

Performance Evaluation of Machine Learning Algorithms for Supply Chain Data Classification

Manuscript received October 30, 2025; revised January 12, 2026

Maniah

Department of Magister Logistic Management
Universitas Logistik dan Bisnis Internasional
Bandung, Indonesia
maniah@ulbi.ac.id

Abstract—Forecasting systems that are data-driven are of great importance in streamlining industrial and business processes during the digital transformation age. Supply chain management (SCM) is among the most significant processes for enhancing operational efficiency and supporting strategic decision-making. This study seeks to evaluate the performance of two machine learning-based classification algorithms, namely Naive Bayes and the k-Nearest Neighbours (K-NN) algorithm, using data in the supply chain. Some of the most valuable operational attributes, including payment method, customer segment, shipment status, profit per transaction, and customer location, are stored in the database. The data were first cleaned and then normalised and label-encoded, after which they were split into training and test sets with a ratio of 80:20. The performance of the two algorithms was assessed using accuracy, precision, recall, and F1-score. The findings of the research indicate that Naive Bayes is the most promising algorithm; its accuracy and precision are 99.75%, and its recall rate is close to 100% in the majority of the classes. These findings show that Naive Bayes is a probabilistic algorithm that better fits the data distribution than a distance-based K-NN algorithm.

Keywords—Data Classification, K-Nearest Neighbours, Machine Learning, Naive Bayes, Predictive Analytics, Supply Chain.

I. INTRODUCTION

In the digital transformation and growing global uncertainty, supply chain management (SCM) has become a strategic potential that directly impacts organizational sustainability and competitiveness [1], [2]. Supply chains in the modern world do not rely anymore on physical streams of materials, production, and distribution, but rather on processes that are based on data to help in timely, accurate, and efficient decision-making [3]. Transactional, operational, and logistics data are growing at a very rapid rate, which has made data analytics and machine learning essential enablers to enhance supply chain performance and responsiveness. The empirical findings of micro, small, and medium enterprises (MSMEs) in Indonesia indicate that the operational digitalization has a strong positive impact on sustainable business performance, which supports the importance of analytical technologies in contemporary SCM activities [4].

Data classification is among other machine learning jobs that are central to SCM analytics [5]. Classification helps organizations to organise heterogeneous and complex data into meaningful groups, including customer segments, demand trends, reliability of suppliers, and logistics risk

profile [6]. More informed strategic and operational decisions, such as inventory planning, distribution prioritization, and differentiation of customer service, can be made with the help of accurate classification [7]. It is stated in previous research that forecasting and classification systems based on analytics have the potential to reduce forecast error by up to 50 percent [8] and reduce maintenance and operating costs by a factor of about 30 percent [9], underscoring their potential to enhance supply chain efficiency and profitability.

The recent disturbances worldwide have further underscored the need for sound, informed SCM strategies. COVID-19 revealed the weaknesses in the global supply chains, such as shipping delays, shortages of raw materials, and sudden changes in demand [10]. Approximately 90% of manufacturing companies worldwide are estimated to have been hit by the pandemic in their supply chains [11]. Such realities have forced organizations to redesign supply chains to become more resilient, flexible, and adaptive. Concurrently, a faster pace of adopting digital technologies (including the Internet of Things (IoT), big data analytics, and artificial intelligence (AI)) has made supply chain data volumes and complexity grow exponentially [12]. As per the projections of the industry, most of the large businesses will add AI-enabled analytics to the supply chain processes within a few years [13], which adds yet again to the significance of effective data classification methods.

The need to support customer demands for quicker, more transparent, and more individualized services has heightened the value of customer- and transaction-level data classification in industrial supply chains [14]. The correct identification of customer features and behavioral patterns helps firms develop the best distribution strategies, enhance service levels, and minimize resource inefficiencies. Nonetheless, with the growing topicality of data-driven SCM, the bulk of the current literature focuses solely on the technical performance of classification algorithms [15], [16], [17], [18], [19]. The relationship between classification performance and subsequent operational insights or strategic decision support in the actual SCM context has received comparatively few studies [20]. The gap shows that there is a necessity to organize the applied empirical research to assess the methods applied to classification, taking into consideration the nature of the data and its applicability to decisions in the supply chain.

Two of the most popular classification algorithms are naive Bayes and k-Nearest neighbors (K-NN) because of their simplicity of concepts and their efficiency in computing, as well as their interpretability [21]. Naive Bayes is a probabilistic classifier based on the Bayes theorem and the assumption that features are independent given a particular class, whereas K-NN is a distance-based clustering algorithm that assigns a class label based on the similarity of observed objects. The two algorithms have found wide use in fields like healthcare analytics, text classification, and pattern recognition. However, relatively few comparative analyses of these techniques in the context of SCM-related classification tasks have been performed, and the results presented by the extant literature are usually not rigorously validated nor in a context of interpretation.

It is against this light that this research paper will empirically conduct a comparison of the performance of the Naive Bayes and K-NN algorithms in the context of supply chain data classification. The analysis is based on classification accuracy and error patterns, while keeping the nature of SCM data in mind. Instead of presenting a new algorithm, this study frames the comparison as an applied assessment to emphasize methodological issues, performance resilience, and possibilities of use of this data in SCM-based decision support. The results will add value to the literature by shedding light on the situational capabilities and shortcomings of popular classifiers currently used in supply chains, as well as offering useful insights for organizations on the path to becoming digital in supply chain management.

II. RESEARCH METHOD

A. Research Design

The research design adopted in this study is a quantitative, which is a comparative approach, using supervised machine learning to classify data in supply chain management (SCM). The aim is to compare the performance of two popular classification algorithms, Naive Bayes and k-Nearest Neighbors (K-NN), when used with SCM-related data. The quantitative methodology is used to have objective, replicable, and statistically measurable outcomes, and the comparative framework provides an opportunity to have systematic performance evaluation in a variety of evaluation metrics [20].

This study is not a proposal for a new algorithm; rather, it presents an applied methodological assessment of the appropriateness and strength of algorithms in a supply chain environment. The research flow is divided into data collection, preprocessing, model training, validation, and performance comparison. To minimize evaluation bias and enhance the reliability of the results, model performance is evaluated using cross-validation rather than a train-test split. The flow of the research is shown in Fig. 1.

B. Data Collection Stage

In the study, secondary structured datasets are used to cover the supply chain activity, and the variables are the type of product, demand features of the customers, delivery time, shipping conditions (on-time or delayed), the location of the warehouse, and the categories of risk associated with the supply chain (low, medium, high). This is a set of variables widely examined in SCM analytics and customer segmentation research. The accessible and popular repositories that were used to get datasets are Kaggle, the UCI

Machine Learning Repository, and OpenML, which are commonly used in machine learning studies to guarantee transparency and reproducibility [21].

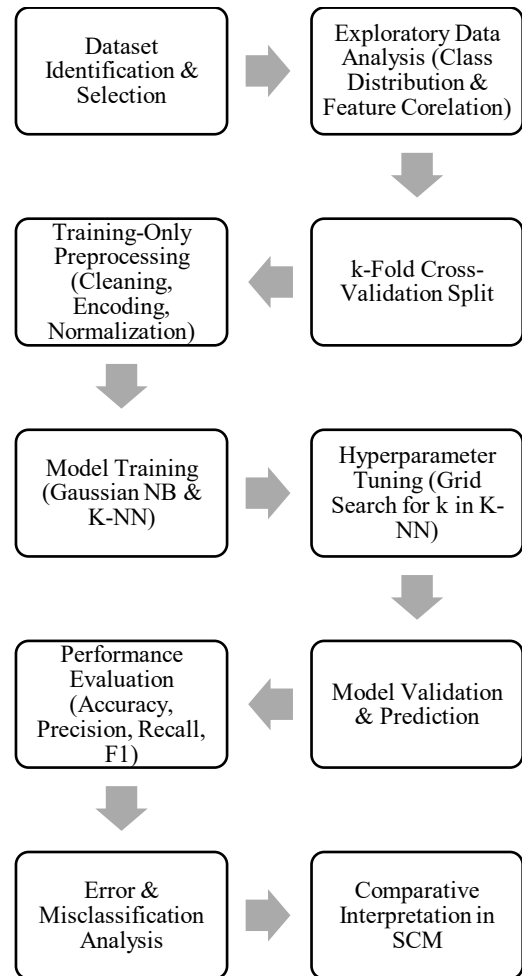


Fig. 1. Research Flow

The selection of datasets was based on relevance to SCM processes, the presence of all attributes, and applicability to multi-class classification tasks. Data preprocessing was performed in multiple steps. To begin with, initial data were cleansed by removing duplicate entries, treating gaps, and detecting extreme data by applying statistical cutoffs [22]. All preprocessing operations involving parameter estimation were performed on the training data only to avoid data leakage, and the learned parameters were subsequently applied to the validation and test data. Categorical variables were encoded with the help of label encoding, whereas numerical variables were brought to the norm, where needed, especially to facilitate distance-based classification in K-NN.

In the case of model evaluation, k-fold cross-validation was used to partition the dataset to give more robust estimates of performance than a single 80:20 split and to make the model less sensitive to random variations in sampling [23].

C. Model Implementation

In the case of probabilistic classification, the research uses the Gaussian Naive Bayes, which applies to continuous numerical properties that are typical in SCM data sets [24]. The model approximates class-conditional probability distribution under the assumption of normality and

conditional feature independence. Though this is not necessarily a match in real-world SCM data, Naive Bayes has performed well empirically on high-dimensional, more modularly correlated data. To determine the plausibility of this assumption, the exploratory correlation analysis of the numerical features was performed before training the model.

The k-Nearest Neighbors (K-NN) algorithm was deployed as a distance classifier based on the Euclidean distance metric [25]. K-NN, in contrast to Naive Bayes, does not go by feature scale, but normalizes all numeric attributes. The k was not randomly chosen, but rather, a grid search process was performed on several candidate values (k = 3, 5, 7, 9, 11) in the (cross-validation) context to determine the best neighborhood size [25]. This model can enhance the fairness of the model and make sure that the difference in performance is not based on poor parameter selection.

D. Evaluation and Comparison Stage

The main part of the research is the evaluation of model performance. Naive Bayes and K-NN classifiers were assessed based on several performance metrics, such as Accuracy, Precision, Recall, and F1-score, that are common in the realm of multi-class classification research [26]. Accuracy is used as the general performance measure, and precision and recall are used as class-specific predictive behaviour. F1-score helps to strike a balance between the precision and the recall, especially when there is an imbalance in the classes. To strengthen the findings, the averages of metric values across cross-validation folds were provided. As much as possible, performance variability was also investigated to check model stability. In addition to comparing the numbers, confusion matrices were used to analyze patterns of misclassification, providing insight into classification errors and their application in supply chain decision-making. This multi-metric, multi-fold assessment model enables a more valid and interpretable comparison of different algorithms for classifying SCM data.

III. RESULTS AND DISCUSSION

A. Dataset Description and Preprocessing

This paper refers to a structured supply chain dataset retrieved from publicly available, reputable archives, namely the Kaggle platform (<https://www.kaggle.com>) and the UCI Machine Learning Repository. The data are based on an open supply chain analytics set that has been extensively applied to empirical research on delivery performance, customer behavior, and operational efficiency. The choice of this data has been based on its completeness, applicability to supply chain management, an adequate sample size for classifying these data, and the licensing terms under which the data are free for academic use, thereby facilitating transparency and reproducibility.

The final dataset is made of 500 observations, 52 predictor variables, and one target variable. After successful data cleaning and preprocessing, it includes both numerical and categorical variables. Some of the numerical characteristics are days to shipping (scheduled and actual), number of sales per customer, number of dollars to make a sale, and number of days to make a sale, and some of the categorical characteristics are type of payment, status of delivery, where to deliver, where not to deliver, and customer segment. The target variable- customer segment- has three classes of consumers: Corporate and Home Office. Table 1 provides a

comprehensive overview of the data set, including feature definitions, target definition, and preprocessing plan.

TABLE I. CHARACTERISTICS OF THE SUPPLY CHAIN DATASET

Aspect	Description
Data domain	Supply chain and logistics analytics
Data source	Kaggle (https://www.kaggle.com); UCI Machine Learning Repository
Dataset type	Structured, tabular dataset
License	Publicly available, permitted for academic use
Number of observations	500
Number of predictor variables	52
Target variable	Customer segment
Target classes	Consumer, Corporate, Home Office
Feature types	Numerical and categorical
Numerical features (examples)	Days for shipping (scheduled and actual), sales per customer, profit per transaction
Categorical features (examples)	Payment type, delivery status, customer location
Class distribution	Moderately imbalanced, with Consumer as the majority class
Missing values	Present in a limited number of records; removed during preprocessing
Duplicate records	Identified and removed
Outlier handling	Detected using statistical thresholds and excluded
Feature engineering	Delivery lead time, temporal features (day and month), and regional aggregation
Preprocessing strategy	Encoding and normalization are applied within training folds only
Validation framework	k-fold cross-validation to prevent data leakage

A class distribution analysis shows a moderate imbalance, with the Consumer segment the largest, followed by the Corporate and Home Office segments. This imbalance is explained by the interpretation of classification results, as shown in Table 1, because a performance measure such as overall accuracy can be biased toward the dominant class. In order to enhance the quality of analysis and reproducibility, the main features of the dataset are presented both in text and in a table form. Namely, Table I replaces the data set presented above, and it provides a more compact, analytically useful depiction of dataset characteristics, without creating non-informative visual objects.

Data quality analysis revealed minor irregularities, including record duplication, missing records, and extreme values. Standard preprocessing methods, such as duplicate removal, filtering of incomplete records, and statistical outlier identification, were used to address these problems. It was then engineered to capture operational and behavioral patterns that are more pertinent to supply chain management. These derived features are delivery lead time (calculated as the difference between the order date and delivery date), temporal features (day of week and month), and aggregated geographical groupings of customer locations to minimize categorical sparsity, as shown in Table 1. Preprocessing was applied across all training folds of the k-fold cross-validation, including categorical encoding and numerical normalization. Table I summarizes this strategy because it was adopted to

eliminate data leakage and maintain the rigor and validity of the model performance evaluation.

B. Implementation of the Naïve Bayes Algorithm and K-Nearest Neighbors (K-NN)

Naive Bayes and k-Nearest Neighbors (K-NN) have been used as base supervised classification models in this research to assess their appropriateness for customer segment classification in a supply chain management (SCM) setting. These algorithms were chosen because of their popularity, efficiency in computations, and the opposite approach to methodology that enables them to be significantly compared to each other in the same data and preprocessing settings [27].

To classify probabilistically, the research uses Gaussian Naive Bayes, which applies to SCM data with continuous numerical values like profit per transaction and delivery lead time [26]. The categorical variables, such as payment method and delivery status, were coded before model training. The model was trained only on the training folds, using a cross-validation scheme as a means to avoid data leakage. Naive Bayes makes the assumption of conditional independence of features; however, similar experiments by other researchers have demonstrated that the algorithm can still perform well even in high-dimensional and moderately correlated data, which only satisfies this assumption approximately. The empirical research in the study showed a weak-to-moderate inter-feature dependence, which explained the applicability of Gaussian Naive Bayes in this context.

The k-Nearest Neighbors (K-NN) classifier was present in the form of a distance-based classifier that was implemented using the Euclidean distance metric, which is usually applied in the numeric feature space [28]. In contrast to probabilistic models, K-NN is not a model that builds a clear training model; it classifies instances based on their similarity to labeled training examples. Since K-NN depends on feature magnitudes, normalization was applied only to the training data. To prevent the arbitrary choice of the parameter, a grid search procedure that is part of the cross-validation process was employed to test different candidate values to find the configuration that produced the most consistent performance [29], [30].

The accuracy, precision, recall, and F1-score were used to evaluate model performance and provide a comprehensive account of predictive performance and the individual behavior of a particular class in multi-class classification tasks [31]. Since a medium level of imbalance between the classes was present, a confusion matrix analysis was used to further investigate misclassification patterns and enhance interpretability in the SCM scenario [32]. The multi-metric evaluation strategy will help ensure that performance comparisons are not based solely on majority-class accuracy.

Fig. 2 and Fig. 3 illustrate the high-resolution workflows of the Naïve Bayes and K-Nearest Neighbors (K-NN) classification processes, respectively. Both figures explicitly represent the sequential stages of the experimental pipeline, including dataset retrieval, model training, validation using unseen data, and performance evaluation. While these stages appear structurally similar, the figures are intentionally designed to highlight the conceptual and computational differences underlying each classification approach. As shown in Fig. 2, the Naïve Bayes workflow emphasizes probabilistic inference based on conditional feature

distributions, reflecting the algorithm's reliance on the assumption of feature independence.

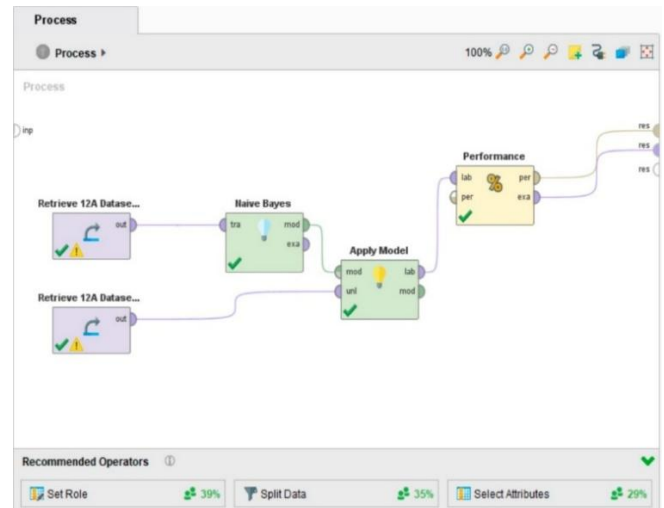


Fig. 2. Naive Bayes Model Result

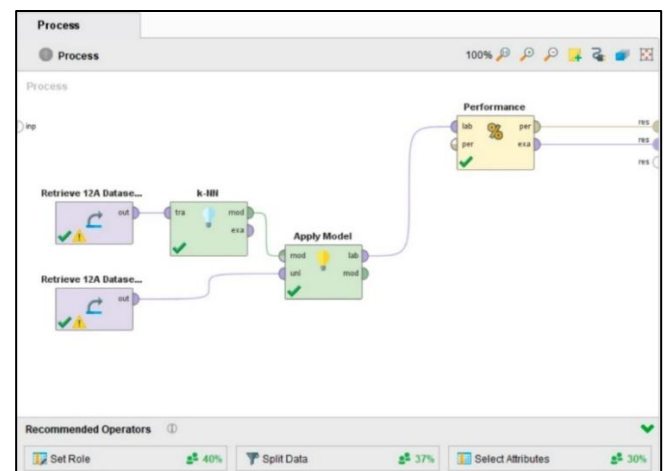


Fig. 3. K-NN Model Result

This visualization clarifies how probability estimation and posterior class selection are central to the decision-making process, thereby providing insight into why Naïve Bayes may perform efficiently on high-dimensional data with relatively limited training samples [33], [34].

In contrast, Figure 3 depicts the K-NN workflow, in which classification decisions are driven by similarity measurements in the feature space. The figure highlights the role of distance computation and neighborhood selection, illustrating how instance-based learning influences classification outcomes, particularly in datasets containing mixed numerical and categorical features. These figures were redrawn with publication-quality resolution to meet journal standards and are explicitly referenced to enhance methodological transparency. Rather than introducing novel pipeline structures, Figures 2 and 3 contextualize the experimental results by visually linking algorithmic principles to the observed performance differences discussed in the Results section. By making the internal logic of each classifier explicit, the figures support a deeper interpretation of model behavior beyond a generic machine learning workflow depiction [35], [36].

TABLE II. CONFUSION MATRIX OF NAIVE BAYES CLASSIFIER

Actual/Predicted	Customer	Home Office	Corporate
Customer	211	0	0
Home Office	0	54	0
Corporate	0	0	134

TABLE III. CONFUSION MATRIX OF K-NN CLASSIFIER

Actual/Predicted	Customer	Home Office	Corporate
Customer	171	16	31
Home Office	12	31	10
Corporate	28	8	93

TABLE IV. PERFORMANCE COMPARISON OF CLASSIFICATION MODELS

Model	Accuracy (%)	Precision (Macro)	Recall (Macro)	F1-score (Macro)
Naïve Bayes	99.75	98.75	99.39	99.07
K-NN	73.75	69.67	68.93	69.28

Comparative output from the two algorithms, Naive Bayes and k-Nearest Neighbors (K-NN), as presented in Tables 1, 2, and 3, reveals significant differences in their performance in customer segmentation of the supply chain dataset. Naive Bayes achieved a total accuracy of 99.75%, compared to 73.75% for K-NN at the aggregate level. As much as this numerical difference can be strong in the case of a multi-class classification task, it is important to take that numerical difference with caution and with reference to the data features, the modeling assumptions, and the context of evaluation, and not to be considered as unconditional algorithmic superiority.

The analysis of the confusion matrix reveals that the Naive Bayes model correctly classified almost all observations across the three customer groups (Consumer, Home Office, and Corporate), with only one misclassification in the Corporate segment. The class-wise precision and recall were consistently above 96%, indicating highly predictive behavior. The feature correlation analysis performed in the previous preprocessing stage also supports these findings: the majority of pairwise correlations among the numerical features were weak to moderate, with most falling within a representative range of 0.3 to 0.4, and no strong patterns of multicollinearity were observed. This general correlation pattern indicates that, although full independence of features is not achieved, the conditional independence assumption of Naive Bayes is approximated, as is congruent with previous findings that Naive Bayes can work well even when there are partial violations of this assumption when discriminative features are known to be separated with good probabilities between classes.

However, the high accuracy recorded in this study cannot be taken at face value. It has been stressed in previous literature that high accuracy in multi-class problems can be due to class imbalance, feature separability, or latent regularities in the data, but not necessarily due to good model generalization. The Consumer segment is the most prevalent in this dataset, which could explain the inflated overall accuracy if the majority of classes are correctly classified. In this connection, the balanced interpretation should be based on auxiliary performance indicators, including analyses of precision and recall by class and a confusion matrix. Although Naive Bayes' high sensitivity to correct class instances is indicated by its near-perfect recall, it is important

to note that further testing on new datasets is needed to assess robustness.

However, the K-NN algorithm showed less heterogeneous performance across classes. Even though the Consumer segment recorded a moderate accuracy of 78.44%, there was a great deal of misclassification between the Home Office and Corporate segments, and the recall values of 56.36% and 69.40%, respectively. These findings indicate that K-NN could not differentiate between similar patterns in the features, consistent with previous findings that, in distance-based classification algorithms, feature scaling, inter-class distance, and data density affect classification performance, especially in heterogeneous samples with partially overlapping distributions.

Although parameter tuning can enhance K-NN stability, the algorithm is sensitive when class boundaries are not clear in the feature space. The misclassification between Home Office and Corporate customers, as witnessed, is especially informative for supply chain management, as it suggests that the two share attributes of the transactional or logistical process, such as order frequency, delivery lead times, and profit margins. Under these circumstances, proximity-based segmentation can result in suboptimal differentiation; hence, it is important to ensure that the methodologies used to classify datasets align with the data structure and organizational decision-making goals.

The success of Naive Bayes in this experiment indicates that probabilistic modelling is a suitable solution for customer segmentation problems when feature distributions exhibit observable conditional trends. This advantage, however, should be viewed contextually, not in a universal way. Existing comparative studies have consistently highlighted that no single algorithm is superior across all fields, and that data characteristics, interpretability needs, and computational budget should be used to select an algorithm. In that regard, Naive Bayes offers significant computational and interpretability efficiency, making it suitable for large-scale or resource-intensive SCM settings.

In spite of these strengths, several limitations should be mentioned. To start, the analysis is conducted on a medium-sized dataset, which limits the generalizability of the findings. Second, the use of k-fold cross-validation to improve evaluation robustness was not accompanied by an assessment of performance variability across folds (e.g., standard error of the mean or confidence intervals), thereby limiting a more in-depth evaluation of model stability. Third, the paper focuses on two base classifiers and fails to consider superior models that can capture complex non-linear relationships, such as Random Forest and Support Vector Machine. Also, ethical aspects of customer segmentation, such as bias and equity implications, were not studied and should be considered in future studies.

In general, the findings justify applying Naive Bayes as a powerful baseline classifier for supply chain customer segmentation, provided probabilistic relationships can be exploited, and computational efficiency is a concern. Although K-NN is flexible and intuitive, it requires careful parameter tuning and feature engineering to achieve competitive results on complex SCM data. An extension of this comparative framework to other algorithms, larger and more diverse data sets, real-time data streams, and IoT-enabled logistics systems in the future would enhance the use of data as a driver of decision-making and digital

transformation efforts in the modern supply chain management framework.

IV. CONCLUSION

This research paper has made a comparison analysis of Naive Bayes and k-Nearest Neighbors (K-NN) in the context of customer segmentation in a supply chain management setting using a structured transactional dataset. The findings show that Naive Bayes performs better on the specific characteristics of the analyzed data, and it is more stable and consistent across customer segments than K-NN. These results do not suggest that all algorithms are universally better. Still, they point out that the probabilistic assumptions of Naive Bayes are rather consistent with datasets where the distributions of features have relatively clear conditional structures. Therefore, the research has added empirical evidence on the situational appropriateness of baseline machine learning classifiers for supply chain analytics, which helps make wiser model choices for data-driven SCM decision-making.

In spite of these findings, there are a series of limitations that should be noted. The study relies on only one medium-sized dataset analysis and two baseline classifiers, restricting the extrapolation value of the results. Future studies ought to improve this comparative framework by adding new algorithms, such as Decision Trees, Random Forests, and Support Vector Machines, and by testing their performance on larger, more varied data sets that are more representative of real supply chain settings. Additional research on parameter optimization, performance variation, and ethical issues in customer segmentation would also strengthen and enhance the usability of machine learning solutions in digital supply chain management.

REFERENCES

- [1] A. Belhadi, S. Kamble, A. Gunasekaran, and V. Mani, "Analyzing the Mediating Role of Organizational Ambidexterity and Digital Business Transformation on Industry 4.0 Capabilities and Sustainable Supply Chain Performance," *Supply Chain Management: An International Journal*, vol. 27, no. 6, pp. 696–711, Nov. 2022, doi: 10.1108/SCM-04-2021-0152.
- [2] L. V. Lerman, G. B. Benitez, J. M. Müller, P. R. de Sousa, and A. G. Frank, "Smart Green Supply Chain Management: A Configurational Approach to Enhance Green Performance through Digital Transformation," *Supply Chain Management: An International Journal*, vol. 27, no. 7, pp. 147–176, Dec. 2022, doi: 10.1108/SCM-02-2022-0059.
- [3] G. Wilson, O. Johnson, and W. Brown, "The Role of Machine Learning in Predictive Analytics for Supply Chain Management," Aug. 05, 2024, doi: 10.20944/preprints202408.0343.v1.
- [4] S. Munambar, A. W. Yuniasih, and A. Prayoga, "Design and Implementation of Functional Drink Product Inventory Applications at Kulon Progo MSMEs," *AJARCODE (Asian Journal of Applied Research for Community Development and Empowerment)*, pp. 228–235, Sep. 2024, doi: 10.29165/ajarcde.v8i3.501.
- [5] D. S. Manik, Nazaruddin, and N. Panjaitan, "Literatur Review: Penggunaan Algoritma Machine Learning untuk Optimalisasi Rantai Pasokan," in *Talenta Conference Series: Energy and Engineering (EE)*, 2025, pp. 1038–1047, doi: <https://doi.org/10.32734/ee.v8i1.2671>.
- [6] A. Asrul *et al.*, "Pemanfaatan Big Data Analytics dalam Proses Manajemen Teknologi untuk Prediksi Permintaan Pasar," *Jurnal Minfo Polgan*, vol. 13, no. 2, pp. 2433–2438, Jan. 2025, doi: 10.33395/jmp.v13i2.14516.
- [7] V. Pasupuleti, B. Thuraka, C. S. Kodete, and S. Malisetty, "Enhancing Supply Chain Agility and Sustainability through Machine Learning: Optimization Techniques for Logistics and Inventory Management," *Logistics*, vol. 8, no. 3, p. 73, Jul. 2024, doi: 10.3390/logistics8030073.
- [8] R. Dubey, A. Gunasekaran, and S. J. Childe, "Big data analytics capability in supply chain agility," *Management Decision*, vol. 57, no. 8, pp. 2092–2112, Sep. 2019, doi: 10.1108/MD-01-2018-0119.
- [9] S. Andradóttir and J. S. Lee, "Pareto set estimation with guaranteed probability of correct selection," *Eur J Oper Res*, vol. 292, no. 1, pp. 286–298, Jul. 2021, doi: 10.1016/j.ejor.2020.10.021.
- [10] D. Sari, N. Harahap, Y. Sasmita, S. A. Sitepu, and Arsyadana, "Manajemen Risiko dalam Rantai Pasok Global pada Masa Pandemi: Studi Kasus Ekspor Produk UMKM di Indonesia," *Jurnal Rumpun Manajemen dan Ekonomi*, vol. 2, no. 4, pp. 115–123, 2025, doi: <https://doi.org/10.61722/jrme.v2i4.5331>.
- [11] G. Ambrogio, L. Filice, F. Longo, and A. Padovano, "Workforce and supply chain disruption as a digital and technological innovation opportunity for resilient manufacturing systems in the COVID-19 pandemic," *Comput Ind Eng*, vol. 169, p. 108158, Jul. 2022, doi: 10.1016/j.cie.2022.108158.
- [12] D. Teh and T. Rana, "The Use of Internet of Things, Big Data Analytics and Artificial Intelligence for Attaining UN's SDGs," in *Handbook of Big Data and Analytics in Accounting and Auditing*, Singapore: Springer Nature Singapore, 2023, pp. 235–253, doi: 10.1007/978-981-19-4460-4_11.
- [13] H. Wu, J. Liu, and B. Liang, "AI-Driven Supply Chain Transformation in Industry 5.0: Enhancing Resilience and Sustainability," *Journal of the Knowledge Economy*, vol. 16, no. 1, pp. 3826–3868, Jun. 2024, doi: 10.1007/s13132-024-01999-6.
- [14] E. Nurmianti and A. Nashikha, "OPTIMALISASI E-CRM PADA STARTUP DIGITAL UNTUK MENINGKATKAN RETENSI PELANGGAN: SYSTEMATIC LITERATURE REVIEW," *JURNAL PERANGKAT LUNAK*, vol. 7, no. 2, pp. 133–143, Jun. 2025, doi: 10.32520/jupel.v7i2.4126.
- [15] D. Nasien *et al.*, "Perbandingan Implementasi Machine Learning Menggunakan Metode KNN, Naive Bayes, dan Logistik Regression Untuk Mengklasifikasi Penyakit Diabetes," *JEKIN - Jurnal Teknik Informatika*, vol. 4, no. 1, pp. 10–17, Feb. 2024, doi: 10.58794/jekin.v4i1.640.
- [16] S. Sahar, "Analisis Perbandingan Metode K-Nearest Neighbor dan Naive Bayes Classifier Pada Dataset Penyakit Jantung," *Indonesian Journal of Data and Science*, vol. 1, no. 3, Dec. 2020, doi: 10.33096/ijodas.v1i3.20.
- [17] P. D. Rinanda, B. Delvika, S. Nurhidayarnis, N. Abror, and A. Hidayat, "Perbandingan Klasifikasi Antara Naive Bayes dan K-Nearest Neighbor Terhadap Resiko Diabetes pada Ibu Hamil," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 2, no. 2, pp. 68–75, Sep. 2022, doi: 10.57152/malcom.v2i2.432.
- [18] S. Baadel, K. Attai, B. Bassey, E. Attai, A. Majeed, and F.-M. Uzoka, "Comparative Analysis of Machine Learning Classifiers for Yellow Fever Diagnosis Using Causative Data: Evaluating Naïve Bayes, KNN, RIPPER, and PART," 2025, pp. 160–169, doi: 10.1007/978-3-031-92178-0_13.
- [19] S. Ganesh and R. Baskar, "Comparative analysis of early detection of lung cancer through breath analysis by comparing K-nearest neighbours (KNN) with Naive Bayes," 2025, p. 020169, doi: 10.1063/5.0262284.
- [20] S. Saberi, M. Koushizadeh, J. Sarkis, and L. Shen, "Blockchain technology and its relationships to sustainable supply chain management," *Int J Prod Res*, vol. 57, no. 7, pp. 2117–2135, Apr. 2019, doi: 10.1080/00207543.2018.1533261.
- [21] M. M. Queiroz, D. Ivanov, A. Dolgui, and S. F. Wamba, "Impacts of epidemic outbreaks on supply chains: mapping a research agenda amid the COVID-19 pandemic through a structured literature review," *Ann Oper Res*, vol. 319, no. 1, pp. 1159–1196, Dec. 2022, doi: 10.1007/s10479-020-03685-7.
- [22] X. Yu, L. Tang, L. Long, and M. Sina, "Comparison of deep and conventional machine learning models for prediction of one supply chain management distribution cost," *Sci Rep*, vol. 14, no. 1, p. 24195, Oct. 2024, doi: 10.1038/s41598-024-75114-9.
- [23] R. W. Ahmad, H. Hasan, I. Yaqoob, K. Salah, R. Jayaraman, and M. Omar, "Blockchain for aerospace and defense: Opportunities and open research challenges," *Comput Ind Eng*, vol. 151, p. 106982, Jan. 2021, doi: 10.1016/j.cie.2020.106982.
- [24] S. S. A. and K. K. T. G., "Comparative Study of Naive Bayes, Gaussian Naive Bayes Classifier and Decision Tree Algorithms for Prediction of Heart Diseases," *International Journal for Research in Applied Science & Engineering Technology (IJRASET)*, vol. 9, no. 3, pp. 475–486, 2021, Accessed: Sep. 19, 2025. [Online]. Available: https://www.researchgate.net/profile/Sushma-Sa-2/publication/350525962_Comparative_Study_of_Naive_Bayes_Gaussian_Naive_Bayes_Classifier_and_Decision_Tree_Algorithms_for_Prediction_of_Heart_Diseases/links/6363dce8431b1f5300689ddc/Comparative-Study-of-Naive-Bayes-Gaussian-Naive-Bayes-Classifer-and-Decison-Tree-Algorithms-for-Prediction-of-Heart-Diseases.pdf

-
- [25] P. Cunningham and S. J. Delany, “k-Nearest Neighbour Classifiers - A Tutorial,” *ACM Comput Surv*, vol. 54, no. 6, pp. 1–25, Jul. 2022, doi: 10.1145/3459665.
- [26] L. M. Sinaga, Sawaluddin, and S. Suwilo, “Analysis of classification and Naïve Bayes algorithm k-nearest neighbor in data mining,” *IOP Conf Ser Mater Sci Eng*, vol. 725, no. 1, p. 012106, Jan. 2020, doi: 10.1088/1757-899X/725/1/012106.
- [27] G. Naidu, T. Zuva, and E. M. Sibanda, “A Review of Evaluation Metrics in Machine Learning Algorithms,” 2023, pp. 15–25. doi: 10.1007/978-3-031-35314-7_2.
- [28] A. R. Deore, N. Shafighi, and A. A. Hajibashi, “Impact of Machine Learning on Supply Chain Optimization,” *Archives of Business Research*, vol. 12, no. 12, pp. 112–126, Dec. 2024, doi: 10.14738/abr.1212.18101.
- [29] A. S. Ahmed and H. A. Salah, “A comparative study of classification techniques in data mining algorithms used for medical diagnosis based on DSS,” *Bulletin of Electrical Engineering and Informatics*, vol. 12, no. 5, pp. 2964–2977, Oct. 2023, doi: 10.11591/eei.v12i5.4804.
- [30] F. Iseri, H. Iseri, N. J. Chrisandina, E. Iakovou, and E. N. Pistikopoulos, “AI-based predictive analytics for enhancing data-driven supply chain optimization,” *Journal of Global Optimization*, Jul. 2025, doi: 10.1007/s10898-025-01509-1.
- [31] S. Borrohou, R. Fissoune, and H. Badir, “Data cleaning survey and challenges – improving outlier detection algorithm in machine learning,” *Journal of Smart Cities and Society*, vol. 2, no. 3, pp. 125–140, Oct. 2023, doi: 10.3233/SCS-230008.
- [32] S. Bu, “Logistics engineering optimization based on machine learning and artificial intelligence technology,” *Journal of Intelligent & Fuzzy Systems*, vol. 40, no. 2, pp. 2505–2516, Feb. 2021, doi: 10.3233/JIFS-189244.
- [33] M. B. Saputro and A. Alamsyah, “Comparison of Naive Bayes Classifier and K-Nearest Neighbor Algorithms with Information Gain and Adaptive Boosting for Sentiment Analysis of Spotify App Reviews,” *Recursive Journal of Informatics*, vol. 2, no. 1, pp. 37–44, Mar. 2024, doi: 10.15294/rji.v2i1.68551.
- [34] R. Fauzan Mu’taz and A. Rosita, “Comparison of KNN and Naïve Bayes Classification Algorithms for Predicting Stunting in Toddlers in Banjaran District,” *Journal of Applied Informatics and Computing*, vol. 9, no. 5, pp. 2711–2717, Oct. 2025, doi: 10.30871/jaic.v9i5.9703.
- [35] M. Nasrullah, B. Surarso, and O. D. Nurhayati, “Analysis of Naïve Bayes and K-Nearest Neighbors Algorithms for Classifying Fishermen Aid Eligibility,” *Jurnal Penelitian Pendidikan IPA*, vol. 10, no. 10, pp. 7652–7664, Oct. 2024, doi: 10.29303/jppipa.v10i10.8818.
- [36] M. C. Ramadhan, J. Wiratama, and A. A. Permana, “A Prototype Model on Development of Web-Based Decision Support System for Employee Performance Assessments with Simple Additive Weighting Method,” *JSiI (Jurnal Sistem Informasi)*, vol. 10, no. 1, pp. 25–32, Mar. 2023, doi: 10.30656/jsii.v10i1.6137.
-