

See discussions, stats, and author profiles for this publication at: <https://www.researchgate.net/publication/381690706>

Sentiment Polarity Analysis of Bangla Food Reviews Using Machine and Deep Learning Algorithms

Conference Paper · April 2024

DOI: 10.1109/ICAEEE62219.2024.10561876

CITATIONS

3

READS

26

6 authors, including:



Al Amin

American International University-Bangladesh

15 PUBLICATIONS 30 CITATIONS

SEE PROFILE



Anik Sarkar

National Institute Of Technology Silchar

3 PUBLICATIONS 3 CITATIONS

SEE PROFILE



Md Mahamodul Islam

University of Oslo

6 PUBLICATIONS 12 CITATIONS

SEE PROFILE



Asif Ahammad Miazee

Maharishi International University

11 PUBLICATIONS 189 CITATIONS

SEE PROFILE

Sentiment Polarity Analysis of Bangla Food Reviews Using Machine and Deep Learning Algorithms

Al Amin¹, Anik Sarkar², Md Mahamodul Islam³, Asif Ahammad Miaze⁴, Md Robiul Islam⁵, Md Mahmudul Hoque⁶

^{1,5} Department of Computer Science and Engineering, Uttara University, Dhaka, Bangladesh

² National Institute of Technology, Silcar, India

³ Department of Informatics, University of Oslo, Oslo, Norway

⁴ Department of Computer Science, Maharishi International University, Iowa, USA

⁶ Department of Computer Science and Engineering, CCN University of Science & Technology, Comilla, Bangladesh
{alaminbhuyan321, aniksarkar.cs, mdmahamodul1998, asifahammad7,robiul.cse.uu, cse.mahmud.evan}@gmail.com

Abstract—The Internet has become an essential tool for people in the modern world. Humans, like all living organisms, have essential requirements for survival. These include access to atmospheric oxygen, potable water, protective shelter, and sustenance. The constant flux of the world is making our existence less complicated. A significant portion of the population utilizes online food ordering services to have meals delivered to their residences. Although there are numerous methods for ordering food, customers sometimes experience disappointment with the food they receive. Our endeavor was to establish a model that could determine if food is of good or poor quality. We compiled an extensive dataset of over 1484 online reviews from prominent food ordering platforms, including Food Panda and HungryNaki. Leveraging the collected data, a rigorous assessment of various deep learning and machine learning techniques was performed to determine the most accurate approach for predicting food quality. Out of all the algorithms evaluated, logistic regression emerged as the most accurate, achieving an impressive 90.91% accuracy. The review offers valuable insights that will guide the user in deciding whether or not to order the food.

Index Terms—Online Review, Deep Learning, Machine Learning, Online Food Order, Sentiment Analysis

I. INTRODUCTION

Text-based opinion analysis is a branch of NLP that uses computational methods to analyze subjective information in text. This method, applied to unstructured text data, enables organizations to extract meaningful and socio-emotional information. As we approach the Internet era and the next phase of technological advancement, it is evident that text-based reviews shared across different online platforms, particularly restaurants via online meal delivery platforms, are prevalent. Text-based reviews found online express what consumers think about a specific product. A user can easily share his or her thoughts by reviewing the food. Nowadays, people are so smart that when they are ready to order any food, they first check the reviews of the food. If most of the users gave positive feedback, then users can easily be sure that the food is good in taste. So, it is very crucial for businessmen or

restaurant owners to make sure the food is of good quality. People are now less conscious of money and more concerned about the healthiness of food. If the food is good, then they are ready to pay a large amount of money to buy it. Its essential for the owners to provide good food to the user so that they will be satisfied and give a good review because the next customer will depend on the previous customer as the internet usability is growing so rapidly and almost everyone has access to the internet. The market position of a particular restaurant can be ascertained in order to gauge user opinion. New user reviews will be categorized by the model according to their target level and the total number of reviews. positive or negative [1]. Nowadays, the majority of restaurants have pages on Facebook or other review websites where patrons can leave comments. Once again, accurate review analysis is essential to a business's success. If the restaurant's management is unable to resolve the issue raised by customers in reviews, it will be difficult for the authority to remedy the problem.

According to the source [2], the assessed number of web clients in Bangladesh by January 2023 is 66.94 million. They generally write reviews in English, Bangla, or a combination of the two, with the latter being more common among smartphone users. There are no clear guidelines for posting reviews, so it is impossible to carefully review them, and there are frequently a ton of remarks made. Therefore, it may be useful to have A tool that can measure the emotional valence of text. Both positive and negative aspects of the user's affective state can be found in Bangla. Over 1484 reviews were collected for our study from numerous online food delivery websites, such as Food Panda. We strive to design a machine learning architecture that employs sentiment analysis and natural language processing (NLP) to automatically detect and display the ratio of negative to positive reviews. We preprocessed the data by eliminating stop words, tokenization, stemming, superfluous punctuation, and other symbols. We employ CountVectorizer, Term Frequency-Inverse Document

Frequency vectorizer, and N-gram to vectorize text into feature vectors, which is a means of evaluating deep learning and machine learning techniques. Finally, we preferred the Random Forest Classifier, Linear SVM, Naïve Bayes, Decision Tree Classifier, Logistic Regression, Multinomial Naïve Bayes, and LSTM for classifying positive and negative evaluations. In conclusion, we collected over 1484 reviews of Bengali food for our study without converting or altering the original content, which the public will have access to for upcoming studies. These evaluations were gathered by hand from online meal delivery services. In order to achieve greater accuracy than just machine learning techniques, we employed deep learning techniques. Overall, this research project's contribution is-

- We gathered diverse Bengali data from several well-known online food delivery services in order to apply sentiment analysis of food reviews in Bengali text
- The gathered dataset was carefully annotated with two sentiment polarities: positive and negative.
- This much data on online Bangla food reviews in Bengali text has never been used before.
- Our suggested strategy outperforms existing approaches employed by researchers in terms of accuracy.

The rest of the work is separated into four segments. Segment II keeps works that are correlated. Our recommended methodology is presented in Segment III. Segment IV displays the findings of the experiment. Lastly, in Segment V, we came to a conclusion and suggested some next steps.

II. RELATED WORKS

In this age of advanced technology, online food ordering systems are very common. Its becoming popular with all kinds of customers because of the easy ordering process and because people can easily pay the bill through an online bank. Sentiment analysis, also referred to as opinion mining, is a natural language processing (NLP) method employed to ascertain the emotional tone or sentiment expressed in a piece of text. Sentiment analysis seeks to comprehend and categorize the sentiment or opinion conveyed by the text as positive or negative. It can provide insights into the overall sentiment of a document, sentence, or even individual words. Several researchers have worked on this not only in food but also in other domains such as restaurants, books, movies, etc. Junaid et al. [3] proposed an LSTM model after trying several machine learning and deep learning algorithms, where their best accuracy was 90.89% with word2sequence feature extraction. Hossain et al. [4] proposed a model after collecting 1,000 restaurant reviews from the Priyao website. Theyve used POs tagged and TF-IDF vectorizers for feature extraction. After using several machine learning algorithms such as multinomial Naive Bayes, logistic regression, K-Nearest Neighbor, and support vector machines, 77% accuracy was obtained in the validation dataset by Asiful et al. [5]. The Facebook page Food Bank was used for collecting the dataset. Hasan et al. [6] used several machine learning algorithms on 5000 restaurant reviews, where several feature extraction methods were used, like bag of words, term frequency-inverse term

frequency (TF-IDF), and skim-gram. Skip-gram gives the best accuracy among the. Following the application of LSTM and the CNN attention mechanism to the English dataset, Bhuiyan et al. achieved 98.4% accuracy [7]. The majority of these interviews were utilized for Bangla review sentiment analysis. After collecting data from several translations or groups and translating from English to Bangla with two classes, such as positive and negative Sharif et al. [8] account for the Naive Bayes in the validation set. After translating online shopping reviews, Rahman et al. [9] achieved 83% accuracy for the CNN model and 78% correctness for the Support Vector Machine (SVM). Nayan Banik et al. [10] developed a polarity detection system for movie reviews. They used a precision of 0.86 after using the stemmed unigram feature extraction technique with support vector machines and Naive Bayes. Another work has been done with restaurant reviews. Fableha et al. [11] developed a model with various feature extraction approaches such as N-gram techniques, count vectorizers, and TF-IDF. After using several machine learning algorithms, they got 75.58% accuracy on SVM. Sentiment analysis not only relies on food or restaurants; it is also used in the Bangla Cricket Commentary. Mahtab et al. [12] proposed a model where they used an SVM classifier and a TF-IDF vectorizer. They couldnt achieve significant accuracy. Opinion analysis used not only Bangla or English but also the native language as well. In their [13] work, Rohini et al. present a model for identifying emotion in Kannada, an Indian native language. Machine learning algorithms are applied for classification; they show the accuracy comparison between English and Kannada datasets. Despite their inability to collaborate with other Indian languages, Animesh and Pintu [14] suggested a model in which they used Mutual Information to carry out feature selection [15]. They utilized MNB for classification and demonstrated that the Bangla dataset is more accurate than the English dataset. Nonetheless, they were unable to process actual Bangla reviews. Another work has been done with Twitter posts; they perform their work on a specific application called Traveloka. After collecting 12 thousand tweets and performing numerous machine learning algorithms, SVM got the highest accuracy [16]. The TripAdvisor website was used for collecting reviews of the restaurants in Manoda City. Theyve collected overall 640 reviews from that website, and they got 76.20% accuracy for Naive Bayes [17].

III. METHODOLOGY

Within this section, an exhaustive account of the comprehensive approach and the suite of tools employed for the execution of the sentiment analysis task on Bengali Foods reviews is delineated. Before commencing the training process, it is imperative to elaborate on the dataset pre-processing steps. Subsequently, we will introduce our proposed model for sentiment analysis. The workflow, as elucidated in Figure 1, provides a visual representation of the entire process, with further intricacies expounded upon in the subsequent sections.

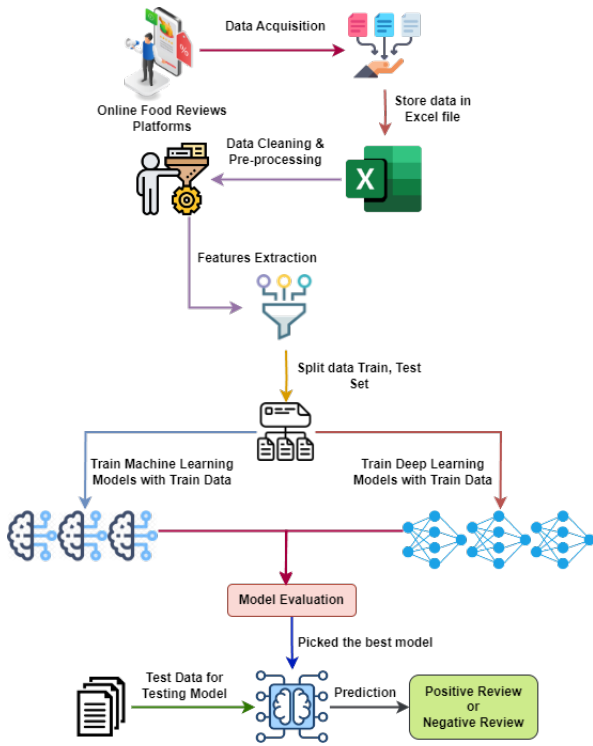


Fig. 1. Proposed Methodology Workflows.

A. Data Acquisition

Convocating, calculating, and checking exact observations for research based on recognized systems constitutes the process of data gathering. Researchers can derive their hypotheses from the information they have collected. The importance of data collection has a direct bearing on the upkeep of research unity and the making of business judgments. Our data was gathered from online reviews found on several online meal delivery services, including FoodPanda, Shohoz Food, Pathao Food, HungryNaki, and others. We collected this data manually from different restaurant and food review sections under the online food delivery services without changing or violating the original review. Then we stored those data in an Excel file and their corresponding sentiments. A total of 1400+ Bengali food reviews have been gathered, of which 718 have been classified as positive and 766 as negative. Later, our dataset contains two columns reviews and sentiment. The task is a binary classification problem to detect whether the review is classified as positive 1 or negative 0. We can see some examples of data in Table I, and Table II shows the distribution for both the positive and negative classes of our dataset.

B. Data Cleaning and Pre-processing

Data wrangling is the process of converting raw data into a usable format for machine learning algorithms. Cleaning the data is far too crucial in order to extract high-quality data for our experiment, which will enable us to minimize test errors. Textual representations of online food reviews are not

TABLE I
EXAMPLE OF DATA

NO.	Reviews	Sentiment
1	এক বারে বাজে খাবার এত বাজে খাবার আমি ভাবতেও পারি নাই।	0
2	ওয়েল ফুডের খাবার মজার ছিল এবং আচরণ ও ভাল ছিল।	1

TABLE II
DISTRIBUTION OF CLASSES

Review Sentiment Polarity	Number of Data
Positive	718
Negative	766
Total	1484

sufficient for machine learning. We must make sure that the data has been cleansed before we train the model. Our dataset has to be pre-processed using the Bangla text pre-processing approach in order to get the textual representations ready for model building. First, include contractions in our dataset, and then extract the contractions' brief form. Using regular expressions, we first eliminated punctuation marks, emoticons, pictorial icons, emojis, random English words, and alphabets from our dataset. There are also various Unicode characters in Bangla text that we must precisely put back. Second, we employ the BNL Python library to tokenize our assessment transform phrases into individual words and retrieve Bangla language-specific exclusionary terms. Furthermore, we extract the root words from each word token using Bangla word stemming methods. Finally, all related words and, to improve feature extraction, word stems are amalgamated into sentences. Below Table III is a demonstration of raw data and cleaned data differences.

TABLE III
COMPARISON OF RAW REVIEWS AND CLEANED REVIEWS

NO.	Raw Reviews	Cleaned Reviews
1	ওয়েটার দের ব্যবহার একদম ফালতু,,,, ইটা তাদের কাসে আসা সিল না,,,, পুরাই জঘন্য ব্যবহার..... রাবিশ	ওয়েটার দের ব্যবহার একদম ফালতু ইটা তাদের কাসে আসা সিল না পুরাই জঘন্য ব্যবহার রাবিশ
2	অরিজিনাল বারবিকিউ চিকেন এবং হানি মাস্টার্ড সস !!! খাওয়ার জন্য নিয়েছিলাম স্বাদটা চমৎকার ছিল এবং পরিবেশও খুব ভাল ছিল	অরিজিনাল বারবিকিউ চিকেন এবং হানি মাস্টার্ড সস খাওয়ার জন্য নিয়েছিলাম স্বাদটা চমৎকার ছিল এবং পরিবেশও খুব ভাল ছিল

C. Features Distillation

To feed the classifier models, several features were taken out of the review text. As previously indicated, we classified our meal review dataset using both deep learning models and conventional machine learning techniques. Textual representations cannot be classified by either of these models.

They require numerical representations of numerous textual and sentence-level qualities and attributes to analyze and forecast distinct feelings. To explore and obtain better results, we employed different pieces for different classifiers. In this part, we'll quickly review those features.

Counter Vectorizer: Count Vectorizer facilitates the conversion of text data into a matrix format, where each row represents a document and each column represents a word, with the matrix elements representing the word counts. The count vectorizer creates a sparse matrix of all the unique words in the dataset. In a document, every word has an index number that indicates how often that word appears in that document. Then, we use the vector delineation as a feature for diverse models. We applied count vectors to our dataset and tested them on abundant imitation.

Let $\mathbf{D}$ be the collection of documents and t be a term (word). The count vectorization of term t in document d is denoted as $\text{Count}(t, d)$. The matrix representation $\mathbf{X}$ is created by counting the occurrences of terms in documents.

$$X_{t,d} = \text{Count}(t, d) \quad (1)$$

Here, $X_{t,d}$ represents the count of term t in document d .

TF-IDF: TF-IDF stands for term frequency-inverse document frequency. It is a statistical measure that evaluates how important a word is to a document in a collection of documents. This is done by multiplying two metrics. Term frequency (TF): This is the number of times a word appears in a document. Inverse document frequency (IDF): This is a measure of how rare a word is across a set of documents. The rarer a word is, the more important it is to a document. For TF-IDF vectorizing, we employed bigram and unigram, with contiguous two words serving as additional feature input. N-grams are helpful in creating context between words, which helps to convey better sentence qualities. The "*tf-idf*" statistic we're using is,

$$tf-idf(w, r) = tf(w, r) \frac{N}{|\{r \in R : w \in r\}|} \quad (2)$$

Here,

- $tf-idf(w, r)$ = value of word w in the review r ,
- $tf(w, r)$ = frequency of word w in review r ,
- N = total amount of reviews,
- $|\{r \in R : w \in r\}|$ = number of reviews containing w .

GloVe-Vectors: This is another approach to represent a word in vector space. GloVe is a technique for unsupervised learning that generates vector representations of words. The training process is carried out using combined global word-word co-occurrence statistics from a corpus. The outputs demonstrate intriguing linear substructures inside the word vector space. In our study, we use Bengali pre-trained GloVe vector that is trained on Wikipedia+crawl news articles (39M (39055685) tokens, 0.18M (178152) vocab size. This model is built for BNLP package. BNLP is a natural language processing toolkit for Bengali Language. This tool will help

you to tokenize Bengali text, Embedding Bengali words, Embedding Bengali Document, Bengali POS Tagging, Bengali Name Entity Recognition, Bangla Text Cleaning for Bengali NLP purposes. This model is available on the HuggingFace platform.

D. Sentiment Polarity Models

In our study, we split the dataset training and test set as 80% and 20% ratio. We applied different well-known machine learning algorithms such as Logistic-Regression (LR), DecisionTreeClassifier, RandomForestClassifier, MultinomialNB, KNeighborsClassifier, Support Vector Classification (SVC), SGD Classifier or Stochastic Gradient Descent (SGD). We also tried deep learning approach to get good results for that we used the Recurrent Neural Network (RNN), Long Short-Term Memory (LSTM), and Gated Recurrent Unit (GRU) model. We used some hyperparameters to tune the model to get better accuracy. The hyperparameter value is given in the below table for machine learning and deep learning models.

E. Model Implementation

We've gathered various attributes such as Tf-Idf, glove vector, count vector, and more from the refined raw text information. Following this, we input the data into our deep learning and machine learning models. The machine learning model, also referred to as the logistic regression model, has produced results that are reasonably excellent. We trained our Logistic Regression model with random state 123 using the unigram features vectors. Therefore, our accuracy is 90.91%.

IV. EXPERIMENTAL RESULTS

In this research for the training model, data was split into 80% for training and 20% for testing. We have examined several machine learning and deep learning models. To assess our model's effectiveness, we've employed four distinct categories of metrics. These metrics are outlined below.

Accuracy: is one crucial factor in evaluating the classifier's performance. This metric indicates the proportion of correct classifications among all predictions made by the classifier. Accuracy calculation is defined in Equation-3.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (3)$$

Precision: can be described as the count of accurate positive outputs generated by the model. The precision calculation is shown in Equation-4.

$$\text{Precision} = \frac{TP}{TP + FP} \quad (4)$$

Recall: measures the effectiveness of a model in recognizing positive instances among all the genuine positive instances within the dataset. Recall calculation is defined in Equation-5.

$$\text{Recall} = \frac{TP}{TP + FN} \quad (5)$$

F1-Score: assists in assessing both recall and precision concurrently. The F1-score reaches its peak when recall matches precision. It can be calculated using Equation-6.

$$F1\text{-Score} = \frac{2 * Precision * Recall}{Precision + Recall} \quad (6)$$

The subsequent parameters are employed in computing the aforementioned metrics:

- **True Positive (TP):** The model has predicted a sentiment as positive, and the actual sentiment was positive.
- **False Positive (FP):** The model has predicted a sentiment positive, but the actual sentiment was negative.
- **True Negative (TN):** The model has predicted a sentiment negative, and the actual sentiment was also negative.
- **False Negative (FN):** The model has predicted a sentiment was negative and actually sentiment was positive.

A. Machine Learning Model Evaluation

Table-V displays the machine learning model's output. Tf-Idf with unigram was utilised in this experiment as a feature extraction method to feed the machine learning models. It is evident that Logistic Regression performs fairly well when a Tf-Idf is employed.

TABLE IV
RESEARCH COMPARISON WITH OTHERS

No.	Authors	Topic	Data	Model	Accuracy
1	Junaid et al . [2]	Food	1040	LSTM	90.86%
2	O. Sharif et al [7]	Restaurant	1000	MNB	80.48%
3	Proposed Model	Food	1484	LR	90.91%

TABLE V
MODEL EVALUATION TABLE

No.	Algorithm	Accuracy	Precision	Rcall	F1-Score
1	Logistic Regres-sion	90.91	96.40	85.90	90.85
2	DecisionTree Classifier	85.52	86.93	85.26	86.08
3	RandomForest Classifier	89.56	94.33	85.26	89.56
4	MultinomialNB	84.51	91.04	78.21	84.14
5	KNeighbors Classifier	80.47	81.41	81.41	81.41
6	Linear SupportVec-torClassification	84.90	98.46	82.05	89.51
7	RBF SupportVec-torClassification	89.90	98.46	82.05	89.51
8	SGD Classifier	90.24	95.04	85.90	90.24

Our model's confusion matrix is displayed in Fig-2. Here, 0 and 1 stand for negative and positive values, respectively. In this figure, we can clearly see the false positive and false negative values are very low in comparison to true negative and true positive. This indicates that our model is more accurate.

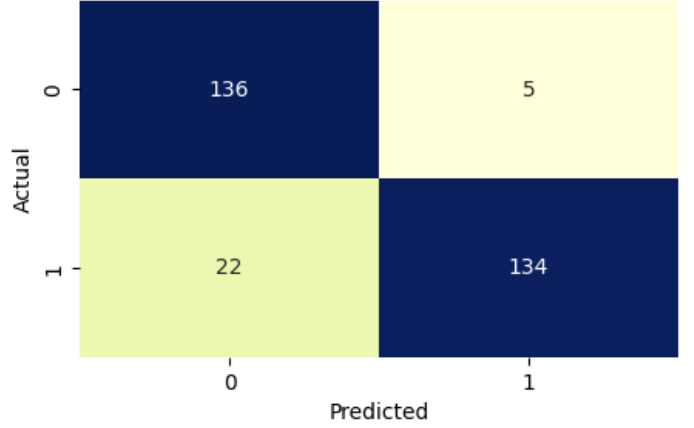


Fig. 2. Confusion Matrix of Proposed Model

Fig-3 showing the ROC curve analysis for unigram features in different machine learning models.

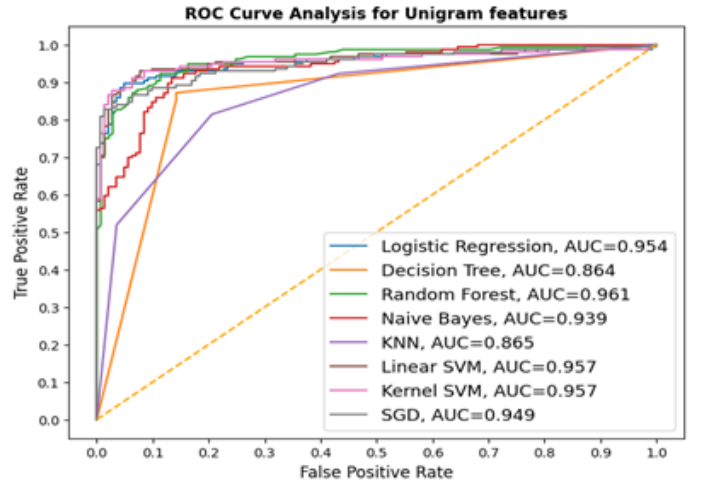


Fig. 3. ROC Curve of Proposed Model

Fig-4 showing the Precision and Recall curve analysis for unigram features in different machine learning models.

V. CONCLUSION AND FUTURE WORK

Our research aimed to examine public sentiment towards Bengali cuisine through the analysis of online reviews. As internet usage has risen in Bangladesh, individuals have increasingly turned to online platforms to express their opinions and experiences regarding food. We collected 1400+ Bengali (Bengali Text) food reviews from social media and other online platforms and trained them using a range of machine learning and deep learning techniques. To make the text reviews easier for comprehension by our algorithm, we

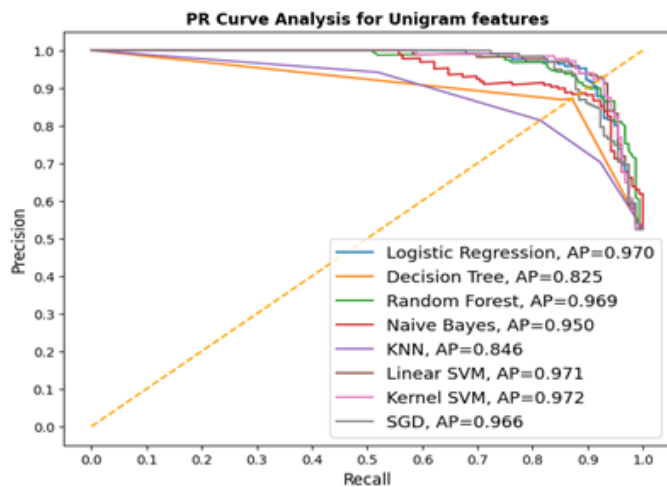


Fig. 4. Precision, Recall Curve of Proposed Model

used some data cleaning and feature extraction techniques before the model training. A comparative evaluation of all the classifiers is carried out using feature extraction methods like TF-IDF. According to the results of our experiment, the Logistic Regression (LR) classifier can reach the greatest accuracy of TF-IDF, at 90.91%.

Even though we received a decent score for the proposed model, expanding the dataset will yield better results. In the future, mobile or web-based applications can be developed using pre-trained models based on BERT or Transformer to obtain greater performance. In this study, we worked on binary class classifications, positive and negative. So, in future work, it can be expanded to include multi-class classification problems like positive, negative, and neutral sentiment polarity analysis. Since our dataset contains only positive and negative classes. So, by adding some neutral reviews, it can be possible to solve a multi-class classification problem.

REFERENCES

- [1] X.-L. Wang and Cloete, "Learning to classify email: a survey," in *2005 International Conference on Machine Learning and Cybernetics*, vol. 9, pp. 5716–5719 Vol. 9, 2005.
- [2] DataReportal, "Digital 2023: Bangladesh." <https://datareportal.com/reports/digital-2023-bangladesh>. [Online; accessed 23-March-2024].
- [3] M. I. Hossain Junaid, F. Hossain, U. S. Upal, A. Tameem, A. Kashim, and A. Fahmin, "Bangla food review sentimental analysis using machine learning," in *2022 IEEE 12th Annual Computing and Communication Workshop and Conference (CCWC)*, pp. 0347–0353, 2022.
- [4] N. Hossain, M. R. Bhuiyan, Z. N. Tumpa, and S. A. Hossain, "Sentiment analysis of restaurant reviews using combined cnn-lstm," in *2020 11th International Conference on Computing, Communication and Networking Technologies (ICCCNT)*, pp. 1–5, 2020.
- [5] S. M. Asiful Huda, M. M. Shoikot, M. A. Hossain, and I. J. Ila, "An effective machine learning approach for sentiment analysis on popular restaurant reviews in bangladesh," in *2019 1st International Conference on Artificial Intelligence and Data Sciences (AiDAS)*, pp. 170–173, 2019.
- [6] T. Hasan, A. Matin, M. Kamruzzaman, S. Islam, and M. O. Faruq Goni, "A comparative analysis of feature extraction methods for human opinion grouping using several machine learning techniques," in *2020 IEEE International Women in Engineering (WIE) Conference on Electrical and Computer Engineering (WIECON-ECE)*, pp. 272–275, 2020.
- [7] M. R. Bhuiyan, M. H. Mahedi, N. Hossain, Z. N. Tumpa, and S. A. Hossain, "An attention based approach for sentiment analysis of food review dataset," in *2020 11th International Conference on Computing, Communication and Networking Technologies (ICCCNT)*, pp. 1–6, 2020.
- [8] O. Sharif, M. M. Hoque, and E. Hossain, "Sentiment analysis of bengali texts on online restaurant reviews using multinomial naïve bayes," in *2019 1st International Conference on Advances in Science, Engineering and Robotics Technology (ICASERT)*, pp. 1–6, 2019.
- [9] M. A. Rahman and E. Kumar Dey, "Aspect extraction from bangla reviews using convolutional neural network," in *2018 Joint 7th International Conference on Informatics, Electronics Vision (ICIEV) and 2018 2nd International Conference on Imaging, Vision Pattern Recognition (icIVPR)*, pp. 262–267, 2018.
- [10] N. Banik and M. Hasan Hafizur Rahman, "Evaluation of naïve bayes and support vector machines on bangla textual movie reviews," in *2018 International Conference on Bangla Speech and Language Processing (ICBSLP)*, pp. 1–6, 2018.
- [11] F. Haque, M. M. H. Manik, and M. Hashem, "Opinion mining from bangla and phonetic bangla reviews using vectorization methods," in *2019 4th International Conference on Electrical Information and Communication Technology (EICT)*, pp. 1–6, 2019.
- [12] S. Arafin Mahtab, N. Islam, and M. Mahfuzur Rahaman, "Sentiment analysis on bangladesh cricket with support vector machine," in *2018 International Conference on Bangla Speech and Language Processing (ICBSLP)*, pp. 1–4, 2018.
- [13] V. Rohini, M. Thomas, and C. A. Latha, "Domain based sentiment analysis in regional language-kannada using machine learning algorithm," in *2016 IEEE International Conference on Recent Trends in Electronics, Information Communication Technology (RTEICT)*, pp. 503–507, 2016.
- [14] A. K. Paul and P. C. Shill, "Sentiment mining from bangla data using mutual information," in *2016 2nd International Conference on Electrical, Computer Telecommunication Engineering (ICECTE)*, pp. 1–4, 2016.
- [15] H. Liu, J. Sun, L. Liu, and H. Zhang, "Feature selection with dynamic mutual information," *Pattern Recognition*, vol. 42, no. 7, pp. 1330–1339, 2009.
- [16] Z. A. Dieksona, M. R. B. Prakosoa, M. S. Qalby, M. S. A. F. S. Putraa, S. Achmada, and R. Sutoyo, "Sentiment analysis for customer review: Case study of traveloka," *Procedia Computer Science*, vol. 216, pp. 682–690, 2023.
- [17] J. P. Matrutty, A. M. Adrian, and A. Angdresey, "Sentiment analysis of visitor reviews on star hotels in manado city," *Journal of Information Technology and Computer Science*, vol. 8, no. 1, pp. 21–32, 2023.