

MACHINE LEARNING-BASED SENTIMENT ANALYSIS OF PRODUCT REVIEWS

Oluwatoyin Agbonifo¹
Victor Olutayo
Oluwaseyi Oluyode

Received 07.09.2024.
Revised 04.10.2024.
Accepted 08.11.2024.

Keywords:

Sentiment Analysis, Machine Learning, Consumer Sentiment, Naive Bayes, Logistic Regression, Random Forest.

Original research



ABSTRACT

In digital commerce, understanding consumer sentiment is crucial for businesses to make informed decisions and enhance customer satisfaction. With millions of reviews generated daily, consumers and manufacturers struggle to process this vast data. Hence, this research uses exploratory data analysis tools for the classification of product reviews into positive, negative, and neutral sentiments. Furthermore, three supervised machine learning techniques (Naïve Bayes, Logistic Regression, and Random Forest) are applied on Samsung Amazon reviews from Kaggle, conducting extensive training, testing, and evaluation. The research findings show that Random Forest outperformed the other models, achieving over 90% accuracy. This research offers an efficient solution for sentiment analysis, providing businesses with valuable insights to support product development and marketing strategies.

© 2025 Journal of Trends and Challenges in Artificial Intelligence |

1. INTRODUCTION

The rapid expansion of e-commerce platforms and social media has revolutionized how consumers make purchasing decisions (Dehghantanha, 2015). User-generated product reviews now serve as crucial resources for potential buyers, significantly influencing their choices and perceptions of products and brands. As businesses increasingly recognize the value of customer feedback, the demand for tools that can efficiently analyze and interpret consumer sentiments has grown. Sentiment analysis, a branch of natural language processing (NLP) and machine learning, plays a pivotal role in this area by automatically classifying written

opinions into positive, negative, or neutral categories (Liu, 2015). This process enables businesses to gauge customer satisfaction, address areas of concern, and make data-driven decisions to enhance both products and overall customer experiences.

In this research, we applied machine learning techniques—Naïve Bayes, Logistic Regression, and Random Forest—to Amazon product reviews sourced from Kaggle. These models were used to classify sentiments and compared based on their performance. The results demonstrate that Random Forest outperformed the other models, achieving an accuracy exceeding 90%. By automating sentiment analysis, this study provides an efficient solution for businesses

¹ Corresponding author: Victor Olutayo
Email: vaolutayo@futa.edu.ng

seeking to extract actionable insights from customer feedback, ultimately improving product development and marketing strategies.

2. LITERATURE REVIEW

Sentiment analysis is a particular sub-domain within the field of opinion mining, wherein the primary objective lies in the extraction of individuals' emotions and opinions pertaining to a specific subject matter, derived from either structured, semi-structured, or unstructured textual data (Liu, 2015). This technology enables machines to discern the emotional tone behind words, allowing for the classification of opinions as positive, negative, or neutral. The application of sentiment analysis spans various domains, including social media monitoring, customer feedback analysis, and product review evaluation. Figure 1 shows different types of sentiment analysis.

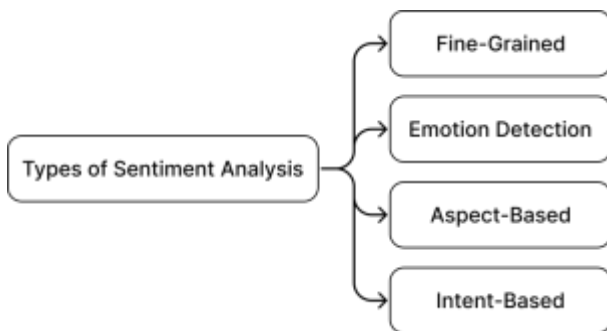


Figure 1. Types of sentiment analysis (Chandra Kala & Sindhu, 2012)

Sentiment analysis utilizes a variety of techniques to extract and classify sentiments from text. These techniques can be broadly categorized into two main approaches: lexicon-based and machine learning-based techniques. The Lexicon-based technique relies on predefined sentiment lexicons or dictionaries that assign sentiment scores to words or phrases. For example, the well-known Affective Norms for English Words lexicon assigns a sentiment score to each word, enabling the calculation of a sentiment score for a given text by summing the scores of its constituent words. (Pang & Lee, 2008) demonstrate the effectiveness of lexicon-based methods in sentiment analysis tasks. Machine learning-based techniques, on the other hand, have gained prominence due to their adaptability and scalability. (Bharathi et al., 2022) research paper employs three machine learning models for sentiment analysis of Twitter text sentiment analysis of Amazon unlocked mobile reviews, showcasing the efficacy of machine learning techniques in this field.

2.1 The Application of Sentiment Analysis

Sentiment analysis is utilized in diverse fields and sectors, fundamentally transforming the comprehension and reaction of enterprises toward customer feedback, brand perception monitoring, and public sentiment

evaluation. A notable employment is observed in the realm of social media monitoring, where sentiment analysis empowers organizations to observe and examine the sentiments conveyed on social media platforms regarding their offerings, services, or brand (Cambria et al., 2013). Several fundamental applications of sentiment analysis comprise:

- i. **Customer Service:** Sentiment analysis aids businesses in enhancing customer service through the automated classification and prioritization of customer inquiries and feedback. By discerning adverse sentiment in customer interactions, establishments can elevate concerns, offer prompt responses, and enhance overall customer gratification (Mudambi & Schuff, 2010).
- ii. **Marketing:** Businesses can use sentiment analysis to understand how customers feel about their products or brands. This can help in designing more effective marketing strategies. Research published in the "Journal of Marketing Research" by (Wang & Kim, 2017) underscores the role of sentiment analysis in market research by exploring its implications in predicting product sales.
- iii. **Politics:** Sentiment analysis is utilized in political campaigns and public opinion polling for the purpose of examining the sentiment of the general public towards political candidates, parties, and policies. Through the analysis of sentiment conveyed in posts on social media, articles in the news, and speeches made in public, analysts are able to evaluate the sentiment of voters, recognize important matters, and determine the impact of political messaging (O'Connor et al., 2010).

2.2 Challenges of Sentiment Analysis

Despite the numerous applications and usefulness of sentiment analysis, it encounters a multitude of challenges that have a significant impact on its accuracy, reliability, and effectiveness (Liu, 2012; Liu, 2015). Some of the principal hurdles in sentiment analysis encompass:

- i. **Domain-specific:** sentiment analysis models might not exhibit satisfactory performance when employed on domain-general datasets, like technical documents or industry-specific jargon, owing to variations in language usage and the expression of sentiment (Liu, 2012).
- ii. **Context and sentiment shift:** Sentiment analysis algorithms may encounter challenges in accurately comprehending the intricate context of language, comprising elements such as sarcasm, irony, slang, or cultural allusions. Consequently, these algorithms may result in the misclassification of sentiments (Pang & Lee, 2008).
- iii. **Quality data:** Accurately labeling sentiment in data is frequently intricate and expensive to

obtain (Astya, 2017). This arises from the demanding process of annotating sentiment labels in the acquired data. The annotation of sentiment labels faces challenges posed by different sentence types, such as sarcasm, ridicule, rhetorical questions, and negation. Nevertheless, researchers navigate these challenges by employing semi-supervised and unsupervised machine learning methods.

2.3 Machine Learning-Based Sentiment Analysis

Machine learning refers to systems that adapt to perform tasks associated with artificial intelligence (AI), including recognition, prediction, diagnosis, and more. These adaptations may either enhance existing systems or generate new ones (Nilsson, 2005). Machine learning algorithms are generally categorized into three types: supervised learning, unsupervised learning, and reinforcement learning (Dhumale et al., 2019).

Supervised learning involves recognizing patterns between variables using labeled datasets. Data with known outputs ("y") and features ("X") are provided to the machine, which creates models to predict outcomes based on new data. Algorithms include regression analysis, decision trees, neural networks, and support vector machines (Deva et al., 2020).

- i. Naïve Bayes Algorithm: A probabilistic classification method based on Bayes' theorem, often used in text classification and sentiment analysis. It assumes predictors are independent and calculates the probability of different classes given the observed features (Webb, 2011).
- ii. Logistic Regression: Used for binary classification, logistic regression estimates the probability of an event based on predictor variables. It uses the sigmoid function to map predictions between 0 and 1 (Hosmer et al., 2013).
- iii. Random Forest Algorithm: A decision tree-based method that creates multiple trees from random data subsets. This approach reduces overfitting and improves accuracy by combining the predictions of individual trees (Breiman, 2001).

Unsupervised learning works without labeled data, seeking to identify hidden patterns, cluster data points, or reduce dimensionality in datasets (Naeem et al., 2023). While reinforcement learning distinguishes itself by continuously improving a model through feedback from prior iterations. Unlike supervised and unsupervised learning, this method evolves dynamically with ongoing learning processes (Wang & Zhan, 2011).

2.4 Review of Related Works

Several studies have leveraged both traditional machine learning and deep learning techniques for sentiment analysis of product reviews, often comparing the performance of different models. Singla et al. (2017) and Anusuya et al. (2018) both focused on using traditional

machine learning models like Naïve Bayes, Support Vector Machines (SVM), and Decision Trees. Singla et al. (2017) found that SVM performed best with an accuracy of 81.75%, while Anusuya et al. (2018) discovered that both Linear SVC and Logistic Regression achieved high accuracy rates (88.11%). A notable limitation in both studies was dataset-related—Singla et al. (2017) highlighted mismatches between ratings and sentiments, and Anusuya et al. (2018) pointed out that their datasets were outdated, impacting the relevance of the results.

Similarly, Sobia et al. (2021) and Bharathi et al. (2022) employed supervised learning techniques such as Random Forest and Logistic Regression to classify product reviews. Their findings reinforced the effectiveness of SVM and Random Forest in large-scale sentiment analysis. However, both studies were limited by the need for sophisticated hardware due to the large datasets involved.

Research involving deep learning models such as Long Short-Term Memory (LSTM) and Convolutional Neural Networks (CNN) has gained traction in recent years. Kia et al. (2021), Shaozhang et al. (2020) and Salma et al. (2023) all explored the use of LSTM and CNN for sentiment analysis. Shaozhang et al. (2020) found LSTM to be more effective than Naïve Bayes and Logistic Regression, while Kia et al. (2021) demonstrated that LSTM outperformed models like Multilayer Perceptron and SVM. Salma et al. also confirmed that deep learning models provided better results for product reviews but faced limitations due to the high computational resources required to process large datasets.

Studies by Li et al. (2020) and Maas et al. (2011) went a step further, combining deep learning with more advanced models such as Bidirectional Gated Recurrent Units (Bi-GRU) and BERT. Li *et al.* showed that deep learning models outperformed traditional models in sentiment classification, while Maas et al. (2011) found that BERT achieved the highest accuracy (90%) among machine learning and deep learning models. The limitations of both studies again revolved around the need for advanced hardware and computational resources due to the sheer volume of reviews analyzed.

Building on these studies, this research focuses on applying Naïve Bayes, Logistic Regression, and Random Forest to Amazon product reviews from Kaggle. It aims to automate sentiment analysis while comparing the efficiency of these three methods.

3. METHODOLOGY

This section covers the system architecture, mathematical modeling, data preprocessing techniques, feature extraction methods, and model training processes.

3.1 System Architecture

The system architecture as depicted in Figure 2 demonstrates the arrangement of parts and the connections between different components in the system.

The dataset, sourced from Kaggle.com, focuses on Samsung smartphone reviews from Amazon, providing the basis for sentiment analysis. After acquiring the data in CSV format, preprocessing steps are undertaken to clean and standardize the text, removing noise such as punctuation and slang. Feature extraction techniques, including TF-IDF and Count Vectorizer, are used to convert the text into numerical data for machine learning. Three models—naive bayes, logistic regression, and random forest—are trained to classify sentiments as neutral, positive, or negative. Finally, the models' performance is evaluated using metrics like accuracy, precision, recall, and F1-score to assess their effectiveness in sentiment analysis.



Figure 2. System Architecture

The project's system development tools are essential in shaping its functionality and efficiency. Integrated Development Environments (IDEs) and data processing tools contributed to the seamless implementation of sentiment analysis. Google Colab, a cloud-based interactive platform, provided an accessible coding environment with support for GPUs and TPUs, enhancing performance for computational tasks. Pandas and Numpy are utilized for data manipulation and numerical computations, while NLTK facilitated text mining and natural language processing. Scikit-learn supported machine learning models, and Matplotlib and Seaborn are used for data visualization, presenting findings clearly during exploratory data analysis (EDA).

3.2 Mathematical Modeling

The sentiment analysis system for product reviews employs Naive Bayes, Logistic Regression, and Random Forest models. These models are accompanied by the technique of TF-IDF feature extraction. The representation of these models and techniques involves a combination of mathematical notations that describe the algorithms of Naive Bayes, Logistic Regression, and Random Forest, as well as the calculations associated with TF-IDF.

- a. **Naïve Bayes Model:** Naive Bayes is a classifier that uses Bayes' theorem and assumes independence between features. The computation of the posterior of class y can be achieved when provided with a collection of features, namely $x_1, x_2, \dots, x_n$ and a target class y :

$$P(x_1, x_2, \dots, x_n) = \frac{P(y) \times P(x_1|y) \times P(x_2|y) \times \dots \times P(x_n|y)}{P(x_1) \times P(x_2) \times \dots \times P(x_n)} \quad (1)$$

where:

$P(x_1, x_2, \dots, x_n)$ is the posterior probability of class y given the features.

$P(y)$ is the prior probability of class y .

$P(y)$ is the likelihood of feature x_i given class y .

$P(x_i)$ is the marginal probability of feature x_i .

- b. **Logistic Regression:** Logistic Regression models the probability that a given input x belongs to a particular class. It uses the logistic function (sigmoid) to map the output to the range $[0, 1]$:

$$P(x) = \frac{1}{1 + e^{-z}} \quad (2)$$

where:

$P(x)$ is the probability that the output y is 1 given input x .

$z = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n$ is the linear combination of input features weighted by coefficients β .

The logistic regression model learns the optimal coefficients β by minimizing the logistic loss function using techniques like gradient descent.

- c. **Random Forest:** Random Forests are ensemble learning methods based on decision trees. Each tree in the forest is constructed using a subset of the training data and a subset of the features. The model averages or aggregates predictions from multiple decision trees to make final predictions. The mathematical model behind a single decision tree involves recursively partitioning the feature space into regions and assigning a class label to each region.

$$h(x) = \sum_{i=1}^N w_i \times I(x \in R_i) \quad (3)$$

where:

$h(x)$ is the prediction made by the decision tree for input x .

N is the number of regions (leaves) in the tree.

w_i is the weight (class probability) associated with region R_i .

$I(x \in R_i)$ is an indicator function that returns 1 if x belongs to region R_i , and 0 otherwise.

Random Forests combine predictions from multiple decision trees to produce the final output.

- d. **Term Frequency-Inverse Document Frequency:** The Term Frequency-Inverse Document Frequency (TF-IDF) is a numerical statistic that reflects the importance of a term in a document relative to a collection of documents. The mathematical model for TF-IDF can be expressed as follows:

$$TF(t, d) = \frac{ft, d}{\sum_{t' \in d} ft', d} \quad (4)$$

where:

$TF(t, d)$ is the Term Frequency of term t in document d .

ft, d is the frequency of term t in document d .

$\sum_{t' \in d} ft', d$ is the total number of terms in document

ii. **Inverse Document Frequency (IDF):**

$$IDF(t, D) = \log \left(\frac{N}{|\{d \in D: t \in d\}|} \right) \quad (5)$$

where:

$IDF(t, D)$ is the Inverse Document Frequency of term t in document D .

N is the total number of documents in document set D .

$|\{d \in D: t \in d\}|$ is the number of documents containing term t .

iii. **TF-IDF Weighting:**

$$TFIDF(t, d, D) = TF(t, d) \times IDF(t, D) \quad (6)$$

where:

$TFIDF(t, d, D)$ is the TF-IDF weight of term t in document d with respect to document set D

TF-IDF provides greater significance to terms that are commonly found within a specific document but are infrequently observed throughout the entirety of the document collection. Consequently, this approach proves to be a highly efficient technique for representing the importance of terms within a given document.

3.3 The Dataset

Figure 3 displays a selected sample of Samsung smartphones reviews dataset from the Kaggle.com website, providing data that encompasses the following areas:

- “Product Name” is the name of the product
- “Brand Name” is the name of the brand
- “Price” is the amount of the product
- “Rating” is the star rating of the product
- “Reviews” is the title of the review
- “Reviews Votes” is the vote of the review

For the experiment, the “text” field from the datasets is utilized.

Product Name	Brand Name	Price	Rating	Reviews	Reviews Votes
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10
Samsung Galaxy S21	Samsung	1099.99	4.5	I love my S21! It's a great phone with a great camera.	10

Figure 3. Sample of Samsung Reviews Dataset

3.4 Data Preprocessing

This is a systematic approach that transforms raw data into a clean and structured format suitable for analysis or model training. This process involves addressing inconsistencies such as punctuation, hyperlinks, and special characters, ensuring the dataset is both accurate and complete. Additionally, advanced text optimization techniques like lemmatization and stopword removal are applied to reduce redundancy, focusing the content on meaningful information for improved model performance.

- Text Cleaning:** This is the process to remove unwanted characters, URLs, stopwords, special symbols, and numbers. Converting text to lowercase for uniformity.
- Tokenization:** This is to break down the reviews into individual words or tokens for analysis.
- Lemmatization:** This text preprocessing technique reduced words to their root forms to identify similarities. For instance, the lemma of "grains" is "grain," and the lemma of "saw" is "see."
- Data Transformation:** This is designed to convert target classes into integer representations. The function assigned integer values to sentiment classes: 4 and 5 stars are represented as 1, 3 stars as 0, and 1 and 2 stars as -1. This conversion facilitated the application of machine learning algorithms that required numerical target values, enabling seamless integration with the sentiment analysis model.

3.5 Feature Extraction and Model Training

The dataset is split into two segments: a training set for model training and a test set for evaluating the final performance of the trained model on unseen data. The split ratio is 75% for training and 25% for testing. Three machine learning models are employed in this research: Random Forest, Logistic Regression, and Naive Bayes. These models utilize features generated from the TfidfVectorizer, capturing various aspects of the text. Each model is trained on the training set, and hyperparameters are tuned to optimize performance.

4. PERFORMANCE EVALUATION

This section discusses the findings of the exploratory data analysis and the three selection models of Naive Bayes, Logistic Regression and Random Forest performed on the dataset obtained from Kaggle on Samsung Amazon reviews.

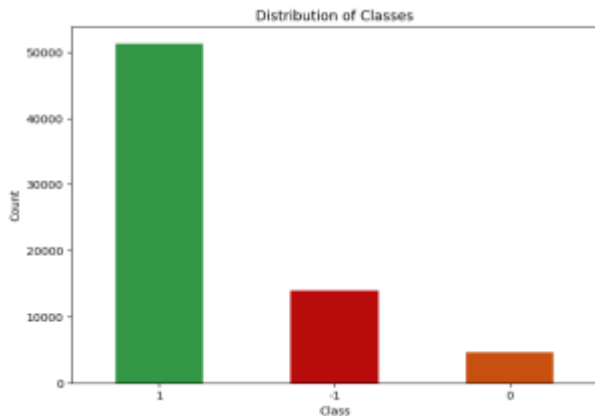
4.1 Data Analysis

An exploratory data analysis is conducted on Samsung Amazon reviews dataset using Python tools like Pandas, Matplotlib, and Seaborn for analysis and visualization. Exploratory data analysis findings from the dataset are discussed as follows:

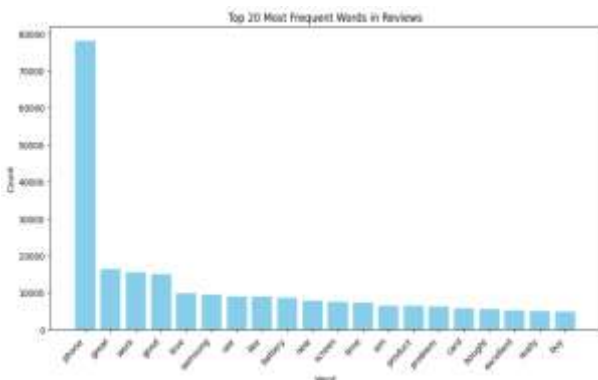
- Sentiment Class Distribution:** Table 1 shows insight of the balance or imbalance of sentiment classes such as positive, negative and neutral within the dataset.
- Distribution of Class in Bar Chart:** Figure 4 illustrates the distribution of sentiment classes in the dataset. The dataset predominantly comprises positive reviews, outnumbering the combined count of negative and neutral reviews.

Table 1. Sentiment Class Distribution of the Dataset

1 (Positive)	-1 (Negative)	0 (Neutral)
51285	13977	4655

**Figure 4.** Bar Chart showing the Sentiment Classes Distribution of the Dataset

- iii. **Top 20 Most Frequent Words:** Figure 5 provides an overview of the most frequently occurring words in the dataset, offering insight into the terms most closely associated with the Samsung phones.

**Figure 5.** Top 20 Most Frequent Words in the Dataset

- iv. **Wordcloud of the Positive Class:** Figure 6 illustrates that the positive dataset prominently features words such as love, great, good, work, nice, and feature—terms commonly associated with positivity.

**Figure 6.** Wordcloud of the Positive Dataset

- v. **Wordcloud of the Negative Class:** Figure 7 reveals prevalent words in the negative dataset,

including screen, problem, return, bad, charger, never, and issue. In addition to conventionally negative terms like frustrating and annoying, the presence of words such as screen, problem, charger, bad, return, and issue provides valuable insights into specific concerns voiced by Samsung phone customers.

**Figure 7.** Wordcloud of the Negative Dataset

4.2 Performance Evaluation

Several metrics and techniques are employed to evaluate models across different types of machine learning tasks.

- a. **Accuracy:** Accuracy serves as a metric employed to assess machine learning classification models, quantifying the ratio of accurately predicted model outcomes to the overall number of model predictions. The Accuracy formula is given below:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} (\times 100)$$

- b. **Precision:** Precision can be defined as the ability of a classification model to correctly identify relevant instances. It is measured by calculating the ratio of true positives to the sum of true positives and false positives. The formula for precision is given below:

$$Precision = \frac{TP}{TP + FP} (\times 100)$$

- c. **Recall:** Recall evaluates the proportion of relevant instances accurately detected by a classification model. It is determined by dividing the number of true positives by the sum of true positives and false negatives. The formula for recall is given below:

$$Recall = \frac{TP}{TP + FN} (\times 100)$$

- d. **F1 Score:** F1 Score aims to find a balance between Precision and Recall in a classification model. It assigns equal importance to both of these metrics in order to provide a comprehensive evaluation. The formula for f1 score is given below:

$$F1\ score = 2 * \frac{Precision * Recall}{Precision + Recall}$$

Table 2 showcases the performance metrics for the machine learning models, highlighting the exceptional performance of the Random Forest model across all metrics. Random Forest model outperforms Naïve Bayes and Logistic Regression model with an accuracy of 91%.

Table 2.Performance metric scores of the ML models

Models	Accuracy (%)	Precision (%)	Recall (%)	F1 Score (%)
Naïve Bayes	85	89	85	86
Logistic Regression	87	90	88	88
Random Forest	91	90	89	89

Performance Metrics Comparison: Figure 8 depicts a graph illustrating the performance comparison of the three machine learning models employed in the sentiment analysis task.

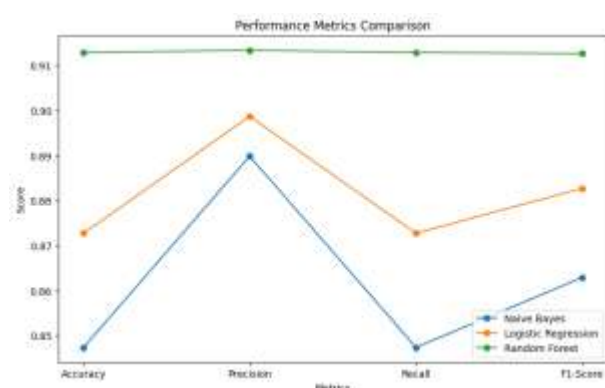


Figure 8.Graph comparing the performance of the Machine Learning Models

The x-axis of the graph represents the three machine learning models used in the experiment, while the y-axis denotes the accuracies of the models.

5. CONCLUSION

Analyzing product reviews to determine customer satisfaction is indeed a huge task, primarily due to the vast amount of data generated and stored daily across the internet. Furthermore, reviews are characterized with

complexity since they are in different formats, languages and tones, requiring flexible techniques to understand them. Therefore, this research made a careful selection models and preprocessing techniques, including three machine learning methods (Naïve Bayes, Logistic Regression, and Random Forest), the Random Forest algorithm proved to be the most accurate, achieving an accuracy rate of 91%. The findings highlight how sentiment analysis can provide useful insights for businesses, allowing them to improve brand perception and engage customers effectively. By using insights from sentiment analysis, businesses can adjust their strategies to better meet consumer preferences and drive positive outcomes. This research offers opportunities for businesses to shape their brand image and enhance customer engagement based on these insights.

Addressing grammatical errors in textual data is important to ensure the accuracy of sentiment analysis results. This can be done through preprocessing techniques like text normalization and error correction, using natural language processing (NLP) tools such as language correction libraries. Regularly updating models is essential to adapt to changing trends in language use and expressions. Implementing these methods can significantly improve the performance of sentiment analysis systems, enhancing decision-making and the quality of customer interactions. Future studies should explore more advanced approaches, such as deep learning, to achieve further improvements in performance.

References:

- Anusuya, D., Arkadeb, S., Sourish, S., & Pranita, B. (2018). Sentiment Analysis of Product-Based Reviews Using Machine Learning Approaches. *International Journal of Computer Sciences and Engineering*, 6(5), 132-136.
- Astya, P. (2017, May). Sentiment analysis: approaches and open issues. In *2017 International Conference on computing, Communication and automation (ICCCA)* (pp. 154-158). IEEE.
- Bharathi, R., Bhavani, R., & Priya, R. (2022). Twitter Text Sentiment Analysis of Amazon Unlocked Mobile Reviews Using Supervised Learning Techniques. *Journal of Data Science*, 20(1), 45-58.
- Breiman, L. (2001) Random Forests. *Machine Learning*, 45, 5-32. DOI: 10.1023/A:1010933404324
- Cambria, E., Schuller, B., Xia, Y., & Havasi, C. (2013). New avenues in opinion mining and sentiment analysis. *IEEE Intelligent Systems*, 28(2), 15-21.
- Chandra Kala, S., & Sindhu, C. (2012). Opinion mining and sentiment classification: a survey, *International Journal of Computer Applications*, 3(1), 420-427.
- Dehghantanha, A. (2015). Mining the Social Web: Data Mining Facebook, Twitter, LinkedIn, Google+, GitHub, and More, by Matthew A. Russell. *Journal of Information Privacy and Security*, 11, 137-138, 10.1080/15536548.2015.1046287.
- Deva, K. A., Prem, K. J., Prakash, V. S., & Divya, K. S. (2020). Supervised Learning Algorithms: A Comparison. *Kristu Jayanti Journal of Computational Sciences (KJCS)*, 1(1), 1-12. DOI: 10.59176/kjcs.v1i1.1259

- Dhumale, R. B., Thombare, N. D., & Bangare, P. M. (2019, April). Machine learning: A way of dealing with Artificial Intelligence. In 2019 1st International Conference on Innovations in Information and Communication Technology (ICIICT) (pp. 1-6). IEEE.
- Hosmer Jr, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied logistic regression*. John Wiley & Sons.
- Kia, D., Mandar, G., Ahsan, A., Hadi, L., & Amir, H. (2021). Sentiment Analysis of Persian Movie Reviews Using Deep Learning. *Entropy*, 23(5), 596, DOI: 10.3390/e23050596
- Li, Y., Ying, L., Jin, W., & Simon, S. R. (2020). Sentiment analysis for e-commerce product reviews in chinese based on sentiment lexicon and deep learning. *IEEE Access*, 8, 1-8, DOI: 10.1109/ACCESS.2020.2969854.
- Liu B. (2015). Introduction. In: Sentiment Analysis: Mining Opinions, Sentiments, and Emotions. Cambridge University Press; 1-46.
- Liu, B. (2012). Sentiment analysis & opinion mining. *Synthesis Lectures on Human Language Technologies*, 5(1), 1-167.
- Maas, A. L., Daly, R. E., Pham, P. T., Huang, D., Ng, A. Y., & Potts, C. (2011). Learning word vectors for sentiment analysis. In Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies, 1, 142-150.
- Mudambi, S. M., & Schuff, D. (2010). What makes a helpful online review? A study of customer reviews on Amazon.com. *MIS Quarterly*, 34(1), 185-200.
- Naeem, S., Ali, A., Anam, S., & Ahmed, M. (2023). An Unsupervised Machine Learning Algorithms: Comprehensive Review. *IJCDS Journal*, 13, 911-921, 10.12785/ijcds/130172.
- Nilsson, N. J. (2005). Introduction to Machine Learning. Stanford, CA: Stanford University Press
- O'Connor, B., Balasubramanyan, R., Routledge, B., & Smith, N. (2010, May). From tweets to polls: Linking text sentiment to public opinion time series. In Proceedings of the international AAAI conference on web & social media, 4(1), 122-129.
- Pang, B., & Lee, L. (2008). Opinion mining and sentiment analysis. *Foundations and Trends® in Information Retrieval*, 2(1-2), 1-135.
- Salma, E., Wael, A., Nagham, M., & Mohammed, E. (2023). Sentiment Analysis for E-commerce Product Reviews: *Current Trends and Future Directions*. Preprints, 2023051649. DOI: 10.20944/preprints202305.1649.v1
- Shaozhang, X., Hexiang, W., Zhi, L., Lanfang, W., & Zhaoxia, T. (2020). Sentiment analysis for product reviews based on deep learning. *Journal of physics: conference series*, 1651, 012103. DOI: 10.1088/1742-6596/1651/1/012103.
- Singla, Z., Randhawa, S., & Jain, S. (2017, June). Sentiment analysis of customer product reviews using machine learning. In 2017 international conference on intelligent computing and control (I2C2) (pp. 1-5). IEEE.
- Sobia, W., Xi, C., Tian, S., Muhammad, W., & NZ, J. (2021). "Amazon Product Sentiment Analysis using Machine Learning Techniques. 30(1), 695, DOI:10.24205/03276716.2020.2065
- Wang, Z., & Kim, H. G. (2017). Can Social Media Marketing Improve Customer Relationship Capabilities and Firm Performance? Dynamic Capability Perspective. *Journal of Interactive Marketing*, 39, 15-26.
- Wang, Q., & Zhan. Z. (2011) Reinforcement learning model, algorithms and its application, 2011 International Conference on Mechatronic Science, Electric Engineering and Computer (MEC), Jilin, China, 1143-1146, doi: 10.1109/MEC.2011.6025669.
- Webb, G.I. (2011). Naïve Bayes. In: Sammut, C., Webb, G.I. (eds) Encyclopedia of Machine Learning. Springer, Boston, MA, 713-714, DOI: 10.1007/978-0-387-30164-8_576

Oluwatoyin Agbonifo

Department of Information Systems,
Federal University of Technology,
Akure, Nigeria.
ocagbonifo@futa.edu.ng
ORCID: 0000-0001-8888-1137

Victor Olutayo

Department of Computer Science,
Federal University of Technology,
Akure, Nigeria.
vaolutayo@futa.edu.ng
ORCID: 0000-0003-0090-8994

Oluwaseyi Oluyode

Department of Information Systems,
Federal University of Technology,
Akure, Nigeria.
oluyedeoluwaseyi123@gmail.com
