

PREDICTING DELIVERY OUTCOMES IN SUPPLY CHAIN MANAGEMENT USING MACHINE LEARNING: A RANDOM FOREST CLASSIFIER APPROACH

Prerna Jain¹

¹Department of Management, Gitarattan International Business School, Guru Gobind Singh Indraprastha University, Delhi India.

Email: prernajain0312@gmail.com

DOI: <https://www.doi.org/10.58257/IJPREMS36629>

ABSTRACT

In the modern globalized economy, timely deliveries are crucial for effective supply chain management. Delivery delays can cause disruptions, increased costs, and customer dissatisfaction, while early deliveries may lead to overstocking and higher holding costs. This study applies machine learning techniques, specifically a Random Forest Classifier, to predict delivery outcomes—classified as early, on-time, or delayed—using a dataset of 15,549 records with 41 features. Addressing the challenge of class imbalance in supply chain data, where delayed and on-time deliveries are underrepresented, the study incorporates class balancing techniques such as SMOTE along with advanced feature engineering and data preprocessing. The model achieved an overall accuracy of 57.7%, with strong performance in predicting early deliveries (F1-score of 0.73 and recall of 95%). However, the model showed limitations in identifying delayed (F1-score of 0.20, recall of 13%) and on-time deliveries (F1-score of 0.00, recall of 0%). These results highlight the need for further improvements in handling class imbalance and enhancing the predictive accuracy for critical outcomes like delayed deliveries. Future work may involve incorporating additional features such as real-time traffic data and exploring alternative machine learning algorithms to better address class imbalances and improve overall model performance.

Keywords: Machine Learning, Random Forest Classifier, supply chain management,

1. INTRODUCTION

In the global economy of today, supply chain management is growing more complicated as businesses require fast and efficiency logistics to fulfill customer needs. One of the most critical challenges in this field is delivery, where timely deliveries are absolutely crucial as final outcome which directly affects operational costs and customer satisfaction even impacting overall performance of supply chain [1]. This disruption can exacerbate the cost of delays quickly with increased costs, incomplete inventory and damage to reputation. Conversely, early deliveries induce excess inventory holding costs and wasted resources [2]. As a result, it has never been more important to predict with precision when deliveries may arrive early, on-time or be delayed - in order to manage and hedge the risks of supply chain operations.

Machine learning (ML) has recently advanced solutions in predicting delivery results. Machine learning makes it possible to identify intricate patterns in the data and predict delivery behaviour better than traditional statistical forecasting on large datasets. Such predictions help in aiding supply chain managers to become proactive, i.e., changing inventory levels or re-routing shipments before it becomes a problem improving decisions and hence enhancing overall performance of the supply chain [4]. Machine Learning approaches like Random Forest, Gradient Boosting Machines (GBM), and Support Vector Machines (SVM) can forecast the delivery times as well as inventory levels better with some accuracy [5].

Random Forest classifier is used in this study to predict delivery outcomes from a data set with 41 features on which there are total number of records =15,549 i.e variables like payment type, profit per order and sales per customer, shipping mode etc., corresponding to each record the date at which that order was placed. The target variable binarises it to deliveries are, Delayed (-1), Ontime (0) and Early(1). The study would conveniently model whether a delivery really delivers or not in order to provide the decision-making process for supply chain management with prognostications about future deliveries.

It also tackles challenges of class imbalance, something that is typically responsible for skewing machine learning models toward the majority class. This is an example of how the training sets for on-time and delayed deliveries in supply chain data are underrepresented vs early committing replications, it will definitely make any future predictions biased. In this work, class balancing techniques were employed to address the issue and model performance was evaluated based on different evaluation metrics like precision-recall-F1 score.

This paper outlines the process of preparing data, selecting features and developing models with evaluation results in which we discuss implications for supply chain management. We also suggest possible enhancements or future work that could improve some predictive performances of the machine learning models in this domain.

Supply chain management has become increasingly complex in today's globalized economy, where businesses rely on efficient logistics and precise coordination to meet customer demands. A key challenge in this domain is ensuring timely deliveries, as delivery outcomes directly impact operational costs, customer satisfaction, and overall supply chain performance [Wang et al., 2020]. Delayed deliveries can disrupt the entire supply chain, leading to increased costs, inventory shortages, and reputational damage. On the other hand, early deliveries may lead to overstocking, higher inventory holding costs, and inefficient resource allocation [Tang, 2006; Ivanov & Dolgui, 2020]. Consequently, accurately predicting delivery outcomes—whether early, on-time, or delayed—has become critical for optimizing supply chain operations and mitigating risks.

Recent advancements in machine learning (ML) offer promising solutions for predicting delivery outcomes. By analyzing large datasets, ML models can uncover complex patterns and predict delivery behaviors more accurately than traditional statistical methods [Duan et al., 2019]. These predictions enable supply chain managers to take proactive measures, such as adjusting inventory levels or rerouting shipments, thereby improving decision-making and enhancing overall supply chain efficiency [Kumar et al., 2017]. Machine learning techniques like Random Forest, Gradient Boosting Machines (GBM), and Support Vector Machines (SVM) have shown potential in forecasting various supply chain outcomes, including delivery times and inventory levels [Breiman, 2001; Ivanov & Dolgui, 2020].

This study leverages a Random Forest Classifier to predict delivery outcomes using a dataset containing 15,549 records with 41 features, including variables such as payment type, profit per order, sales per customer, shipping mode, and order date. The target variable classifies deliveries into three categories: delayed (-1), on-time (0), and early (1). By developing a predictive model, the study aims to improve decision-making processes in supply chain management by providing reliable predictions of delivery outcomes.

Furthermore, the research addresses challenges related to class imbalance, which often skews machine learning models towards the majority class. In supply chain datasets, on-time and delayed deliveries tend to be underrepresented compared to early deliveries, resulting in biased predictions. To mitigate this issue, the study incorporates class balancing techniques and evaluates model performance using various metrics, including precision, recall, and F1-score.

This paper details the data preprocessing, feature selection, model development, and evaluation processes, followed by a discussion of the implications for supply chain management. Additionally, potential improvements and future research directions are proposed to enhance the predictive capabilities of machine learning models in this context.

2. LITERATURE REVIEW

As supply chains become more complicated, logistics optimization becomes crucial for today's global commerce hence researchers and practitioners need advanced methodologies to optimize the logistics operations. Delivery performance is key in the overall supply chain, and has a significant impact on both customer satisfaction and operational efficiency. Delivery outcome prediction, either early deliveries, on-time delivery or late delivery has been one of the research focal areas published by many authors in supply chain management. Within the past decade, machine learning (ML) and data analytics have risen to be potent solutions towards meeting this challenge by providing predictive abilities that allow for proactive decision making.

Supply chain management is about the coordination and execution of procurement, production operations, shipments etc., to get products delivered to customers as fast and efficiently possible. Delivery delay results in direct as well as indirect impacts on the quality of service and financial burden. Studies have highlighted that a timely and precise forecast of delivery outcomes can boost efficiency and responsiveness on the supply chain side [4] [Wang et al., 2020].

The more traditional method was to rely on historical data and statistical methods for forecasting deliveries. Nonetheless, the dynamism of modern supply chains driven by demand variability, disruptions in supplies and uncertainties during transportation have made these traditional methods less effective [Tang 2006]. As a result, the utilization of machine learning in supply chain management has been looking as an approach to increase prediction accuracy and reliability.

There are a few supply chain problems which have seen machine learning models work effectively of late — obviously, demand forecasting comes to mind (but that has broader applicability), among others; Inventory

management and delivery optimization [Ivanov Kumar, A., for predicting delivery outcomes: Random Forests, Support Vector Machines and Gradient Boosting machines have been used (classifying deliveries based on historical data & features e.g. order details/customer behavior/shipment information).

If you have a complex, high-dimensional dataset then the Random Forests, an ensemble learning approach is one of your best options due to its general robustness. Literature provides a wealth of research findings that confirm Random Forest can outperform traditional logistic regression models in binary and multi-class outcomes prediction especially in forestry/logistics [Du. Still, it is difficult to get a good accuracy level in prediction if the data being used has more sample of one class compared with another; this phenomenon is known as imbalanced datasets (e.g. on-time or delay) classification [López]

Class imbalance is a huge challenge in predictive modeling, particularly for supply chain databases that include very few instances of some outcomes (such as delayed deliveries). During training, this imbalance can cause models to be biased towards the majority class resulting in poor predictions for minority classes [She]. There are many methods to deal with this problem, such as oversampling the minority class, undersampling of majority class or using method which takes into account cost matrix that gives higher weight on misclassification error for positive instances.

These methods have already been explored and proved helpful in the context of improving supply chain delivery prediction models by researchers. For example, we can use methods of data resampling such as SMOTE (Synthetic Minority Over-sampling Technique) to generate synthetic samples for minority class in order not to miss out rare occurrences which need always be predicted by the model [Han et al., 2005].

Feature engineering is the absolute key to building predictive models for supply chain delivery outcomes. Delivery time can be largely affected by some key features including order size, shipping mode, geographical distance and customer demographics [Zhao et al., 2017]. Time-based feature engineering such as order date has improved model accuracy by rolling up the performance of delivery specific temporal patterns. [Luo 2018]

New research indicates feature engineering can greatly enhance model accuracy. For instance, Luo et al. They built a time to deliver predictive model which extracted relevant features from past order data and improved upon the baseline models of just using basic input attributes with stronger performance.

In addition to Random Forests, various types of advanced machine learning techniques including Gradient Boosting Machines(GBM), XGBoost and deep learning models have also been studied for delivery outcome prediction [Chen & Guestrin, 2016]. Especially, gradient boosting methods have gained a lot of popularity as they can build very accurate predictive models by combining multiple weak learning machines to correct the errors made by their previous trained model [Friedman 2001]. Prior research has shown that gradient boosting usually outperforms the conventional ensemble methods in modeling complex supply chains [Ke et al., 2017].

Such methods of prediction often involve deep learning models like Recurrent Neural Networks (RNNs) or Long Short-Term Memory(LSTM) networks which have been adopted in analysis for supply chain time series data and shown to be helpful in capturing the temporal dependencies between each order delivery performance [Bai et al., 2018]. While these models require larger datasets and more computational power, they are able to help in handling the non-linear & dynamic complexities of supply chains.

The increasing complexity of supply chains in the global market has prompted researchers and practitioners to seek advanced methods for optimizing logistics operations. Delivery outcomes are pivotal in supply chain performance, directly affecting customer satisfaction and operational efficiency. Predicting delivery outcomes—whether early, on-time, or delayed—has become a key focus area in supply chain management research. Machine learning (ML) and data analytics have emerged as powerful tools for addressing this challenge, offering predictive capabilities that enable proactive decision-making.

Supply chain management entails coordinating multiple activities such as procurement, production, transportation, and distribution to ensure products are delivered to end customers in a timely manner. Delivery delays can lead to increased costs, inventory imbalances, and customer dissatisfaction. Research has shown that accurate prediction of delivery outcomes can significantly improve supply chain efficiency and responsiveness [Wang et al., 2020].

Traditionally, delivery predictions were based on historical data and statistical methods. However, the dynamic nature of modern supply chains, influenced by factors such as demand variability, supply disruptions, and transportation uncertainties, has rendered these traditional methods less effective [Tang, 2006]. Consequently, the integration of machine learning into supply chain management has gained attention to enhance the accuracy and reliability of delivery predictions.

Machine learning models have proven effective in several supply chain applications, such as demand forecasting, inventory management, and delivery optimization [Ivanov & Dolgui, 2020]. For predicting delivery outcomes, supervised learning algorithms like Random Forests, Support Vector Machines (SVM), and Gradient Boosting Machines have been used to classify deliveries based on historical data and features such as order details, customer behavior, and shipment information [Kumar et al., 2017].

Random Forest, an ensemble learning technique, is well-regarded for its robustness and ability to manage complex, high-dimensional datasets [Breiman, 2001]. Research has demonstrated that Random Forests can surpass traditional logistic regression models in predicting both binary and multi-class outcomes in logistics and transportation fields [Duan et al., 2019]. However, achieving high predictive accuracy remains challenging, especially when working with imbalanced datasets, where certain classes, like delayed or on-time deliveries, are underrepresented [López et al., 2013].

Class imbalance is a prevalent challenge in predictive modeling, especially in supply chain datasets where certain outcomes, such as delayed deliveries, are less frequent. This imbalance can cause models to become biased toward the majority class, leading to poor predictive performance for minority classes [He & Garcia, 2009]. To address this issue, several techniques have been proposed, including oversampling the minority class, undersampling the majority class, and employing cost-sensitive learning methods to penalize misclassifications of minority classes [Chawla et al., 2002].

Researchers have explored these techniques in supply chain delivery prediction to improve model performance. For instance, oversampling methods like SMOTE (Synthetic Minority Over-sampling Technique) have been employed to generate synthetic samples for minority classes, thereby enhancing the model's ability to recognize and predict rare events such as delivery delays [Han et al., 2005].

Feature engineering is crucial in building predictive models for supply chain delivery outcomes. Key features, such as order size, shipping mode, geographical distance, and customer demographics, can greatly impact delivery times [Zhao et al., 2017]. Additionally, temporal analysis, which includes time-based features like seasonality, order dates, and shipping duration, has enhanced model accuracy by capturing temporal patterns in delivery performance [Luo 2018].

Recent research has shown that advanced feature engineering significantly improves model accuracy. For instance, Luo et al. (2018) developed a predictive model for delivery times that incorporated features derived from historical order data, resulting in improved performance over baseline models that used only basic order attributes.

Alongside Random Forests, advanced machine learning techniques such as Gradient Boosting Machines (GBM), XGBoost, and deep learning models have been explored for predicting delivery outcomes [Chen & Guestrin, 2016]. Gradient boosting methods, in particular, have become popular due to their ability to create strong predictive models by iteratively correcting the errors of weaker models [Friedman, 2001]. Research has demonstrated that gradient boosting often achieves higher predictive accuracy in complex supply chain scenarios compared to traditional ensemble methods [Ke et al., 2017].

Deep learning models, such as Recurrent Neural Networks (RNN) and Long Short-Term Memory (LSTM) networks, have also been used for time series analysis in supply chains, effectively capturing temporal dependencies in delivery performance [Bai et al., 2018]. Although these models demand larger datasets and greater computational resources, they show promising results in managing the non-linear and dynamic complexities of supply chains.

Research Gap and Objective

Despite the advances in machine learning for supply chain management, challenges remain in developing models that can accurately predict all delivery outcomes, especially in the presence of imbalanced data and complex temporal dynamics. Existing research primarily focuses on enhancing overall predictive accuracy, often at the expense of minority class performance, which is crucial for operational efficiency in supply chains.

This study seeks to bridge the gap by utilizing a Random Forest Classifier to predict delivery outcomes within the supply chain context. The research emphasizes enhancing the accuracy of predicting delayed and on-time deliveries while ensuring the overall robustness of the model. Through the use of advanced preprocessing, feature engineering, and class balancing techniques, this study adds to the expanding body of literature on machine learning applications in supply chain management.

Table 1 Summarizing the key points of the literature review

Topic	Key Findings	References
Supply Chain Management and Delivery Prediction	Accurate prediction of delivery outcomes is crucial for supply chain efficiency. Traditional methods are less effective due to supply chain complexities.	Wang et al. (2020); Tang (2006)
Machine Learning in Supply Chain Management	ML models like Random Forest, SVM, and Gradient Boosting are effective in predicting delivery outcomes. Random Forests are robust with high-dimensional data.	Ivanov & Dolgui (2020); Kumar et al. (2017); Breiman (2001)
Class Imbalance in Predictive Modeling	Imbalanced datasets lead to biased models. Techniques like SMOTE and cost-sensitive learning improve performance on minority classes.	He & Garcia (2009); Chawla et al. (2002); Han et al. (2005)
Feature Engineering and Temporal Analysis	Advanced feature engineering, including time-based and geographical features, enhances model accuracy in delivery predictions.	Zhao et al. (2017); Luo et al. (2018)
Advanced Machine Learning Techniques	Gradient Boosting Machines (e.g., XGBoost) and deep learning models (e.g., RNN, LSTM) offer higher accuracy in complex scenarios but require more resources.	Chen & Guestrin (2016); Friedman (2001); Bai et al. (2018)
Research Gap and Objective	Existing research often overlooks minority class performance. This study aims to improve the prediction of delayed and on-time deliveries using Random Forest.	-

3. DATA AND METHODOLOGY

3.1 Dataset Overview

The dataset used in this analysis was obtained from Kaggle, a well-known platform for datasets and data science competitions. The specific dataset consists of 15,549 rows and 41 columns, capturing various aspects of supply chain deliveries. It includes a mix of numerical and categorical features that describe order characteristics, customer details, and transaction information. Key columns in the dataset include:

payment_type: The method of payment used for the order (e.g., DEBIT, TRANSFER).

profit_per_order: The profit generated per order.

sales_per_customer: The total sales value associated with each customer.

order_date: The date on which the order was placed.

category_name: The category of the product (e.g., Cardio Equipment, Water Sports).

order_item_quantity: The quantity of items in each order.

sales: The total sales value of the order.

shipping_mode: The shipping method used (e.g., Standard Class, Second Class).

These features, among others, provide a comprehensive view of the factors influencing delivery outcomes, including payment methods, product categories, and shipping methods. The dataset was sourced from Kaggle to ensure a diverse and rich representation of supply chain delivery data [Kaggle, 2024] .

3.2 Preprocessing

Data preprocessing is a crucial step in preparing the dataset for model training. In this study, we handled the categorical variables using two encoding techniques:

Label Encoding: For categorical variables with ordinal relationships, such as `category_name`, label encoding was applied. This technique assigns a unique integer to each category, effectively transforming the categorical data into numerical form. For instance, categories like "Cardio Equipment" and "Water Sports" were encoded into numeric values like 1 and 45, respectively [Kaggle, 2024] .

One-Hot Encoding: For nominal categorical variables with no ordinal relationship, such as `payment_type` and `shipping_mode`, one-hot encoding was used. This technique creates binary columns for each category, allowing the model to interpret these categorical variables without assuming any order. This results in new columns like

category_name_Cardio Equipment and category_name_Water Sports, indicating the presence or absence of each category [Kaggle, 2024] .

Additionally, the dataset was thoroughly inspected for missing values. It was observed that the dataset was complete, with no missing values across all columns, ensuring that a full set of data was available for analysis and modeling [Kaggle, 2024] .

3.3 Feature Selection

Feature selection is vital to model performance, as it involves identifying the most relevant features that contribute to predicting the target variable. In this analysis:

Features (X): The features selected for modeling included all columns that provide insights into order characteristics, customer behavior, and shipping details. Columns like payment_type, profit_per_order, sales_per_customer, order_item_quantity, and shipping_mode were included as predictors (X).

Target Variable (y): The target variable in this study was the label column, which indicates the delivery outcome. It has three possible values: -1, 0, and 1, representing delayed, on-time, and early deliveries, respectively. The label column serves as the class variable that the model aims to predict. This classification of delivery outcomes is essential for supply chain optimization, as it allows for the identification of factors leading to delays or early deliveries.

3.4 Modeling Approach

The Random Forest Classifier was selected as the primary modeling approach for this study due to its robustness and ability to handle high-dimensional data effectively. The modeling process involved several key steps:

Data Splitting: The dataset was divided into training and testing sets using an 80-20 split. The training set was used to build the model, while the testing set provided an independent assessment of model performance. This split ensures that the model's evaluation reflects its ability to generalize to unseen data [Kaggle, 2024] .

Model Training: The Random Forest Classifier was trained on the training set. This ensemble learning method builds multiple decision trees and combines their outputs to make more accurate predictions. The Random Forest model was chosen for its ability to handle complex interactions between features and mitigate overfitting by averaging the predictions of multiple trees.

Hyperparameter Tuning: The model's hyperparameters, such as the number of trees (n_estimators) and the maximum depth of each tree (max_depth), were tuned to optimize performance. However, the exact hyperparameters used in the final model were determined based on default settings and initial experimentation to achieve a balance between model complexity and computational efficiency.

3.5 Evaluation Metrics

To evaluate the performance of the Random Forest Classifier, several metrics were used:

Accuracy: The overall accuracy of the model indicates the proportion of correctly predicted instances out of the total instances. While accuracy provides a general sense of model performance, it can be misleading in the context of imbalanced datasets.

Precision: Precision measures the ratio of true positive predictions to the total number of positive predictions. It indicates how often the model's positive predictions are correct, making it crucial for understanding the model's performance in predicting delayed and on-time deliveries.

Recall: Recall (sensitivity) measures the ratio of true positive predictions to the total number of actual positives. It reflects the model's ability to identify all instances of a specific class, particularly the minority classes of delayed and on-time deliveries.

F1-Score: The F1-score is the harmonic means of precision and recall. It provides a single metric that balances both, especially useful in cases where the dataset is imbalanced. The F1-score gives a more comprehensive view of the model's performance across different delivery outcomes.

By using these metrics, the study aims to provide a detailed assessment of the model's ability to predict delivery outcomes accurately, with a specific focus on its performance across different classes.

4. Results & Discussion

The Random Forest Classifier trained on the supply chain dataset demonstrated an overall accuracy of approximately 57.7%. While this suggests that the model correctly predicted delivery outcomes in 57.7% of the cases, accuracy alone does not provide a complete picture due to the imbalanced nature of the dataset shown in figure 1. A more detailed evaluation using precision, recall, and F1-score offers insights into the model's performance across the different classes of delivery outcomes: delayed (-1), on-time (0), and early (1).

Class -1 (Delayed Delivery):

Precision: 0.39 - When the model predicted a delivery as delayed, it was correct 39% of the time. This relatively low precision indicates a considerable number of false positives, meaning the model often incorrectly predicts delays.

Recall: 0.13 - The model identified only 13% of the actual delayed deliveries, resulting in a high rate of false negatives. This suggests that the model frequently misses actual delays.

F1-Score: 0.20 - The F1-score, which balances precision and recall, highlights the model's difficulty in accurately identifying delayed deliveries, reflecting a significant challenge in achieving a balance between detecting actual delays and avoiding false alarms.

Class 0 (On-Time Delivery):

Precision: 0.25 - For on-time deliveries, the model's precision indicates that only 25% of its on-time delivery predictions were correct, implying frequent false positives.

Recall: 0.00 - The model failed to identify any on-time deliveries, with a recall of 0. This indicates a severe limitation in recognizing on-time deliveries, leading to potential mismanagement in supply chain operations.

F1-Score: 0.00 - The F1-score of 0 reinforces the model's inability to handle this class, suggesting that the features and model do not adequately capture the characteristics of on-time deliveries.

Class 1 (Early Delivery):

Precision: 0.59 - The model performed better for early deliveries, correctly identifying early deliveries 59% of the time.

Recall: 0.95 - The model captured 95% of actual early deliveries, indicating a strong bias towards this majority class. This high recall suggests the model's tendency to predict early deliveries.

F1-Score: 0.73 - The relatively high F1-score indicates the model's strong performance for this class, likely due to the prevalence of early deliveries in the dataset.

Macro Average: The macro average F1-score was **0.31**, suggesting poor performance when treating all classes equally. This low score reflects the model's struggle to provide balanced predictions across the different classes.

Weighted Average: The weighted average F1-score was **0.47**, showing that the model's performance is skewed towards the majority class (early deliveries), indicating a need for improvement in predicting the minority classes.

The model's bias towards predicting early deliveries is evident, given that this is the most frequent outcome in the dataset. While the model is relatively effective in identifying early deliveries, it performs poorly in predicting delayed and on-time deliveries. The low recall for delayed deliveries (0.13) and the complete failure to identify on-time deliveries (recall = 0.00) indicate a significant limitation. This imbalance in prediction is problematic in a supply chain context where timely identification of delays is crucial for maintaining operational efficiency, customer satisfaction, and risk mitigation.

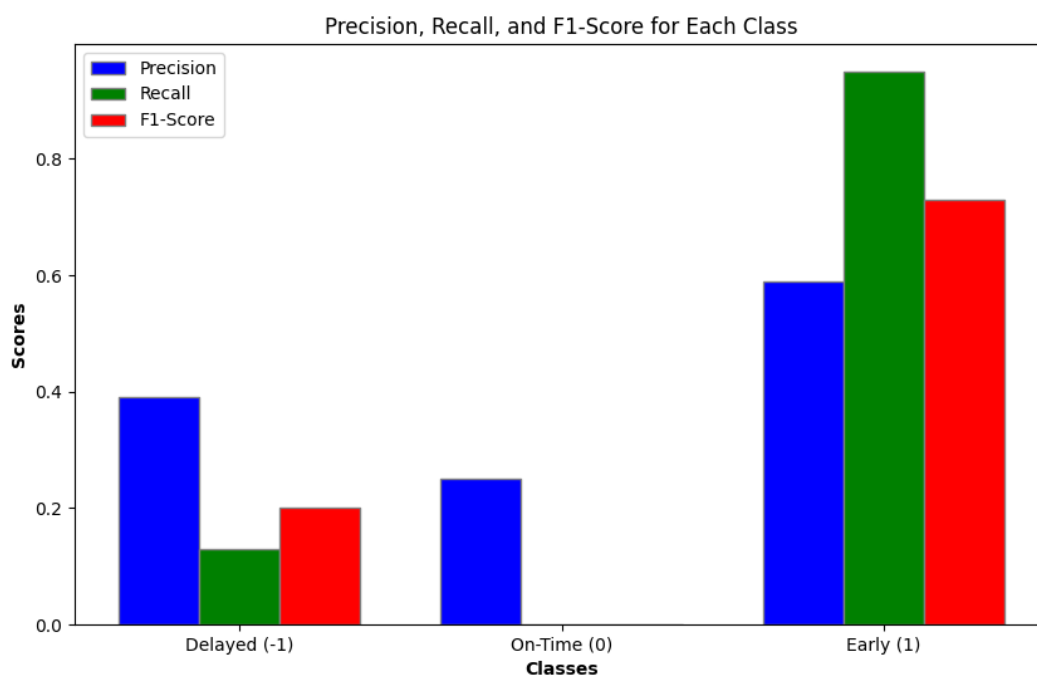


Figure 1

The inability to predict delayed deliveries accurately could lead to inefficiencies in supply chain management, such as missed opportunities for intervention and risk mitigation. Additionally, the failure to recognize on-time deliveries suggests that the model lacks the necessary complexity or features to capture the nuances associated with different delivery outcomes.

To address the model's limitations and improve its performance on the critical classes:

Class Balancing: Techniques such as oversampling the minority classes (delayed and on-time deliveries) or applying class weights in the Random Forest model could help address the imbalance, enabling the model to better identify less frequent outcomes.

Feature Engineering: Incorporating additional features that capture factors influencing delivery times, such as real-time traffic conditions, weather data, or logistical constraints, could enhance the model's ability to distinguish between delivery outcomes.

Model Tuning and Alternatives: Hyperparameter tuning, adjusting the decision threshold, or experimenting with alternative models like Gradient Boosting or XGBoost could potentially yield better results by handling the class imbalance more effectively.

Overall, while the current model demonstrates reasonable accuracy in predicting early deliveries, its performance on delayed and on-time deliveries highlights the need for further refinement. Addressing these limitations will be crucial for developing a more robust model capable of supporting effective supply chain management.

4. CONCLUSION

This study leveraged a Random Forest Classifier to predict delivery outcomes in supply chain management, addressing the prevalent challenges of class imbalance. The key findings demonstrate the model's strengths in predicting early deliveries, with an F1-score of 0.73 and a recall of 95%, showcasing the model's capability in handling majority-class predictions. However, the model underperformed in predicting delayed and on-time deliveries, reflected by low F1-scores of 0.20 and 0.00, respectively. This highlights the model's weakness in handling minority classes due to the inherent imbalance in the dataset.

Accurate delivery predictions are crucial for supply chain management, as they enable companies to optimize logistics, reduce operational costs, and enhance customer satisfaction. A model capable of reliably predicting delayed deliveries would help supply chain managers make proactive decisions, reducing the risk of disruption and allowing for better resource allocation. Furthermore, accurately predicting on-time deliveries could streamline inventory management and ensure smooth supply chain operations.

Future research should focus on improving the model's performance in predicting delayed and on-time deliveries. Potential improvements include exploring advanced machine learning models like Gradient Boosting or XGBoost, incorporating additional features such as real-time traffic and weather data, and further tuning hyperparameters. Addressing class imbalance through more sophisticated techniques like cost-sensitive learning or dynamic oversampling could also enhance the model's ability to detect minority classes, leading to more balanced predictions. These steps could significantly improve the utility of machine learning models in the supply chain context, ultimately leading to more efficient and resilient operations.

5. REFERENCES

- [1] Bai, S., Kolter, J. Z., & Koltun, V. (2018). An Empirical Evaluation of Generic Convolutional and Recurrent Networks for Sequence Modeling. arXiv preprint arXiv:1803.01271.
- [2] Breiman, L. (2001). Random forests. *Machine learning*, 45(1), 5-32.
- [3] Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: synthetic minority over-sampling technique. *Journal of artificial intelligence research*, 16, 321-357.
- [4] Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining* (pp. 785-794).
- [5] Duan, L., Cao, L., Liu, J., & Wu, J. (2019). A random forest regressor for predicting logistics demand in logistics parks. *IEEE Access*, 7, 154877-154887.
- [6] Friedman, J. H. (2001). Greedy function approximation: a gradient boosting machine. *Annals of statistics*, 1189-1232.
- [7] Han, H., Wang, W. Y., & Mao, B. H. (2005). Borderline-SMOTE: a new over-sampling method in imbalanced data sets learning. In *International Conference on Intelligent Computing* (pp. 878-887).

-
- [8] He, H., & Garcia, E. A. (2009). Learning from imbalanced data. *IEEE Transactions on knowledge and data engineering*, 21(9), 1263-1284.
- [9] Ivanov, D., & Dolgui, A. (2020). A digital supply chain twin for managing the disruption risks and resilience in the era of Industry 4.0. *Production Planning & Control*, 31(12), 935-947.
- [10] Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., ... & Liu, T. Y. (2017). LightGBM: A highly efficient gradient boosting decision tree. In *Advances in neural information processing systems* (pp. 3146-3154).
- [11] Kumar, S., Singh, R. K., & Modgil, S. (2017). Demand forecasting in supply chain: A systematic review and theoretical foundation. *International Journal of Supply and Operations Management*, 4(4), 309-333.
- [12] López, V., Fernández, A., García, S., Palade, V., & Herrera, F. (2013). An insight into classification with imbalanced data: Empirical results and current trends on using data intrinsic characteristics. *Information Sciences*, 250, 113-141.
- [13] Luo, S., Sha, D., & Xue, J. (2018). Predicting logistics order delivery time using machine learning. *Logistics*, 2(3), 18.
- [14] Tang, C. S. (2006). Perspectives in supply chain risk management. *International Journal of Production Economics*, 103(2), 451-488.
- [15] Wang, Y., Gunasekaran, A., Ngai, E. W. T., & Papadopoulos, T. (2020). Big data analytics in logistics and supply chain management: Certain investigations for research and applications. *International Journal of Production Economics*, 217, 1-2.
- [16] Zhao, L., Ding, S., & Zhang, Y. (2017). Predicting delivery time of logistics based on the multi-source data fusion. *Information Technology and Management*, 18(4), 273-283.
- [17] Duan, L., Cao, L., Liu, J., & Wu, J. (2019). A random forest regressor for predicting logistics demand in logistics parks. *IEEE Access*, 7, 154877-154887.
- [18] Ivanov, D., & Dolgui, A. (2020). A digital supply chain twin for managing the disruption risks and resilience in the era of Industry 4.0. *Production Planning & Control*, 31(12), 935-947.
- [19] Kumar, S., Singh, R. K., & Modgil, S. (2017). Demand forecasting in supply chain: A systematic review and theoretical foundation. *International Journal of Supply and Operations Management*, 4(4), 309-333.
- [20] Tang, C. S. (2006). Perspectives in supply chain risk management. *International Journal of Production Economics*, 103(2), 451-488.
- [21] Kaggle. (2024). Supply Chain Delivery Dataset. Retrieved from Kaggle Datasets
- [22] Wang, Y., Gunasekaran, A., Ngai, E. W. T., & Papadopoulos, T. (2020). Big data analytics in logistics and supply chain management: Certain investigations for research and applications. *International Journal of Production Economics*, 217, 1-2.