

**CLASSIFICATION OF BANGLA BOOK GENRES USING BOOK
COVER AND TITLE**

BY

**SHAHADAT HOSSAIN
ID: 201-15-13593**

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

Mr. Tanvirul islam

Lecturer

Department of Computer Science and Engineering
Daffodil International University

Co-Supervised By

Dr. Arif Mahmud

Associate Professor

Department of Computer Science and Engineering
Daffodil International University



DAFFODIL INTERNATIONAL UNIVERSITY

DHAKA, BANGLADESH

July 2024

APPROVAL

This Project titled “**CLASSIFICATION OF BANGLA BOOK GENRES USING BOOK COVER AND TITLE**”, submitted by **Shahadat Hossain** 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 **July 13, 2024**.

BOARD OF EXAMINERS



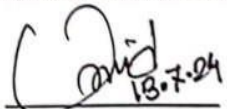
Dr. Md. Fokhray Hossain (MFH)
Professor
Department of CSE
Faculty of Science & Information Technology
Daffodil International University

Board Chairman



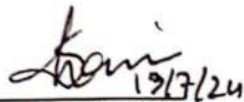
Mr. Shah Md. Tanvir Siddiquee (SMTS)
Assistant Professor
Department of CSE
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner 1



Mr. Md Umaid Hasan (MUH)
Lecturer
Department of CSE
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner 2



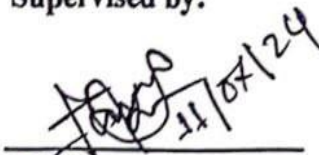
Dr. Shamim H Ripon (DSR)
Professor
Department of Computer Science and
Engineering
East West University

External Examiner

DECLARATION

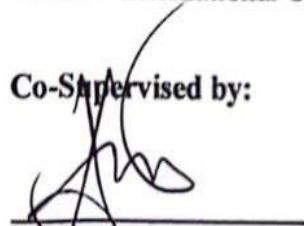
We hereby declare that this project has been done by us under the supervision of Mr. Tanvirul islam, Lecturer, Department of Computer Science and Engineering, Daffodil International University. We also declare that neither this project nor any part of this project has been submitted elsewhere for the award of any degree or diploma.

Supervised by:



Mr. Tanvirul islam
Lecturer
Department of CSE
Daffodil International University

Co-Supervised by:



Dr. Arif Mahmud
Associate Professor
Department of CSE
Daffodil International University

Submitted by:



Shahadat Hossain
ID: 201-15-13593
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 **Mr. Tanvirul islam, Lecturer**, Department of CSE Daffodil International University, Dhaka. Deep Knowledge & keen interest of our supervisor in the field of “*Artificial Intelligence and Machine Learning*” 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 **Dr. Sheak Rashed Haider Noori**, 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

The rapid expansion of e-books and digital libraries has underscored the need for efficient and accurate genre classification systems, particularly for information retrieval, content filtering, and book recommendation purposes. This research "Classification of Bangla Book Genres using Book Cover and Title," aims to develop a reliable system that leverages both linguistic and visual data for the accurate classification of Bangla book genres. A comprehensive dataset of over 15,000 Bangla book covers and titles spanning various genres was collected through web scraping from the Rokomari e-commerce site and annotated for training and evaluating the models. Transfer learning architectures employing pre-trained models like Xception were implemented for feature extraction from images, while the BanglaBERT model was used to obtain contextualized word embeddings for book titles. The core of the research involves a comparative analysis of four distinct machine learning and deep learning models: Logistic Regression, Random Forest, Support Vector Machine, and Neural Network. The findings reveal that the Neural Network model emerges as the front-runner, achieving the highest overall accuracy of 78%. This superior performance underscores the model's ability to generalize well and capture robust features essential for distinguishing between diverse Bangla book genres. This research demonstrates the transformative potential of automated content classification systems in enhancing the accessibility and global reach of Bangla literature, thereby bridging the gap between traditional literary classification and modern digital behaviors.

TABLE OF CONTENTS

Contents	Page No
Board of Examiners	i
Declaration	ii
Acknowledgments	iii
Abstract	iv
 CHAPTER 1: INTRODUCTION	 1-6
1.1 Overview	1-3
1.2 Background and Present State	3
1.3 Problem Statement	3-4
1.4 Objectives	4
1.5 Scope and Limitations	4-5
1.6 Report Organization	5-6
1.7 Summary	6
 CHAPTER 2: LITERATURE REVIEW	 7-10
2.1 Overview	7
2.2 Related Works	7-8
2.3 Comparison between existing works	8-9
2.4 Open Issues	10
2.5 Summary	10
 CHAPTER 3: METHODOLOGY/ REQUIREMENT ANALYSIS & DESIGN SPECIFICATION	 11-19
3.1 Overview	11
3.2 Data Collection Procedure	11-13
3.3 Proposed Methodology	14-15
3.4 Hardware/ Software Requirement	16-17
3.5 Project Management and Financial Analysis	17-19
3.6 Summary	19

CHAPTER 4: IMPLEMENTATION	20-24
4.1 Overview	20
4.2 Train Model	20-22
4.3 Model Evaluation	22-23
4.4 Summary	23-24
 CHAPTER 5: RESULT AND ANALYSIS	 25-37
5.1 Overview	25
5.2 Experimental Result	26-35
5.4 Comparative Analysis	36-37
5.5 Summary	37
 CHAPTER 6: IMPACT ON SOCIETY, ENVIRONMENT AND SUSTAINABILITY	 38-43
6.1 Impact on Life	38-39
6.2 Impact on Society & Environment	39-40
6.3 Ethical Aspects	40-41
6.4 Sustainability Plan	41-42
6.5 Summary	42-43
 CHAPTER 7: CONCLUSION AND FUTURE WORK	 44-46
7.1 Conclusions	44-45
7.2 Further Suggested Works	45-46
7.3 Limitations	46
 REFERENCES	 47-48

**COURSE OUTCOMES, COMPLEX ENGINEERING
PROBLEMS (EP) AND COMPLEX ENGINEERING
ACTIVITIES (EA) ADDRESSING**

LIST OF FIGURES

Figures	Page no
Fig 1.1.1: Inception module	2
Fig 3.2.1: Percentage of data collection from various sources	12
Figure 3.3.1: Work flow Diagram for Bangla Book Classification.	15
Figure 3.5.1: SWOT analysis	18
Figure 5.2.1: Training and validation accuracy over the epochs (Neural Network).	30
Figure 5.2.2: Training and validation loss over the epochs (Neural Network).	31
Figure 5.2.3: CM after Logistic Regression Model	32
Figure 5.2.4: CM after Random Forest model	33
Figure 5.2.5: CM after Support Vector Machine model	34
Figure 5.2.6: CM after Neural Network model	35
Figure 5.3.1: Accuracy comparison among Machine Learning and Deep Learning models.	36

LIST OF TABLES

Tables	Page no
Table 2.3.1: Comparative Analysis Table	9
Table 3.2.1 Label and frequency of Dataset	13
Table 3.5.1 Financial Analysis	19
Table 5.2.1: Comparative Analysis Table	26
Table 5.2.2: Precision, Recall, F1-Score, and Support (n) for Logistic Regression Model (depending on the number of pictures, n=numbers).	27
Table 5.2.3: Precision, Recall, F1-Score, and Support (n) Random Forest Model (depending on the number of pictures, n=numbers).	27
Table 5.2.4: Precision, Recall, F1-Score, and Support (n) Support Vector Machine Model (depending on the number of pictures, n=numbers).	28
Table 5.2.5: Precision, Recall, F1-Score, and Support (n) Neural Network Model (depending on the number of pictures, n=numbers).	28

LIST OF ACRONYMS

ML	Machine Learning
DL	Deep Learning
SVM	Support Vector Machine
CNN	Convolutional Neural Network
OCR	Optical Character Recognition
NLP	Natural Language Processing
CM	Confusion Matrix
TP	True Positive
FP	False Positive
FN	False Negative
TN	True Negative

CHAPTER 1

INTRODUCTION

1.1 Overview

Over the past few years, e-books and digital libraries have grown at an exponential rate, leading to a rise in the need for accurate and efficient genre classification for information retrieval, content filtering, and book recommendation systems. The goal of this study, "Classification of Bangla Book Genres using Book Cover and Title," is to further the creation of a unique and reliable system that uses linguistic and visual data to accurately classify genres. The classification of book genres using elements from the cover and title has been a focus of recent literature. Biradar, Raagini, Varier and Sudhir (2019) have used machine learning methods to combine cover and title information for genre prediction, such as logistic regression. While there isn't research specifically on using cover and title to categorize Bangla book genres in the context of Bangla literature, there is comparable work on the classification of Bangla music genres (Ahmed, Alam, Paul, Hasan and Rab 2022) and the development of a Bengali literature classification system indicates a greater interest in categorizing Bangla cultural content (Panigrahi & Roy 2014). These findings show that the use of computational methods for the classification of Bangla cultural resources is becoming more and more popular. These methods can also be used to classify book genres in Bangla by utilizing information found on book covers and titles. The primary goals of this research are to improve user experience, increase Bangla literature's reach, and demonstrate the potential of automated content classification systems outside of the English-speaking world. Through this investigation, the research aims to close the gap between modern digital behaviors and traditional literary classification, making Bangla literature enthusiasts' experiences smoother and more rewarding.

Xception architecture:

The Xception architecture, introduced by (Chollet 2017), took an "Extreme" approach to the Inception architecture developed by (Szegedy, Liu, Jia, Sermanet, Reed, Anguelov, and Rabinovich 2015). The core component of the Inception model was the Inception module, which performed parallel convolutions, as illustrated in Figure 1.1.1. The primary motivation behind this design was to avoid manually selecting convolution sizes. Instead, by using multiple convolution types simultaneously, the model itself could determine the most effective ones.

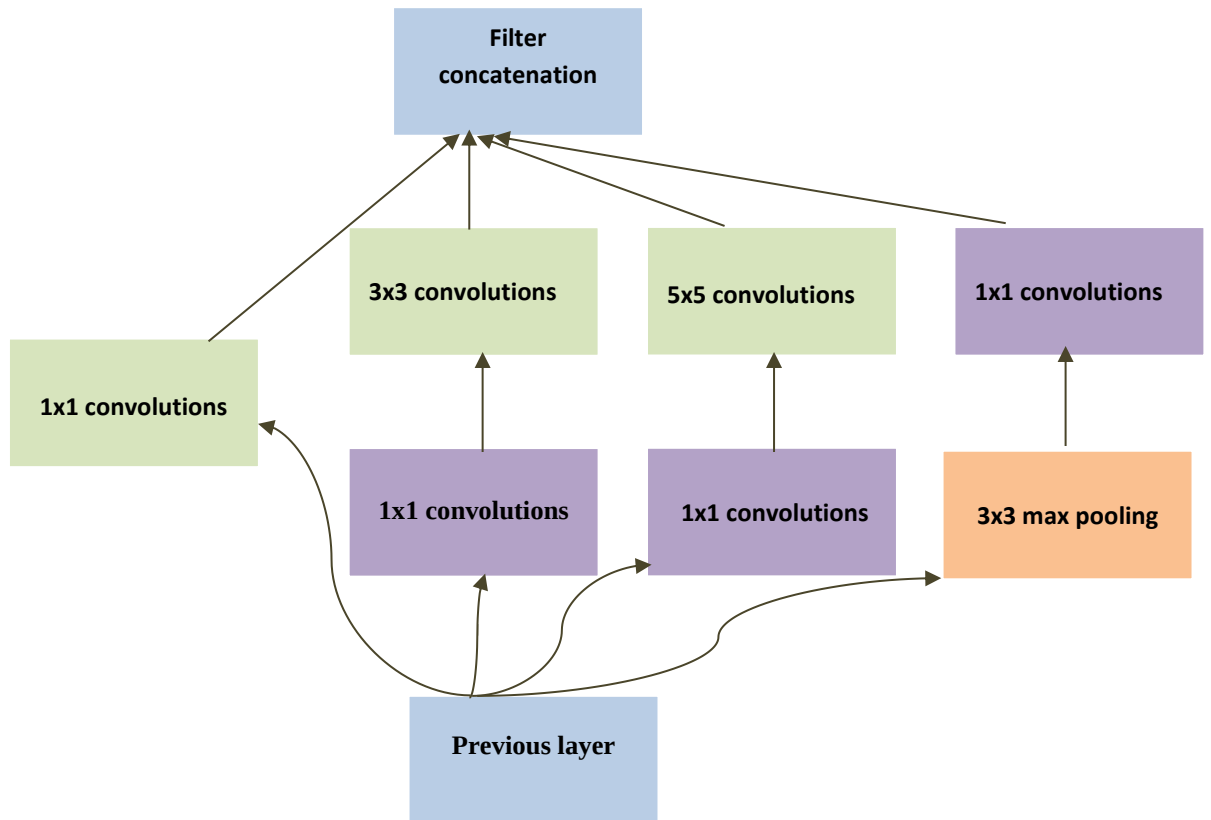


Fig 1.1.1: Inception module

Another key feature of the Inception module was the use of 1x1 convolutions before other convolutions, which significantly reduced computational load and allowed the network to become larger and wider. This feature was based on the idea that spatial correlation and cross-channel correlation could be decoupled.

Xception extended this concept further by assuming that these correlations could be entirely decoupled. Consequently, the proposed architecture exclusively utilized depthwise separable convolutional layers. A depthwise separable convolution involves performing a "depthwise convolution" on each channel separately, followed by a "pointwise convolution" after concatenating these channels. However, in the Xception architecture, these operations are implemented in the reverse order.

BanglaBERT Model:

BanglaBERT is a language model specifically designed by (Bhattacharjee, Hasan, Ahmad, Samin, Islam, Iqbal & Shahriyar 2021) for understanding and generating text in Bangla, one of the most widely spoken languages in the world. It is part of the

family Bidirectional Encoder Representations from Transformers models, which are neural network-based techniques for natural language processing tasks.

The BanglaBERT model was pretrained on a large corpus of Bangla text data, dubbed the "Bangla2B+" dataset, which consists of 27.5 GB of text crawled from 110 popular Bangla websites. The pretraining process utilized the Replaced Token Detection objective, which is similar to the Masked Language Modeling objective used in the original BERT model but more computationally efficient. The pretraining was carried out using the ELECTRA framework, resulting in an ELECTRA discriminator model that serves as the final BanglaBERT model. This approach has been shown to provide better performance and sample efficiency compared to the standard BERT pretraining procedure.

1.2 Background and Present State

The classification of book genres in Bangla literature has become increasingly important and necessary, particularly in the context of book covers and titles. In the digital age, book covers serve as the primary point of interaction between readers and books, acting as the first impression that shapes the reader's expectations and decision to engage with a literary work. The textual and visual elements of a book cover are carefully chosen to convey the book's genre and overall message, making it a crucial factor in the reader's experience. However, manually categorizing book genres based on covers and titles is a labor-intensive and time-consuming task that requires extensive human effort and expertise. This challenge has led to the development of automated systems that leverage advanced technologies like deep learning and text mining to make genre classification more efficient. By using these modern techniques, the aim is to improve how easily people can find and access different books, which will benefit both authors and readers in the Bangla literary community. The motivation behind this project is to change the way we find and interact with books, making the process more efficient, accurate, and enjoyable for everyone involved in the literary world. The development of automated book genre classification systems is driven by the increasing significance of book covers and titles as the primary touchpoints between readers and books in the digital age.

1.3 Problem Statement

The primary focus of the problem statement is the requirement to enhance the accuracy of categorizing Bangla book genres by utilizing the book cover and title. There is a lack of specific research on using book cover and title aspects to categorize

Bangla book genres, even though previous studies have examined genre classification in the context of English literature. Because of the unique characteristics of Bangla literature, such as its variety of genres and complex cultural themes, this gap poses a special problem. Therefore, the problem statement contains developing and applying machine learning, computer vision, and natural language processing techniques to effectively classify Bangla book genres using both visual and textual information. The objective is to address the weaknesses of current genre classification techniques and offer an accurate and culturally appropriate method for classifying literary works written in Bangla

1.4 Objectives

The primary objectives of the research project "Classification of Bangla Book Genres using Book Cover and Title" are as follows:

- To develop and design a machine learning model that can accurately classify Bangla book genres based on book cover images and titles.
- Showcase the possibilities of automated content classification methods outside of English literature to help Bangla literature gain international attention.
- To utilize a comprehensive dataset of over 15,000 Bangla book titles and images from various genres, ensuring a diverse and representative collection for training and evaluating the machine learning models.
- To employ robust and efficient machine learning models that can handle the complexity of Bangla text and images, ensuring accurate and reliable genre classification.

By achieving these objectives, the project aims to overcome the limitations of manual genre classification, enhance user experience, and contribute to the global recognition of Bangla literature through a systematic and innovative approach to genre classification using machine learning, computer vision, and natural language processing techniques.

1.5 Scope and Limitations

The scope of the research project "Classification of Bangla Book Genres using Book Cover and Title" encompasses the development of an accurate and efficient system for categorizing Bangla books into their respective genres using machine learning, computer vision, and natural language processing techniques. The project aims to tackle the challenges associated with manual genre classification, which can be time-consuming, subjective, and inconsistent, especially with the growing number of books available. The project will focus on leveraging visual analysis of book covers

and linguistic analysis of titles to create a robust and reliable genre classification system. The system will be designed to aid readers in discovering books that match their interests and help organize large collections of Bangla literature.

Limitations: Although the project strives for comprehensiveness, it must be understood that it has several limitations.

- **Genre Representation:** Although a wide variety of genres is the project's main focus, it might not cover every subgenre or niche in Bangla literature.
- **Data Source Bias:** Since the Rokomari platform is where the dataset will be collected, the site's catalog may have an impact on how genres are represented.
- **Image Size:** The images obtained are relatively small (240×320 pixels). While this size is considered adequate for feature extraction, fine features in book covers might not be fully captured.

Despite these limitations, the project aims to make significant contributions to the field of Bangla book genre classification by developing an innovative and efficient system that leverages state-of-the-art machine learning techniques. The system will be designed to be scalable and adaptable, allowing for future improvements and expansions as the field of Bangla literature continues to grow and evolve.

1.6 Report Organization

The report layout for the project "Classification of Bangla Book Genres using Book Cover and Title" is meticulously structured to provide a comprehensive and detailed investigation into the genre classification process. The report is organized into several chapters, each focusing on specific aspects of the research endeavor.

Chapter 1: Introduction: This chapter sets the foundation for the research, outlining the background information on the significance of genre classification in Bangla literature, the specific goals of the project, and the challenges faced in the genre classification process.

Chapter 2: Literature Review: This chapter delves into a comprehensive review of existing literature on genre classification based on book covers and titles. It critically analyzes previous studies, methodologies, and findings to establish a solid theoretical framework for the current research.

Chapter 3: Methodology: The methodology chapter details the proposed method for classifying genres using machine learning, computer vision, and natural language

processing techniques. It explains the data collection process, model architecture, feature extraction methods, and the overall research design.

Chapter 4: Implementation: This chapter describes the implementation process, such as model training or prototype design, and discusses the system testing and model evaluation procedures.

Chapter 5: Result and Analysis: In this chapter, the experimental or simulation results are presented, and the performance of the proposed solution is analyzed and compared with existing methods.

Chapter 6: Impact on Society, Environment and Sustainability: This chapter discusses the potential impact of the project on life, society, and the environment, addresses the ethical aspects of the research, and develops a sustainability plan for the project.

Chapter 7: Conclusion and Future Work: The summary chapter serves as a synthesis of the entire report, summarizing key findings, conclusions, and recommendations. It offers insights into the implications of the research and potential avenues for future exploration in genre classification in Bangla literature

1.7 Summary

In response to the growing demand for precise genre classification in Bangla literature, this project employs a unique integration of machine learning, computer vision, and natural language processing. The study intends to overcome the lack of research in Bangla book genre classification by focusing on using book covers and titles. One of the main problems is that there aren't any studies specifically on how to categorize Bangla genres using elements from covers and titles, which makes an automated approach necessary. Improving user experience and helping writers and readers in the digital age are the driving forces behind this. The goals include increasing the accuracy of genre classification, improving user experience, and advancing the international recognition of Bangla literature. The project's scope is comprehensive, covering diverse Bangla literary genres by utilizing a dataset obtained through web scraping the Rokomari e-commerce site. The expected results include the creation of a precise system for classifying genres, and a rise in the stature of Bangla literature worldwide. The introduction, methodology, project scope, expected outcomes, literature review, and organizational structure are all included in the report's chapters. The goal of this research is to provide an accurate and culturally appropriate automated approach to bridge the gap in the classification of Bangla literature.

CHAPTER 2

LITERATURE REVIEW

2.1 Overview

An essential part of any research process is reviewing the literature already in existence since it gives researchers an understanding of the state of the area, identifies knowledge gaps, and influences the direction of their research. A literature review offers an in-depth analysis and overview of the arguments and ideas of existing knowledge in a particular field. In the context of classifying Bangla book genres based on cover and title, examining the literature is important for understanding the prospects and difficulties in this field. A thorough review allows for the identification of gaps and limitations in previous studies. There is a lack of particular research on the use of book cover and title features to identify Bangla book genres, even though earlier studies have looked at genre categorization in the context of English literature and other languages. Because of the unique characteristics of Bangla literature, including its range of genres and rich cultural topics, this gap presents a particular challenge. Analyzing existing literature helps one make sensible decisions on the methods used in the current study. It offers guidelines for choosing appropriate image processing, language processing, and machine learning algorithms. In essence, the literature study is an essential element that influences the project's direction rather than just serving as a precursor. It serves as a compass, directing the research across the existing knowledge landscape, discovering unexplored areas, and opening the path for innovation and improvement in the classification of Bangla book genres.

2.2 Related Works

In the context of classifying Bangla book genres based on cover and title, several relevant studies have been accomplished. Chiang, Ge and Wu (2015) discussed the classification of books purely based on cover image and title, without prior knowledge or context of author and origin. Chiang, Ge and Wu (2015) used a color-based distribution approach and also implemented transfer learning with convolutional neural networks on the cover image along with natural language processing on the title text. Biradar, Raagini, Varier and Sudhir (2019) used a logistic regression model to classify book genres using book Cover and Title and got accuracy of 87.2% by combined features from cover and title and got accuracy of using only cover features and an accuracy of 86.2% using only the title features. By creating a dataset of assembled texts from online book reviews (Scofield, Silva, de Melo-

Gomes, and Moro 2022) presented a model for automatically classifying book genres. The model also employed multiple machine learning algorithms to classify a book into a particular genre. Ozsarfati, Sahin, Saul and Yilmaz (2019) showed algorithmic comparisons for producing a book's genre based on its title. Moreover, several data preprocessing steps were implemented, in which word embeddings were created to make the titles operable by the computer and also different machine learning models were tested throughout the experiment. Each different algorithm was fine-tuned for attaining the best parameter values, while no modifications were conducted on the dataset and got accuracy of 65.58%. Gupta, Agarwal and Jain (2019) proposed method gains knowledge from a large number of words from the books and transforms them into a feature matrix also reduce the size of the initial matrix by using Wordnet and Principle Component Analysis and applied AdaBoost classifier is applied to predict the genres of the books. Optical Character Recognition is the technique of creating machine-readable text from photographs of printed, typewritten, or handwritten text. Among the computer vision fields, optical character recognition (OCR) has been studied the most (Dipu, Shohan and Salam 2021). Furthermore, the majority of the systems in use today are unable to identify compound letters. So, Dipu, Shohan and Salam (2021) tried to resolve this issue by proposing three neural network-based image classification models for Bangla OCR by using models like are Inception V3, VGG16, and Vision Transform. Last few decades, the availability and accessibility of the Bangla document and its content have rapidly increased due to the rapid technological advancement (Pasha, Islam, Rahman, Ahmed, Foysal & Alam 2023). However, Pasha, Islam, Rahman, Ahmed, Foysal and Alam (2023) implemented Deep Learning approaches like Convolutional Neural Network (CNN) and Long Short-Term Memory (LSTM) to classify Bangla text documents and got accuracy 95.42% on the Prothom Alo data set using LSTM. Together, these studies add to our knowledge of how to categorize book genres using their covers and titles, showing the variety and creativity of methods used in this area.

2.3 Comparative Analysis and Summary

The comparative analysis conducted in the provided sources offers a detailed examination of various studies related to the classification of Bangla book genres based on cover and title information, as shown in Table 2.3.1.

Table 2.3.1: Comparative Analysis Table

SL No	Author Name	Used Algorithm	Best Accuracy with Algorithm
1.	Biradar, G. R. et al [1]	Logistic regression	87.2% (combined features from cover and title), 86.2% (only cover features), 86.2% (only title features)
2.	Chiang, H. et al [8]	SVM, Transfer Learning	SVM= 62.4%
3.	Gupta, S. et al [3]	AdaBoost Classifier	Accuracy: 81.18% (Without Unlabeled Data) Accuracy: 92.88% (With Unlabeled Data)
4.	Scofield, C. et al [6]	Naïve Classifier, Decision Tree Classifier, Random Forest, Stochastic Gradient Descent, Gaussian Naive Bayes, K Nearest Neighbor, AdaBoost Classifier	Random Forest = 96%
5.	Ozsarfati, E. et al [2]	Naive Bayes, Convolutional Neural Networks, Recurrent Neural Networks, Gated Recurrent Unit, Long Short-Term Memory, Bi-Directional LSTM	Long Short-Term Memory = 65.58%
6.	Dipu, N. et al [9]	Inception V3, VGG16, Vision Transformer	Inception V3= 98.65%
7.	Pasha, S.T. et al [10]	Logistic Regression, Naïve Bayes (NB), Support Vector Machine (SVM), Long Short-Term Memory (LSTM), Convolutional Neural Network (CNN), Transformer Model	Transformer Model = 87%
8.	Ahmed, T. et al [11]	Artificial Neural Network	Artificial Neural Network= 77.68%

2.4 Open Issues

In the context of classifying book genres based on cover and title, the literature review identifies several relevant research studies that have advanced our knowledge of the subject, several open issues and challenges remain that need to be addressed. While some research focuses on classifying book genres using cover and title information, other studies have investigated the application of deep learning techniques, online text reviews, and OCR. In addition, the majority of the studies that have already been conducted have been on English-language books. Effectively combining visual data from book covers and textual data from titles remains a challenge. Ensuring that both types of data are appropriately preprocessed and weighted in the model is critical for achieving accurate classification.

Overall, the existing literature provides a foundation for further research in this area, and more thorough and integrated methods are required to accurately classify Bangla language books' genres based on information identified on their covers and titles.

2.5 Summary

The literature review provides insightful information about the many methods and strategies for classifying Bangla book genres according to their covers and titles. Many studies have shown how useful deep learning, machine learning, and natural language processing methods are in this field. When predicting book genres from cover and title information, the applications of logistic regression, convolutional neural networks, and recurrent neural networks have demonstrated encouraging outcomes. Furthermore, research has been done on the use of character-level deep learning and optical character recognition for Bangla text classification, which has improved our comprehension of the existing approaches. The review also highlights the importance of natural language processing in text categorization and summarization, emphasizing how it can potentially solve problems related to manual genre classification. The literature also emphasizes the necessity of robust and efficient machine learning models to manage the growing amount of internet text data and book reviews needed for precise genre classification. The knowledge gained from the literature study as an entire assist in overcoming the limitations of conventional manual categorization techniques by assisting in the creation of complex and precise algorithms for categorizing Bangla book genres.

CHAPTER 3

METHODOLOGY/ REQUIREMENT ANALYSIS AND DESIGN SPECIFICATION

3.1 Overview

The research paper's methodology part provides a thorough strategy for categorizing Bangla book genres according to their covers and titles. The methodology enhances the accuracy and reliability of genre categorization by combining visual identification of cover images with linguistic analysis of titles. The methodology has been designed to tackle the various challenges that arise from the complex nature of Bangla literature. In Linguistic analysis of titles, A detailed linguistic analysis of book titles is the initial step in the methodology. Natural language processing techniques are used to extract relevant features from the titles, including sentiment, keywords, and semantic meaning. The goal is to enable an in-depth understanding of the genre by capturing the core of the literary content included in the title. The second part uses computer vision algorithms to identify book covers visually. Convolutional neural networks, one type of deep learning model, are used to extract information from cover photographs. These characteristics include visual cues like color schemes, patterns, and symbols that allow the system to identify distinct visual signals linked to various genres. The fused features are then fed into machine learning and deep learning models for genre prediction. To accomplish accurate and effective genre classification, the methodology part offers a thorough and systematic framework for the categorization of Bangla book genres. This framework integrates deep learning and sophisticated machine learning techniques.

3.2 Data Collection Procedure

A. Datasets

The project "Classification of Bangla Book Genres using Book Cover and Title" utilizes a dataset collected through web scraping the Rokomari e-commerce site. The dataset consists of over 15,000 Bangla book covers and titles, representing a diverse range of genres.

Rokomari is chosen because of its extensive book collection, which covers many different genres. The research aims to create a dataset that reflects different Bangla literary categories, and this diversity is in line with the goal. The platform offers a broad and diverse set of data including categories like academic books, fiction, non-fiction, and children's literature. Rokomari gives the title and an image for every book, which

is essential for the linguistic and visual analysis that is needed for the research. However, the Rokomari images are relatively small, with dimensions of 240×320 pixels. Smaller images can be used since the feature extraction process does not mainly depend on fine details. Smaller images also help to increase processing speed without sacrificing functionality.

It is expected that the web scraping technique for data collecting from Rokomari will offer a strong basis for training and evaluating the model, allowing accurate genre classification based on cover photos and titles. Fig 3.2.1 provides a summary of the dataset statistics.

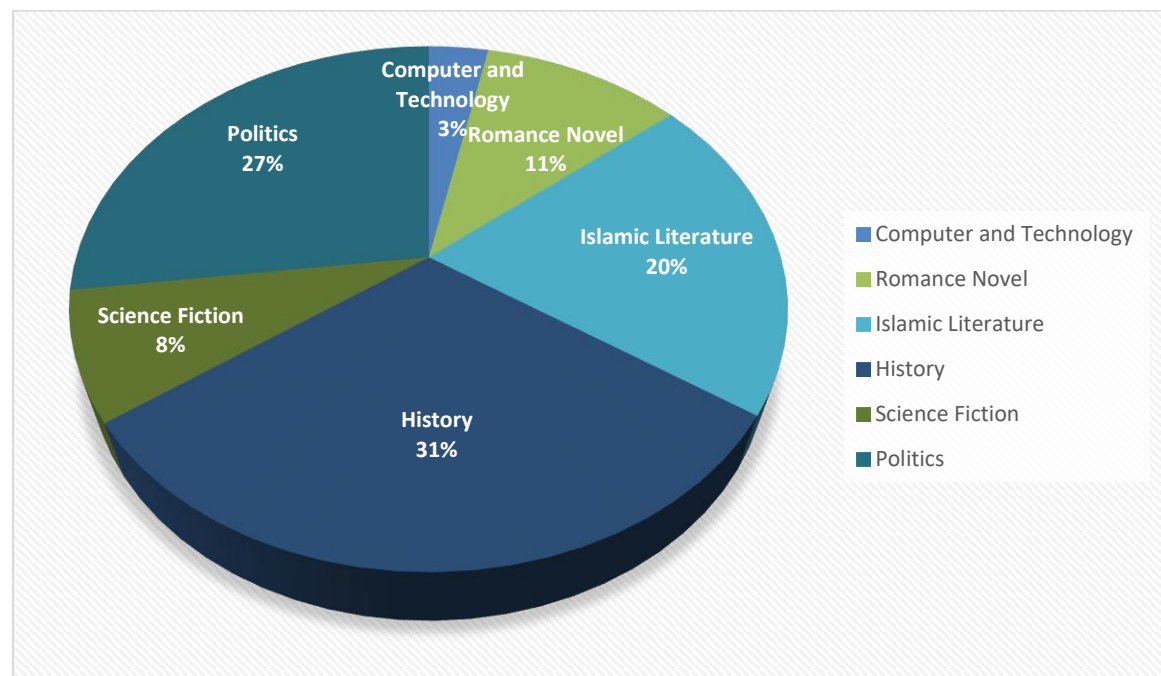


Fig 3.2.1: Percentage of data collection from various sources

The distribution of our dataset is as following, where the number of books in each category, Total Number of Books are shown in the table 3.2.1 below.

Table 3.2.1 Label and frequency of Dataset

Label	Frequency
Computer and Technology	511
Romance Novel	1718
Islamic Literature	3226
History	4863
Science Fiction	1248
Politics	4279
Total	15,845

B. Data Collection Process

Web Scraping:

Web scraping tools are used to collect the book cover images and titles from the Rokomari e-commerce site. The scraping process ensures that a comprehensive dataset is obtained, covering a wide range of Bangla book genres. Python's “BeautifulSoup” library is chosen for web scraping. This is a popular choice for parsing HTML and XML documents, which is essential for extracting data from web pages.

The data collection process is automating the retrieval of images from a specific category on the Rokomari.com website. First, Determining the number of pages in the category by fetching pagination information from the category URL. Extracting unique links to individual book pages within the category. Processing each book page to obtain the book title and corresponding image URL. Sanitizing the book titles to create valid filenames. Downloading the images from the extracted URLs and saving them to a designated folder. Tracking the progress of the data collection process, including the number of pages processed and images downloaded.

The data collection procedure is designed to ensure that the dataset is comprehensive, representative, and suitable for training and evaluating the machine learning models for Bangla book genre classification.

3.3 Proposed Methodology

The proposed methodology for the classification of Bangla book genres using book covers and titles involves several steps. Each step is carefully designed to contribute to the overall goal of accurately classifying book genres. The methodology leverages both image processing and natural language processing techniques to extract features from book covers and titles, respectively. These features are then combined to train a machine learning model that can classify book genres. The following figure 3.3.1 shows methodology for Bangla Book Classification.

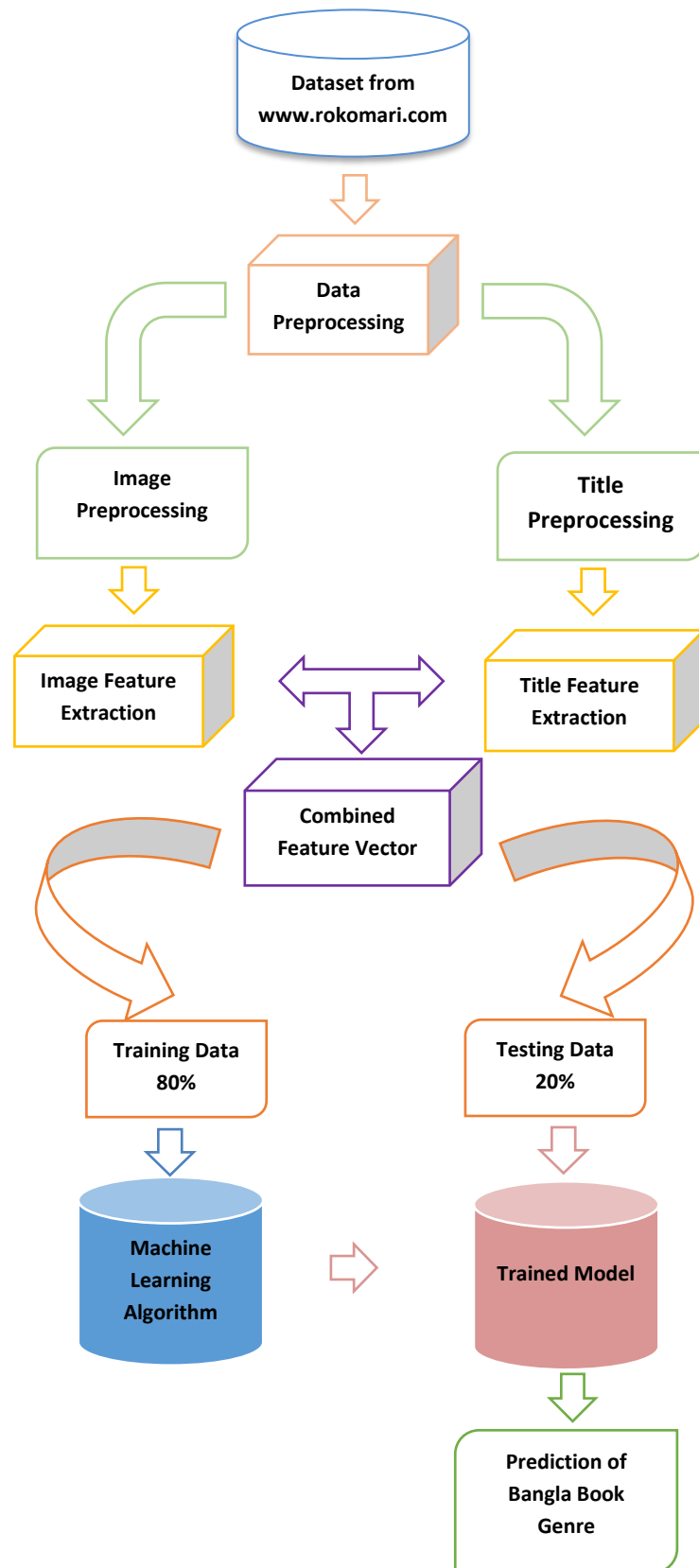


Figure 3.3.1: Work flow Diagram for Bangla Book Classification.

3.4 Hardware/ Software Requirement

The implementation of the project "Classification of Bangla Book Genres using Book Cover and Title" entails specific requirements to ensure the successful execution of the study. These implementation requirements encompass both hardware and software components, as well as considerations related to data collection and model evaluation.

Hardware Requirements

- **High-performance computing resources:** Access to computational infrastructure with sufficient processing power and memory to handle the complex tasks involved in deep learning and image analysis.
- **Graphics Processing Unit (GPU):** A GPU is essential to leverage Keras and PyTorch, both of which benefit from GPU acceleration, particularly for training deep Convolutional Neural Networks.

Software Requirements

- **Python programming language:** Python provides a versatile and widely-used platform for implementing deep learning algorithms and conducting data analysis.
- **Deep learning frameworks:** Utilization of popular deep learning frameworks such as TensorFlow or PyTorch for the development, training, and evaluation of the CNN models.
- **Image processing libraries:** Integration of image processing libraries like the Keras image module for processing book cover images.
- **Natural Language Processing libraries:** Transformers from HuggingFace is Used to tokenize titles and obtain embeddings using a BERT model.
- **BanglaBERT model:** Availability of the BanglaBERT model for obtaining contextualized word embeddings for book titles.

Data Collection

- **Diverse and representative dataset:** Collection of a comprehensive dataset of book cover images and titles, including various genres and authors, to ensure the robustness and generalization of the developed models.

- **Annotated dataset:** Availability of annotated data for training and validating the models, with clear labels indicating the genre of each book.

Model Training and Evaluation

- **Transfer learning models:** Implementation of transfer learning architectures, leveraging pre-trained models such as employs the pre-trained Xception for feature extraction from images.
- **Metrics for evaluation:** Definition and utilization of appropriate evaluation metrics, such as accuracy, precision, recall, and F1 score, to assess the performance of the models in both detection and segmentation tasks.

Documentation and Reproducibility

- **Code documentation:** Thorough documentation of the implementation code to facilitate transparency, reproducibility, and future collaboration.
- **Version control:** Adoption of version control systems to track changes in the codebase and maintain a structured development process.

By meeting these implementation requirements, the research can be conducted systematically and rigorously, ensuring the reliability of the results and contributing to advancements in the field of book genre classification through transfer learning and deep Learning.

3.5 Project Management and Financial Analysis

Project Management:

- **Project Objectives:**
 - A. To develop and design a machine learning model that can classify bangla book genre with precise accuracy.
 - B. To create a user-friendly system that give readers a simple and comfortable way to find books that suit their interests.
- **Project Timeline:**
 - A. **Phase 1:**
 - Project planning (1-2 weeks)
 - Research (2-3 weeks)
 - Data collection (3-4 weeks)

- Data Preprocessing (2 weeks)

B. Phase 2:

- Feature Extraction (2-3 weeks)
- Machine Learning Model Selection (1-2 weeks)
- Training and testing the model (3-4 weeks)
- Evaluating the model (1 weeks)

C. Data and Content:

Data is collected from Rokomari.com which is online ecommerce website. We collect all the book images through web scraping.

D. Training and Skill Development:

- Ensure that the development team has the necessary training and skills in python, web scraping, Machine Learning, Deep Learning, Natural Possessing Language.
- Provide additional training on specific technologies and tools such as feature extraction by pre-trained deep learning models.

E. Risk Management:

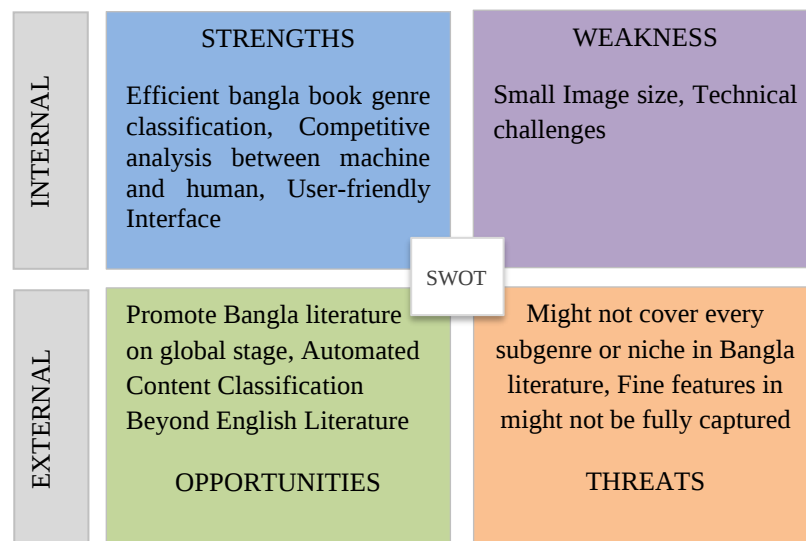


Figure 3.5.1: SWOT analysis

F. Financial Analysis:

Table 3.5.1 Financial Analysis

SN	Components	Estimated Cost (BDT)
01.	Visiting Stakeholders	1000-2000
02.	Software and Tools	2000-3000
03.	Data Collection and Processing	3000-4000
04.	Documentation and Report Writing	500
Total Estimated Cost		6500-9500

3.6 Summary

In summary, the proposed methodology for the classification of Bangla book genres leverages both visual and linguistic features to improve accuracy and reliability. The approach involves a multi-step process starting with data collection from the Rokomari e-commerce site, followed by data preprocessing to standardize book covers and titles. Feature extraction is performed using advanced models like Xception for image features and BanglaBERT for title features, resulting in high-dimensional vectors that capture the essence of the book covers and titles. These features are combined and used to train various machine learning models, including logistic regression, random forest, SVM, and neural networks, to predict book genres accurately. The implementation of this methodology requires robust hardware and software resources, including high-performance computing infrastructure, GPUs, and various deep learning and NLP libraries. Ensuring a diverse and annotated dataset is crucial for training and validating the models. The performance of the models is evaluated using metrics such as accuracy, precision, recall, and F1 score, ensuring comprehensive assessment and validation.

Effective project management, including clear objectives, a structured timeline, resource planning, and risk management, supports the successful execution of the study. Financial analysis ensures budget allocation for essential components like stakeholder visits, software tools, and data processing.

Overall, this methodology provides a comprehensive framework that integrates deep learning and sophisticated machine learning techniques, advancing the field of book genre classification for Bangla literature.

CHAPTER 4

IMPLEMENTATION

4.1 Overview

The implementation chapter of the research project "Classification of Bangla Book Genres using Book Cover and Title" outlines the key steps and processes involved in the development and evaluation of the genre classification system. Building upon the methodological framework established in the previous chapter, this phase focuses on the practical execution of the proposed approach, detailing the steps for training the machine learning models, designing the system architecture, and conducting comprehensive testing and evaluation.

Initially, data preprocessing is carried out to prepare book cover images and titles for feature extraction and model training. This step ensures that the data is in a suitable format for further processing. Following preprocessing, features are extracted from the data to create a comprehensive representation of the book covers and titles. These features serve as the foundation for training the machine learning models.

The next stage involves training various machine learning models to classify Bangla book genres accurately. Each model's performance is optimized to ensure the highest possible accuracy in genre classification. After training, the models are evaluated using a set of performance metrics. This evaluation provides insights into the effectiveness of each model and helps determine the most suitable approach for the genre classification task. Comprehensive system testing is then conducted to ensure the reliability and robustness of the genre classification system. This includes evaluating the system's performance on a test set and conducting real-world user testing to assess usability and user experience. The implementation chapter translates theoretical concepts and methodologies into a functional and deployable genre classification system for Bangla books, forming the backbone of the research project.

4.2 Train Model

The primary objective of this project is to develop a robust classification system for Bangla book genres using a combination of image and text data. The system employs a deep learning-based approach to extract features from book cover images and titles, followed by classification using various machine learning algorithms. This hybrid approach leverages the strengths of convolutional neural networks for image feature

extraction and transformer-based models for text feature embedding, ensuring a comprehensive feature representation.

Data Collection:

The data collection process involves collecting book cover images and titles from the Rokomari e-commerce site.

Data Preprocessing:

All the book covers were pre-processed, which included resizing the images and normalizing the red, green, and blue channels to standardize the color value ranges. All the titles are converted to lowercase and tokenized. Stop words are removed from the titles, and contextualized word embeddings are obtained for each tokenized title using the BanglaBERT model.

Feature Extraction:

- **Image Features:** We used an Xception model pre-trained on the ImageNet dataset to extract image features from the book covers. The model processes images of size 299x299x3, and resizing was done during the pre-processing step. Features were extracted after the average pooling layer, resulting in a 2048-dimensional feature vector
- **Title Features:** Extract features from preprocessed titles using the BanglaBERT model. BanglaBERT word embeddings will depend on the specific model size, but the publicly available BanglaBERT model has 768 dimensions. The average of these word vectors was used as the title vector. This finally gave a 768 sized feature vector,
- **Combined Features:** We concatenated with 2048 features extracted from image and title 768 features into a single feature set.

Data Splitting:

The combined dataset is split into training and testing sets with an 80-20 ratio. This splitting ensures that the model is trained on a substantial portion of the data while retaining a separate test set for unbiased evaluation. Features are flattened to ensure compatibility with classifiers.

Model Training

Four different machine learning models and deep Learning are trained using the

combined feature set:

- **Logistic Regression:** A logistic regression model is trained to classify book genres based on the combined feature set.
- **Random Forest Classifier:** A random forest classifier is trained, leveraging the ensemble learning capabilities of decision trees to improve classification accuracy.
- **Support Vector Machine:** An SVM model is trained to find the optimal hyperplane that separates the book genres in the high-dimensional feature space.
- **Neural Network:** A neural network model is trained, allowing the model to learn complex non-linear relationships between the input features and book genres.

The training process for each model involves the following steps:

1. Initializing the model with appropriate hyperparameters.
2. Feeding the training data to the model.

Hyperparameter Tuning

The hyperparameters of each model, such as the learning rate, regularization strength, and the number of estimators in the case of the Random Forest Classifier, are tuned using techniques like grid search or random search. The goal is to find the optimal configuration that maximizes the classification performance on the validation set.

4.3 Model Evaluation

The trained machine learning models for the "Classification of Bangla Book Genres using Book Cover and Title" project were thoroughly evaluated using a range of performance metrics to assess their effectiveness and identify the most suitable model for the task. We utilized various evaluation metrics to comprehensively assess the performance of the trained models, including accuracy, precision, recall, F1-score, and log loss. The evaluation also involved analyzing confusion matrices to gain deeper insights into the classifiers' performance.

Evaluation Metrics:

- **Accuracy:** Calculates the percentage of correctly identified cases among all

instances.

$$Accuracy = \frac{True\ Positives + True\ Negatives}{True\ Positives + True\ Negatives + False\ Positives + False\ Negatives}$$

- **Precision:** Shows the percentage of actual positive forecasts among all positive predictions.

$$Precision = \frac{True\ Positives}{True\ Positives + False\ Positives}$$

- **Recall:** calculates the ratio of real positives to genuine positive forecasts.

$$Precision = \frac{True\ Positives}{True\ Positives + False\ Negatives}$$

- **F1-Score:** Harmonic mean of memory and accuracy, offering a solitary measure that equalizes the two.

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

- **Confusion Matrix:** Provides insights into the model's classification performance by providing a thorough analysis of true positives, false positives, true negatives, and false negatives. The confusion matrix is represented as:

$$Confusion\ Matrix = \begin{bmatrix} TP & FP \\ FN & TN \end{bmatrix}$$

These evaluation metrics provide a comprehensive assessment of the models' performance, allowing for a comparative analysis and the identification of the most suitable model for the Bangla book genre classification task.

4.4 Summary

The implementation phase of the "Classification of Bangla Book Genres using Book Cover and Title" project encompasses the critical steps of data preprocessing, feature extraction, model training, and evaluation. A reliable and thorough representation of the input data was made possible by the hybrid technique, which combined transformer-based models for text embedding with convolutional neural networks for image feature extraction. Book cover photos and titles were first preprocessed to make sure they were in a state that could be used for feature extraction. For image feature extraction, the pre-trained Xception model on the ImageNet dataset was used,

and for text feature embedding, the BanglaBERT model. A single feature set was created by merging the final features from the two sources to train the model. Four machine learning models Logistic Regression, Random Forest, SVM, and Neural Network were trained on the combined feature set. Hyperparameter tuning was performed to optimize the performance of each model. The evaluation phase involved assessing these models using metrics such as accuracy, precision, recall, F1-score, and log loss, alongside confusion matrix analysis.

CHAPTER 5

RESULT AND ANALYSIS

5.1 Overview

This chapter presents a comprehensive evaluation of the performance of the trained machine learning and deep learning models for the "Classification of Bangla Book Genres using Book Cover and Title" project. The analysis delves into the experimental results, providing a detailed comparative assessment of the models' effectiveness in accurately classifying Bangla book genres.

The evaluation employs a range of performance metrics, including accuracy, precision, recall, F1-score, and confusion matrices. These metrics offer a multifaceted approach to assessing the models' capabilities, shedding light on their strengths and weaknesses across different book genres.

The experimental results are then presented, highlighting the comparative performance of the four models: Logistic Regression, Random Forest, Support Vector Machine and Neural Network. The accuracy, precision, recall, and F1-score for each model are discussed, allowing for a detailed comparison of their classification capabilities across different book genres.

The chapter also includes a visual representation of the models' performance through confusion matrices, which provide a deeper understanding of the classification errors and the distribution of correct and incorrect predictions for each genre.

Furthermore, the chapter explores the training and validation performance of the Neural Network model, the best-performing model, to gain insights into its learning dynamics and potential for overfitting. The trends in training and validation accuracy, as well as loss, are analyzed to understand the model's ability to generalize to unseen data.

The comprehensive evaluation and comparative analysis presented in this chapter serve as a foundation for understanding the strengths and limitations of the various machine learning and deep learning approaches employed in the Bangla book genre classification task. The insights gained from this chapter will inform future improvements and the selection of the most suitable model for real-world deployment.

5.2 Experimental Result

The performance of the four-machine learning and deep learning models, including Logistic Regression, Random Forest, Support Vector Machine, and Neural Network, is compared in this section.

Table 5.2.1: Comparative Analysis Table

Models	Training Accuracy	Model Accuracy
Logistic Regression	89%	76%
Random Forest	95%	64%
Support Vector Machine	94%	72%
Neural Network	87%	78%

The comparative analysis of several machine learning and deep learning models, including Logistic Regression, Random Forest, Support Vector Machine, and Neural Network, is presented in Table 5.2.1, revealing insights into their respective accuracies for Bangla book genre classification. The findings showcase that the Random Forest model achieved the highest training accuracy of 95%, but its test accuracy was lower at 64%. This indicates that the Random Forest model may have overfit the training data, resulting in a decreased performance on the unseen test set. In contrast, the Neural Network model had a lower training accuracy of 87% but performed better on the test set with an accuracy of 78%. This suggests that the Neural Network model was able to generalize better and capture more robust features for Bangla book genre classification. The Logistic Regression model had a training accuracy of 89% and a test accuracy of 76%, while the Support Vector Machine model had a training accuracy of 94% and a test accuracy of 72%. These results demonstrate a more balanced performance between the training and test sets, indicating a moderate level of generalization for these models. The detailed accuracy metrics provided in Table 4.2.1 offer a valuable reference point for understanding the comparative strengths and weaknesses of these models, contributing to the ongoing discourse in the optimization of machine learning methodologies for Bangla book genre classification.

Table 5.2.2: Precision, Recall, F1-Score, and Support (n) for Logistic Regression Model (depending on the number of pictures, n=numbers).

Logistic Regression				
Class	Precision	Recall	F1-Score	Support
Computer and Technology	0.85	0.86	0.85	95
History	0.71	0.71	0.71	877
Islamic Literature	0.92	0.93	0.92	621
Romance Novel	0.74	0.77	0.76	305
Science Fiction	0.75	0.76	0.75	240
Politics	0.67	0.65	0.66	763
Accuracy			0.76	2901
Macro Avg	0.77	0.78	0.78	2901
Weighted Avg	0.75	0.76	0.76	2901

Table 5.2.3: Precision, Recall, F1-Score, and Support (n) Random Forest Model (depending on the number of pictures, n=numbers).

Random Forest				
Class	Precision	Recall	F1-Score	Support
Computer and Technology	1.00	0.38	0.55	95
History	0.54	0.72	0.62	877
Islamic Literature	0.92	0.84	0.88	621
Romance Novel	0.75	0.48	0.59	305
Science Fiction	0.82	0.42	0.56	240
Politics	0.53	0.56	0.55	763
Accuracy			0.64	2901
Macro Avg	0.76	0.57	0.62	2901
Weighted Avg	0.68	0.64	0.64	2901

Table 5.2.4: Precision, Recall, F1-Score, and Support (n) Support Vector Machine Model (depending on the number of pictures, n=numbers).

Support Vector Machine				
Class	Precision	Recall	F1-Score	Support
Computer and Technology	0.85	0.87	0.86	95
History	0.63	0.66	0.65	877
Islamic Literature	0.93	0.93	0.93	621
Romance Novel	0.75	0.76	0.75	305
Science Fiction	0.77	0.72	0.75	240
Politics	0.62	0.59	0.61	763
Accuracy			0.72	2901
Macro Avg	0.76	0.76	0.76	2901
Weighted Avg	0.72	0.72	0.72	2901

Table 5.2.5: Precision, Recall, F1-Score, and Support (n) Neural Network Model (depending on the number of pictures, n=numbers).

Neural Network				
Class	Precision	Recall	F1-Score	Support
Computer and Technology	0.80	0.83	0.81	95
History	0.75	0.74	0.74	877
Islamic Literature	0.91	0.95	0.93	621
Romance Novel	0.71	0.85	0.77	305
Science Fiction	0.79	0.71	0.75	240
Politics	0.73	0.68	0.70	763
Accuracy			0.78	2901
Macro Avg	0.78	0.79	0.79	2901
Weighted Avg	0.78	0.78	0.78	2901

The tables 5.2.2, 5.2.3, 5.2.4, and 5.2.5 provide a detailed comparison of the performance metrics for four distinct machine learning and deep learning models: Logistic Regression, Random Forest, Support Vector Machine, and Neural Network, specifically in the context of Bangla book genre classification. These metrics include precision, recall, F1-score, and support values for each class, allowing for a comprehensive evaluation of each model's strengths and weaknesses.

The Logistic Regression model demonstrates overall balanced performance with a notable accuracy of 76%. It achieves high precision, recall, and F1-scores across various classes, indicating its effectiveness in accurately classifying Bangla book genres. For example, it shows a precision of 0.85, recall of 0.86, and F1-score of 0.85 for the "Computer and Technology" genre, with similarly robust metrics across other genres such as History (precision of 0.71, recall of 0.71, and F1-score of 0.71) and Islamic Literature (precision of 0.92, recall of 0.93, and F1-score of 0.92).

The Random Forest model, while part of the comprehensive comparison, yields a comparatively lower identification rate with an overall accuracy of 64%. It shows significant variability in performance across different classes, such as a high precision of 1.00 for "Computer and Technology" but a lower recall of 0.38, resulting in an F1-score of 0.55. Other genres like Islamic Literature also show a high precision of 0.92 but a lower recall of 0.84, leading to an F1-score of 0.88.

The Support Vector Machine model offers a balanced performance with an overall accuracy of 72%, demonstrating reasonable precision, recall, and F1-scores across all classes. It achieves a precision of 0.85, recall of 0.87, and F1-score of 0.86 for "Computer and Technology." For other genres, such as Islamic Literature, it maintains a precision of 0.93, recall of 0.93, and F1-score of 0.93, indicating consistent and reliable performance.

The Neural Network model, however, shows the highest overall accuracy at 78%, suggesting superior performance in general classification tasks. This model excels in certain genres, particularly Islamic Literature, with a precision of 0.91, recall of 0.95, and F1-score of 0.93. It also performs well in History (precision of 0.75, recall of 0.74, and F1-score of 0.74) and Romance Novel (precision of 0.71, recall of 0.85, and F1-score of 0.77).

Among the four models compared, the Neural Network model stands out as the best performer due to its highest overall accuracy and slightly better macro averages for precision, recall, and F1-score. This makes it the most reliable choice for Bangla book genre classification, effectively handling a diverse set of genres with varying levels of complexity.

Training and Validation Accuracy and loss of Neural Network:

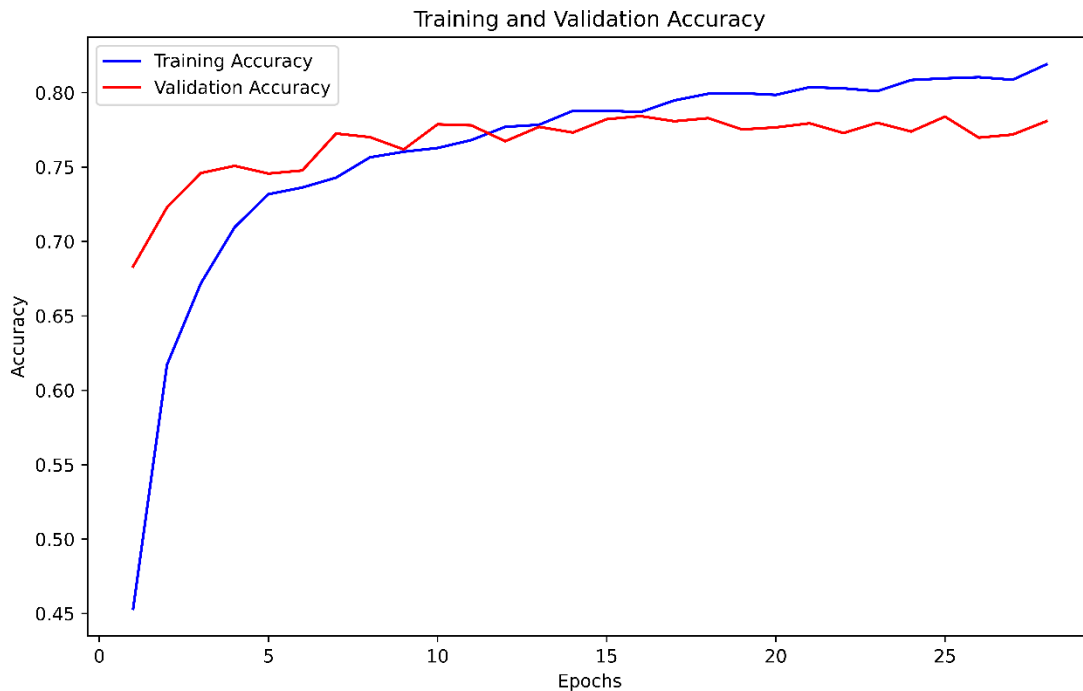


Figure 5.2.1: Training and validation accuracy over the epochs (Neural Network).

The graph shows the training and validation accuracy of a neural network model over 28 epochs. The x-axis represents the number of epochs, which indicates the number of complete passes through the training dataset. The y-axis represents the accuracy, which is a measure of how well the model is performing, with values ranging from 0 to 1 or 0% to 100%. The blue line represents the training accuracy, showing how well the model is performing on the training data over time. The red line represents the validation accuracy, showing how well the model is performing on unseen validation data. Initially, both accuracies increase rapidly, but after about 10 epochs, the validation accuracy begins to fluctuate around 0.75 to 0.78 while the training accuracy continues to rise more steadily, reaching about 0.82. This suggests that the model is learning and generalizing well, but there may be slight overfitting as the training accuracy is higher than the validation accuracy towards the end of the training period.

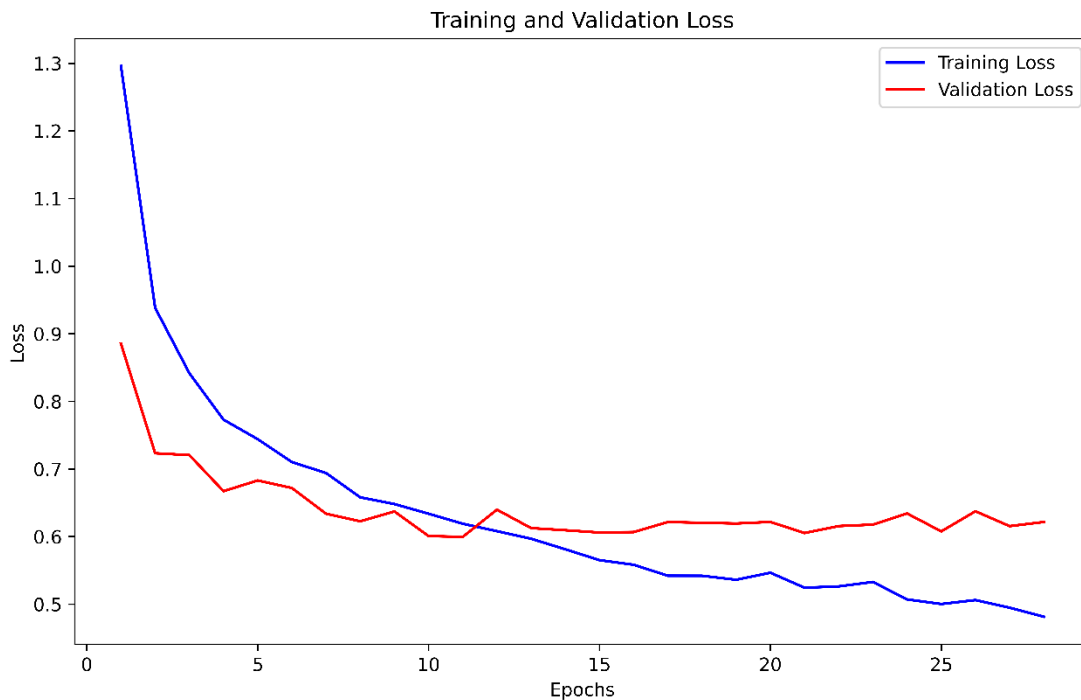


Figure 5.2.2: Training and validation loss over the epochs (Neural Network).

The graph depicts the training and validation loss of a neural network model over 28 epochs. The x-axis represents the number of epochs, indicating the number of complete passes through the training dataset. The y-axis represents the loss, a measure of how well the model's predictions match the actual data, with lower values indicating better performance. The blue line represents the training loss, showing how the model's performance on the training data improves over time. The red line represents the validation loss, indicating how well the model generalizes to unseen validation data. Initially, both losses decrease rapidly, indicating that the model is learning. The training loss continues to decrease steadily, while the validation loss decreases but starts to plateau around 0.6 after about 10 epochs, with minor fluctuations. This pattern suggests that the model is effectively learning and generalizing, though the decreasing training loss combined with the plateauing validation loss hints at potential overfitting as training progresses.

Confusion Matrix after Models:

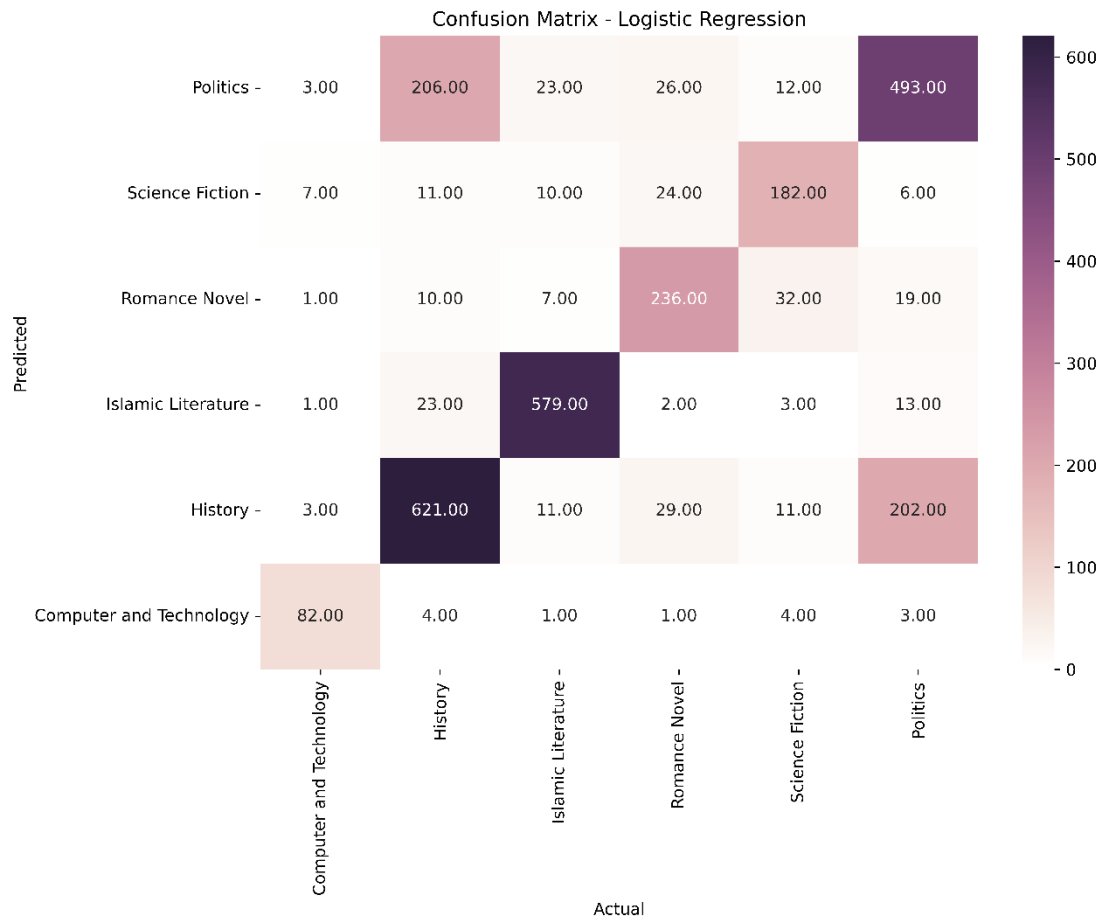


Figure 5.2.3: CM after Logistic Regression Model



Figure 5.2.4: CM after Random Forest model



Figure 5.2.5: CM after Support Vector Machine model

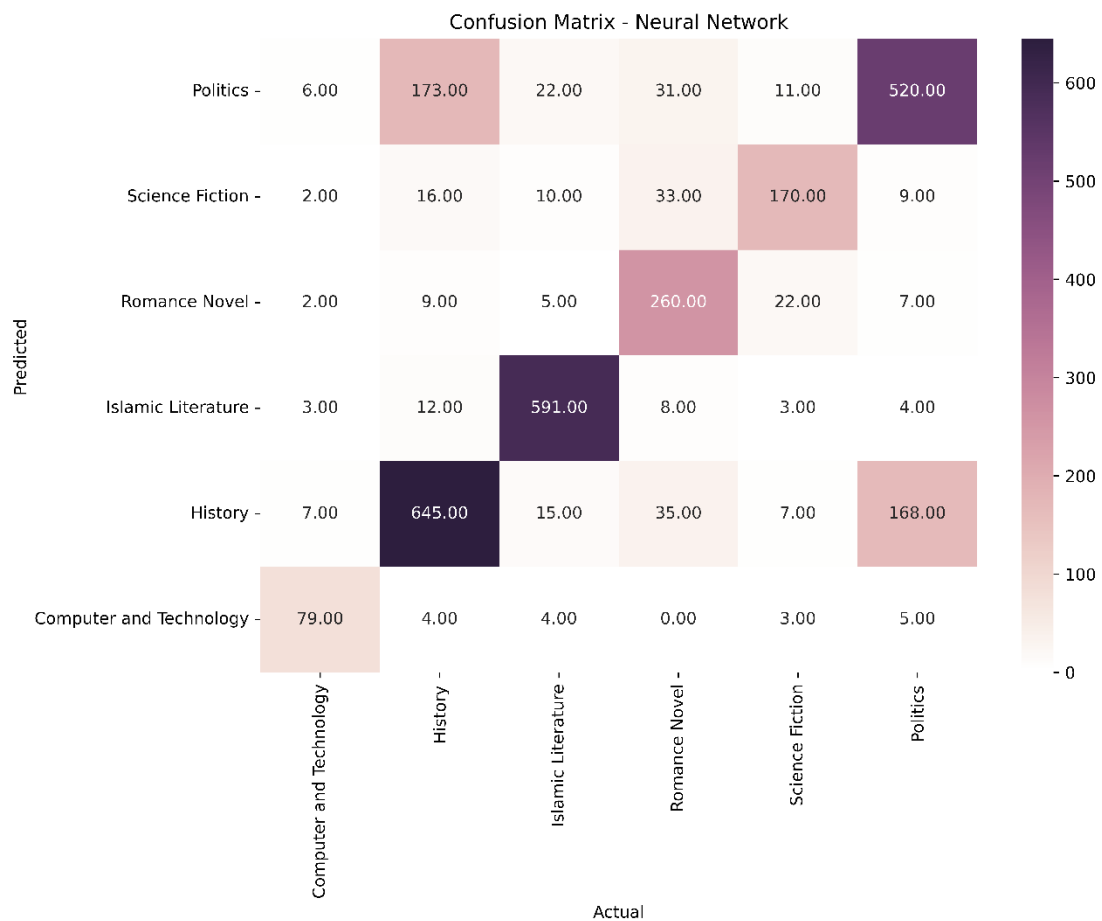


Figure 5.2.6: CM after Neural Network model

Overall, the Neural Network model stands out as the best performer in Bangla book genre classification, achieving an overall accuracy of 78%. It demonstrates robust performance across various genres, with notable precision, recall, and F1-scores. The model's ability to generalize well and capture more robust features is evident in its higher test accuracy compared to the training accuracy. While the other models, such as Logistic Regression and Support Vector Machine, show balanced performance, the Neural Network model's superior performance makes it the most reliable choice for this task. However, it is essential to note that the performance of the models may vary depending on the dataset they are trained on, and other metrics such as precision, recall, and F1-score can provide a more comprehensive understanding of their performance. The confusion matrix provides a helpful visual representation of the model's performance, highlighting both true positives and false negatives, which can inform strategies to improve the model's accuracy.

5.4 Comparative Analysis

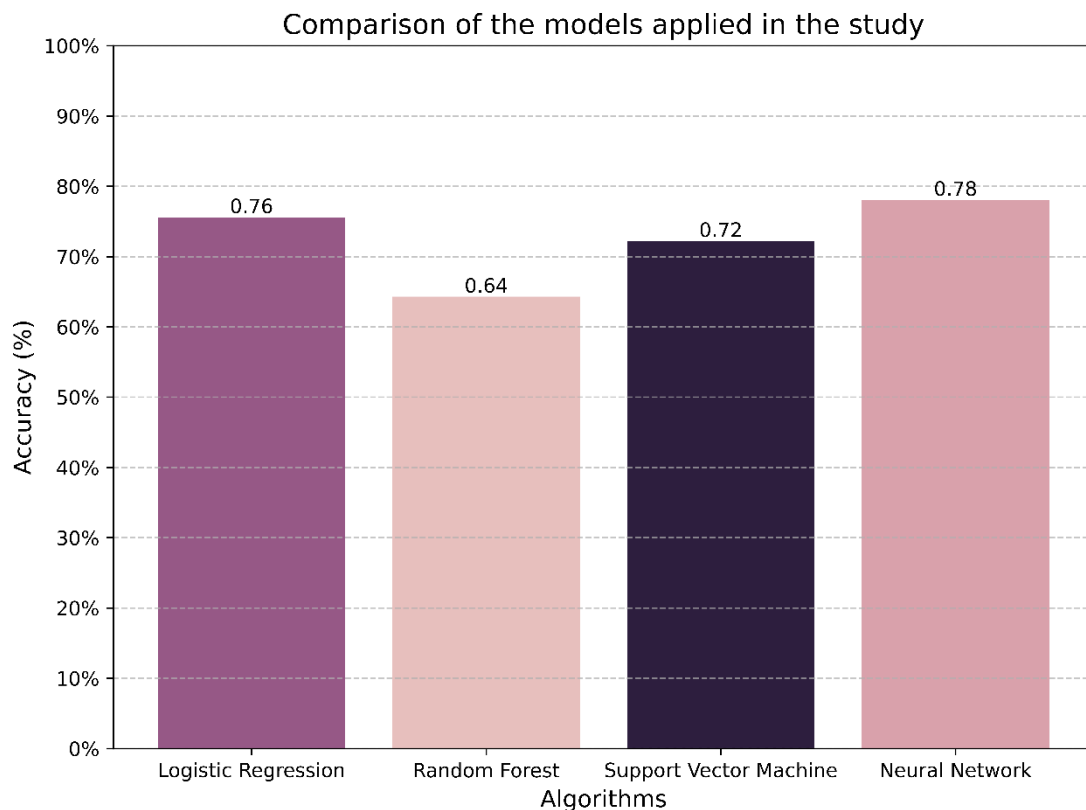


Figure 5.3.1: Accuracy comparison among Machine Learning and Deep Learning models.

This comprehensive study delves into the efficacy of various machine learning and deep learning models in classifying Bangla book genres. The research leverages a substantial dataset, employing sophisticated techniques to compare the performance of four distinct models: Logistic Regression, Random Forest, Support Vector Machine (SVM), and Neural Network. Considering the procedure for determining classification accuracy, the study acknowledges that it is a crucial indicator of each model's effectiveness. The results demonstrate that, with a profound accuracy score of 78%, the Neural Network model comes out on top. This improved performance highlights how well the model generalizes and captures strong characteristics necessary for classifying between different Bangla book genres. The research sheds light on the nuanced impacts of overfitting, particularly evident in the Random Forest model, which, despite a high training accuracy of 95%, suffers from a significant drop to 64% in model accuracy. In contrast, the Neural Network demonstrates a more balanced and effective learning process.

However, the research shows that the less complex Logistic Regression model also does well, with a 76% model accuracy. This indicates that even less complex models

can be effective in certain classification tasks, providing a good balance between interpretability and performance. The Support Vector Machine, with a model accuracy of 72%, showcases its robustness in handling the complexity of the dataset. The study emphasizes how well-suited Neural Networks are for text categorization tasks and how well-handled vast and complicated datasets are. It also shows that more straightforward models, such as SVM and logistic regression, can offer a significant level of accuracy, especially when computational resources are a consideration.

Essentially, this study offers a detailed examination of several deep learning and machine learning models for the classification of Bangla book genres, providing insightful information about the advantages and disadvantages of each approach. The results indicate that while Neural Networks are the most accurate overall, a balanced strategy that takes into account various model designs may be helpful, depending on the particular needs and limitations of the classification task.

5.5 Summary

This chapter provided an evaluation of four machine learning and deep learning models Logistic Regression, Random Forest, Support Vector Machine, and Neural Network for the classification of Bangla book genres using book cover and title. Performance was assessed using key metrics such as accuracy, precision, recall, and F1-score, supplemented by confusion matrices for a visual representation of classification errors. The Neural Network model demonstrated the highest overall test accuracy, indicating its strong generalization capabilities and robustness in capturing relevant features for genre classification. In contrast, the Random Forest model, despite achieving the highest training accuracy, exhibited signs of overfitting, resulting in a lower test accuracy. Logistic Regression and SVM models showed balanced performance, with moderate accuracy and consistent metrics across various genres. The comparative analysis highlighted the significance of choosing the appropriate model based on specific requirements and constraints of the classification task. While simpler models like Logistic Regression and SVM offer balanced performance, the Neural Network model stands out as the most reliable choice for this task. Overall, the insights gained from this comprehensive evaluation will guide future improvements and inform the selection of the most suitable models for real-world deployment in Bangla book genre classification.

CHAPTER 6

IMPACT ON SOCIETY, ENVIRONMENT AND SUSTAINABILITY

6.1 Impact on Life

The "Classification of Bangla Book Genres using Book Cover and Title" project has a profound impact on various aspects of life, enhancing the reader experience and contributing to cultural and educational growth.

- **Enhanced Reader Experience:** By improving book discovery and providing personalized recommendations, the project enriches the reading experience, making it more engaging and tailored to individual preferences. This can lead to increased reading satisfaction and a deeper appreciation for literature.
- **Educational Benefits:** Easier access to a wide range of genres and titles can support educational goals, promoting literacy and encouraging lifelong learning. Readers, especially students, can find books that match their academic interests and personal tastes more efficiently.
- **Cultural Preservation:** By categorizing and promoting Bangla literature, the project helps preserve and promote the rich cultural heritage of the Bengali-speaking community. This can foster a greater sense of cultural identity and pride.
- **Increased Accessibility:** The project makes it easier for readers to find books that are relevant to their interests, which can be particularly beneficial for those new to a particular genre or looking for books that match specific tastes. This can make literature more accessible to a wider audience, including marginalized communities.
- **Support for Authors and Publishers:** Improved classification and discovery systems can lead to increased visibility for authors and publishers, potentially boosting book sales and encouraging the production of more diverse and high-quality content.
- **Mental Health and Well-being:** Reading is known to have numerous mental health benefits, including stress reduction, increased empathy, and improved cognitive function. By facilitating access to a variety of books.

Overall, the project has the potential to positively impact the lives of readers, authors, and the broader literary community by making literature more accessible, engaging, and diverse.

6.2 Impact on Society & Environment

The project "Classification of Bangla Book Genres using Book Cover and Title" has the potential to significantly impact society in several ways.

- **Improved Book Discovery:** The system can help readers discover new books that match their interests and preferences, leading to a more engaging and personalized reading experience.
- **Enhanced Book Recommendations:** The system can provide personalized book recommendations to readers, which can lead to increased book sales and a more vibrant book culture.
- **Increased Accessibility:** The system can make it easier for readers to find books that are relevant to their interests, which can be particularly beneficial for readers who are new to a particular genre or are looking for books that match their specific tastes.
- **Improved Book Classification:** The system can help improve book classification by providing a more accurate and comprehensive classification system, which can be beneficial for authors, publishers, and readers alike.
- **Increased Efficiency:** The system can help reduce the time and effort required to classify books, which can be beneficial for authors, publishers, and readers who are looking for a more efficient way to discover and classify books.
- **Increased Book Sales:** The system can help increase book sales by providing personalized book recommendations to readers, which can lead to increased sales and a more vibrant book culture.

By providing personalized book recommendations and improving book classification, the system can have a significant impact on society by making it easier for readers to discover and enjoy books that match their interests and preferences.

Impact on Environment:

The "Classification of Bangla Book Genres using Book Cover and Title" project significantly impacts the environment by encouraging digital book discovery and classification, which reduces the demand for physical book production and its associated environmental impact. By streamlining the publishing process with an automated classification system, the project promotes resource efficiency and reduces waste from traditional categorization and distribution methods. Additionally, the promotion of digital libraries minimizes the need for physical transportation, storage, and disposal of books, thereby reducing the overall environmental footprint of the publishing industry. This shift supports sustainable reading habits by facilitating access to digital books and contributing to a significant reduction in paper waste, aligning with broader environmental sustainability goals.

Moreover, the project emphasizes collaboration with eco-conscious stakeholders, such as environmentally conscious publishers and green initiatives, to collectively minimize ecological footprints. These partnerships can drive the use of recycled paper for printed books and eco-friendly packaging solutions. Lastly, by raising environmental awareness among readers and industry players about the impacts of their choices, the project fosters a culture of responsibility and promotes sustainable literature consumption patterns. Overall, the project plays a vital role in advancing environmental sustainability within the publishing industry.

6.3 Ethical Aspects

The research on "Classification of Bangla Book Genres using Book Cover and Title" is underpinned by a commitment to ethical considerations throughout its methodology.

1. **Data Collection and Use:** The project involves collecting and using data from various sources, including book cover images and titles. It is essential to ensure that the data is collected and used in an ethical manner, respecting the privacy and intellectual property rights of the authors and publishers.
2. **Intellectual Property:** The project involves using pre-trained models and datasets, which may be protected by intellectual property rights. It is essential to ensure that the project complies with all applicable laws and regulations regarding intellectual property.
3. **Transparency and Accountability:** The project should be transparent about its methods, data, and results. It is essential to ensure that the project is accountable for its actions and decisions, and that any errors or biases are addressed in a

timely and effective manner.

4. **Transparency in Model Development:** The project should be transparent about the development and training of the classification model, including the selection of algorithms, hyperparameters, and datasets.
5. **Continuous Monitoring and Improvement:** The project should continuously monitor and improve its ethical practices, ensuring that the project remains ethical and responsible throughout its lifecycle.

By considering these ethical aspects, the project can ensure that it is conducted in an ethical and responsible manner, respecting the privacy, intellectual property, and dignity of all individuals involved.

6.4 Sustainability Plan

The sustainability plan for the project "Classification of Bangla Book Genres using Book Cover and Title" aims to ensure the long-term viability and impact of the developed system. It outlines key strategies to maintain relevance, adaptability, and continued contribution to the book classification domain.

- **Continuous Data Acquisition and Curation:** Implement a mechanism for frequently updating the dataset with new book cover images and titles. Partner with publishers, libraries, and other stakeholders to acquire diverse and representative data.
- **Model Maintenance and Improvement:** Continuously monitor the performance of machine learning models, identifying areas for improvement. Use feedback from users, industry experts, and domain specialists to enhance model accuracy and robustness. Incorporate new advancements in deep learning and natural language processing to improve classification capabilities.
- **Scalability and Infrastructure Optimization:** Design the system architecture to be scalable, allowing for efficient handling of increasing data volumes and user demands. Leverage cloud-based infrastructure and serverless computing solutions to ensure cost-effective and resource-efficient deployment.

- **User Engagement and Feedback Incorporation:** Develop a user-friendly interface or API that allows seamless integration of the book classification system into various applications and platforms. Regularly analyze user feedback and incorporate it into the system's development roadmap to enhance user experience and satisfaction.
- **Partnerships and Collaborations:** Engage with publishers, libraries, and other stakeholders in the book industry to foster partnerships and explore opportunities for integration and co-development. Explore opportunities for open-sourcing the project's codebase or providing it as a service to the broader community, enabling wider adoption and contributions.

By implementing this comprehensive sustainability plan, the project "Classification of Bangla Book Genres using Book Cover and Title" can maintain its relevance, adaptability, and impact in the book classification domain, contributing to the ongoing advancement and accessibility of Bangla literature.

6.5 Summary

The "Classification of Bangla Book Genres using Book Cover and Title" project offers significant benefits across various domains, including societal, environmental, and ethical aspects. The project enhances the reader experience through improved book discovery and personalized recommendations, supports educational goals by promoting literacy, and aids in cultural preservation by categorizing and promoting Bangla literature. It also increases accessibility for readers, supports authors and publishers by improving visibility and sales, and contributes to mental health and well-being by facilitating access to a diverse range of books. The project positively influences society by enhancing book discovery and classification, leading to a more vibrant book culture and increased efficiency for authors, publishers, and readers. Environmentally, it promotes digital book usage, reducing the demand for physical book production and its associated resource consumption. The project's emphasis on digital libraries and partnerships with eco-conscious stakeholders further minimizes the environmental footprint of the publishing industry, promoting sustainable reading habits and environmental awareness. The project adheres to ethical standards in data collection and use, respects intellectual property rights, and maintains transparency and accountability in its methodologies. Continuous monitoring and improvement ensure that the project remains ethical and responsible throughout its lifecycle. The

sustainability plan outlines strategies for maintaining the system's relevance and adaptability. Continuous data acquisition and curation, model maintenance and improvement, scalability optimization, user engagement, and partnerships with stakeholders are key components. These strategies ensure the long-term viability and impact of the project, contributing to the ongoing advancement and accessibility of Bangla literature. However, the "Classification of Bangla Book Genres using Book Cover and Title" project demonstrates a comprehensive approach to enhancing literature accessibility and engagement while promoting sustainability and ethical practices within the publishing industry.

CHAPTER 7

CONCLUSION AND FUTURE WORK

7.1 Conclusions

This comprehensive research on the classification of Bangla book genres using book cover images and titles successfully achieved its primary objective: to develop an effective machine learning and deep learning-based system capable of accurately categorizing Bangla books into distinct genres such as Computer and Technology, History, Islamic Literature, Romance Novel, Science Fiction, and Politics. The study began by curating and meticulously organizing a substantial dataset of Bangla book cover images and titles. This involved extracting visual features from the book covers using state-of-the-art deep learning models and preprocessing the book titles to obtain meaningful textual representations. The core research focused on a comparative analysis of four distinct machine learning and deep learning models: Logistic Regression, Random Forest, Support Vector Machine, and Neural Network. The findings revealed that the Neural Network model emerged as the most effective, achieving the highest overall accuracy of 78%. This model demonstrated superior generalization capabilities and was particularly adept at capturing robust features essential for distinguishing between the diverse Bangla book genres. In contrast, the Random Forest model, despite its high training accuracy of 95%, suffered from significant overfitting, with a model accuracy drop to 64%. The Logistic Regression and SVM models demonstrated a more balanced performance, with test accuracies of 76% and 72%, respectively, indicating moderate generalization capabilities.

Comprehensive performance measures, such as accuracy, retention, and F1-score for every category, also underscored the advantages and drawbacks of the different models. Certain genres, like Islamic literature, demonstrated the Neural Network model's proficiency in managing intricate and varied information. This study highlights the tremendous promise of deep learning and machine learning methods for accurately recognizing Bangla literature genres. The comparison study shed light on the advantages and disadvantages of each strategy, with the Neural Network model showing the most promise in terms of overall accuracy and generalization. The research moreover underscores the paramount significance of many pivotal elements in guaranteeing the enduring prosperity and influence of the established categorization scheme. These include continuous data acquisition and curation, model maintenance and improvement, scalability, user engagement, and collaboration with eco-conscious stakeholders.

Overall, this study makes a significant contribution to the current discussion on the best machine-learning and deep-learning techniques for genre classification in Bangla literature. The knowledge gained from this research can help direct future developments, resulting in the creation of categorization systems that are more complex, precise, and easy to use. Moreover, the findings highlight the importance of integrating technological innovation with strategic planning to achieve sustainable and impactful results in digital literature management.

7.2 Further Suggested Works

This research has established a strong foundation for Bangla book genre classification using machine learning and deep learning techniques, revealing several promising directions for future investigation. Expanding the dataset to include a wider array of sources beyond Rokomari is essential to reduce bias and enhance the generalizability of the models. Additionally, utilizing higher resolution images could improve the extraction of visual features, providing a more detailed basis for genre differentiation. The application of advanced natural language processing techniques, such as transformer-based models like BERT and GPT, could capture deeper semantic relationships within the text and further improve classification accuracy. Moreover, the integration of contextual information and metadata, including author and publication year, could enhance the model's performance. Future research should also explore hybrid models that effectively combine visual and textual data through multi-modal learning techniques. Experimenting with ensemble methods could leverage the strengths of different models, potentially leading to significant performance gains. Implementing real-time classification systems will present technical challenges, necessitating optimization for speed and scalability to ensure practical deployment. Integrating these systems into existing digital libraries and e-commerce platforms is crucial for seamless operation. The ultimate goal of this research is to enable the development of user-friendly applications or APIs that integrate the genre classification system, allowing seamless integration into various platforms and enhancing the accessibility of Bangla literature. Future studies should focus on designing and implementing such applications, gathering user feedback, and continuously improving the system based on user needs and preferences.

Extending these methodologies to other languages and cultural contexts will also be beneficial, contributing to a more global understanding of literature classification using machine learning and deep learning. Comparative studies across different linguistic datasets could reveal unique challenges and opportunities, informing more generalized

approaches to genre classification. By addressing these implications and exploring new avenues for further research, the project "Classification of Bangla Book Genres using Book Cover and Title" can continue to evolve, contributing to the advancement and accessibility of Bangla literature in the digital age. The insights gained from this study serve as a solid foundation for future researchers and practitioners to build upon, driving innovation and progress in the field of book genre classification.

7.3 Limitations

Limitations

Despite the successful implementation and promising results of the "Classification of Bangla Book Genres using Book Cover and Title" project, several limitations were identified throughout the research process:

1. **Dataset Constraints:** The dataset used in this study, although comprehensive with over 15,000 Bangla book covers and titles, was collected from a single e-commerce platform, Rokomari. This could limit the generalizability of the model to other datasets or platforms with different book collections and genre distributions.
2. **Model Generalization:** While the neural network model achieved the highest accuracy, the model's performance may vary when applied to different or more diverse datasets. The training data was specific to the Bangla language and may not fully capture the nuances of book covers and titles from other languages or cultures.
3. **Feature Extraction Limitations:** The use of pre-trained models like Xception and BanglaBERT, while effective, may not fully leverage domain-specific features unique to Bangla book covers and titles. Custom models trained specifically on a larger and more varied Bangla literature dataset might improve classification accuracy.
4. **Computational Resources:** The deep learning models used in this study required significant computational power for both training and inference. This might limit the practical application of the model in resource-constrained environments or for real-time genre classification tasks.

By acknowledging these limitations, we aim to provide a transparent and comprehensive overview of the research, ensuring the integrity and credibility of the findings.

References:

- [1] Biradar, G. R., Raagini, J. M., Varier, A., & Sudhir, M. (2019, June). Classification of book genres using book cover and title. In 2019 IEEE International Conference on Intelligent Systems and Green Technology (ICISGT) (pp. 72-723). IEEE.
- [2] Ozsarfati, E., Sahin, E., Saul, C. J., & Yilmaz, A. (2019, February). Book genre classification based on titles with comparative machine learning algorithms. In 2019 IEEE 4th International Conference on Computer and Communication Systems (ICCCS) (pp. 14-20). IEEE
- [3] Gupta, S., Agarwal, M., & Jain, S. (2019, January). Automated genre classification of books using machine learning and natural language processing. In 2019 9th International Conference on Cloud Computing, Data Science & Engineering (Confluence) (pp. 269-272). IEEE. [4] Hertzog, M.A., Pozehl, B. and Duncan, K., 2010. Cluster analysis of symptom occurrence to identify subgroups of heart failure patients: a pilot study. *Journal of Cardiovascular Nursing*, 25(4), pp.273-283.
- [5] Panigrahi, P., & Roy, B. (2014). Developing a Classification System for Bengali Literature: Some Thoughts. *Advances in Library & Information Science*, 2(1), 1-8.
- [6] Scofield, C., Silva, M. O., de Melo-Gomes, L., & Moro, M. M. (2022, March). Book genre classification based on reviews of portuguese-language literature. In *International Conference on Computational Processing of the Portuguese Language* (pp. 188-197). Cham: Springer International Publishing.
- [7] Rahman, A. (2012). Bangla Music Genre Classification.
- [8] Chiang, H., Ge, Y., & Wu, C. (2015). Classification of Book Genres By Cover and Title.
- [9] Dipu, N. M., Shohan, S. A., & Salam, K. M. A. (2021, December). Bangla optical character recognition (ocr) using deep learning based image classification algorithms. In 2021 24th International Conference on Computer and Information Technology (ICCIT) (pp. 1-5). IEEE
- [10] Pasha, S.T., Islam, A., Rahman, M.M., Ahmed, E., Foysal, M.F., & Alam, M.Z. (2023). Genre Classification of Bangla Poem Using Machine Learning and Deep Learning Techniques. 2023 IEEE World AI IoT Congress (AIIoT), 0236-0242.
- [11] 1. Ahmed, T., Alam, M. A., Paul, R. R., Hasan, M. T., & Rab, R. (2022, February). Machine Learning and Deep Learning Techniques For Genre Classification of Bangla Music. In 2022 International Conference on Advancement in Electrical and Electronic Engineering (ICAEEE) (pp. 1-6). IEEE.
- [12] Narendra, M., Rayudu, K. M., Sai, T. S., Rajshekar, K., & Lingala, V. (2020). A Survey on Book Genre Classification System using Machine Learning. *Mathematical Statistician and Engineering Applications*, 69(1), 147-160.
- [13] Chollet, F. (2017). Xception: Deep learning with depthwise separable convolutions. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 1251-1258).
- [14] Szegedy, C., Liu, W., Jia, Y., Sermanet, P., Reed, S., Anguelov, D., ... & Rabinovich, A. (2015). Going deeper with convolutions. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 1-9).
- [15] Bhattacharjee, A., Hasan, T., Ahmad, W. U., Samin, K., Islam, M. S., Iqbal, A., ... & Shahriyar, R. (2021). BanglaBERT: Language model pretraining and benchmarks for low-resource language understanding evaluation in Bangla. *arXiv preprint arXiv:2101.00204*.
- [16] Pasha, S. T., Islam, A., Rahman, M. M., Ahmed, E., Foysal, M. F. A., & Alam, M. Z. (2023, June). Genre Classification of Bangla Poem Using Machine Learning and Deep Learning Techniques. In 2023 IEEE World AI IoT Congress (AIIoT) (pp. 0236-0242). IEEE.
- [17] Afroze, A., Dutta, K., Sadik, S., Khanam, S., Rab, R., & Rahim, M. A. (2024). Empirical Text Analysis for Identifying the Genres of Bengali Literary Work. *Journal of Advances in Information Technology*, 15(5).

- [18] Al Mamun, M. A., Kadir, I., Rabby, A. S. A., & Al Azmi, A. (2019, November). Bangla music genre classification using neural network. In *2019 8th international conference system modeling and advancement in research trends (SMART)* (pp. 397-403). IEEE.
- [19] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner, "Gradient-based learning applied to document recognition," *Proceedings of the IEEE*, vol. 86, no. 11, pp. 2278–2324, 1998.
- [20] J. Zujovic, L. Gandy, S. Friedman, B. Pardo, and T. N. Pappas, "Classifying paintings by artistic genre: An analysis of features & classifiers," in *Multimedia Signal Processing, 2009. MMSP'09. IEEE International Workshop on*. IEEE, 2009, pp. 1–5.
- [23] J. Worsham and J. Kalita, "Genre identification and the compositional effect of genre in literature," in *Proceedings of the 27th International Conference on Computational Linguistics*, 2018, pp. 1963–1973.
- [22] P. Buczkowski, A. Sobkowicz, and M. Kozłowski, "Deep learning approaches towards book covers classification." in *ICPRAM*, 2018, pp. 309–316.
- [23] C. Szegedy, W. Liu, Y. Jia, P. Sermanet, S. Reed, D. Anguelov, D. Erhan, V. Vanhoucke, and A. Rabinovich, "Going deeper with convolutions," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2015, pp. 1–9.
- [24] J. Pennington, R. Socher, and C. Manning, "Glove: Global vectors for word representation," in *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*, 2014, pp. 1532–1543.
- [25] Chollet, F. (2017). Xception: Deep learning with depthwise separable convolutions. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 1251-1258).

Appendix A

Course Outcomes, Complex Engineering Problems (EP) and Complex Engineering Activities (EA) Addressing

**Title: CLASSIFICATION OF BANGLA BOOK GENRES USING BOOK
COVER AND TITLE**

Student ID: 201-15-13593

CO Description for FYDP

CO	CO Descriptions	PO
Phase -I		
CO1	Integrate recently gained and previously acquired knowledge to identify a Bangla book classification problem 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 these 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 CO (1 to 8), Knowledge Profile (K), Attainment of Complex Engineering Problems (EP), and Attainment of Complex Engineering Activities (EA)

Addressing CO (1 to 8), Knowledge Profile (K), 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, CO3, CO5, CO6, CO7 and CO8	This project demonstrates fundamental engineering (K3) principles by employing machine learning and deep learning techniques, transfer learning architectures, and natural language processing for genre classification tasks. The project demonstrates specialist knowledge (K4) by conducting a comparative analysis of models such as Logistic Regression, Random Forest, Support Vector Machine, and Neural Network, enhancing the accuracy of Bangla book genre classification based on book covers and titles, crucial for improving information retrieval, content filtering, and book recommendation systems.	Page no: [20-22] Section: [4.2] Page no: [36-37] Section: [5.4]
				The project applies engineering practice & design (K5) by the figure of process of experiments. The project addresses engineering	Page no: [14-15]

				practice & technology (K6) by employing a variety of machine learning and deep learning models, including Logistic Regression, Random Forest, Support Vector Machine, and Neural Network, as well as transfer learning techniques.	Section: [3.3] Page no: [20-22] Section: [4.2]
				This project ensures to K8 (Research Literature) by synthesizing insights from recent studies, to advance the classification of Bangla book genres using book cover and title, showcasing a comprehensive understanding of current methodologies.	Page no: [7-9] Section: [2.2, 2.3]
2.	EP2: Range of Conflicting Requirements	Yes	CO2, and CO7	This project addresses EP-2 by recognizing the hurdles in classifying Bangla book genres, including facing challenges to collect high quality images, relatively small (240 × 320 pixels). Thought this size is considered adequate for feature extraction, but fine features in book covers might not be fully captured.	Page no: [4-5] Section: [1.5]

3.	EP3: Depth of analysis required	Yes	CO2, and CO6	This project addresses EP-3 by finding a clear-cut solution for project. Which is challenging, primarily because of the image and title processing and feature extraction. To process and extract features from images and title, we have to use a pretrained deep learning model and natural language processing model then train the data by using a machine learning algorithm and deep learning algorithm. We conducted an in-depth analysis to carefully choose a specific deep learning pre-trained model and algorithm from numerous alternatives.	Page no: [1-3] Section: [1.1] Page no: [20-22] Section: [4.2]
4.	EP4: Familiarity of Issues	Yes	CO8	To understand and study book Bangla genre classification for our project, we need to look into various aspects we were not very familiar with. Specifically, we have to figure out how to process and extract data from images and texts which indicates EP-4.	Page no: [7-8] Section: [2.2] Page no: [1-3] Section: [1.1]
5.	EP5: Extends of application codes	No	CO5	N/A	N/A

6.	EP6: Extends of stakeholders involved and	No	CO8	N/A	N/A
7.	EP7: Interdependence	Yes	CO5	This project addresses EP-7 by demonstrating the interdependence between various components of the classification system for Bangla book genres. The success of the project relies on the seamless integration of multiple elements such as data collection and preprocessing, feature extraction, model training, evaluation, system performance.	Page no: [11-15] Section: [3.2, 3.3] Page no: [20-22] Section: [4.2]

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	In order to ensure systematic research and make progress toward the goal of classifying Bangla book genres using book covers and titles through deep learning techniques, our project makes use of a variety of resources, including GPUs, deep learning frameworks, annotated datasets, and ethical considerations.	Page no: [16-17, 40-41] Section: [3.4, 6.3]

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		Our project addresses EA-4 by considering the consequences for society and the environment. By developing an effective system for classifying Bangla book genres, we aim to enhance the accessibility and discoverability of Bangla literature, fostering cultural preservation and promoting literacy. Additionally, the use of digital classification methods reduces the need for physical resources, thereby minimizing the environmental impact associated with traditional classification and cataloging processes.	Page no: [38-40] Section: [6.1, 6.2]

5.	EA-5: Familiarity	Yes		Our project addresses EA-5 by demonstrating a high level of familiarity with the latest advancements in machine learning , deep learning and natural language processing. The project leverages state-of-the-art deep learning frameworks and methodologies, incorporating best practices from current research to effectively classify Bangla book genres. This familiarity ensures that the project is built upon a solid foundation of existing knowledge and techniques, allowing for the successful implementation and potential future improvements of the system.	Page no: [7-9] Section: [2.1, 2.2, 2.3]
----	-------------------	-----	--	---	--

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

SN	COs	Attainment	Justification	References
1	CO4	Yes	Through the integration of efficient project management and financial control, this project tackles CO4 by guaranteeing careful planning, resource allocation, and budget estimate for the best possible resource utilization throughout the study lifetime.	Page no: [17-19] Section: [3.5]
2	CO9	Yes	The By making sure that data is handled properly, getting the required rights to use the data, and openly recording the study process, the project displays commitment to ethical norms. This complies with CO9 and guarantees ethical application of cutting-edge machine learning technology for the purpose of societal well-being and responsible information distribution.	Page no: [40-41] Section: [6.3]
3	CO10	No	N/A	N/A
4	CO12	Yes	The extensive data collecting, careful data pretreatment, methodical methodology development, and in-	Page no: [11-15, 20-22, 26-

			depth experimental findings and analysis all demonstrate the project's commitment to continuous learning (CO12) and adaptability within the ever-changing technological context. This demonstrates a dedication to keeping up with the latest developments and perfecting methods to tackle contemporary difficulties in the machine learning and deep learning technologies used to classify Bangla book genres.	37] Section: [3.2, 3.3, 4.2, 5.2, 5.3]
--	--	--	--	--

FYDP Final Report

ORIGINALITY REPORT

13%

SIMILARITY INDEX

9%

INTERNET SOURCES

6%

PUBLICATIONS

2%

STUDENT PAPERS

PRIMARY SOURCES

1	dspace.daffodilvarsity.edu.bd:8080 Internet Source	4%
2	Ganeshprasad R Biradar, Raagini JM, Aravind Varier, Manisha Sudhir. "Classification of Book Genres using Book Cover and Title", 2019 IEEE International Conference on Intelligent Systems and Green Technology (ICISGT), 2019 Publication	2%
3	Submitted to Daffodil International University Student Paper	2%
4	Mosab. A. Hassan, Alaa. H. Ali, Atheer A. Sabri. "Recent Progress in Arabic Sign Language Recognition: Utilizing Convolutional Neural Networks (CNN)", BIO Web of Conferences, 2024 Publication	1%
5	doctorpenguin.com Internet Source	1%
6	www.journaltocs.ac.uk Internet Source	1%