

Title of Thesis

**Sentiment Polarity Detection on Bengali Book Reviews By Word2vec
and Hybrid Machine Algorithm**

Submitted By

Saidur Rahman

ID: 201-16-499

Supervised By

Sonia Nasrin

Lecturer

Computing & Information System
Daffodil International University



DEPARTMENT OF COMPUTING AND INFORMATION SYSTEM
DAFFODIL INTERNATIONAL UNIVERSITY
DHAKA, BANGLADESH

Submission date: 23-09-2024

APPROVAL

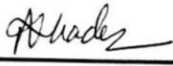
This thesis titled “Sentiment Polarity Detection on Bengali Book Reviews By Word2Vec and Hybrid Machine Algorithm”, submitted by **Saidur Rahman**, ID No: **201-16-499** to the Department of Computing & Information Systems, Daffodil International University, has been accepted as satisfactory for the partial fulfillment of the requirements for the degree of B.Sc. in Computing & Information Systems and approved as to its style and contents. The presentation has been held on- 23-09-2024.

BOARD OF EXAMINERS



Mr. Md. Sarwar Hossain Mollah
Associate Professor and Head
Department of Computing & Information System
Faculty of Science & Information Technology
Daffodil International University

Chairman



Md. Nasimul Kader
Assistant Professor
Department of Computing & Information System
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner



Mr. Mehedi Hasan
Lecturer
Department of Computing & Information System
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner



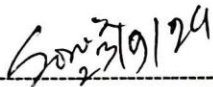
Dr. Saifuddin Md. Tareeq
Professor
Department of Computer Science and Engineering
University of Dhaka, Dhaka

External Examiner

Declaration

I hereby declare that; this thesis has been done by me under the supervision of **Sonia Nasrin, Lecturer**, Department of Computing and Information System (CIS) of Daffodil International University. I am also declaring that this thesis or any part of there has never been submitted anywhere else for the award of any educational degree like, B.Sc., M.Sc., Diploma or other qualifications.

Supervised By

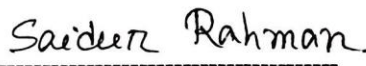


Sonia Nasrin

Lecturer

Department of CIS
Daffodil International University

Submitted By



Saidur Rahman

ID: 201-16-499

Department of CIS
Daffodil International University

ACKNOWLEDGEMENT

First, we express our heartiest thanks and gratefulness to almighty for His divine blessing making it possible for us to complete the final year project/internship successfully.

I am grateful and wish our profound indebtedness to **Sonia Nasrin, Lecturer**, Department of Computing and Information System (CIS) Daffodil International University, Dhaka. Deep Knowledge & keen interest of our supervisor in the field of *Artificial Intelligence* to carry out this project. His endless patience, scholarly guidance, continual encouragement, constant and energetic supervision, constructive criticism, valuable advice, reading many inferior drafts, and correcting them at all stages have made it possible to complete this project.

I would like to express our heartiest gratitude to the **Mr. Sarwar Hossain Mollah, Head of the Department of Department of Computing and Information System (CIS)**, for his kind help in finishing our project and also to other faculty members and the staff of the Department of Computing and Information System (CIS), Daffodil International University.

I would like to thank our entire course mate in Daffodil International University, who took part in this discussion while completing the course work.

Finally, I must acknowledge with due respect the constant support and patience of our parents.

ABSTRACT

Despite sentiment analysis's significance in gauge public opinion, current methods frequently fail to account for Bangla and other languages sufficiently. Specifically, this work develops a Sentiment Polarity Detection on Bengali Book Reviews By Word2vec and Hybrid Machine Algorithm. A total of 2012 reviews, including both good and unfavorable assessments, were compiled from various sources. The study assesses several ML models, including Word2Vec, for feature extraction, RF, SVM, and a hybrid of the two. With superior recall, accuracy, and precision, the best model was the one that combined RF with SVM. Results showed that the most effective RF and SVM hybrid models scored 0.92%. On top of that, they took first place in the categorization report with an exceptional accuracy rate of 0.93%. These outcomes prove that it can easily handle complex Bangla text. Improved processing of the Bangla natural language is the goal of the suggested effort. It paves the way for more advancements, such as creating an AI-driven mobile app that can analyze sentiment in real time and make it more accessible than a web app.

TABLE OF CONTENTS

CONTENTS	PAGE NO.
Declaration	i
Abstract	ii
List of Tables	v
List of Figures	vi
 CHAPTER	
 CHAPTER 1: INTRODUCTION	 1-4
1.1 Overview	1
1.2 Research Objective	2
1.3 Motivation	2
1.4 Rational of Study	3
1.5 Research Question	3
1.6 Project Outcome	4
 CHAPTER 2: BACKGROUND STUDY	 5-8
2.1 Introduction	5
2.2 Related Work	5
2.3 Perspective of Bangladesh	8
2.4 Scope of this problem	8
2.5 Challenge of this work	8
	9-13

CHAPTER 3: Methodology

3.1 Overview	9
3.2 Data Collection	9
3.3 Data Preprocessing	10
3.4 Classification	12
3.5 Implementation Requirement	12

CHAPTER 4: RESULT ANALYSIS

4.1 Experimental Result	14
5.2 Classification report	15
5.4 Decision	18
5.5 Evaluation	18

CHAPTER 5: CONCLUSION

CHAPTER 6: FUTURE WORK

REFERENCES

PLAGIARISM REPORT

LIST OF TABLES	
Table	PAGE NO.
Figure 1 Tokenization table	11
Figure 2 Accuracy Table	15
Figure 3 RF and SVM Classification Report	16
Figure 4 DT, NN, and SVM Classification Report	17
Figure 5 DT, RF, and SVM Classification Report	18
Figure 6 Evaluation table of Rf and SVM	19

LIST OF FIGURES	
FIGURES	PAGE NO.
Figure 3.1 Mythology Diagram	9
Figure 3.2 Processing Step	10
Figure 3.3 Dataset Representation	12
Figure 4.1 Confusion matrix	19
Figure 4.2 Evaluation graph	21
Figure 4.3 ROC Curve	22

LIST OF ACRONYMS

SVC Support Vector Classifier

RF Random Forest

NN Neural Networks

DT Decision Trees

KNN K-Nearest Neighbors

LR Logistic Regression

CHAPTER 1

INTRODUCTION

1.1 Overview

The study of book reviews is a promising and innovative area of application. To ascertain a book's opinions and level of success, it is advisable to peruse its reviews. A comprehensive grasp of the tone of these evaluations is essential for readers, publishers, and writers seeking to ascertain public sentiment toward literary works, enabling them to make informed judgments. The complex morphology, wide range of vocabulary, and cultural nuances of the Bangla language provide challenges for conventional sentiment analysis systems in comprehending these reviews. Thanks to the proliferation of online review sites, more people than ever before can share their thoughts and feelings about books with the world. There are a lot of reviews of Bangla literature on literary forums, social media, and websites that sell books. By reading these reviews, we can get a sense of the consensus among readers regarding the works of literature. However, processing and interpreting the sentiment in Bangla text is challenging due to the language's rich and complex morphological structure. The absence of natural language processing (NLP) resources tailored to Bangla sentiment analysis is a big problem in this area. Bangla has not been studied nearly as much as English. Hence, there is a lack of data to build reliable sentiment analysis models for the former language. Review sentiment analysis is a win-win for all authors, readers, and publishers. Authors and publishers can use it to track reader reactions, which improves their publications and allows for more personalized recommendations. Identifying positive and negative reviews becomes challenging without effective machine learning techniques tailored to Bangla text. This method for analyzing Bangla book reviews is being developed to address this issue. It incorporates feature extraction using Word2Vec and classification using hybrid machine learning techniques. Using a dataset of 2012 reviews sourced from widely-read Bangla book review groups, online book-selling platforms, and forums like Rosemary, we investigate various hybrid models, such as Decision Trees (DT) with Neural Networks (NN), Random Forests (RF)

with Support Vector Machines (SVM), and others. The RF/SVM hybrid model achieved the best results, with a test accuracy of 92% and a classification accuracy of 93%. The success of this integrated machine learning strategy in processing Bangla text and sentiment analysis is evident from the results. This paper offers a framework for efficient sentiment classification of Bangla text, filling a gap in the literature on Bangla NLP. Not only do the proposed hybrid models—especially those that merge RF and SVM—deliver remarkable accuracy, but they also point to avenues for further investigation into using this approach to various forms of Bangla text, such as news articles and social media posts. The study concludes that the field of Bangla sentiment analysis has grown thanks to the innovative hybrid machine-learning approach used in this research. By building on the results of this work, I hope to improve automated sentiment analysis in languages other than English and advance Bangla's natural language processing.

4.2 Research Objective

- Implement Word2Vec and other hybrid machine-learning techniques to construct a sentiment analysis model for Bengali book reviews.
- Evaluate the recall, accuracy, and precision of hybrid models that integrate Random Forest (RF) and Support Vector Machine (SVM) neural network algorithms.
- Discover the hybrid approach that yields the most optimal classification of Bangla texts.
- Consider morphology and context as distinct challenges specific to artificial intelligence processing the Bangla natural language.
- Achieve a significant impact on the Bangla natural language processing (NLP) resources.
- Adapt it to serve purposes beyond book evaluations in the Bangla language.
- Initiate the exploration of Bangla sentiment analysis in the prospective future.

1.3 Motivation

The exponential growth of online information, mainly reviews and opinions, creates a

significant demand for practical sentiment analysis tools compatible with numerous languages. The absence of similar advancements in Bangla natural language processing (NLP) has created a gap in our knowledge of sentiment in Bangla text compared to the extensive research conducted for languages such as English. The increasing number of Bangla-speaking individuals engaging in online discussions about literature and other topics has created an urgent need for robust sentiment analysis models. Book assessments, which provide keen insights into reader experiences, can significantly influence literary works' popularity. However, comprehending Bangla reviews is challenging due to its intricate morphology and restricted natural language processing capabilities. This project aims to address the knowledge gap by employing advanced techniques such as Word2Vec and hybrid machine learning models to develop a sentiment analysis system designed explicitly for Bangla book reviews. Furthermore, it will contribute to the broader advancement of Bangla natural language processing.

1.4 Rational of Study

This article fills the most significant gaps in Bangla natural language processing and sentiment analysis. There aren't many advanced methods for accurately detecting sentiment in the growing amount of Bangla content online, mainly reviews of books. Because of Bangla's complex morphology and unique linguistic features, classical sentiment analysis approaches for English don't always work. To determine if evaluations of Bangla books are favorable or unfavorable, this study employs feature extraction with Word2Vec and hybrid machine learning techniques such as Random Forest (RF) and Support Vector Machine (SVM). The objectives are to improve Bangla natural language processing and to supply valuable data for writers, publishers, and readers. In addition to contributing to multilingual sentiment analysis, this research aims to enhance Bangla sentiment analysis and set a standard for future efforts in this neglected area.

1.4 Research Question

- How effective is Word2Vec feature extraction for Bangla sentiment analysis?
- In terms of accuracy, which RF/SVM hybrid model best categorizes reviews of Bangla books?
- Why is it so hard to do sentiment analysis on Bangla text?
- How do hybrid models for sentiment categorization in Bangla stack up regarding recall and accuracy?
- Could this idea work with any Bangla content, not just reviews of books?

1.4 Research Outcome

- Use Word2Vec and hybrid ML to build a model for sentiment analysis of Bangla book reviews.
- Please take a look at RF and SVM, two hybrid models, and see how they do.
- Determine the most effective model combination for accurately classifying sentiment.
- Take on the contextual and morphological challenges of Bangla natural language processing systems.
- Incorporate a robust sentiment analysis framework with Bangla natural language processing tools and resources.
- Extend the method to Bangla literature other than book reviews.
- Prepare the way for a study on Bangla sentiment analysis.

CHAPTER 2

BACKGROUND STUDY

2.1 Introduction

The field of sentiment analysis within Natural Language Processing (NLP) is expanding rapidly, providing valuable insights into consumer behavior and public opinion. This chapter of the sentiment analysis literature examines various approaches for various languages and topics. Compared to languages like Bangla, which have less robust natural language processing resources, there is a dearth of English sentiment analysis studies. The literature review delves into several sentiment analysis methods, including Word2Vec for feature extraction and classic algorithms like SVM, RF, and NN for sentiment analysis. Next, we delve into hybrid models incorporating techniques to enhance precision and deal with intricate language factors. The review covers issues and developments in Bangla natural language processing (NLP) tokenization, sentiment representation, and morphology.

2.2 Related work

M.N.M Adnan et al. [1] state that fuzzy logic utilizes "degrees of truth" instead of the binary "true or false" representation of Boolean logic (1 or 0). Recommender systems curate reading and browsing material by leveraging user preferences. Furthermore, they elaborated on how e-commerce and data access efficiently analyze extensive data sets to assist consumers in making well-informed product choices. Several recommendation techniques exist, encompassing content-based, collaborative, and knowledge-based methodologies. The application of fuzzy logic enables the identification and suggestion of content similar to that of readers. A fundamental objective exists for fuzzy logic. Employing the numerical values of 0 and 1 poses a difficulty for news articles. The mere assertion of the relationship between 'X' and 'Y' is inadequate. The researchers embarked on developing a fuzzy algorithm that, based on several characteristics of news items, would ascertain the suitability of the piece for user recommendations.

A study by C. Feng et al. [2] showed that news organizations have redirected their attention from print newspapers to online platforms and mobile applications. News recommendation systems can autonomously analyze comprehensive articles and provide relevant goods according to criteria specified by the user. News article The ongoing effort aims to categorize and document challenges in the recommender space, territorial strategies based on application space, assessment datasets, and sources, and existing addressed problems. Eighty-one interconnected issues were selected and categorized into six categories for this study, covering the period from 2001 to 2019. Considerable debate ensued about datasets, with 60% of the participants employing a hybrid methodology. Various news proposal frameworks were developed to tackle many issues inherent in the news domain. This paper represents the initial comprehensive analysis of news ideas and their corresponding metrics—in-depth research results in superior recommendations for news articles.

Tan el. At.[3] The significance of e-learning recommendation systems arises from the phenomenon of information overload, as they empower students to make informed choices even in the absence of adequate practical experience. The primary objective of our study is to investigate collaborative filtering techniques that leverage user input and remote learning. An online e-learning recommendation system consists of five stages: data collecting, data ETL, model implementation, strategy setup, and service deployment. Also included is a plan for future growth. Among the seven components of this architecture, the databases utilized by recommendation models and systems, together with the modules responsible for data management and suggestions, hold paramount significance. The Goynai assembly [4]. Given the abundance of data, online higher education institutions are employing recommendation algorithms to assist students in making informed choices. Our research suggests that recommendation algorithms and online education systems are the optimal settings for implementing user-based collaborative filtering approaches. An online e-learning recommendation system consists of five stages: data collection, Extreme Value Linearization (ETL), model development, strategy implementation, and service delivery. Furthermore, you receive a modular development framework. Of their seven components, four are indispensable: a repository of model recommendations,

recommendation systems, recommendation management, and data/model management. Alejandro et al. [5] created an e-learning recommendation system to aid students in decision-making within the current age of copious information. The primary objective of our study is to investigate collaborative filtering techniques that leverage user input and remote learning. An online e-learning recommendation system consists of five stages: data collecting, data ETL, model implementation, strategy setup, and service deployment. We also provide architecturally designed building blueprints that are structurally sound. The four essential components of this seven-part framework are recommendation systems, data/model management, recommendation administration, and databases for reference models. The Hindi evaluation approach developed by Mittal et al. achieves an accuracy of 82.89% for positive results and 76.59% for negative scores [6]. In order to attain consistency, they assessed emotions and augmented the database. The objectives of this study are to examine the perspectives of Roman Urdu speakers on politics, culinary arts, computer science, athletics, and computer programming. This presentation comprises 10,021 sentences extracted from 566 online talks. The objective of this project is to develop a Roman Urdu corpus that humans have annotated and to investigate sentiment analysis techniques for processing emotional language utilizing Rule-based and N-gram (RCNN) models. Aspect-based opinion evaluation, a subset of assumption analysis, can reveal people's opinions on a topic. Rahman and colleagues [7] investigated Bangladesh using this strategy. Bengali estimate is a popular academic topic. Due to resource constraints, Bengali data collection, corporate dialect research, and voice tagger vocabulary creation are difficult. Their main goals were to audit the restaurant and collect cricket perspectives using aspect-based research. The highest accuracy for recognizing and removing food-related pests is 77% for Support Vector Machines (SVM) and 71% for Classification Trees. In their investigations on emotion perception in Bangladesh, Tuhin et al. [8] proposed two methodologies. A spectrum of feelings was encountered, encompassing enthusiasm, rage, sorrow, dread, and sensitivity. The Naive Bayes algorithm employs segmentation and topical solutions. By employing a topical approach, the researchers attained a 90% accuracy rate with 7,400 sentences originating from Bangladesh. One paper achieved a

Support Vector Machine (SVM) score of 93%, while the other piece had a document frequency score of 83%. This paper presents three distinct inquiries into different aspects of the human emotional experience.

2.3 Scope the problem

This project focuses on developing and testing Bangla book review sentiment analysis tools. The scope includes:

- Discusses Bangla's language difficulties.
- Only book reviews, no news or social media.
- Examines Word2Vec and hybrid models like SVM and Random Forest.
- Punjabi context and morphological issues.
- Model evaluation emphasizes recall, precision, and accuracy.

2.4 Bangladesh Perspective

It is pertinent and could have an impact on Bangladesh by analyzing the sentiment of Bangla text, mainly book reviews. The internet presence and powerful literary culture of Bangla have led to the blossoming of Bangla text online, which includes social media posts, forum debates, and book reviews. While English has state-of-the-art natural language processing and sentiment analysis tools, Bangla has not. Bangladesh brings unique challenges and opportunities to this domain because of its storied literary heritage and burgeoning digital economy. Specialized sentiment analysis is necessary for Bangla due to its varied vocabulary and rich morphology. Publishers, writers, and readers can benefit from accurate emotion classification to better grasp literary trends and engage readers. Improving tools for Bangla sentiment analysis can also aid in cultural studies and marketing.

2.5 Challenge of this problem

The intricate syntax and morphology of the Bangla language are one of the primary hurdles to accurate sentiment extraction in this quest. The complex contextual clues and variable vocabulary of the Bangla language necessitate feature extraction and classification approaches. The lack of pre-existing Bangla-specific natural language processing (NLP) resources and annotated datasets further limits the accessibility of training data and the robustness of models. One approach to addressing this challenge is to develop techniques tailored to Bangla sentiment analysis, while another is to enhance the capacity of hybrid machine learning models to accurately classify emotions in a context with a wide range of languages.

CHAPTER 3

METHODOLOGY

3.1 Methodology

This chapter provides a comprehensive description of the development and evaluation of a sentiment analysis model for reviews written in Bangla. It describes all stages of the methodology, including data collection, preprocessing, feature extraction, model construction, and evaluation. Using cutting-edge machine learning methods, the goal is to create a sentiment categorization system that is both efficient and responsive to the nuances of the Bangla language.

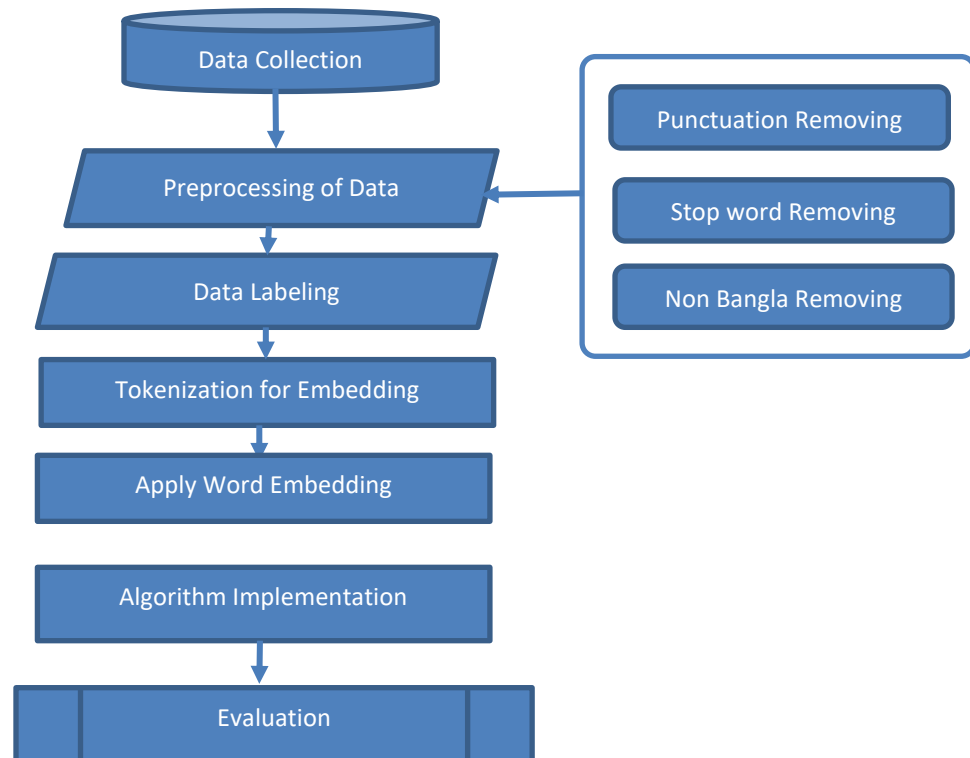


Figure 3.1 Methodology diagram

3.2 Data collection

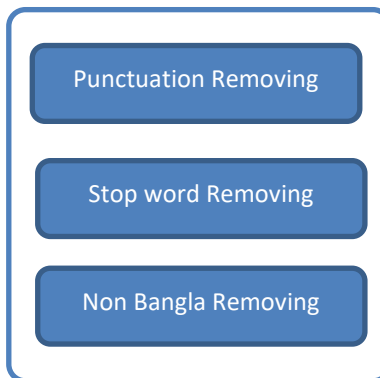
In order to guarantee diversity and representation, we gathered 2012 comments on Bangla books from various sources for this study. This enabled us to capture many viewpoints. Online literary forums, book review clubs, and book marketplaces offered evaluations of

literary works. The wide range of perceptions and sentiments expressed by Bangla literary reviewers on these platforms is remarkable. To guarantee the dataset is appropriate for training and testing sentiment analysis models, every review is meticulously categorized as either positive or negative based on its emotional content. By incorporating evaluations from several genres and authors, this extensive compilation enhances the analysis and provides a comprehensive understanding of the mood of Bangla literary readers.

3.3 Data Preprocessing

Data preparation is a method in data mining that involves the cleansing and organization of raw data before carrying out additional processing. The acquisition of knowledge requires the careful and thorough preparation of information. The function is implemented utilizing the KDD framework. Kamiran et al. [9] identify elimination, data massaging, weighting, and same poling as the four fundamental data pre-processing procedures. The principal approach we used to generate publicly accessible data sets was through data messaging systems. The removal of punctuation, stop words, and non-Bangla words is an essential component of our preprocessing routine. Furthermore, we have eliminated any superfluous words or characteristics from the Bangla stop. The decision has been made to incorporate our revised recommendations into the procedures to a higher degree. Figure 3.2 illustrates our methods for preprocessing.

Figure 3.2 Processing steps



Dataset Labelling

Positive and negative were the two distinct categories into which the data was given to me. The instructional design approach considers the user's emotional state. This component will be highly rated if the novel analysis is comprehensive and rigorous. There are various ways to classify negative reviews. The dataset's labeling is shown in Table 3.1. Negative data appears 1,133 times and positive data 1712 times in this table, according to a thorough examination. About 4,300 instances make up the training dataset.

Tokenization

Tokenization, as described by Pinto et al. [10], is a technique used to decompose flag phrases into their constituent words or signals. Our database encompasses a vast array of terms. Our team has determined that using a phrase mark instead of a standard word label will be more effective in accomplishing our objective. Furthermore, tokenization is essential. Tokenization is the process of breaking down a sentence into its constituent words. Refer to Table 1 for the tokenization procedure.

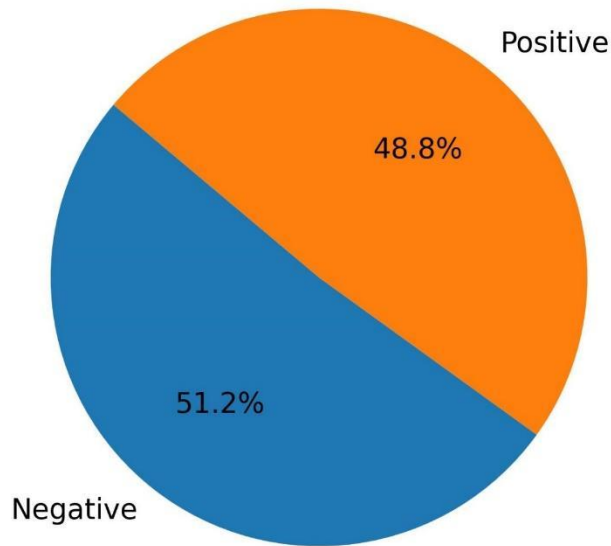
Table 1 Tokenization Table

Raw Data	Type	Tokenized data
বইটি আসলেই চমৎকার	Positive	'বইটি' 'আসলেই' 'চমৎকার'
ফাউল একটা বই।	Negative	'ফাউল' 'একটা' 'বই'
ভালো বই ছোট বড় সবার পড়া উচিত	Positive	'ভালো' 'বই' 'ছোট' 'বড়' 'সব' 'পড়া' 'উচিত'
লেখক কি লিখেছে নিজেই বোঝেনাই	Negative	'লেখক' 'কি' 'লিখেছে' 'নিজেই' 'বোঝেনাই'
কোথায় হুমায়ন আহমদ আর কোথায় এই লেখক	Negative	'কোথায়' 'হুমায়ন' 'আহমদ' 'ও' ' 'আর' 'কোথায়' 'এই' 'লেখক'

3.4 Dataset classification

After all, the Preprocessing steps are completed. I classify my data set. In this project data set, it is seen that most of the equal data present this dataset. It is also essential for a balanced dataset to work well in a project. This data set shows in figure 3.2 is 48.8% positive reviews and 51.2% negative reviews. It is clearly said that this data set is purely balanced and suitable for this project.

Figure 3.2 Dataset representation



3.5 Algorithms Implementation

Implementing algorithms for sentiment analysis involves several sequential steps. A Word2Vec model is trained on the Bangla book reviews dataset to provide word embeddings for feature extraction. Vectors generated by text embeddings achieve the numerical representation of word semantics. The implementation of Support Vector Machine (SVM) and Random Forest (RF) models is presented here. Random Forest (RF) trains its ensemble of decision trees using feature vectors, while Support Vector Machine

(SVM) trains its hyperplane-based categorization. Furthermore, investigated is the feasibility of a hybrid model combining Random Forest (RF) and Support Vector Machines (SVM). The performance and generalizability of each model are improved by systematic training and rigorous cross-validation. [11] The usefulness of a model for categorizing the sentiment of Bangla book reviews is assessed by evaluating its accuracy, precision, and recall properties. Integrating sci-kit-learn and Gensim in the implementation enhances the efficiency of data processing and analysis.

CHAPTER 4

Result Analysis

4.1 Overview

The accuracy table of the algorithms employed in combination is shown in Table 2. I utilized the matrices in this table to assess the level of satisfaction with my outcomes. The results indicate that hybrid models significantly enhance the accuracy of sentiment analysis compared to traditional approaches relying on a single algorithm. The deep learning, neural network, and support vector machine models have superior recall and precision metrics, indicating their ability to accurately detect and categorize emotions in book reviews with fewer false positives and negatives. [12] The initial integration of DT, RF, and SVM occurred at this point. Subsequently, SVM and RF have collaborated. Furthermore, I incorporated the three algorithms: Decision Tree, Neural Network, and Support Vector Machine. By amalgamating my three methodologies, we successfully attained almost indistinguishable degrees of test precision. Delta, Random Forest, and Support Vector Machines yield a recall rate of 0.92% and a test accuracy of 0.92%. Random Forest (RF) and Support Vector Machine (SVM) obtained test accuracy and recall of 0.91% and 0.93%, respectively. Based on a precision value of 0.90 and an accuracy of 0.86% in the test, the models DT, NN, and SVM demonstrate the highest level of accuracy. Book review enterprises in Bangladesh can utilize these findings. Organizations can improve their decision-making process and gain a more comprehensive understanding of customer sentiment by implementing the RF and SVM hybrid models. Potentially, this might result in improved book selections, enhanced customer service, and increased customer satisfaction and loyalty. [13] A statistical study of the table reveals that the integration of RF and SVM achieved the optimal test outcome and three algorithms. The findings indicate that the hybrid method is the most optimal for sentiment categorization in terms of efficiency, robustness, and accuracy.

Table 2 Accuracy table

Algorithm name	Test Accuracy	Precision	Recall	F1-score
RF, and SVM	0.92	0.92	0.93	0.92
DT, RF, and SVM	0.88	0.86	0.86	0.87
DT, NN, and SVM	0.91	0.91	0.90	0.90

4.2

Classification report

RF and SVM

The hybrid model incorporating SVM and RF achieves an outstanding rate of 0.93 in sentiment classification, as shown in Table 3. Class 0, which stands for negative mood, has an accuracy of 0.92, indicating that the model does a decent job of identifying negative situations. A Class 0 recall rate of 0.94 suggests that the model disregards certain adverse events. Class 1 has a very optimistic outlook, with a recall of 0.91 and a precision of 0.94. An outstanding example of the model's recall rate is its perfect recognition of all optimistic scenarios. Due to its inherent inaccuracy, the model's projections can produce outlandish numbers. Class 0 and class 1 F1-scores of 0.93 and 0.92, respectively, show that the model is very accurate. The class appears to have little effect on performance, as indicated by a weighted average F1-score of 0.93. Surprisingly, the hybrid model that combines RF and SVM does very well when recognizing pleasant emotions. Using this technology for sentiment analysis can significantly assist online shops. This hybrid algorithm gained very healthy values for all combined algorithms.

Table 3 RF and SVM classification Table

	Precisio n	Recall	F1-score	Support
0	0.92	0.94	0.93	310
1	0.94	0.91	0.92	290
Accuracy			0.93	600
Macro avg	0.93	0.93	0.93	600
Weighte d avg	0.93	0.93	0.93	600

DT, NN, and SVM

Table 4 displays the remarkable sentiment classification rate of 0.91 achieved by the hybrid model that combines DT, NN, and SVM. A class 0 accuracy of 0.91 indicates that the model adequately identifies negative situations, where 0 represents negative mood. An incorrect class 0 recall rate of 0.93 suggests that the model ignores specific negative occurrences. Class 1 looks forward to a bright future thanks to their excellent memory and precision. Accurately identifying every favorable scenario is a prime illustration of the model's recall rate. The model's estimates can yield absurd figures because of its intrinsic unreliability. The model's high accuracy is demonstrated by its Class 0 F1 score of 0.92 and Class 1 F1 score of 0.92. Given a weighted average F1-score of 0.91, the class doesn't do much to improve performance. Remarkably, the mixed model incorporating DT, NN, and SVM performs admirably when identifying positive emotions. Online stores can significantly benefit from this technology's sentiment analysis capabilities. For all combined algorithms, this hybrid algorithm achieved excellent values.

Table 4 DT, NN, and SVM classification Table

	Precision	Recall	F1-score	Support
0	0.91	0.93	0.92	310
1	0.92	0.90	0.91	290
Accuracy			0.91	600
Macro avg	0.91	0.91	0.91	600
Weighted avg	0.91	0.91	0.91	600

DT, RF, and SVM

Table 5 shows that the hybrid model incorporating DT, RF, and SVM produced an impressive sentiment classification rate of 0.88. The model correctly recognizes adverse conditions (class 0 accuracy = 0.87), where 0 denotes a negative mood. With an inaccurate class 0 recall rate of 0.82, the model disregards certain unpleasant events. Class 1 is confident in their future because of how well they remember and how precise they are. An excellent example of the recall rate of the model is its ability to identify all favorable scenarios correctly. Due to its inherent inaccuracy, the model's estimations can produce ludicrous values. Class 1 F1 score of 0.87 and Class 0 F1 score of 0.89 show that the model is very accurate. The class doesn't affect performance much, with a weighted average F1 score of 0.88. To everyone's surprise, the hybrid model that uses DT, RF, and SVM to identify positive emotions does a fantastic job. The ability to analyze customer opinion is a massive boon to online retailers. This hybrid algorithm outperformed all of the combined methods.

Table 5 DT, RF, and SVM classification Table

	Precision	Recall	F1-score	Support
0	0.87	0.90	0.89	310
1	0.89	0.86	0.87	290
Accuracy			0.88	600
Macro avg	0.88	0.88	0.88	600
Weighted avg	0.88	0.88	0.88	600

4.3 Decision

Among the three hybrid models, RF and SVM Showed the best test results, which were 0.93%. They also scored the best in the classification report, showing 0.89% accuracy, which was a remarkable score. This algorithm is also nicely suitable for my project dataset. Since they performed best, I selected this hybrid algorithm for future implementation.

4.4 Evaluation

This part is an essential part of a project. It can be said that It is also called the real-life implementation of a project. In this section, I applied my selected RF and SVM hybrid algorithm to evaluate my project with accurate data. [14]

Table 6 shows the real-life data detection table of my selected algorithm. Hybrid RF and SVM show excellent value as they detect review.[15] For label 0, it gained 0.84% precision, and for label 1, it also gained 1.00% precision. This algorithm performed 0.94% accuracy when the test data was 103.

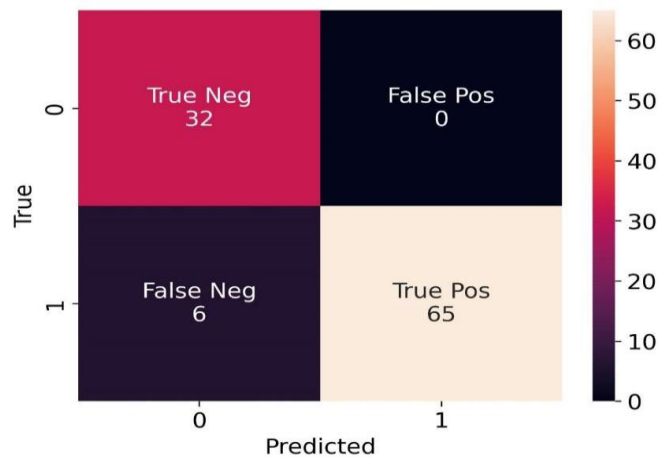
Table 6 Evaluation table of Rf and SVM

	Precision	Recall	F1-score	Support
0	0.84	1.00	0.91	32
1	1.00	0.92	0.96	71
Accuracy			0.94	103
Macro avg	0.94	0.94	0.94	103
Weighted avg	0.94	0.94	0.94	103

Confusion matrix

Viewing the confusion matrix, which shows the outcomes in terms of TP, TN, FN, and FP, is one way to assess the performance of a classification model. Here is how to understand the confusion matrix: for this project in figure 4.1.

figure 4.1 Shows the confusion matrix



accuracy calculation

$$= \frac{TP + TN}{TP + TN + FP + FN} = \frac{26 + 21}{26 + 21 + 5 + 4} = \frac{47}{56} \approx 0.839 (83.9\%)$$

Precision calculation

$$= \frac{TP}{TP + FP} = \frac{26}{26 + 5} = \frac{26}{31} \approx 0.839 (83.9\%)$$

Recall calculation

$$= \frac{TP}{TP + FN} = \frac{26}{26 + 4} = \frac{26}{30} \approx 0.867 (86.7\%)$$

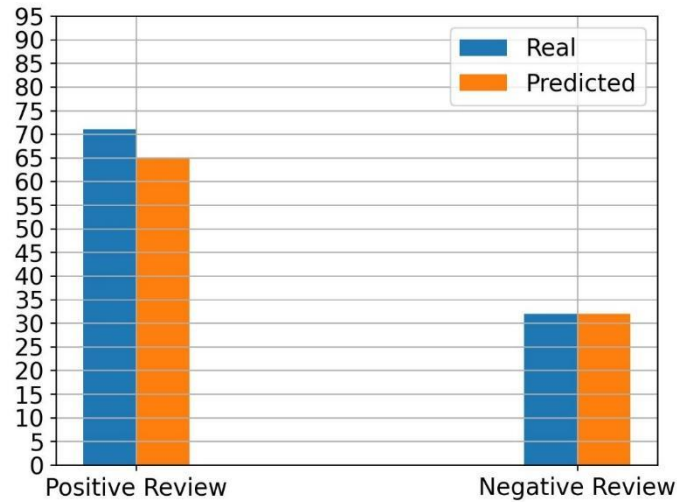
F1score calculation

$$= 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} = 2 \times \frac{0.839 \times 0.867}{0.839 + 0.867} \approx 0.853$$

This confusion matrix illustrates the model's equilibrium between recall and precision. However, there exists potential for improvement, especially in reducing the incidence of false positives and false negatives.

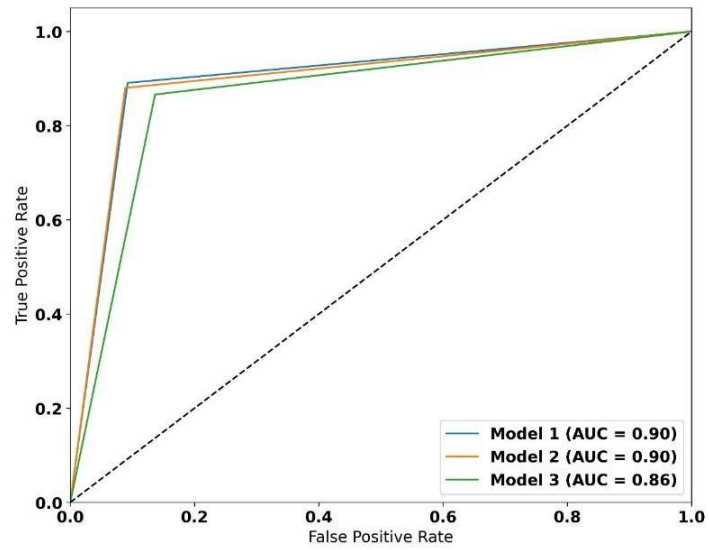
The evaluation graph for the RF and SVM algorithms is illustrated in Figure 3. The blue bar in this graph signifies the data that had not been previously trained before to the commencement of my project, while the orange bar denotes the projected data. [16] This RF and SVM provide outstanding performance in detecting. Despite the presence of thirty positive data points, my model failed to find three fewer data points. [17] The total count of hazardous data was 31; however, the methodology employed in my research failed to uncover five less data points. This signifies that the methodology I employed for training trains operates with considerable efficacy.

Figure 4.3 Evaluation Graph



Each line of the ROC curve shows in figure 4.4, three models' performance. At varying thresholds, ROC curves compare True Positive Rate (TPR) and False Positive Rate (FPR). Sensitivity, or True Positive Rate (TPR), is the percentage of genuine positives the model correctly identifies, while FPR is the percentage of true negatives misclassified as positives. Region Model efficacy is measured by the area beneath the curve. The model's classification accuracy is assessed. AUCs of 0.90 indicate excellent performance for Models 1 and 2, while Model 3 performs well but poorly with 0.86. An ideal classification model has an AUC of 1.0, indicating a TPR of 1.0 and no FPR. When AUC approaches 1, the model better separates positive and negative classes. Model 3 performs poorly on all criteria compared to Models 1 and 2.

Figure 4.4 ROC Curve



CHAPTER 5

Conclusion

This work constructed and assessed a predictive model for sentiment analysis of Bangla book reviews using cutting-edge machine-learning techniques. To evaluate Bangla's intricate structure and extensive lexicon, the study employed Word2Vec to extract features and investigated various hybrid models, including Random Forest (RF) and Support Vector Machine (SVM). Integrating Random Forest (RF) and Support Vector Machines (SVM) into a unified model enhanced sentiment categorization, particularly in terms of accuracy, precision, and recall. This study presents a sentiment analysis framework for many domains of Bangla text, thereby improving the processing of the Bangla natural language. By emphasizing specific methods for addressing linguistic complexity, this work establishes a foundation for applying sentiment analysis to languages that are not well represented.

One potential avenue for expanding this research is to develop a web application that employs the sentiment analysis model to examine Bangla book reviews interactively. Another advantage of an artificial intelligence-driven app is the enhanced availability of real-time sentiment analysis and insights obtained from mobile devices. Implementing these applications would improve the accessibility of Bangla literature, therefore benefiting publishers, writers, and readers alike.

CHAPTER 6

Future Work

Future research in this area should focus on expanding the existing sentiment analysis model's practical applicability. Creating a web app that uses the sentiment analysis method is a big step in the right direction. Those interested in reading and analyzing internet evaluations of Bangla novels will have a place to do it with this app. With features like sentiment trend tracking, user-friendly dashboards, and real-time sentiment research, this online tool has something for everyone: publishers, writers, and users. On top of that, creating an AI-powered mobile app can make sentiment analysis tools far more accessible and easier to use. With this software, users may engage with Bangla literature and gain insights from sentiment analysis without ever having to leave their mobile devices. Among the potential features is the capacity to constantly assess the tone of reviews, alert users to major patterns in tone, and provide personalized suggestions according to their preferences. This work's sentiment analysis approach, if implemented, would make online and mobile apps more accessible and engage users. This study's influence would extend to real-world applications in Bangla text analysis thanks to the aforementioned innovations, which would make sentiment analysis more accessible and usable.

Reference:

- [1]. M. N. M. Adnan, M. R. Chowdury, I. Taz, T. Ahmed and R. M. Rahman, "Content based news recommendation system based on fuzzy logic," 2014 International Conference on Informatics, Electronics & Vision (ICIEV), 2014, pp. 1-6
- [2]. C. Feng, M. Khan, A. U. Rahman and A. Ahmad, "News Recommendation Systems - Accomplishments, Challenges & Future Directions," in IEEE Access, vol. 8, pp. 16702-16725, 2020
- [3]. H. Tan, J. Guo and Y. Li, "E-learning Recommendation System," 2008 International Conference on Computer Science and Software Engineering, 2008, pp. 430-433
- [4]. Goyani, Mahesh; Chaurasiya, Neha. "A Review of Movie Recommendation System". ELCVIA: electronic letters on computer vision and image analysis, [online], 2020, Vol. 19, Num. 3, pp. 18-37, <https://raco.cat/index.php/ELCVIA/article/view/373942> [View: 3-02-2022]
- [5]. Alejandro Montes-García, Jose María Álvarez-Rodríguez, Jose Emilio Labra-Gayo, Marcos Martínez-Merino, Towards a journalist-based news recommendation system: The Wesomender approach, Expert Systems with Applications, Volume 40, Issue 17, 2013, Pages 6735-6741
- [6]. M. Rahman and E. Kumar Dey, "Datasets for Aspect-Based Sentiment Analysis in Bangla and Its Baseline Evaluation," Data, vol. 3, no. 2, p. 15, May 2018.
- [7]. N. Mittal, B. Agarwal, G. Chouhan, N. Bania, and P. Pareek, "Sentiment analysis of hindi reviews based on negation and discourse relation," in Proceedings of the 11th Workshop on Asian Language Resources, 2013, pp. 45-50.
- [8]. R. A. Tuhin, B. K. Paul, F. Nawrine, M. Akter and A. K. Das, "An Automated System of Sentiment Analysis from Bangla Text using Supervised Learning Techniques," 2019 IEEE 4th International Conference on Computer and Communication Systems (ICCCS), Singa-pore, 2019, pp. 360-364, doi: 10.1109/CCOMS.2019.8821658.
- [9]. R. R. Chowdhury, M. Shahadat Hossain, S. Hossain and K. Andersson, "Analyzing Sentiment of Movie Reviews in Bangla by Applying Machine Learning Techniques," 2019 International Conference on Bangla Speech and Language Processing (ICBSLP), 2019, pp. 1-6, doi: 10.1109/ICBSLP47725.2019.201483.
- [10]. Fang, X., Zhan, J. Sentiment analysis using product review data. Journal of Big Data 2, 5 (2015). <https://doi.org/10.1186/s40537-015-0015-2>.

- [11]. M. H. Alam, M. Rahoman and M. A. K. Azad, "Sentiment analysis for Bangla sentences using convolutional neural network," 2017 20th International Conference of Computer and Information Technology (ICCIT), Dhaka, Bangladesh, 2017, pp. 1-6, doi: 10.1109/ICCITECHN.2017.8281840
- [12]. C. Chauhan and S. Sehgal, "Sentiment analysis on product reviews," 2017 International Conference on Computing, Communication and Automation (ICCCA), Greater Noida, India, 2017, pp. 26-31, doi: 10.1109/CCAA.2017.8229825.
- [13]. R. A. Tuhin, B. K. Paul, F. Nawrine, M. Akter and A. K. Das, "An Automated System of Sentiment Analysis from Bangla Text using Supervised Learning Techniques," 2019 IEEE 4th International Conference on Computer and Communication Systems (ICCCS), Singapore, 2019, pp. 360-364, doi: 10.1109/CCOMS.2019.8821658.
- [14]. J. M. Keller, M. R. Gray, J. A. J. I. t. o. s. Givens, man,, and cybernetics, "A fuzzy k-nearest neighbor algorithm," no. 4, pp. 580-585, 1985.
- [15]. S. R. Safavian, D. J. I. t. o. s. Landgrebe, man,, and cybernetics, "A survey of decision tree classifier methodology," vol. 21, no. 3, pp. 660-674, 1991.
- [16]. Logistic Regression available at <<https://www.javatpoint.com/logistic-regression-in-machine-learning>> last accessed on 4-08-2021 at 11AM.
- [17]. <https://github.com/eftekhar-hossain/Bengali-Book-Reviews/tree/master>

PLAGIARISM REPORT

Sentiment Polarity Detection on Bengali Book Reviews By Word2vec and Hybrid Machine Algorithm

ORIGINALITY REPORT

10%	10%	6%	4%
SIMILARITY INDEX	INTERNET SOURCES	PUBLICATIONS	STUDENT PAPERS

PRIMARY SOURCES

1	dspace.daffodilvarsity.edu.bd:8080 Internet Source	8%
2	Md. Hamidur Rahman, Md. Saiful Islam, Md. Mine Uddin Jewel, Md. Mehedi Hasan, Ms. Subhenur Latif. "Classification of Book Review Sentiment in Bangla Language Using NLP, Machine Learning and LSTM", 2021 12th International Conference on Computing Communication and Networking Technologies (ICCCNT), 2021 Publication	1%
3	Rely Das, Md Forhad Hossain, Taufiq Ahmed, Ananyana Devanath, Shahnaz Akter, Abdus Sattar. "Classification of Product Review Sentiment by NLP and Machine Learning", 2022 Second International Conference on Advances in Electrical, Computing, Communication and Sustainable Technologies (ICAECT), 2022 Publication	1%
4	ir.juit.ac.in:8080 Internet Source	1%

Exclude quotes On Exclude matches < 1%
Exclude bibliography Off