techniques on model performance.

The findings revealed that XGBoost outperformed other algorithms in terms of both accuracy and computational efficiency, making it the most effective choice overall. Random Forest also demonstrated high accuracy, presenting itself as a viable alternative for accuracy-focused applications. Meanwhile, AdaBoost and Gradient Boosting Machines excelled in computational efficiency, highlighting their utility in scenarios requiring rapid execution. Feature selection using correlation analysis enhanced computational efficiency by eliminating irrelevant variables but introducing slight forecast accuracy reductions. The study also found that hyperparameter tuning played a critical role in optimizing algorithm performance, with Random Search outperforming Grid Search regarding accuracy and execution time.

The research provided a structured framework for selecting the most suitable machine learning model based on specific objectives such as maximizing accuracy, minimizing computational time, or balancing both. XGBoost and Random Forest were identified as top performers for tasks prioritizing accuracy. Random Search further enhanced model performance by strategically evaluating hyperparameter combinations, proving superior to traditional Grid Search. In scenarios prioritizing computational efficiency, XGBoost and AdaBoost stood out, with AdaBoost being particularly effective in minimizing execution time. Feature selection was shown to be beneficial for reducing

computational overhead, making it ideal for applications where speed is critical, even at the cost of minor reductions in precision.

Overall, the study's most significant contribution lies in identifying XGBoost combined with Random Search as the optimal solution for achieving a balance between accuracy and computational efficiency. This comprehensive framework addressed the challenges of demand forecasting in the Philippine food and beverage industry and introduced a scalable, adaptable approach for improving supply chain practices in other contexts.

Acknowledgements

The team would like to acknowledge the Department of Science and Technology - Science Education Institute (DOST-SEI) of the Philippines for providing financial support. Their sponsorship significantly assisted the team in the completion of their research study.

References

- Aci, M., and Yergok, D. Demand forecasting for food production using Machine Learning Algorithms: A case study of University Refectory, *Tehnički Vjesnik*, vol. 30, no. 6, pp. 1683–1691, 2023.
- Aguiar, G. J., and Cano, A. Enhancing concept drift detection in drifting and imbalance data streams through meta-learning. *2023 IEEE International Conference on Big Data (BigData)*, 2023.
- Azmi, S. S., and Baliga, S. An overview of boosting decision tree algorithms utilizing AdaBoost and XGBoost boosting strategies, *International Research Journal of Engineering and Technology*, vol. 7, no. 5, pp. 6867–6870, 2020.
- Bujisic, M., Bogicevic, V., and Parsa, H. G. The effect of weather factors on restaurant sales. *Journal of Foodservice Business Research*, vol. 20, no. 3, pp. 350–370, 2016.
- Brassington, G.B. Mean absolute error and root mean square error: which is the better metric for assessing model performance? *Geophysical Research Abstracts*, 19, 2017.
- Chan, J. Y.-L., Leow, S. M. H., Bea, K. T., Cheng, W. K., Phoong, S. W., Hong, Z.-W., and Chen, Y.-L. (2022). Mitigating the multicollinearity problem and its machine learning approach: A review, *Mathematics*, vol. 10, no. 8, pp. 1283, 2022.
- Collins, S. E. Associations between socioeconomic factors and alcohol outcomes, *Alcohol Research*, vol. 38, no. 1, pp. 83–94. PMID: 27159815; PMCID: PMC4872618, 2016.
- Du, X. F., Leung, S. C.H., Zhang, J. L., and Lai, K.K. Demand forecasting of perishable farm products using support vector machine, *International Journal of Systems Science*, vol. 44, no. 3, pp. 556–567, 2013.
- Falatouri, T., Darbanian, F., Brandtner, P., & Udokwu, C. Predictive analytics for demand forecasting: A comparison of SARIMA and LSTM in retail SCM, *Procedia Computer Science*, vol. 200, pp. 993–1003, 2022.
- Gianey, H. K., and Choudhary, R. Comprehensive Review on Supervised Machine Learning Algorithms, *2017 International Conference on Machine Learning and Data Science*, 2017.
- Groene, N., and Zakharov, S. Introduction of AI-based sales forecasting: How to drive digital transformation in food and beverage outlets. *Discover Artificial Intelligence*, vol. 4, no. 1, 2024.
- Hallak, R., Onur, I., and Lee, C. Consumer demand for healthy beverages in the hospitality industry: Examining willingness to pay a premium, and barriers to purchase, *PLoS ONE*, vol. 17, no. 5: e0267726, 2022.
- Hirche, M., Haensch, J., and Lockshin, L. Comparing the day temperature and holiday effects on retail sales of alcoholic beverages – A time-series analysis. *International Journal of Wine Business Research*, 2021.
- Hodson, T. O. Root-mean-square error (RMSE) or mean absolute error (MAE): When to use them or not. *Geoscientific Model Development*, vol. 15, pp. 5481–5487, 2022.
- Holidays | Official Gazette of the Republic of the Philippines, Available: <https://www.officialgazette.gov.ph/nationwide-holidays/2021/>. Accessed on December 9, 2024.
- Holidays | Official Gazette of the Republic of the Philippines, Available: <https://www.officialgazette.gov.ph/nationwide-holidays/2022/>. Accessed on December 9, 2024.
- Holidays | Official Gazette of the Republic of the Philippines, Available: <https://www.officialgazette.gov.ph/nationwide-holidays/2023/>. Accessed on December 9, 2024.
- Holidays | Official Gazette of the Republic of the Philippines, Available: <https://www.officialgazette.gov.ph/nationwide-holidays/2024/>. Accessed on December 9, 2024.
- Holidays -- Quezon City. Senate of the Philippines Legislative Reference Bureau, Available: <https://issuances-library.senate.gov.ph/subject/holidays--quezon-city>. Accessed on December 9, 2024.

- Jiang, L., Rollins, K. M., Ludlow, M., and Sadler, B. Demand forecasting for alcoholic beverage distribution. *SMU Data Science Review*, vol. 3, no. 1, Article 5, 2020.
- Kumar, V., and Garg, M. L. Predictive Analytics: A Review of Trends and Techniques, *International Journal of Computer Applications*, vol. 182, no. 1, 2018.
- Lasek, A., Cercone, N., and Saunders, J. Restaurant sales and customer demand forecasting: Literature survey and categorization of methods, *Institute for Computer Sciences, Social Informatics and Telecommunications Engineering*, vol. 166, pp. 479-491, 2016.
- Liashchynskiy, P., and Liashchynskiy, P. Grid search, random search, genetic algorithm: A big comparison for NAS, 2019.
- Mantovani, R. G., Horváth, T., Cerri, R., Vanschoren, J., and de Carvalho, A. C. P. L. F. Hyper-parameter tuning of a decision tree induction algorithm. *2016 5th Brazilian Conference on Intelligent Systems (BRACIS)*, 2016.
- Nassibi, N., Fasihuddin, H., and Hsairi, L. Demand forecasting models for food industry by utilizing machine learning approaches, *International Journal of Advanced Computer Science and Applications*, vol. 14, no. 3, pp. 892-898, 2023.
- Posch, K., Truden, C., Hungerländer, P., and Pilz, J. A Bayesian approach for predicting food and beverage sales in staff canteens and restaurants, *International Journal of Forecasting*, vol. 38, no. 1, pp. 321-338, 2022.
- Prajwala, T. R. A comparative study on decision tree and random forest using R tool, *International Journal of Advanced Research in Computer and Communication Engineering*, vol. 4, no. 1, pp. 196-199, 2015.
- Priyadarshi, R., Panigrahi, A., Routroy, S., and Garg, G. K. Demand forecasting at retail stage for selected vegetables: a performance analysis, *Journal of Modelling in Management*, vol. 14, no. 4, pp. 1042-1063, 2019.
- Restrepo, B. J., Rabbit, M. P., and Gregory, C. A. The effect of unemployment on food spending and adequacy: Evidence from coronavirus-induced firm closures, *Applied Economic Perspectives and Policy*, vol. 43, no. 1, 2021.
- Robeson, S. M., and Willmott, C. J. Decomposition of the mean absolute error (MAE) into systematic and unsystematic components, *PLoS ONE*, vol. 18, no. 2, 2023.
- Rodrigues, M., Miguéis, V., Freitas, S., and Machado, T. Machine learning models for short-term demand forecasting in food catering services: A solution to reduce food waste. *Journal of Cleaner Production*, 435, 2024.
- Sadik-Zada, E. R., and Niklas, B. Business cycles and alcohol consumption: Evidence from a nonlinear panel ARDL approach, *Journal of Wine Economics*, vol. 16, no. 4, pp. 429-438, 2021.
- Slováčková, T., Birčiaková, N., and Stávková, S. Forecasting alcohol consumption in the Czech Republic. *Procedia - Social and Behavioral Sciences*, 220, pp. 472-480, 2016.
- Tanizaki, T., Hoshino, T., Shimmura, T., and Takenaka, T. Demand forecasting in restaurants using machine learning and statistical analysis, *Procedia CIRP*, 79, pp. 679-683, 2019.
- Tirkeş, G., Güray, C., and Çelebi, N. Demand forecasting: A comparison between the Holt-Winters, trend analysis, and decomposition models, *Tehnički vjesnik*, vol. 24, no. 2, pp. 503-509, 2017.
- Tsoumakas, G. A survey of machine learning techniques for food sales prediction, *Artificial Intelligence Review*, 52(1), pp. 441-447, 2019.
- Venkatesh, B., and Anuradha, J. A review of feature selection and its method, *Cybernetics and Information Technologies*, vol. 19, no. 1, pp. 3-26, 2019.
- Wiyanti, D. T., Kharisudin, I., Setiawan, A. B., and Nugroho, A. K. Machine learning algorithm for demand forecasting problem, *Journal of Physics: Conference Series*, 1918, 042012, 2021.
- Yerragudipadu, S., Gurram, V. R., Rayapudi, N. S., Bingi, B., Gollapalli, L., and Peddapatlolla, U. An efficient novel approach on machine learning paradigms for food delivery company through demand forecasting in societal community, *E3S Web of Conferences*, 391, 01089, 2023.
- Yoo, T. W., and Oh, I. L. Time series forecasting of agricultural products' sales volume based on seasonal long short-term memory, *Appl. Sci.*, 10, 8169, 2020.
- Zebari, R. R., Abdulazeez, A. M., Zeebaree, D. Q., Zebari, D. A., and Saeed, J. N. A comprehensive review of dimensionality reduction techniques for feature selection and feature extraction, *Journal of Applied Science and Technology Trends*, vol. 1, no. 1, pp. 56-70, 2020.

Biographies

Ray Baltazar Alunen is a 4th year undergraduate student taking up BS in Industrial Engineering at the University of Santo Tomas. He is currently specializing in Operations Research and Analytics. He also worked as a Supply Chain Management Intern in a retail company, working in demand-supply operations such as demand forecasting.

Cyrene Franchesca Molina is a 4th year Industrial Engineering student of the University of Santo Tomas, specializing in Quality Engineering as her Professional Elective course. She worked on a Process Flow Optimization Project as a Commissary Intern and was the former Executive Vice President of the department's mother organization.

Raven Francheska Quesada is a 4th-year Industrial Engineering student at the University of Santo Tomas with a professional elective in Operations Research and Analytics. She held a role as a Global Supply Chain Intern, providing administrative support in purchase order management and supplier communication.

Chloe Nicole Reyes is a 4th-year Industrial Engineering student at the University of Santo Tomas, specializing in Production Engineering. Her experience includes roles as a Health, Safety, and Environment intern and a Systems Improvement intern, where she focused on optimizing operations and enhancing workplace safety.

Engr. Delfin R. Jacob is the President of e2 Consulting Inc. He is a Professional Industrial Engineer (PIE) certified by the Philippine Institute of Industrial Engineers (PIIE). He holds a Master of Science in Industrial Engineering degree from the University of the Philippines – Diliman and completed the academic requirements of the Doctor of Philosophy in Human Resource Management at the UST Graduate School.