

**BANGLA PRODUCT REVIEW SENTIMENT ANALYSIS BY ML AND DL
APPROACH**

**BY
SADMAN SADIK**

ID: 193-15-13419

AND

RASHEDUL ISLAM

ID: 193-15-13568

FINAL YEAR DESIGN PROJECT REPORT

This Report Presented in Partial Fulfillment of the Requirements for the
Degree of Bachelor of Science in Computer Science and Engineering

Supervised By

Ms. Naznin Sultana
Associate Professor
Department of CSE
Daffodil International University

Co-Supervised By

Ms. Sharmin Akter
Lecturer
Department of CSE
Daffodil International University



DAFFODIL INTERNATIONAL UNIVERSITY

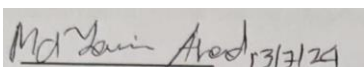
DHAKA, BANGLADESH

July 2024

APPROVAL

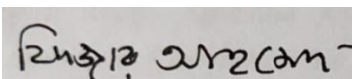
This Project titled **BANGLA PRODUCT REVIEW SENTIMENT ANALYSIS BY ML AND DL APPROACH**, submitted by **SADMAN SADIK** and **RASHEDUL ISLAM** to the Department of Computer Science and Engineering, Daffodil International University has been accepted as satisfactory for the partial fulfillment of the requirements for the degree of B.Sc. in Computer Science and Engineering and approved as to its style and contents. The presentation has been held on **13 July 2024**.

BOARD OF EXAMINERS



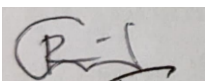
Dr. Md. Taimur Ahad (MTA)
Associate Professor & Associate Head
Department of CSE
Faculty of Science & Information Technology
Daffodil International University

Board Chairman



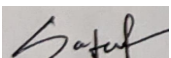
Dr. Fizar Ahmed (FZA)
Associate Professor
Department of CSE
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner 1



Mr. Rahmatul Kabir Rasel Sarkar (RKR)
Lecturer
Department of CSE
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner 2



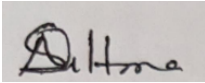
Sadat Hasan
Data Scientist (Senior Principal Officer)
Risk Management Division
BRAC Bank

External Examiner

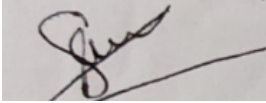
DECLARATION

We hereby declare that this project has been done by us under the supervision of **Ms. Naznin Sultana, Associate Professor, Department of Computer Science and Engineering, Daffodil International University**. We also declare that neither this project nor any part of this project has been submitted elsewhere for the award of any degree or diploma.

Supervised by:

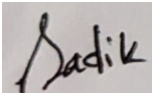


Ms. Naznin Sultana
Associate Professor
Department of CSE
Daffodil International University
Co-Supervised by:

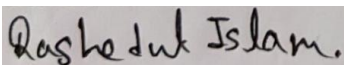


Ms. Sharmin Akter
Lecturer
Department of CSE
Daffodil International University

Submitted by:



SADMAN SADIK
ID: 193-15-13419
Department of CSE
Daffodil International University



RASHEDUL ISLAM
ID: 193-15-13568
Department of CSE
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.

We are grateful and wish our profound indebtedness to **Ms. Naznin Sultana, Associate Professor**, Department of CSE 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.

We would like to express our heartiest gratitude to the **Head of the Department of CSE**, for his kind help in finishing our project and also to other faculty members and the staff of the Department of CSE, Daffodil International University.

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

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

ABSTRACT

In this age of internet technology, e-commerce and online marketing companies in Bangladesh were already doing well. As a result of the widespread quarantines caused by the COVID-19 pandemic, online shopping has surpassed all other methods as the preferred method of making purchases. As a result, Businesses were able to get online much more quickly. While an increase in online product service providers does many good things, it also brings up questions about the reliability and quality of these offerings. Consequently, new customers are easy prey for internet scammers. Our goal is to develop an AI system that can automatically analyze customer reviews of e-commerce products and provide a set of positive and negative comments made by previous customers in the Bangla language using Natural Language Processing (NLP) and Artificial Intelligence algorithms. From different e-commerce websites, we collected 1016 Bangla product review comments for analysis. After completing preprocessing, we used two different domains of AI called Machine Learning(ML) and Deep Learning(DL) algorithms. There are five ML algorithms like KNN, Random Forest, Logistic Regression, SVM, and Gaussian Naïve Bayes used and for deep learning we used CNN and LSTM.

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 Problem Statement	1
1.3 Motivation and Objectives	2
1.4 Project Scope	3
1.5 Project Outcome	3
1.6 Report Organization	3
1.7 Summary	4
 CHAPTER 2: BACKGROUND STUDY	5-7
2.1 Overview	5
2.2 Related Work	5
2.3 Comparative Analysis and Summary	6
2.4 Gap Analysis	7
2.5 Summary	7
 CHAPTER 3: RESEARCH METHODOLOGY	8-13
3.1 Overview	8

3.2 Requirement analysis	8
3.3 Proposed Methodology	9
3.4 Data Collection	11
3.5 Project Management and Financial Analysis	12
3.6 Summary	13
CHAPTER 4: IMPLEMENTATION	14-23
4.1 Overview	14
4.2 Train Model	14
4.3 Model Evaluation	23
4.4 Summary	23
CHAPTER 5: RESULT AND ANALYSIS	24-31
5.1 Overview	24
5.2 Experimental Result	24
5.4 Comparative Analysis	29
5.5 Summary	31
CHAPTER 6: IMPACT ON SOCIETY, ENVIRONMENT AND SUSTAINABILITY	32-33
6.1 Impact on Life	32
6.2 Impact on Society & Environment	32
6.3 Ethical Aspects	33
6.4 Sustainability Plan	33
6.5 Summary	34
CHAPTER 7: CONCLUSION AND FUTURE WORK	34
7.1 Conclusions	34
7.2 Further Suggested works	34

7.3 Limitations	34
REFERENCES	14
APPENDIX	15
PLAGIARISM REPORT	18

LIST OF FIGURES

Figures	PAGE NO.
Figure 3.1: Proposed methodology	10
Figure 3.2: Data collection procedure	11
Figure 4.1: LSTM Model Architecture	21
Figure 4.2: CNN architecture	22
Figure 4.3: Confusion matrix for best model	23

LIST OF TABLES	
Tables	PAGE NO.
Table 1.1 Comparative analysis with previous work	6
Table 3.1 Data collection source	11
Table 3.2 Collected sample data	12
Table 3.3 Project management and finance analysis	12
Table 3.4 Estimated Cost for Machine Learning-based Project	13
Table 4.1: KNN's Best Hyperparameters	16
Table 4.2 Random Forest's Best Hyperparameters	17
Table 4.3 Logistic Regression's best hyperparameters	18
Table 4.4 SVM's Hyperparameters	19
Table 4.5 Gaussian Naïve Bayes's Hyperparameters	19
Table 5.1 Classification report of KNN algorithm	25
Table 5.2 Classification Report of Random Forest	25
Table 5.3 Classification Report of Logistic Regression	26
Table 5.4 Classification Report of SVM	27
Table 5.5 Classification report of Naïve Bayes	27
Table 5.6 Classification Report of LSTM	28
Table 5.7 Classification Report of CNN	29
Table 5.8 Comparative result analysis	30
Table 5.9 Comparative result analysis for balance dataset	31

CHAPTER 1

INTRODUCTION

1.1 Overview

Over the last decade, Bangladesh has made great strides in technology, reflecting its progressive nature. As high-speed internet became more accessible, many new businesses, now known as e-commerce spread online. Particularly traditional businesses were able to extend their operations over the whole nation as a result of the explosion of e-commerce. There may not have been a suitable online payment gateway in the late 1990s, but the e-commerce industry grew rapidly from the early 2000s through 2008. [1] E-commerce user numbers rose drastically with the advent of SSL COMMERZ's payment aggregator in 2010 with WiMAX, which boosted internet speed. In any e-commerce, there are so many ratings and comments, so it's hard for users to choose the best product. Reviews written in Bangla tend to be more genuine since people feel more at ease expressing themselves in their tongues. And this is the main strength point of our work. We are on the lookout for both positive and negative product evaluations written in the Bangla language. Public sentiment is an important factor in identifying shady sellers and faulty products. When there are several comments and evaluations for a product, it could be difficult for purchasers to pick the best one. We want to find out how people feel about different things by looking at reviews and comments written in Bangla. We take a different strategy on this. Whether the number of participants is positive or negative, why is it necessary to do this review analysis? It seems to reason that individuals' evaluations of the numbers could vary. By using our natural language processing it could be very easy to quantify positive and negative sentiment replies. To achieve the goal, CNN and LSTM algorithms are used for deep learning, and for machine learning five different algorithms like KNN, Random Forest, Logistic Regression, Support Vector Machine, and Gaussian Naïve Bayes are used in our work.

1.2 Problem Statement

Buying things is where most people in Bangladesh today spend their time. More and more people in Bangladesh are buying things online as e-commerce grows in popularity.

More and more people are making purchases via online stores. We conducted this study to save time and provide the highest quality of service. Machine learning and natural language processing were used. Some problems have arisen because of this endeavor. Due to the sensitive nature of our inquiry, data collection has proven to be challenging. We visited many online marketplaces to collect data. Every single product review was also mentally noted. We have good and negative feedback. When creating a feature, this is the data we utilize. We received 1016 comments in Bangla. This is the unprocessed data we have, and it's full of mistakes like emojis, additional punctuation marks, and multiple words. In order for our algorithm to learn correctly, we eliminated all of this noise during the preprocessing stage. We transformed the string to integers using TFIDF after preprocessing. To provide numerical representations for our categorization assignment, we used numerous categorization Machine Learning algorithms. All comments are positive or negative. Untrained data is used to evaluate our progress after training. Evaluation-wise, our technique trumps rivals. For deep learning algorithms, we used a tokenizer and word embedding approach.

1.3 Motivation and Objective

In Bangladesh nowadays eCommerce sites are increasing day by day Besides increasing eCommerce people have been interested in buying something online from previous decades. When a person visits eCommerce to buy something at first he or she monitors the review of the product and if the product review is better then he wants to buy it other hand if the preview is most of negative then he did not want to buy it. But it is a very time-consuming process. Customers lost lots of time to analyze product reviews. And this is the main part of our motivation. We want to work here to reduce customer time by introducing an effective Artificial intelligence approach. The objective of our work is given below:

- To use classification methods to analyze customer sentiment.
- To Develop an accurate model for identifying positive and negative product sentiment.
- To assist consumers in choosing the right product.

- To create a product review sentiment analysis program using AI engineering and machine learning.
- To reduce customer time by providing an appropriate percentage of reviews of each product.

1.4 Project Scope

The project scope includes the study's background, objectives, population of interest, and methodology. In Bangladesh there is no effective and accurate product review sentiment analysis product available so we want to build an assistant who can predict positive and negative reviews more accurately without any bias. For this purpose, we tried to collect 1016 rows of Bangla comments that contain negative and positive sentiments. Then we processed it and applied two different domains for experimental results.

1.5 Project Outcome

The expected outcome of our work is given below:

- At first, we want to build a system that is biased and free.
- The system can accurately detect positive and negative comments on any product.
- When the system can separate positive and negative reviews, we hope it can save customers time.
- We built a robust web API that shows the result of any review sentiment.

1.6 Report Organization

Each chapter of this paper describes itself as:

Chapter 1 The preliminary inquiry would not be complete without this first component. We also explain our rationale for doing this study in this chapter. The issue description is the most crucial part of this chapter. Here we cover the research problem and the difficulty of product review.

Chapter 2 An input analysis summarizing previous work in the field makes up this. It can be discovered more about the related machine learning work here.

Chapter 3 the methodology for the study will be detailed. A more in-depth explanation of the process is available [here](#). For more on the report's methodology, see the "Methodology" section.

1.7 Summary

The purpose of this chapter is to introduce the fundamental ideas that underpin our work. As well as the obstacle that stands in its way of our work, are also discussed along with potential solutions. In addition to that, we spoke about the ways in which this initiative might assist those working in the e-commerce business.

CHAPTER 2

BACKGROUND STUDY

2.1 Overview

The debate in the article is set in context by the background information on the study. Therefore, the background of the investigation is both fascinating and necessary for the research. Through background study, we can identify which analyses have been done in previously and what we need to do for new analysis. Discrepancies in academic performance between students of different socioeconomic backgrounds may be investigated in the context of a research project. Several specialists in the preceding field have made fruitful efforts, which are summarized in this chapter.

2.2 Related Works

Safin et al. found that the majority of Bangladeshi internet users speak Bangla. Labeling them verbally is hard. Classifying social media posts is difficult. Bangla writing is tougher to analyze. They categorize social media comments by tone to make searching, filtering, and organizing easier. They evaluated the posts' persuasiveness using Sentiment Analysis. The Bangla post-classification model utilizes Logistic Regression, Decision Tree, SVM, and K-Nearest Neighbour. Their Bangla fluency allowed them to apply an algorithm that properly detects social media content. The logistic regression method maybe 88% accurate. [2]

Reviews and ratings of books are useful for readers. They want to review the Bangla language assessments conducted by Hamidur et al.[3] and provide insightful feedback on books and online stores so that readers may find and enjoy more relevant content. The algorithm is trained using six 2,281 input data points. In order to determine if a sentiment was favourable or negative, this research used Machine Learning and Deep Learning. Natural Language Processing (NLP) texts are improved. Attitudes are classified using Decision Tree, KNN, Logistic Regression, Random Forest, SVM, and LSTM. We achieved an accuracy of 97.49% by using LSTM.

The way individuals feel about a subject is investigated using aspect-based sentiment analysis. Rahman et al. [4], who were researchers from Bangladesh, used this method. The field of Bengali sentiment analysis is seeing a surge in academic interest. Business language analysis, voice tagging, and vocabulary development are all made more difficult by the absence of resources like annotated data in Bengali. For this study, they employed aspect-based research and restaurant reviews to understand cricket fans' perspectives. When it comes to identifying restaurants based on their polarity or extracting polarity from insects, SVM is the most valid model (71% validity).

There is an 82.89% positive validity and a 76.59% negative validity for the Hindi analysis technique by Mittal et al. [5]. Emotions were evaluated, and the database was augmented to ensure consistency. This show takes a look at Roman Urdu feelings via the lens of software, cuisine, recipes, politics, and sports. From 566 online chats, 10,021 phrases were culled. Sentiment analysis using Roman Urdu corpora annotated by humans is the main emphasis of this research. Second, to compare sentiment analysis approaches, employ Rule-based, N-gram (RCNN) models.

2.3 Comparative Analysis and Summary

Table 1 represents a comparative analysis with previous work. From this table, we can see that some work has been done previously and the authors got the expected accuracy. Only one author uses RCNN approach and they got accuracy about 82.89%. but our CNN model got about 91% accuracy.

TABLE 1.1 COMPARATIVE ANALYSIS WITH PREVIOUS WORK

SL No	Author Names	Algorithm Used	Best Score with Algorithm Name
1	Mittal et al.	Rule-based, N-gram (RCNN)	82.89% RCNN
2	Rahman et al.	SVM	77% with SVM
3	Hamidur et al.	LSTM	97.49% with LSTM
4	Safin et al.	Logistic Regression	88% with Logistic Regression

2.4 Gap Analysis

The gap analysis is a way of proper planning that enables the researcher to focus their work more effectively and it also indicates weak points that need improvements. In previous work, we can see that most of the work algorithms are less than 90% only one algorithm achieved about 97% accuracy achieved by LSTM algorithm but there is no ML algorithm used there and only accuracy is measured in this work. But accuracy is not the only way of measurements, from the confusion matrix we can see that accuracy only worked with true positives and true negatives, but false positives and false negatives also be an important factor which rate is calculated by precision and recall score, this research did not address these points. RCNN was used in another work but the accuracy was 82.89% whereas our CNN model achieved 91% accuracy. From the above discussion, we can say that our work is very updated to all other's work because we used ML and DL in both domains and we applied all measurement metrics call precision, recall, and f1 score.

2.5 Summary

In this chapter, we discussed some previous work that has been done in a similar field. From his chapter, we pointed out the performance of previous work by comparison table. We also find out previous work gaps and where we need to upgrade in our work.

CHAPTER 3

METHODOLOGY

3.1 Overview

The methodology makes it simple for us to choose the best approach to tackling an issue. We tackled the challenge by dividing the process into five primary stages data collection and annotation, data preprocessing, dataset making, algorithm implementation, and evaluation. Most of the primary stages are divided into sub-stages. Our entire methodology stages are integrated as a result if we bypass any stage entire system will fail to generate appropriate results. We will go over each stage in detail in this chapter, and we will also include information on how we overcame the obstacles we encountered to solve each difficulty.

3.2 Requirements Analysis

The specifications and features outline the necessary conditions for the product review sentiment system. It must accurately analyze and interpret sentiment within reviews. In our work we divided our requirements into 3 different parts one is Data requirements, second is technical requirements, and finally Maintenance and support.

3.2.1 Data requirements

Data is the heart of the AI system as our entire system stands with historical data. Our required data is Bangla product review comments the comment can be anything but it must be in Bangla language.

3.2.2 Technical requirements

The technical requirements are divided into three different categories. One is software based the is hardware-based and programming requirements.

Software-based requirements:

1. Any eCommerce platform
2. Jupiter notebook ID

3. Microsoft Excel

Hardware-based requirements:

1. Processor core i5
2. Ram 8 GB minimum
3. 100 GB free space in C drive

Programing Requirements

1. Python 3.6 minimum for programming language
2. Library requirements is given below:

`joblib==1.1.0`

`matplotlib==3.3.4`

`numpy==1.19.5`

`openpyxl==3.0.9`

`pandas==1.1.5`

`pyxlsx==1.1.3`

`pymzmq==22.3.0`

`scikit-learn==0.24.2`

`seaborn==0.11.2`

`xlrd==2.0.1`

In library requirements, joblib 1.1.0 version is very important because in future we will use this library for the rest Api integration.

3.3 Proposed Methodology

Through the use of methodology, we are able to easily choose the most appropriate strategy for addressing a problem. Our strategy for overcoming the obstacle consisted of breaking the procedure down into five basic parts. In this chapter, we will discuss each step in great depth, and we will also include information on how we overcame the challenges that we faced in order to find solutions to each issue. Figure 3.1 represents the methodology diagram of our project. Some main part are divided into subparts like data processing contains two subsections one is non Bangla removing another is punctuation

removing. Same way dataset is split into two categories on his test data set another is train data set. and in our work, we apply to different domains of artificial intelligence called deep learning and machine learning. As the result of the experimentation we used different performance Matrix to find out which algorithm was based for our problem then this best problem we evaluated our work. After finding the best model we saved this model by joblib library and finally by using this model we can make web web-based rest api.

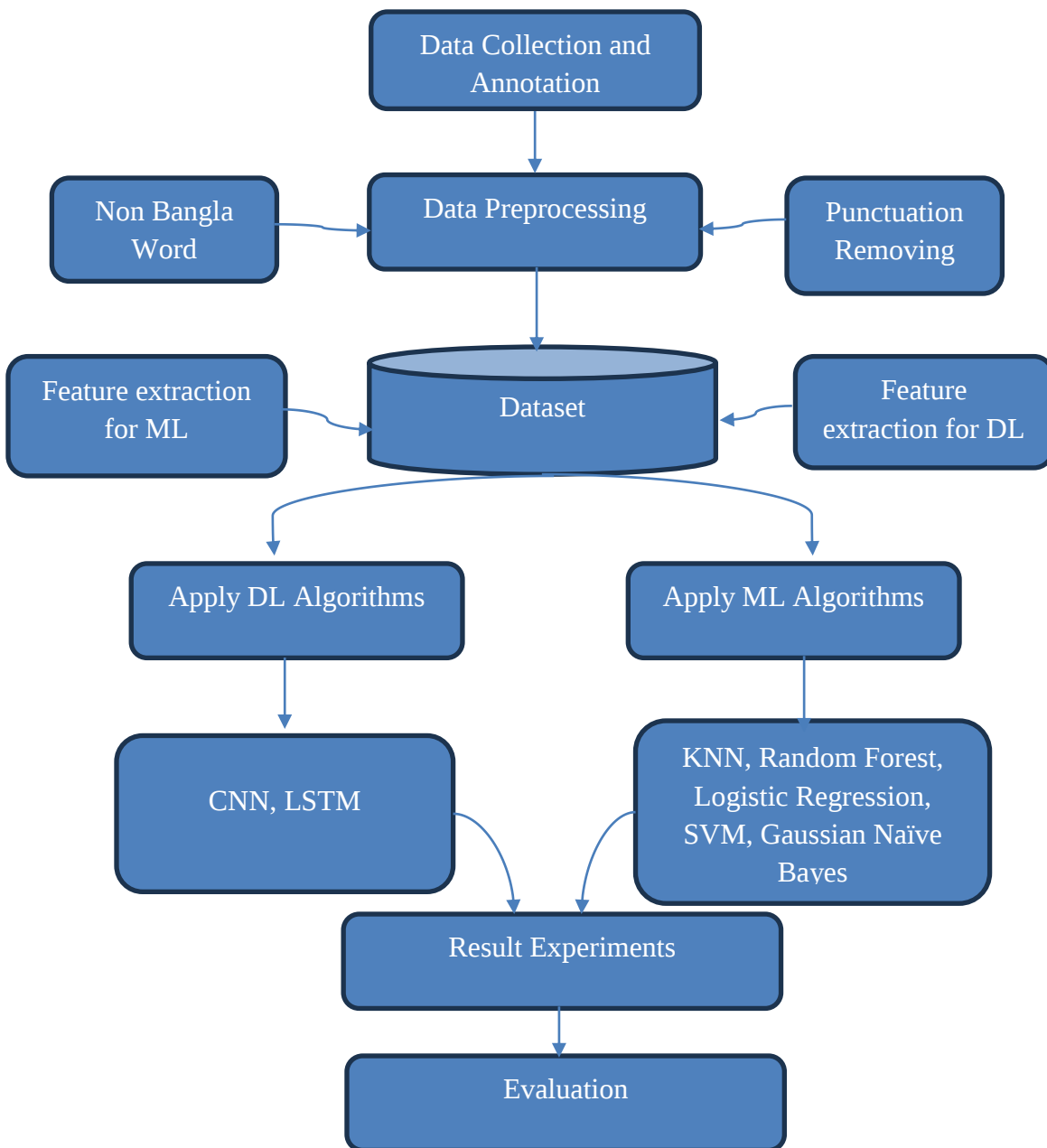


Figure 3.1 Proposed methodology

3.4 Data collection

As we work with product review sentiment analysis in the Bangla language our main target to collect data is different E-commerce websites like daraz and Chaldal.

TABLE 3.1 DATA COLLECTION SOURCE

E-commerce name	Collected data
Daraz	600
Chaldal	250
aarong	600

we collected comments on different products from this website after collecting the website we annotated the data set as positive reviews and negative reviews. Figure 3.2 represents the data collection procedure. Sometimes it is very difficult to collect comments directly. In such cases, we bypass by inspecting the element. This idea is given by the stakeholders.

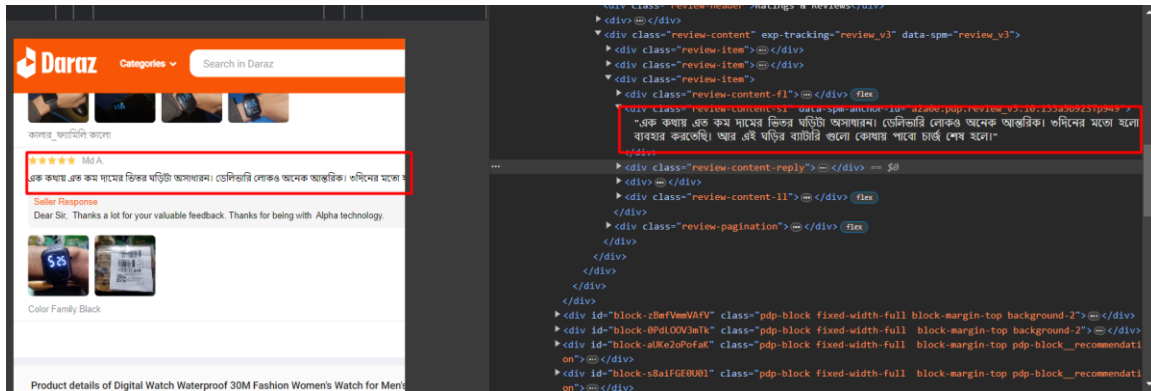


Figure 3.2: data collection procedure

The data annotation table is represented by Table 3.1. after data collection, we noticed that some of the comments contained foreign language we built a function for this purpose which can remove language without Bangla. After manually reading every comment set labels as positive and negative which is represented by 0 and 1 correspondingly while model train.

TABLE 3.2 COLLECTED SAMPLE DATA

Comment	Annotation
একদম হালকা আর নিম্নমানের	Negative
পাওয়ার বাটন প্রেস করলেই শুধু পোড়া গন্ধ বেরোতে থাকে ফালতু ;	Negative
প্যাকেজ টা একদম বাজে ছিল -	Negative
ডেলিভারিতে আমি অনেক খুশি প্যাকেজিং ভালো ছিলো	Positive
অরিজিনাল টা দেওয়ার জন্য ধন্যবাদ	Positive
চা অরিজিনাল ছিল কোন পিনিক না থাকলেও খেতে দারুন	Positive

3.5 Project Management and Finance Analysis

The management of projects is shown in table 3.2. According to the data shown in this table, the majority of the time we carried out work in parallel.

TABLE 3.3 PROJECT MANAGEMENT AND FINANCE ANALYSIS

	Oct 2023	Nov 2023	Dec 2023	Jan 2023	Feb 2023
Title defense					
Data collection					
Data annotation					
Report Writing					
Phase I					

In our work, we selected a minimum core i5 13th-gen processor. For motherboard, we chose an Asus exb660m motherboard and for ram, we chose 8 GB for HDD we chose 256gb team SSD. The project cost and finance analysis are given in Table 3.3.

TABLE 3.4: ESTIMATED COST FOR MACHINE LEARNING-BASED PROJECT

Sl. No	Item Quantity & Unit	Unit Price	Total Price
1	1. CPU: Intel 13th Gen Core i5 13400 Raptor Lake Processor.	25,000	44,349
	2. Motherboard: ASUS EX-B660M-V5 D4 M-ATX Motherboard.	14,000	
	3. RAM: OCP XT II DDR4 3200MHz 8GB Desktop RAM.	2,500	
	4. Storage: TEAM CX2 2.5" SATA 256GB SSD	2,849	
3	Visiting ML engineers		3000.00
4	Thesis/Project Typing, Printing and Binding		2000.00
5	Miscellaneous: (Maximum amount: Tk. 3,000)		3000.00
Total (In Figure):			52,349 BDT
In words: Fifty-two thousand three hundred forty-nine taka only			

3.6 Summary

Throughout this chapter, we spoke about the process and strategy that we take to our work. In addition to this, we go into great detail about each stage, and we will also provide information on how we overcame the obstacles that we encountered in order to discover answers to each problem.

CHAPTER 4

IMPLEMENTATION

4.1 Overview

In our work, we are going to construct an AI system for reading the Bangla customers' evaluations automatically, using NLP and AI algorithms. To meet this goal, we analyzed the 1016 comments in Bangla about the products that were sold on several online marketplaces. In the implementation phase, magnetic and deep learning algorithms, as well as data pretreatment, are used. These five machine learning techniques were utilized: KNN, Random Forest, Logistic Regression, Support Vector Machine (SVM), and Gaussian Naïve Bayes are examples of the techniques. Second, we made use of two DL models: CNN and LSTM are CNN and LSTM's twins at the party of technology. The main idea of this chapter is that in it, the whole implementation process from data preprocessing to prototype design to model evaluation, and a summary of the important findings are shown.

4.2 Train Model

At first, our system started with the careful gathering and cleaning of 1016 Bangla review comments from several e-commerce sites to make sure the data was of high quality. The preprocessing methods were the normalization, the tokenization, and the elimination of the irrelevant data (like the stop words), which made the data ready for the next steps. For the feature extraction, we used advanced methods like Term Frequency-Inverse Document Frequency (TF-IDF) and word embeddings which helped in the transformation of the textual information into a format that is machine learning algorithms can process. The basic structure of the system has an arrangement of five machine learning models—KNN, Random Forest, Logistic Regression, SVM, and Gaussian Naïve Bayes—and two deep learning approaches, namely Convolutional Neural Networks (CNN) and Long Short-Term Memory networks (LSTM). Every model was chosen depending on its success in the tasks of text classification and it was used with the current AI frameworks like TensorFlow and Scikit-learn.

4.2.1 Data Preprocessing

The process of data preprocessing is of great importance because it involves filtering of incoming data with the aim of generating data sets of higher quality at input into the feature extraction and model training. This step is divided into three main categories.

4.2.1.1 Removing Punctuation

Punctuation marks do impact such types of text analysis and those might be displaced by word frequencies. The main preprocessing step is to demutualize the text which eliminates all the punctuation symbols.

4.2.1.2 Eliminating Stop Words

The stop words are the quintessence of the marginal (like এবং, আছে, and থেকে in Bangla), the words which don't convey the particular tone rather than the distribution of the words if it is eliminated. It assists in rejecting the non-valuable words and selecting the more relevant ones such that machine learning may be able to extract better predictions.

4.2.1.3 Filtering Out Non-Bangla Sentences

Taking the assumption into account that the dataset might have noise – in the form of non-Bangla text, e.g. numbers or sentences of English words or numerical data – elimination of any non-Bangla text is of a crucial importance. This ensures that a model is sensitive to the linguistic features that are linguistically pinpointed in the Bengali language and not general language features during the training period which enhances the performance of the analysis.

The preprocessing steps were made into work in Python libraries like the NLTK for stop word removal and regular expressions for tokens removal and non-Bangla text.

4.2.2 Feature Extraction

As I worked with two different Domin(ML and DL) so I have to extract feature different way. That is why I used TFIDF for the machine learning algorithm and Word Embedding for Deep learning algorithm.[6]

4.2.2.1 TF-IDF

Via TF-IDF which is integrated in the machine learning model, each word in a document acquires an importance that is related to its frequency across the whole set of documents.

They allow for the word-significance distinction and the model algorithms like SVM and Random Forest to be used to make a correct prediction.[7]

4.2.2.2 Keres Word Embeddings

By deploying CNN, LSTM, and Keras such as deep learning models, the word embedding was also used. The models are tweaked so that words get a dense vector representation whereby semantics and relationships of the words are encoded, which in return, the AI language data is understood and processed more effectively.

Consequently, it becomes easier to organize the data according to the requirements of ML or DL, and the models will perform well in follow-up model runs.[8]

4.2.3 Machine Learning Model

In this section, I used 5 traditional machine-learning approaches. Each model design architecture and configuration are given below:

4.2.3.1 KNN algorithm

The k-Nearest-Neighbors (KNN) algorithm is a simple non-parametric and lazy learning technique which is used to solve classification and regression cases. In KNN classification is determined from the majority vote of its k nearest neighbors. The object is assigned to the class most common among the k closest to the neighborhood.[9]

Table 4. The selection of the highest suitable settings for the K-nearest point algorithms (KNN) applied to this step is called "KNN Best Hyperparameters." The first row shows that the Euclidean distance was taken as a metric to calculate the distance between events. Adjoining neighbors (k nearest neighbors, knn) now serve as the basis for classifying and computation of a value of a data point. The weights will be the probability vectors with values given by distance, and the neighbors with a more immediate impact on classification or prediction will be more important.

TABLE 4.1: KNN'S BEST HYPERPARAMETERS

Hyperparameters	Values
metric	euclidean
N neighbors	7
weights	distance

4.2.3.2 Random Forest Algorithm

A Random Forest algorithm could be utilized a mixture of ensemble learning techniques addressing both classification and regression problems. Initially, it has a few decision trees that become more entwined in order to avoid overfitting it but anyway they achieve high accuracy.[10] For every tree, the forest was partitioned into random data chunks, as well as taking a random subset among the features, Therefore, our model has to be adjusted in such a way that it better relates to the single decision trees and be reliable instead of erroneous.

Table 4. 2, called "Random Forest Best Hyperparameters" is the one which depicts the best settings for the Random Forest algorithm. The table illustrates that the maximum depth of the trees (Max depth) is set to None, meaning that the trees can continue to grow until all leaves are pure or contain less than the minimum number of samples that can be used in splitting. The minimum number of samples is set at exactly 1 for each leaf node which is represented by Min samples leaf. The Minimum to cut an internal node (Min samples split) is set to 2. Moreover, the number of trees in the forest (N estimators) is fixed to 200 as well. These hyperparameters are determined to improve the performance of the Random Forest model.

TABLE 4.2 RANDOM FOREST'S BEST HYPERPARAMETERS

Hyperparameters	Values
Max depth	None
Min samples leaf	1
Min samples split	2
N estimators	200

4.2.3.3 Logistic regression

The logistic regression is a statistical algorithm, which is for two classifications.[11] It builds up the probability that the given input belongs to a given class by means of a logistic function which converts any real number into a value between 0 and 1. The algorithm determines the influence of the dependent binary variable and one or more independent variables by finding the logistic curve fitting of the data points. It is most

often used for its ease of use and the different applications in the areas of machines and statistics.

Table 4. I will write two papers on the subject, namely "Best Hyperparameters for our Logistic Regression" and "Logistic Regression Best Practices to Optimize Model Performance," which explain the optimum settings for logistic regression. The corresponding parameter (Rho) is set at 10 to be less regularized. The penalty is L2, which does Ridge regularization by bringing the sum square of coefficients instead of the loss function. For this soup recipe, the solver is the 'liblinear' which would be appropriate for the smaller datasets and also supports both L1 and L2 regularization. They are tweaked in such a way and fine-tuned through each experiment until the predictive ability of the model is improved.

TABLE 4.3 LOGISTIC REGRESSION'S BEST HYPERPARAMETERS

Hyperparameters	Values
c	10
penalty	L2
solver	liblinear

4.2.3.4 Support Vector Machine

Support Vector Machines (SVM), is a modular supervised machine learning algorithm capable of both classification and regression. It does this by creating a plane of the optimum hyperplane which has enough space to separate the data points that belong to different classes in the beyond dimension of the space.[12] Support vector machines achieve this by turning off the closest points of different classes, which consequently are the kernel's relevant support vectors and, therefore we put up a strong classifier. This kernel is suited for decision function training as it deals with the data samples being linearly and nonlinearly through the kernel functions with abilities in addition to getting input transformed.

Table 4. Scanning item 4, "SVM's Hyperparameters", explores SVM's most efficient settings. The regularization parameter (c) will be set at 1 as its task is to be balanced such as a low training error and a big margin. Gamma is set to 'scale' value, that means gamma is automatically being adjusted based on the received data. The 'linear' nature of the

kernel used, implies linear relationship between the input features and the target variable. These hyperparameters, in particular, are the one that will be used to increase the Sensitivity, Specificity and Correlation Ratio of the SVM model.

TABLE 4.4 SVM'S HYPERPARAMETERS

Hyperparameters	Values
c	1
gamma	scale
Kernel	linear

4.2.3.5 Gaussian Naïve Bayes

Gaussian Naive Bayes (Gaussian NB) is a probabilistic classifier that uses Bayes' theorem and assumes that the observable variables (x) are conditionally normally distributed with zero mean and unit variance. It is assumed that the features are each located by the normal distribution, and there are no associations among them with a specific class label. Probably, the main advantage of Gaussian NB is the plain structure and strong assumption of complete independence of its features.[13] Undoubtedly, the independence assumption when not contradicted to the actual case is a good tool applied to many real-world applications. It is proven to be a solid tool for high-dimensional data as well as for learning and predicting phase what is performed quick and efficiently.

Table 4. 5. We also, discuss here the optimal values for contrary parameters for the Gaussian Naive Bayes Machine learning Algorithm. Create a table if the smoothing variance is given as 0. 1. The explanation of this variance is the factor that evokes this displacement style to make each feature variance and by a small value to prevent dividing to zero and stabilize these calculations, as a result, the algorithm variance does not become too small thereby improves its performance. The education reform is result-based, it gives students the ability to learn and become smart as a result of sustainability.

TABLE 4.5 GAUSSIAN NAÏVE BAYES'S HYPERPARAMETERS

Hyperparameters	Values
Var smoothing	0.1

4.2.4 Deep Learning Model

Deep learning of natural language processing (NLP) is the study that uses artificial neural networks to understand and generate human linguistics. It creates an intelligent interaction with applications such as machine translation, sentiment analysis and chatbots on the basis of learned complex text patterns in the data.

4.2.4.1 LSTM model

The following graph that represents a neural network built with the Keras sequential API to solve binary classification in Natural Language Processing (NLP).[14] It originates with an embedding layer that transforms integer encoded words into a dimensional vector of size 128 with 10,000 unique words in vocabulary and fixed input sequence length. Thereafter, an LSTM layer with 128 units is imposed, for which we used dropout and recurrent dropout rates equal to 0.2 as a further contribution which the authors take to improve model performance by adding L2 regularization to overcome overfitting. After that, a dropout layer is put up, for which 30% of its input units are initialized to zero randomly, and it is updated for every epoch to stop overfitting. There is one multiplying dense layer with 64 units and a ReLU activation followed by L2 regularization and a dropout layer of 0. Because of this, the model should be free to find its own balanced rate. The last is a dense layer with a node and sigmoid activation function having an L2 regularization to solve binary classification problem. The Adam optimizer is used for model compiling, and the model is trained on binary cross-entropy loss and validated on the accuracy metric.

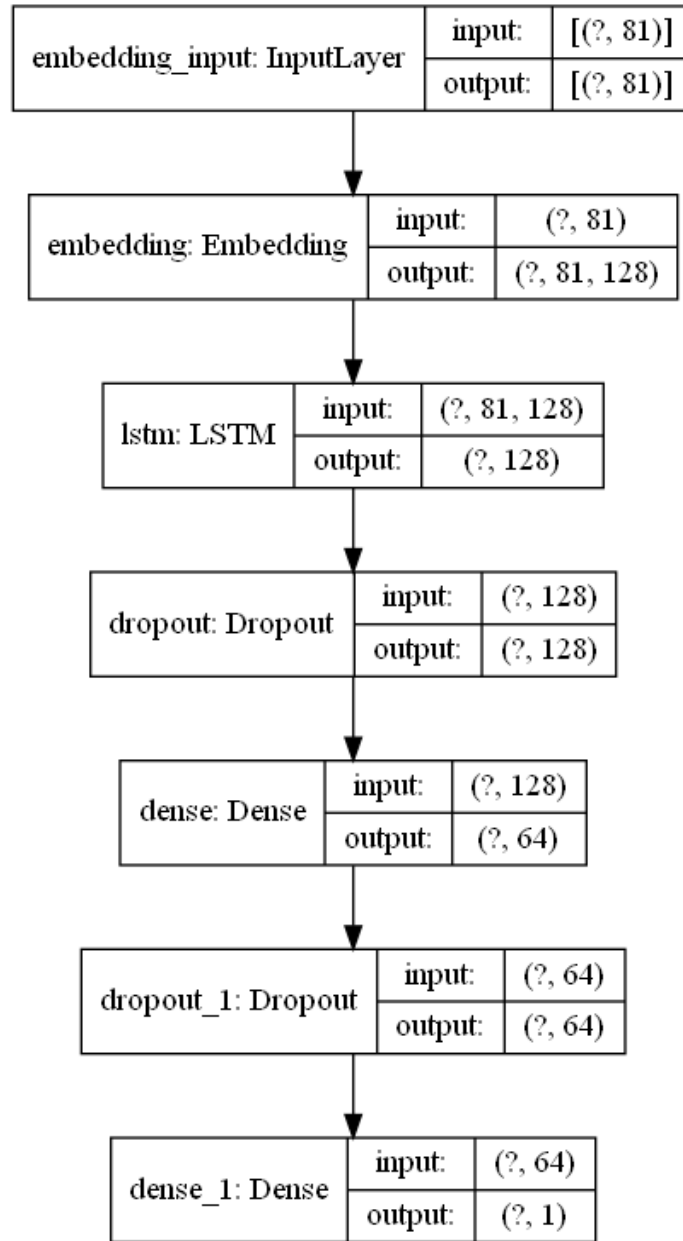


Figure 4.1: LSTM Model Architecture

4.2.4.2 CNN model

Figure 4.2 represents a CNN model for a binary classifier used in NLP and implemented using the Keras Sequential API. The setup begins with a 128-dimensional embedding layer that translates integer-encoded words into 128-size vectors which is sufficient for up to 10 thousand words in vocabulary.[15] Its structure differs from the original one in the addition of the dropout layer with the rate of 0.2 which serves the purpose of preventing overfitting, followed by two 1D convolutional layers with 32 and 64 filters,

respectively, with the kernel size of 5 which is associated with the ReLU as an activation parameter and same padding. The normalization batch layer is geared on the stabilizing and boosting learning, and there we use the global max pooling which lets us reduce feature maps to their topmost values. The architecture is built to help the network by adding in a dense layer of 128 nodes and L2 regularization, a dropout layer with a probability of 30% for their occurrence, a dense layer of 64 nodes and L2 regularization, and another dropout layer with a 20% probability. As the other layer has only one unit, it may use the sigmoid function for classification. The mentioned model compiles with the Adam optimizer -the speed of learning rate is 0.001, gradient descent algorithm (NAG), bi-cross-entropy loss function, and accuracy as the performance metric.

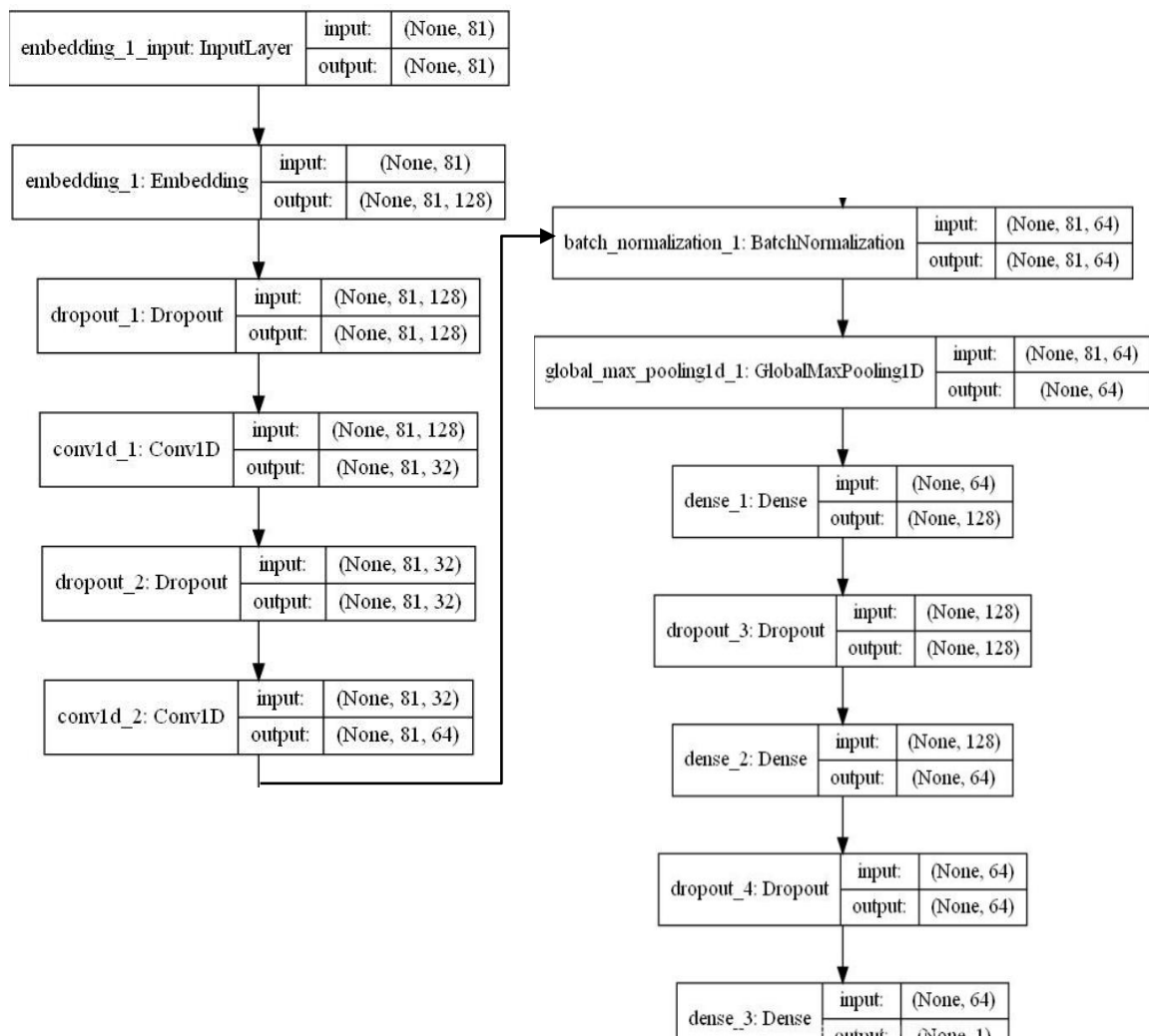


Figure 4.2: CNN architecture

4.3 Model Evaluation

Model evaluation process by which appropriate methods are chosen for estimation of the designed model and to measure its accuracy. The head-to-head would be the performance of metrics that include Accuracy, Precision, Recall, and F1- score. Figure 4.3 represents the Confusion matrix of the best model Logistic Regression. Calculation represents recall for positive is better than recall for negative. That means our model is best for positive comments.

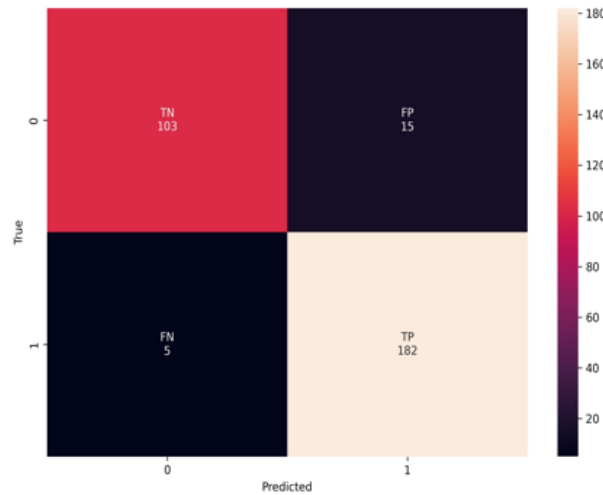


Figure 4.3: Confusion matrix for best model

$$\text{Accuracy} = \frac{TP+TN}{(TP + FP+FN+TN)} = \frac{182+103}{(182+15+5+103)} = 0.933 = 93.44\%$$

$$\text{Error} = 1 - 0.9344 = 0.0656$$

$$\begin{aligned}\text{The recall rate for positive: } TP / (TP + FN) &= 182/182+5 \\ &= 0.9733*100 = 97.33\%\end{aligned}$$

$$\begin{aligned}\text{The recall rate for Negative: } TN / (TN+FP) &= 103/103+15 \\ &= 0.8729*100 = 87.29 \%\end{aligned}$$

4.4 Summary

In this chapter, we gave a demonstration of the K-nearest Neighbors, Random Forest, Logistic Regression, SVM, and Gaussian Naive Bayes Algorithms that they could use for classifying data tasks. Two NLP techniques are used in CNN and LSTM is used, they includes embedding, convolution, pooling, and dense layers. They are later connected together with a linear classifier, dropout, and L2 regularization.

CHAPTER 5

RESULT AND ANALYSIS

5.1 Overview

In this chapter, the impact and match of this project are discussed. Among various machine learning and deep learning models on a dataset of Bangla product reviews, we analyze the performance of models. The main target is to determine which models lead to the best sentiment classification and which is the most precise. We have the apparatus section which has the data preprocessing and model training sections. Additionally, a study into the results of different algorithms, and an appraisal of their pros and cons, also is part of the process. These far-reaching adjustments also add sturdiness and facilitate the achievement of the desired outcomes.

5.2 Experimental Result

The experimental data is a set of obtained data and results that are obtained from carrying out benchmark or systematics tests. The "Experimental Results" section of the "Bangla Product Review Sentiment Analysis by ML and DL Approach" project is an accumulation of the metrics that are used to find out the actual accuracy of different machine learning and deep learning models on the test dataset, and those metrics are usually accuracy, precision, recall, F1-score. Our two-step experimental result is given below:

5.2.1 KNN algorithm

Table 5. 1 depicts the classifier report of the knn model with two classes. For class 0, precision is 0. 92, and recall is 0. Accuracy is 0.8, and F1-score is 0. 85 Voters with a support of 118 instances. Class 1 made an accuracy of 0. 88, recall of 0. 0.96 and F1-score of 0. 92 accompanied with 187 examples. The accuracy of the model is overall 0. 90, through the use of macro and weighted averages for precision, recall, and F1-score across all movements. 89-0.90, implying that the classes are balancing each other.

TABLE 5.1 CLASSIFICATION REPORT OF KNN ALGORITHM

	Precision	Recall	F1-score	Support
0	0.92	0.80	0.85	118
1	0.88	0.96	0.92	187
accuracy			0.90	305
macro avg	0.90	0.88	0.89	305
weighted avg	0.90	0.90	0.89	305

5.2.2 Random Forest

Examples of the measurements between the two classes are illustrated in Table 5. 2 of this classification report. Class 0 achieves an accuracy of 0%. 94, recall of 0. 84, F1-score of 0. The R2 value at 89 showed a high degree of predictability, while the support 118 points toward a considerable reliability that can be relied on. Class 1 displayed a recall as close to 0 as possible precision. 97 with a 0 precision. 0.91, which presents an F1-score of 1. We predicted the accuracy of the model by setting three parameters of mean=0 and standard deviation=0.5, and it yielded an identification result of 94, with a larger support of 187, which implied that it could complete this task with more precision. The general accuracy for the model is considered as 0. 92. The results are averaged with macro and weighted averages and all scores are either for precision, recall, and F1-score. 92. It means that the model's performance is not just a mimicry of stereotypes but can be applied to both specific groups of patients, which gives better results.

TABLE 5.2 CLASSIFICATION REPORT OF RANDOM FOREST

	Precision	Recall	F1-score	Support
0	0.94	0.84	0.89	118
1	0.91	0.97	0.94	187
accuracy			0.92	305
macro avg	0.92	0.90	0.91	305
weighted avg	0.92	0.92	0.92	305

5.2.3 Logistic Regression

This table 5.3 shows a classification report for a model having two classes of Logistic regression model. A Class 0 has a precision of 0. 95, recall of 0. 87, and F1-score 0. 90, sharing 118 occurrences, thus showing a remarkable skill to accurately classify positive cases. Class 1 accuracy is 0. 92 or 0, a recall of. 0.97 and the F1-score is 0. 95, 70% of them with a better support (187), emphasizing the sensitivity. The model will have an accuracy of 0 overall. 93. The variances for precision, recall, and F1-score of the macro and weighted methods are nearly identical at 0. 92 to 0. 94, which means steady and precise function in different classes.

TABLE 5.3 CLASSIFICATION REPORT OF LOGISTIC REGRESSION

	Precision	Recall	F1-score	Support
0	0.95	0.87	0.91	118
1	0.92	0.97	0.95	187
accuracy			0.93	305
macro avg	0.94	0.92	0.93	305
weighted avg	0.94	0.93	0.93	305

5.2.4 Support Vector Machine

Table 5 is an illustration of a classification report. 4, which summarizes the performance of a model for two different classes of SVM algorithms. The reporter said that the class 0 precision is 0. 93, the recall rate = 0. 89. 5, and the precision rate is 0. 92% of the 118 samples. This shows that it is able to impart actual good consequences in a high level of accuracy that it has. One can observe that the class equal to one has the maximum precision (0. 94), and the F1-score is 0. 95 shows the confidence interval in which we can state without doubt that the standard error will fall somewhere around 187 samples when the sample size is increased. Regarding the precision, the model is 0 94. The both unweighted and weighted averages and the values of precision, recall, and F1-score are similar to each other and almost always consistent with the last two methods as well. The grade "94" represents that the student has got good grades in all the subjects.

TABLE 5.4 CLASSIFICATION REPORT OF SVM

	Precision	Recall	F1-score	Support
0	0.93	0.91	0.92	118
1	0.94	0.96	0.95	187
accuracy			0.94	305
macro avg	0.94	0.93	0.93	305
weighted avg	0.94	0.94	0.94	305

5.2.5 Naïve Bayes

The accuracy metrics for the classification report are depicted in Table 5.5, as well, that can be eluded to the naïve bayes method that includes two classes. 118 cases are Class 0 that contain predictive, recall, and F1 rates which are all equal to 0. 91. Conversely to the most negative score of Class 2, all of the metrics for Class 1 are close to each other at the value of 0. 94 Referencing this statement with 187 proofs. The model obtained an accuracy of 93%. The values incline only nearly to 0. 92 and 0. thus, being 93,7 and 93,6 respectively, both the conventional and weighted averages show a good consistency all over the classes too. This happens to be the cause of the balanced performance in terms of precision, recall, and F1 score.

TABLE 5.5 CLASSIFICATION REPORT OF NAÏVE BAYES

	Precision	Recall	F1-score	Support
0	0.91	0.91	0.91	118
1	0.94	0.94	0.94	187
accuracy			0.93	305
macro avg	0.92	0.92	0.92	305
weighted avg	0.93	0.93	0.93	305

5.2.6 LSTM

Table 5.6 , reported are the metrics for the classification report of LSTM algorithms. The precisions, recalls and F1-scores of Non-Cancer class are 0. 82, and the accuracy level (65) of the film signifies that it is a movie with a somewhat predictable plot. Class 1 metrics the performance is higher, with precision, recall, and F1-score which are all at 0. 91. This was confirmed by 139 cases since Class 1 was more precise in prediction. The model is 100% accurate. And, at the individual level, this figure shoots up 88%. According to macro-average the values of precision, recall, and F1 are all zero. 86 and 0 respectively, the weighted average is higher at 0. 88. This discrepancy can be explained by the additional support of class 1 that is added to the formula.

TABLE 5.6 CLASSIFICATION REPORT OF LSTM

	Precision	Recall	F1-score	Support
0	0.82	0.82	0.82	65
1	0.91	0.91	0.91	139
accuracy			0.88	204
macro avg	0.86	0.86	0.86	204
weighted avg	0.88	0.88	0.88	204

5.2.7 CNN

Finally, a classifying report in Table 5 is outlined. 7, which is concerned with the presentation of the performance measures of the model which is applied to two distinct groups. The class 0 precision is equal to 0. 86 where the recall is 0 0.93 and the score of F1 mean 0. 89, the cases making up the 82 exemplary classifications of benign nature. In 122 cases Class 1 has shown its highest precision score which is 0. 95, the re-election of 0. Accuracy of 0.90 and an F1-score of 0.0. 92. Such validated information is given through the data The model in general has an accuracy of 0. 91. Both macro and weighted averages for precision, recall, and F1 score are always 0. The commonly known 91, indicating that the performance of all the classes is equally distributed.

TABLE 5.7 CLASSIFICATION REPORT OF CNN

	Precision	Recall	F1-score	Support
0	0.86	0.93	0.89	82
1	0.95	0.90	0.92	122
accuracy			0.91	204
macro avg	0.91	0.91	0.91	204
weighted avg	0.91	0.91	0.91	204

5.3 Comparative Analysis

Table 5.8 uses four metrics, namely test accuracy, precision, recall and F1-score, employing which values of different models can be compared. The model Logistic Regression has now a score of 93, so it is on the top position. 44% in terms of test accuracy, which is the portion of right decisions among all decisions. The “Precision” metric measure the percentage of true positives predicted out of all the positives the CNN prediction has, with a score of 94. 83%. Recollection presents capacity to serve all relevant occurrences, and Logistic Regression has the highest performance rate which is equivalent to 97%. 33%. Logistic Regression was the best exact score with during the event which was balancing the Precision and Recall with a score 94. The overall accuracy including TP Rate, FP Rate, and Accuracy attained by Logistic Regression is 79%, making it appear the most effective model on all the evaluation metrics. As a result, it is a good choice for jobs that require a high degree of extremely effective recognition of positive cases and overall prediction accuracy. Long story short, Logistic Regression is the most accurate.

TABLE 5.8 COMPARATIVE RESULT ANALYSIS

Model	Test Accuracy	Precision	Recall	F1-Score
KNN	0.8951	0.8818	0.9572	0.9179
Random Forest	0.9180	0.9050	0.9679	0.9354
Logistic Regression	0.9344	0.9239	0.9733	0.9479
SVM	0.9377	0.9421	0.9572	0.9496
Naive Bayes	0.9279	0.9412	0.9412	0.9412
CNN	0.9118	0.9483	0.9016	0.9244
LSTM	0.8824	0.9137	0.9137	0.9137

Table 5.8 represents the comparative result analysis for the balance dataset. As our dataset was slightly imbalance we decided to balance this and test the performance of our selected algorithm. SMOTE library is used for balancing datasets, and we used up sampling techniques. If we compare the balance and imbalance datasets, we can see that most of the algorithms drop their performance. For example, KNN produced 82.95% accuracy but earlier accuracy was about 89%. For random forest, accuracy was 88.52% but earlier it was 91%. We can see only Logistic Regression produced better accuracy than the imbalance dataset but accuracy was never exceed than SVM imbalance dataset model. If we look deep learning part we can also see that accuracy has dropped. So we decided to use Imbalance models for final results like classification reports, test performance, etc.

TABLE 5.9 COMPARATIVE RESULT ANALYSIS FOR BALANCE DATASET

Model	Test Accuracy	Precision	Recall	F1-Score
KNN	0. 8295	0. 9789	0. 7394	0. 8424
Random Forest	0. 8852	0. 8768	0. 9468	0. 9105
Logistic Regression	0.9377	0. 9471	0. 9521	0. 9496
SVM	0. 9180	0. 9179	0. 9521	0. 9347
Naive Bayes	0. 9148	0. 9765	0. 8830	0. 9274
CNN	0. 8775	0. 8671	0. 9538	0. 9084
LSTM	0. 8824	0. 9717	0. 8306	0. 8957

5.4 Summary

Once this part is done, they present the result of the experiment and the comparison of different ML and DL models that are used to classify the product reviews sentiment. Accordingly, the ML libraries that are used, their accuracy and the datasets that are used are concretely examined and analyzed. We discover that the deep learning algorithms achieves the higher accuracy and faster processing speed than others in our study. During this comparison, the complexity model of the existing solution versus performance was noted for the given problem. The analysis of outcome will be the basis for the improvement of sentiment analysis concept as a precondition of its transfer to real language applications in the context of Bangla.

CHAPTER 6

IMPACT ON SOCIETY, ENVIRONMENT, AND SUSTAINABILITY

6.1 Impact on Life

"Bangla Product Review Sentiment Analysis by ML and DL Approach" has a considerable implication on daily living via consumers' decision-making based on product reviews. The local customers will be confident in reviews if sentiments are properly identified in Bangla, which will, in effect, steer them to make better buying decisions. They help in the evaluation of products in a shorter time and without spending too much time and effort. Besides, businesses can deliver precise knowledge about the level of customer satisfaction and preferences which additionally helps to improve products and services. Thus, the project at the same time helps the non-English-speaking communities by the provision of sentiment analysis tools in their native language that help to improve their online shopping experience.

6.2 Impact on Society & Environment

Through sentiment analysis in Bangla product reviews, people will know what the best products to buy and will thereby build a level of trust and transparency in e-commerce sites. The main achievement of this project is it can save customers time by providing appropriate suggestions. Customers no need to read each and every comment of a product. It is prone to consumers' rights by allowing them to get feedback that is analyzed accurately and represented properly. Ecologically, the project is able to affect the process of consumers making an informed decision, which will decrease the amount of waste that is produced as a result of people not being satisfied with a purchased product. Thus, this provides avenues for sustainable consumption habits. Hence businesses will be able to optimize their supply chains by listening to customer needs which will result in less waste of resources.

6.3 Ethical Aspects

These ethical question will be position on this project such as privacy of data and consistence sentiment analysis algorithms. It is essential to encrypt communication with an extra layer of passwords between user data and information in order to avoid information leakage. It is vital to instill thorough job training on different types of datasets on the algorithms in order to prevent the inequality or discrimination effects that may be brought forth by opinion analysis results. The manners for analysis being implemented should be provided transparently and the result should be given in the non partial way for keeping the user belief. Besides, customers` opinions must be gathered, impartial and the negative feedback must be accounted for constructively in order to avert unfairly damage to business or personal reputation.

6.4 Sustainability Plan

The implementation of the project sustainability plan includes the constant revision and enhancement of the machine learning and deep learning models to ensure the highest accuracy and relevance. Interacting with the people in the user community to get feedback and understand emergent needs may help in maintaining the tool relevant and usable. E-commerce collaborations would provide consistent data for the model training and the model validation. On the other side, making this technology available in different languages and regions can possibly extend its impacts, making the world wide green goals to be supported by responsible consumption patterns.

6.5 Summary

The "Bangla Product Review Sentiment Analysis by ML and DL Approach" program is a huge milestone in the field of life, society, and the environment. It brings consumer thinking to a different level, creates accountability and builds trust as far as the online systems are concerned. Ethically It deals with the topic of data privacy and algorithmic fairness, which means that the algorithm is not biased and safe. The sustainability plan proposes constant enhancement, involving the community and expanding coverage. In the end, the project is good in informing consumption, and in turn promoting sustainable development objectives by diminishing waste and establishing resource efficiency.

CHAPTER 7

SUMMARY, CONCLUSION AND FUTURE WORK

7.1 Conclusion

In the era of the general buying of products online by the masses because of the COVID-19 pandemic, our AI-based system has shown a 94 percent accuracy in the analysis of the Bangla product review comments from different e-commerce websites. The system uses a mixture of Machine Learning and Deep Learning techniques to process the 1016 review comments, and, as a result, sorts them into the positive or the negative sentiment categories. Of the machine learning algorithms, the Support Vector Machine (SVM) was the most accurate one and it outperformed the other algorithms such as K-Nearest Neighbors (KNN), Logistic Regression, Random Forest as well as Deep Learning methods such as CNN and LSTM. This way of tackling the difficulties of consumer sentiment in Bangla implies its suitability for application in worldwide e-commerce, thus providing a scalable answer for the enhancement of customer experience and business operations.

7.2 Further Suggested Works

Some really outstanding recommendations exist for future work:

- Training data needs to be vast in size to generate test data with higher accuracy.
- LSTM, CNN, and Bangla Bert are examples of deep learning algorithms that can be employed while working with large datasets.
- For deployment, we can use either the Django Rest Framework or the Flask Framework.

7.3 Limitations

Here is a roadmap for where this effort is going from here:

- ❑ Our methodology only included good and negative feedback; it did not consider neutral or caustic remarks. Going forward, we will make neutral or caustic comments.
- ❑ The sentiment expressed in the summary section seems to be typical when sarcasm is present. In this particular analysis, though, the remark has the reverse impact as sarcasm is hard to predict on a computer. This leads us to our next project: developing an algorithm to detect statements containing sarcasm.
- ❑ To that end, we are developing an application programming interface (API) for the web that will define the review of analyses.
- ❑ In order to finish the work, we employed a machine learning and deep learning approach. Eventually, we will build an AI system with LLM based model in future
- ❑ All of our work is conducted in Bangla. Conversely, Banglish is the language used by customers to write comments. Consequently, we will soon be teaching Banglish Sentence to enable our system to understand and respond to comments in Banglish.

Reference:

- [1] "Literature of Bangladesh, culture," Bangladesh.com, 2020. [Online]. Available: <https://www.bangladesh.com/culture/literature/>. [Accessed: Oct-10-2021].
- [2] M. R. Alam, A. Akter, M. A. Shafin, M. M. Hasan, and A. Mahmud, "Social Media Content Categorization Using Supervised Based Machine Learning Methods and Natural Language Processing in Bangla Language," in Proc. 2020 11th International Conference on Electrical and Computer Engineering (ICECE), 2020, pp. 270-273.
- [3] M. H. Rahman, M. S. Islam, M. M. U. Jowel, M. M. Hasan, and M. S. Latif, "Classification of Book Review Sentiment in Bangla Language Using NLP, Machine Learning and LSTM," in Proc. 2021 12th International Conference on Computing Communication and Networking Technologies (ICCCNT), 2021, pp. 1-5.
- [4] M. Rahman and E. Kumar Dey, "Datasets for Aspect-Based Sentiment Analysis in Bangla and Its Baseline Evaluation," Data, vol. 3, no. 2, pp. 15, May 2018.
- [5] N. Mittal, B. Agarwal, G. Chouhan, N. Bania, and P. Pareek, "Sentiment analysis of Hindi reviews based on negation and discourse relation," in Proc. 11th Workshop on Asian Language Resources, 2013, pp. 45-50.
- [6] Y. Ma, K. Zhan, and Z. Wang, "Feature Extraction," in *Applications of Pulse-Coupled Neural Networks*, Berlin, Heidelberg: Springer Berlin Heidelberg, 2010, pp. 111-145. doi: 10.1007/978-3-642-13745-7_7.
- [7] W. Zhang, T. Yoshida, and X. Tang, "A comparative study of TF*IDF, LSI and multi-words for text classification," *Expert Systems with Applications*, vol. 38, no. 3, pp. 2758-2765, 2011. doi: 10.1016/j.eswa.2010.08.066.
- [8] B. Wang, A. Wang, F. Chen, Y. Wang, and C.-C. J. Kuo, "Evaluating word embedding models: Methods and experimental results," *APSIPA Transactions on Signal and Information Processing*, vol. 8, p. e19, 2019.
- [9] A. Moldagulova and R. B. Sulaiman, "Using KNN algorithm for classification of textual documents," in *2017 8th International Conference on Information Technology (ICIT)*, 2017, pp. 665-671.
- [10] S. J. Rigatti, "Random forest," *Journal of Insurance Medicine*, vol. 47, no. 1, pp. 31-39, 2017.
- [11] M. P. LaValley, "Logistic regression," *Circulation*, vol. 117, no. 18, pp. 2395-2399, 2008.
- [12] H. Wang and D. Hu, "Comparison of SVM and LS-SVM for regression," in *2005 International Conference on Neural Networks and Brain*, vol. 1, pp. 279-283, 2005.
- [13] A. Mootooyaloo, "Bayesian model selection with applications to radio astronomy," 2017.
- [14] Y. Luan and S. Lin, "Research on Text Classification Based on CNN and LSTM," 2019 IEEE International Conference on Artificial Intelligence and Computer Applications (ICAICA), Dalian, China, 2019, pp. 352-355,
- [15] S. Moriya and C. Shibata, "Transfer Learning Method for Very Deep CNN for Text Classification and Methods for its Evaluation," 2018 IEEE 42nd Annual Computer Software and Applications Conference (COMPSAC), Tokyo, Japan, 2018, pp. 153-158, doi: 10.1109/COMPSAC.2018.10220.

Appendix A

Complex Engineering Problems (EP) and Complex Engineering Activities (EA) Analysis

Title: BANGLA PRODUCT REVIEW SENTIMENT ANALYSIS BY ML AND DL APPROACH

Students ID: 193-15-13419, 193-15-13568

CO Description for FYDP-Phase-I

CO	CO Descriptions	PO
Phase -I		
CO1	Integrate recently gained and previously acquired knowledge to Product Review Sentiment Analysis for the Final Year Design Project (FYDP)	PO1
CO2	Analyze different aspects of the goals in designing a solution for this FYDP	PO2
CO3	Explore diverse problem domains through a literature review, delineate the issues, and establish this goals for the FYDP	PO4
CO4	Perform economic evaluation and cost estimation and employ suitable project management procedures throughout the development life cycle of the FYDP	PO11
Phase -II		
CO5	Design and develop technical solutions and system components or processes that meet specified requirements, ensuring compliance with public health and safety standards, as well as considering cultural, socioeconomic, and environmental factors in this FYDP	PO3
CO6	Choose and apply appropriate methodologies, resources, and contemporary engineering and IT technologies to address complex engineering processes, encompassing prediction and modeling, while adhering to relevant constraints in this FYDP	PO5
CO7	Analyze societal, health, safety, legal, and cultural considerations, along with associated responsibilities, in the context of professional engineering practice and the resolution of this problem, employing logical reasoning guided by contextual understanding.	PO6
CO8	Comprehend and evaluate the enduring sustainability and impact of professional engineering endeavors in addressing intricate engineering challenges within social and environmental frameworks.	PO7

CO9	Implement ethical principles and adhere to professional standards and norms in this FYDP	PO8
CO10	Capable of operating proficiently both individually and as a team member or leader across diverse teams and interdisciplinary settings in this FYDP.	PO9
CO11	Proficiently communicate with the engineering community and broader society regarding complex engineering endeavors, including the ability to comprehend and generate comprehensive reports and design documentation, as well as provide and receive clear instructions throughout this FYDP.	PO10
CO12	Acknowledge the importance of self-directed and life-long learning within the evolving landscape of technology, and possess the readiness and capability to engage in lifelong learning endeavors.	PO12

Addressing of COs, Knowledge Profile (K), and Attainment of Complex Engineering Problems (EP):

SN	EP Definition	Attainment	CO	Justification (with Knowledge Profile)	References
1.	EP1: Depth of Knowledge required	Yes	CO1, CO2, and CO3	In our project, we needed to collect data from different e-commerce websites after collecting we needed to annotate them as Positive and Negative Review. Finally, we applied NLP to prepare a dataset for Machine Learning and Deep Learning Algorithms which covered K3 and K4	Page no: 11 Section: [3.4]
				After training and evaluation, the model can be deployed by the website or any rest API system that covers K5 and K6	Page no: 9 Section: [3.3]
				We compared my work with recent existing work which covers K8	Page no: 6 Section:

					2.3
2.	EP2: Range of Conflicting Requirements	Yes	CO2	During project work, we faced some problems like Data unavailability. Most of the comments are in mixed languages like Bangla and English. Another problem is data processing for the two domains(DL, ML) is different so we tried to solve this	Page no: 10 Section: [3.3]
3.	EP3: Depth of analysis required	Yes	CO2	We have tried to mitigate these challenges by programming. Initially, we collected all comments then we built a system that can detect foreign languages and remove them. For data preprocessing tidif is used for the ML part and Tokenizer and Embedding are used in the DL part.	Page no: 14 Section: [4.2]
4.	EP4: Familiarity of Issues	Yes	CO3	In the data processing section, we need to learn some new technology. Most e-commerce keep comments in such a way that no one can collect them. So, we tried to learn how to bypass this rule.	Page no: 14 Section [4.2]
5.	EP5: Extends of application codes	No	CO5	N/A	N/A
6.	EP6: Extends of stakeholders involved and conflicting requirements	Yes	CO3	To overcome the obstacles, we discussed with one of the Machine Learning Engineers with him we have been in touch, and the solution to my issue was explained to us by him.	Page no: 10, 11 Section: [3.4]

7.	EP7: Interdependence	Yes	CO3	We divided my problem into different subcategories like data collection, data preprocessing, and algorithm implementation. We tried to solve each category separately. In this project there is two domains are used so we carefully monitor this issue.	Page no: 9-10 Section: [3.3]
----	---------------------------------	-----	-----	--	--

Addressing CO11 with Complex Engineering Activities (EA) [Some or all of the following]:

SN	EA Definition	Attainment	CO	Justification	References
1.	EA1: Range of resources	Yes	CO11	Our research makes substantial use of high-performance computing gear, GPUs, and deep learning frameworks like Tensorflow and Keras.	Page no: 12, 13 Section: [3.5]
2.	EA2: Level of interaction	No		N/A	N/A

3.	EA3: Innovation	No		N/A	N/A
4.	EA4: Consequences for society and the environment	Yes		This project contributes to society by saving time for customers. If any e-commerce integrates our system it can automatically give a percentage of positive reviews and negative reviews as a result customers do not need to read every comment by wasting time.	Page no: 32 Section: [6.2]
5.	EA-5: Familiarity	Yes		This study extends on previous research by investigating a fresh strategy to apply Machine Learning and deep learning algorithms together, as well as a full comparison analysis, bringing new insights into the subject.	Page no: [9-10] Section: [3.3]

Addressing CO (4, 9, 10, and 12):

SN	COs	Attainment	Justification	References
1	CO4	Yes	Excellent project management and financial control, which offer careful planning, resource allocation, and budget estimation for optimal resource utilization throughout the research lifespan, are combined in this project to address CO4.	Page no: [8] Section: [3.2]
2	CO9	Yes	By prioritizing the privacy of e-commerce customers and reviewers and openly recording the research process, the project demonstrates adherence to ethical	Page no: [32, 33]

			principles. It also ensures responsible knowledge dissemination and societal well-being through the ethical use of advanced e-commerce that complies with CO9.	Section: [6.2, 6.3]
3	CO10	No	N/A	N/A
4	CO12	Yes	The project's dedication to continuous learning (CO12) and flexibility in an ever-changing technological environment is evident in the thorough experimental findings and analysis, technique development, data preparation, and data collection.	Page no: [9,10,11,14] Section: [3.3, 3.4, 4.2]

PLAGIARISM REPORT

ORIGINALITY REPORT

17%

SIMILARITY INDEX

15%

INTERNET SOURCES

10%

PUBLICATIONS

11%

STUDENT PAPERS

PRIMARY SOURCES

1

dspace.daffodilvarsity.edu.bd:8080

Internet Source

5%

2

Submitted to Daffodil International University

Student Paper

3%

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

blog.csdn.net

Internet Source

1%

5

"Advanced Computing", Springer Science and Business Media LLC, 2024

<1%