

Prediction of Movie Box Office Success using YouTube Trailer Comments and Wikipedia Data

by

Tasin Mohammad
20341021

Maliha Binta Islam
20101587

Humayra Binte Jamal
20101591

A Thesis submitted to the Department of Computer Science and Engineering
in partial fulfillment of the requirements for the degree of
B.Sc. in Computer Science

Department of Computer Science and Engineering
Brac University
May 2024

© 2024. Brac University
All rights reserved.

Declaration

It is hereby declared that

1. The thesis submitted is our own original work while completing degree at Brac University.
2. The thesis does not contain material previously published or written by a third party, except where this is appropriately cited through full and accurate referencing.
3. The thesis does not contain material which has been accepted, or submitted, for any other degree or diploma at a university or other institution.
4. We have acknowledged all main sources of help.

Student's Full Name & Signature:

Tasin Mohammad

20341021

Maliha Binta Islam

20101587

Humayra Binte Jamal

20101591

Approval

The thesis titled “Prediction of Movie Box Office Success using YouTube Trailer Comments and Wikipedia Data” submitted by

1. Tasin Mohammad (20341021)
2. Maliha Binta Islam (20101587)
3. Humayra Binte Jamal (20101591)

Of Spring 2024 has been accepted as satisfactory in partial fulfillment of the requirement for the degree of B.Sc. in Computer Science on May 2024.

Examining Committee:

Supervisor:
(Member)

Najeefa Nikhat Choudhury

Lecturer
Department of Computer Science and Engineering
Brac University

Program Coordinator:
(Member)

Md. Golam Rabiul Alam, PhD

Lecturer
Department of Computer Science and Engineering
Brac University

Head of Department:
(Chair)

Sadia Hamid Kazi, PhD

Chairperson and Associate Professor
Department of Computer Science and Engineering
Brac University

Abstract

The movie industry is one of the most highly competitive and profitable businesses in the entertainment world, where predicting a movie's success or failure of a film is a difficult task. Along with the movie trailer, key factors such as - the popularity of actors and actresses, comments, ratings, marketing budgets, critical receptions, likes, and dislikes play a significant role in determining a movie's box office performance. Day by day, the use of social media is increasing. People express their opinions and feelings about anything on social media platforms such as Facebook, Instagram, Twitter, and YouTube. This research paper aims to predict movie box office success using opinions expressed in comments on movie trailers and Wikipedia data through the use of machine learning algorithms.

Keywords: Prediction, Movie box office success, YouTube trailer comments, Wikipedia data, Web scraping, Sentiment analysis, Machine learning, Natural Language Processing

Table of Contents

Declaration	i
Approval	ii
Abstract	iii
Table of Contents	iv
List of Figures	vi
List of Tables	vii
1 Introduction	1
1.1 Leveraging Online Data for Predicting Box Office Success	1
1.2 Problem Statement	1
1.3 Research Objectives	2
2 Literature Review	4
3 Dataset	9
3.1 Overview	9
3.2 Wikipedia Data	9
3.3 YouTube Trailer Comments	11
4 Data Analysis	12
4.1 Correlation	12
4.2 Average Box Office Revenue	15
4.3 Most Number of Successful Movies	16
4.4 Highest Average Budgets	16
4.5 Most Number of Movies Released	16
5 Data Preprocessing	18
5.1 Data Cleaning	18
5.2 Success Score Calculation	18
5.3 Data Encoding	19
5.4 Data Splitting	19
6 Methodology	21
6.1 Performance Metrics	21
6.2 Model Description	22

6.2.1	Sentiment Analysis Model	22
6.2.2	Machine Learning Models	22
6.3	Model Training	23
6.3.1	Random Forest Classifier Training	23
6.3.2	Decision Tree Classifier Training	24
6.3.3	K Nearest Neighbors Classifier Training	25
6.3.4	Logistic Regression Training	27
6.3.5	Support Vector Classifier Training	27
7	Results	29
7.1	Individual Results Analysis	29
7.1.1	Random Forest Classifier Results	29
7.1.2	Decision Tree Classifier Results	29
7.1.3	K Nearest Neighbors Classifier Results	30
7.1.4	Logistic Regression Results	30
7.1.5	Support Vector Classifier Results	31
7.2	Results Visualization	31
7.3	Combined Results Analysis	35
7.4	Comparative Results Analysis	35
8	Conclusion	36
8.1	Limitations	36
8.2	Future Work	37
	Bibliography	39

List of Figures

4.1	Scatter Plot of Budget, Runtime, Release Year and Box Office	14
4.2	Scatter Plot of Success Scores and Box Office	14
4.3	Scatter Plot of Success Scores and Box Office	14
4.4	Average Box Office Revenue by Release Year, Release Month, and MPAA Rating	15
4.5	Top 10 Actors, Directors, and Production Companies with Highest Average Box Office Revenue	15
4.6	Top 10 Actors, Directors, and Production Companies with Most Number of Successful Movies	16
4.7	Top 10 Production Companies with Highest Average Budgets	16
4.8	Number of Movies Released by Release Year, Release Month, and MPAA Rating	17
5.1	Label Encoding vs. One Hot Encoding	20
6.1	Feature Importance for Random Forest Classifier	24
6.2	Feature Importance for Decision Tree Classifier	26
7.1	Confusion Matrix for Random Forest Classifier on Test Data	32
7.2	Confusion Matrix for Decision Tree Classifier on Test Data	32
7.3	Confusion Matrix for K Nearest Neighbors Classifier on Test Data . .	33
7.4	Confusion Matrix for Logistic Regression on Test Data	33
7.5	Confusion Matrix for Support Vector Classifier on Test Data	34
7.6	Receiver Operating Characteristic (ROC) for All Models	34

List of Tables

3.1	Dataset Before Preprocessing	10
4.1	Correlation between Budget, Runtime, Release Year and Box Office .	13
4.2	Correlation between Success Score and Box Office	13
4.3	Correlation between Sentiment Score and Box Office	13
5.1	Dataset After Preprocessing	20
6.1	Results for Different Number of Trees for Random Forest Classifier On Validation Set	23
6.2	Results from Different Criteria for Decision Tree On Validation Set	25
6.3	Results from Different Number of Neighbors for KNN Classifier On Validation Set	26
6.4	Results from Different Encodings for Logistic Regression On Valida- tion Set	27
6.5	Results from Different Kernels for SVCs On Validation Set	28
7.1	Random Forest Classifier Performance Metrics on Test Data	29
7.2	Random Forest Classifier Performance Metrics on Test Data	30
7.3	K Nearest Neighbors Classifier Performance Metrics on Test Data . .	30
7.4	Logistic Regression Performance Metrics on Test Data	31
7.5	Support Vector Classifier Performance Metrics on Test Data	31
7.6	Performance Metrics for all Classifiers on Test Data	35
7.7	Accuracy Comparison with Multiple Papers	35

Chapter 1

Introduction

1.1 Leveraging Online Data for Predicting Box Office Success

The movie-going experience for many people has long been the preferred choice when it comes to spending quality time with friends and family. This makes the film industry one of the most rewarding businesses in the entertainment world. The success of a movie can depend on many factors, such as the cast, budget, director, and trailer reception. This research will use Wikipedia data and YouTube trailer comments to predict a movie's success. YouTube is among the most popular platforms for movie studios to upload trailers for upcoming films. Moviegoers engage in passionate discussions regarding a particular film and express their opinions and feelings in the comment section. On the other hand, Wikipedia is one of the most used sites to access information about movies, actors, directors, ratings, and other film industry news. This study aims to use YouTube trailer comments and Wikipedia data to predict the box office success of a movie. Sentiment analysis will be used to categorize YouTube comments and gain insight into the overall reaction of audiences to a film's trailer. This, combined with other factors such as the cast, director, distributor, and production studio, will predict whether a film will succeed at the box office or not.

1.2 Problem Statement

Predicting a movie's box office success is a complex and challenging task with significant implications for the film industry, such as allowing production companies to allocate resources effectively and make informed decisions. Despite the availability of various data sources, such as YouTube trailer comments and Wikipedia data, there is a lack of comprehensive research exploring their potential to accurately forecast the commercial performance of movies, with most existing approaches primarily relying on traditional marketing strategies and limited data sources, leading to sub-optimal predictive models. Therefore, the problem addressed in this thesis is the absence of a robust predictive model that leverages YouTube trailer comments and Wikipedia data to predict the box office success of movies effectively. This study aims to develop a reliable model that can provide valuable insights to filmmakers, distributors, and investors to enable them to make more informed decisions regard-

ing the types of movies they make and the cast and crew they hire for a particular film, backed up by audience approval.

This research aims to address the following challenges:

- **Neglecting the potential of YouTube trailer comments:** YouTube has become a major platform for movie promotion, and user comments on movie trailers reflect audience reactions and expectations. However, most existing prediction models overlook the valuable information embedded in these comments. Therefore, leveraging YouTube trailer comments as a data source and extracting meaningful insights can enhance the accuracy and granularity of box office predictions.
- **Insufficient integration of Wikipedia data with contemporary approaches:** Wikipedia provides much information about movies, including director, release date, and cast details. While previous prediction models utilized Wikipedia data, they often lack integration with newer data sources or advanced machine-learning techniques. This study aims to explore innovative ways to combine Wikipedia data with YouTube trailer comments and employ advanced predictive algorithms to improve the accuracy and reliability of box office predictions.

1.3 Research Objectives

This research aims to develop a predictive model to predict the box office success of movies using a combination of YouTube trailer comments and Wikipedia data. The main aim is to study the relationship between audience feedback from YouTube comments on movie trailers and the film's subsequent box office performance. This research further aims to analyze the traditional influences on a movie's box office, release date, budget, writer, director and cast members from Wikipedia.

The research objectives are as follows:

1. Explore the potential of YouTube trailer comments as a valuable source of information for predicting movie box office success. By analyzing the sentiment of comments, this study aims to determine if there are significant correlations between user-generated responses and the commercial performance of films.
2. Investigate the role of traditional movie attributes from Wikipedia in conjunction with YouTube trailer comments. By combining these two data sources, the research aims to enhance the prediction accuracy of box office success.
3. Develop a predictive model using machine learning algorithms to predict the box office performance of a movie. The study will evaluate classification techniques, such as Random Forest, Decision Tree, and Logistic Regression, to identify the most suitable approach for predicting movie revenues based on the given features.
4. Evaluate the performance of the proposed algorithms through extensive experimentation and validation. The model will be trained and tested on a diverse

dataset of movies spanning different release dates and production budgets. The accuracy, precision, recall, and other relevant metrics will be measured to assess the reliability and effectiveness of the model.

5. Provide insights and recommendations based on the research findings to assist movie studios, distributors, and marketing professionals make informed decisions regarding resource allocation, promotional strategies, and release planning. The aim is to offer actionable recommendations that can improve the chances of achieving box office success for future movie releases.

By accomplishing these objectives, this research aims to contribute to the field of movie industry analysis by leveraging user-generated content on YouTube and traditional movie attributes from Wikipedia to build an accurate and robust predictive model to estimate the box office success of movies.

Chapter 2

Literature Review

In the study by Apala, Krushikanth & Jose, Merin & Motnam, Supreme & Chan, Chien-Chung & Liszka, Kathy & de Gregorio, and Federico [2], data mining is applied to forecast the achievement of the movies. Data is collected from social media like IMDb, Twitter, and Youtube. This paper used a code snippet in PHP to collect data from Youtube, which ran by the Wamp server. From IMDb, they collected the top 50 genres of the movie; from Twitter, they collected the popularity of directors, actors, and actresses; from Youtube, they collected the likes, dislikes, views, and comments. Sentiment analysis is applied to the comments that have been extracted from Youtube. It used a simple min-max method for normalizing the training data, a K-means tool from Weka for clustering, and then created a training set for generating the predictive model. This study showed that uniform and non-uniform weighted approaches produced by Weka (Naive Bayesian classifier and J-48 classifier) could be used to predict a movie's success. The potential gap of this study is that it did not demonstrate which algorithm will give the highest accuracy and relied on a single machine learning algorithm. The study divided the movies into three categories: hit, flop, and neutral but didn't clarify which movie was a hit or flop. Overall, we can conclude that the study only considers the popularity of an actor/actress but does not count views, likes, dislikes, and comments to predict the favorable outcome of a movie.

Vasu Jain [3] worked on sentiment analysis to determine the sentiment expressed in tweets related to movies. In this study, to predict the movie's box office success, sentiment analysis results were used from tweets during a film's release. The sentiments were classified into four categories, which were positive, negative, neutral, and irrelevant. A classifier was trained using the Lingpipe sentiment analyzer, which utilizes an 8-gram language model on character sequences. 48 movies were correctly classified as hits. However, no other factors except Tweet sentiments were considered when making predictions, which makes the study very limited and leaves room for improvements.

M. S. Neethu and R. Rajasree [4] focused on sentiment analysis, particularly on analyzing sentiments expressed in Twitter posts about electronic products using a machine learning approach. The study identified some of the challenges with analyzing the sentiments of Tweets, which include the presence of slang words, misspellings, and the constraint of a 140-character limit per Tweet. The study compares the

performance of three types of basic classifiers (SVM, Naive Bayes, and Maximum Entropy) and an ensemble classifier for sentiment classification. SVM and Naive Bayes classifiers were implemented using built-in functions in Matlab, while the Maximum Entropy classifier was implemented using MaxEnt software. The Naive Bayes classifier exhibited better precision than the other three classifiers but slightly lower accuracy and recall. On the other hand, SVM, Maximum Entropy Classifier, and ensemble classifiers demonstrated similar accuracy, precision, and recall. These classifiers achieved an accuracy of 90%, while Naive Bayes achieved an accuracy of 89.5%. The study concludes that the chosen feature vector for the product domain plays a crucial role in sentiment analysis. Despite the variation in classifier performance, the quality of the feature vector enables better sentiment analysis results. They emphasize that the feature vector's effectiveness does not depend on the specific classifier.

M. S. Usha and M. Indra Devi (2013) [5] propose a Combined Sentiment Topic (C.S.T.) model, which offers a novel approach to sentiment analysis by simultaneously detecting sentiments and topics in the text. Unlike existing supervised and semi-supervised learning methods that rely on labeled documents for classification, the C.S.T. model is purely based on unsupervised learning techniques and utilizes unlabeled documents for classification. This unsupervised nature of the C.S.T. model enhances its portability to different domains. The paper reports that the C.S.T. model outperforms existing semi-supervised approaches, demonstrating its effectiveness in sentiment analysis tasks. One limitation of the C.S.T. model is its inability to detect neutral opinions, as it currently focuses only on classifying positive and negative sentiments.

Vr, Nithin & Pranav, M & Babu, PB & Lijiya, A. [6] focus on forecasting the movie's success depending on IMDb data. This study collected the dataset from IMDb, Wikipedia, and Rotten Tomatoes. This study only prioritizes the movie which was released between 2000-2012 and in the English language. As items can be duplicated, they use the mean and median methods as central tendencies during the data pre-processing time. This study only worked with a numerical value. That's why the central tendency is used to convert nominal attributes to numerical ones in data integration and transformation time. Different machine learning algorithms, such as - Linear Regression, Logistic Regression, and Support Vector Machine Regression (SVN) models have been run to predict how the movie is going to be. It shows that the linear regression model has the highest accuracy in predicting movie success among other machine learning algorithms. One potential gap in the study is that this study only concentrates on IMDb and Wikipedia data. Other data, such as the YouTube trailer, comments under that trailer, Twitter, and social media comments, were ignored. These factors play an important role in identifying whether the movie is going to be successful or not. Briefly, this study gives an idea of predicting the success of movies by collecting the dataset from IMDb, Wikipedia, and Rotten Tomatoes and solving the dataset using machine learning algorithms. Also, it can be seen that among the 20 features in the dataset, actor1, writer, actor2, director, budget, and reviews played the most significant features. Briefly, it gives an idea of how to forecast whether a movie is going to be successful or not by collecting the dataset from Rotten Tomatoes, IMDb, and Wikipedia and solving

the dataset using machine learning algorithms.

A. Bhave, H. Kulkarni, V. Biramane, and P. Kosamkar [7] highlighted the importance of considering both classical and social factors while predicting a movie's achievement at the box office. Factors such as cast, producer, and director are considered classical, while online responses to a film and a film's engagement on social media platforms are considered social factors. The study argues that integrating classical and social factors can lead to more accurate results when predictions are made regarding the success of a film. The study suggests considering YouTube view count and comments, sentiment analysis of Tweets regarding a film, and Wikipedia view and edit counts to improve the prediction success rate.

R. Dhir and A. Raj [8] analyzed the IMDb dataset and provided a prediction of IMDb scores for movies. The dataset, which includes information on 5,043 movies spanning 100 years and 66 nations, is analyzed using machine-learning methods. Then the outcomes from the methods are compared to decide the most successful one. Various machine learning techniques like K-Nearest Neighbors, Support Vector Machine, Gradient Boost, Ada Boost, and Random Forest are used here to analyze the dataset. Being one of the most informative resources, IMDb covers many variables like movie titles, duration, genre, release time, director names, star popularity, and movie ratings from critics. Additional factors like the number of critics for reviews, the number of user votes, Facebook reactions, movie runtime, budget, and gross collection all significantly impact IMDb scores, making it a great source to consider datasets from. According to the findings and predictions received from all the algorithms, Random Forest had the highest rate of prediction accuracy among the evaluated. Compared to earlier studies, the suggested model predicts movie success with greater accuracy, proving it successful. A possible limitation of this study includes limited exploration of social media data like YouTube, Facebook, and Twitter comments. Fully relying on IMDb scores may result in ignoring other significant factors like box-office revenue, awards, and critical acclaim. Moreover, the limitation over sample size is compromising the proposed model to be a global one. Still, the researchers also expressed their desire to work on it using other learning models.

W. Lu [9] explores how to collect and evaluate online comments to predict the movie box office. For data acquisition, the author gathered comments on individual movies using the Baidu search engine and pre-processed the data to extract subjective criticism and make it more uniform. For emotional orientation analysis, two kinds of approaches: machine learning(using a supervised training method) and semantic orientation(using an unsupervised approach), are introduced here. Finally, the study proposes a K.N.N. (K-Nearest Neighbors) model-based prediction method for box office, which classifies comments as positive or negative, and calculating their correlation coefficient proves the model's great accuracy. The potential limitations in the study include using a specific search engine and a specific platform for collecting data, primarily focusing on the K.N.N. model regardless of other machine learning techniques being available, and excluding multiple additional factors like marketing campaigns, the popularity of casts, release timing, movie genre, music ratings, etc.

Verma, Garima and Verma, Hemraj. (2019) [10] chose a logistic regression model because of its versatility and capacity to handle binary outcomes. This L.R. model can predict with an accuracy of up to 80% whether a Bollywood film will be a "Hit" or a "Flop" before their debut. For data processing and analysis, it uses IBM SPSS 21.0. The authors chose the number of screens, IMDb rating, and MusicRating as predictors, with the movie verdict to be hit or flop as the outcome variable. The data was acquired from online sources like IMDb, bollymoviereviewz.com, planetbollywood.com, boxofficeindia.com, and bollywoodhungama.com through Web Scraping, after extracting data for more than 2000 movies, resulting in the final dataset including 116 movies. The model successfully predicted hit or flop movies with 78.1% accuracy for training data and 80% for test data. The possible shortcomings of this study were that it used limited predictors and was confined to Bollywood movies, which may not give accurate results for movies outside of the Bollywood industry because of Bollywood's distinctive characteristics, and not taking into consideration multiple additional factors like audience preference as well as opinion, marketing strategies, star power, etc.

M. D. Athira and K. S. Lakshmi [11] emphasize the importance of predicting movie success based on a strategy that improves prediction accuracy by combining movie metadata and reviews data. For review data collection, IMDb is used for the dataset labeling positive, negative, and neutral reviews. In contrast, meta-data is collected from Github that includes information such as the number of critical reviews, duration, number of Facebook reactions, release time, IMDb score, aspect ratio, film budget, gross revenue, etc. An ensemble classifier technique is used in the study for predicting the success rate, with algorithms that involve Random Forest, Naive Bayes, Logistic Regression and Max Voting. The max voting approach will identify the movie as a success, failure, or neutral depending on similar outcomes from two or more methods among these classifiers. Finally, the researchers concluded that combining metadata and review data with the results obtained using the Max Voting technique provides us with an almost 90% accuracy rate, which is better than using any features independently. There are some limitations in this study also, though it can provide us with a much higher accuracy rate. The possible gaps are mainly a lack of versatility in the datasets, difficulty in ensuring quality review data, missing usage of some additional evaluation metrics like recall, F1 score, R.O.C. curve, etc., along with some external factors like marketing strategies, audience preference, film budget, critical acclaim and some more.

N. Darapaneni et al. [12] aimed to forecast the movie box office performance using various Machine Learning algorithms, including K-Nearest Neighbours, Random Forest, XGBoost Classifier, Decision Tree, and Deep Neural Network. The authors implemented these algorithms on an IMDB dataset which focused on features such as movie crew, plot, audience, and critics' reviews/ratings, to predict their success rate. The study found that the XGBoost Classifier was the most accurate model, achieving an accuracy rate of nearly 90%. It identified user reviews, cast, and budget as key features for predicting the success of movies. A potential gap in the study is that it does not consider other key features, such as the production studio, the film's director, and the film's budget, which have historically been known to impact the success of a movie significantly. Examples include movies produced by Disney

and movies directed by Christopher Nolan.

Sivakumar, Pirunthavi; Abishankar, Kamalanathan; Rajeswaren, Vithusia Puvaneswaren; Mehendran, Yanusha; Ekanayake, E.M.U.W.J.B. [13] showed how to assess the success and rating of a movie using Random Forest and Naive Bayes models before its release. Using data collected from IMDb, Box Office Mojo, and YouTube comments, this paper highlights various factors that influence the success of a movie. Using YouTube API for scraping 1000 comments per video from YouTube and using a dataset from Kaggle for 200 movies as a substitute for IMDb scraping, they covered the part of data collection. A Random Forest Algorithm is trained using IMDb attributes to predict movie success, while a Naive Bayes model is trained using YouTube user reviews to predict movie ratings. The models perform well, with the Success Prediction model having an overall accuracy of 70%. Still, a problem arises due to YouTube API's tendency to prioritize negative comments over positive ones. Therefore, the study proposes the developed models will work successfully on the data collected from the internet rather than the comments collected from YouTube. The potential gaps in this study comprise the limitation of YouTube's API, which only allows access to a limited number of user comments per video, failing to consider a larger and more diverse dataset required to improve the models' generalizability and being unaware of other factors such as critical acclaim, cultural impact, and long-term profitability that contribute to a film's overall success.

Dewan Muhammad Qaseem, Nashit Ali, Waseem Akram, Aman Ullah, and Kemal Polat (2022) [14] focus on forecasting the success rates of movies by applying a technique to spectators' tweets on movie trailers. The study collected tweets from different movies using the hashtag method. Different machine learning algorithms have been used, such as - Linear S.V.C., K.N.N., Naive Bayes, and decision trees. Also, a lexical-based approach algorithm has been used. The outcomes of this paper showed that Linear S.V.C. has the maximum validity among other machine learning algorithms in predicting movie success by using tweets. Also, the lexical-based approach has used three different dictionaries with three different word counts. Among them, ratings_Warriner dictionary gives a more accurate result. Then the result was compared with other sites, such as IMDb. The potential gap in this study only focuses on tweets. Still, it does not pay attention to the comment section under the YouTube trailers, social media comments, and IMDb ratings and comments. Also, it did not consider important key features such as the popularity of actors, actresses, cast, movie budget, quality, promotion, sound effects, etc. In short, this study only used Twitter to predict a movie's success using different machine learning algorithms ignoring some significant features and a lexical-based approach from which Linear S.V.C. and ratings_Warriner dictionary give the highest accuracy.

Chapter 3

Dataset

3.1 Overview

The primary dataset used in this research is an amalgamation of Wikipedia data and YouTube trailer comments that have been collected firsthand by the authors. The dataset initially comprised of 23 features and 2020 rows of data when it was originally collected from Wikipedia but was left with 18 features and 1780 rows once the data was cleaned. The features of the final dataset are described in Table 3.1.

3.2 Wikipedia Data

The raw movie data utilized in this research was web-scraped from Wikipedia using Python’s beautiful soup library. The scope was limited to major American film productions released theatrically in the United States of America between 2010 and 2022 inclusive. In total, data on 2020 films was collected.

The web scraping methodology produced a dataset that required substantial cleaning and preprocessing before analysis could proceed. Three main data issues were addressed: embedded HTML reference tags, heterogeneous data types for key variables, and non-standardized date variables.

First, the raw Wikipedia data retained various HTML tags used for internal Wiki-media formatting and references. As these provided no analytical value, a find-and-replace script removed all HTML tags across the entire dataset. Second, the running time information for films contained heterogeneous values - some records displayed hours and minutes while others contained only minutes. A simple calculation was scripted to standardize all running times to integer minutes. Similarly, budget and worldwide box office gross values required standardization to an integer data type and domestic currency. Records contained a variety of string formats: written numbers, numerals, and currency symbols. Values were parsed to remove non-numeric characters, then converted to integer type, and scaled to hundreds of millions of USD. Finally, movie release dates in the raw data took several forms, such as “June 23, 2017” and “2017-06-23” across records. The Python datetime library was leveraged to convert all release dates into standardized Python datetime objects.

Feature	Data Type	Description
Title	string	The titles of each movie in string format.
Directed by	list	A list of strings containing the names of directors of the film.
Produced by	list	A list of strings containing the names of producers of the movie.
Cinematography	list	A list of strings containing the names of cinematographers of the movie.
Runtime	int	The movie's total duration in minutes.
Distributed by	list	A list of strings containing the names of distributors of the movie.
Language	string	The primary language spoken in the movie.
Written by	list	A list of strings containing the names of writers of the movie.
Cast	list	A list of strings containing the names of actors in the movie.
Edited by	list	A list of strings containing the names of editors of the movie.
Production companies	list	A list of strings containing the names of production companies who produced the movie.
Release Date	datetime	The release date of movie in %Y-%m-%d %H:%M:%S format.
MPAA Rating	string	The MPAA rating that the movie received. The values are G, PG, PG-13, R, and NC-17.
Sentiment Score	float	The sentiment score calculated from the YouTube trailer comments of each movie.
Budget	int	The movie's production budget.
Box Office	int	The total box office gross of the movie.
Box Office Status	string	An estimation of whether or not the movie was profitable at the box office.

Table 3.1: Dataset Before Preprocessing

With the dataset cleaned, an additional column was engineered called “Box Office Status”, which is a binary indicator of whether a given film achieved financial success. The informal Hollywood rule of thumb states that for a film to break even, it must gross double its production budget. An even stricter criterion was adapted whereby $\frac{2}{3}$ rds of a movie’s box office gross amounted to greater than its budget multiplied by 1.5 were labeled box office successes, while films failing to meet that benchmark were labeled as failures. This criterion was formulated by film journalist John Campea where the total budget of a movie (production budget + marketing budget) is assumed to be 1.5 times its production budget, and the take-home box office revenue for production companies is assumed to be $\frac{2}{3}$ rds of the total box office revenue. The primary assumptions are that the marketing budget of movies is 50% of the production budget and that production companies pay $\frac{1}{3}$ rd of total box office revenue to theaters, which leaves them with $\frac{2}{3}$ rds of the total revenue. These assumptions were necessary because information on marketing budget and take-home box office numbers are not available publicly.

$$\text{Box Office Success} = \left(\frac{2}{3} \times \text{Box Office Revenue}\right) > (1.5 \times \text{Budget}) \quad (3.1)$$

3.3 YouTube Trailer Comments

Building on the movie dataset from Wikipedia, additional data was collected from YouTube utilizing the YouTube Data API. The objective was to gather viewer commentary and discussion around trailers for the movies contained in the Wikipedia dataset.

The YouTube Data API allows programmatic access to metadata on videos uploaded to the platform. This capability was leveraged to search for and identify the official trailers for each of the 2,020 American films released between 2010-2022 in the dataset. The trailer’s video ID was collected using the API’s search method to search for the movie trailers using the movie’s title. The trailer video IDs were then used to request user comments made on that trailer video. For each movie trailer video, the oldest 100 comments were collected. This was done to try and ensure that the comments were older than the movie’s release date in order to avoid biased data. In cases where fewer than 100 comments existed for a given trailer video, all available comments were gathered. Any comments consisting solely of non-alphanumeric characters were filtered out. In total, 160,138 YouTube comments were collected across 2,020 trailers. The discrepancy in numbers arises from the fact that some obscure indie films had very little identifiable trailer comments on YouTube whatsoever. Moreover, some movie trailers had their comments turned off, which prevented users from commenting on those trailers.

Sentiment analysis was then performed on these trailer comments and an average sentiment score was derived for each movie. This additional information was then integrated with the core Wikipedia movie features by matching records on movie titles. This combined multi-modal dataset of both structured and unstructured data formed the foundation for investigating what signal, if any, YouTube trailer commentary provides in predicting the financial performance of movies.

Chapter 4

Data Analysis

4.1 Correlation

Correlation refers to the relationship between two variables. It is a statistical measure of how a change in one variable impacts another variable. In this study we have utilized the correlation coefficient to establish relationships between box office success and other movie features, in order to see how much of an impact each feature has in predicting the financial performance of movies.

Table 4.1 and Figure 4.1 show the correlation coefficient and scatter plot between Budget, Runtime, Release Year, and Box Office respectively. Budget has the highest positive correlation with box office revenue among all the other features with a value of 0.7749. This indicates that movies with higher budgets tend to perform better at the box office. This is also demonstrated in the Budget and Box Office scatter plot in Figure 4.1, which shows a growth in box office revenue with an increase in the budget. There is also a positive correlation between a movie's runtime and its box office performance, though to a lesser degree than budget. Most movies released between 2010 and 2022 have a runtime of around 125 minutes to 150 minutes, as shown by the scatter plot in Figure 4.1. There is also a negligible positive correlation between a movie's release year and its box office revenue. This indicates that box office revenue for movie's have grown over the years, although very little.

Table 4.2 and Figure 4.2 show the correlation coefficient and scatter plot between the Success Score of different features and Box Office respectively. The feature with the highest positive correlation with Box Office Revenue is the Production Company with a correlation coefficient of 0.2972. A close second are the producers of the film with a correlation coefficient of 0.2851. This indicates that a film's box office revenue is greatly dependent on the company and producers that produce it. The scatter plot in Figure 4.2 also shows an increase in box revenue with an increase in the success score of production companies and producers. The success score of directors and writers also seems to have a great deal of impact on a movie's box office revenue with a positive correlation coefficient of 0.2521 and 0.2367 respectively. Every other feature to a lesser degree also has a positive correlation with box office revenue.

Table 4.3 and Figure 4.3 show the correlation coefficient and scatter plot between the Sentiment Score and Box Office revenue respectively. The value being very

small at only -0.0142 indicates that there is very little correlation between box office success and sentiment score. This is unusual as the norm would suggest that audience sentiment plays a key role in the success of movies at the box office. This deviation from the norm may in part be due to the way in which sentiment scores were calculated from YouTube comments in this study.

	Box Office
Budget	0.7749
Runtime	0.4237
Release year	0.0136

Table 4.1: Correlation between Budget, Runtime, Release Year and Box Office

	Box Office
Cast Success Score	0.1895
Production Company Success Score	0.2972
Director Success Score	0.2521
Distributor Success Score	0.1776
Producer Success Score	0.2851
Cinematographer Success Score	0.1399
Writer Success Score	0.2367
Editor Success Score	0.1980

Table 4.2: Correlation between Success Score and Box Office

	Box Office
Sentiment Score	-0.01423

Table 4.3: Correlation between Sentiment Score and Box Office

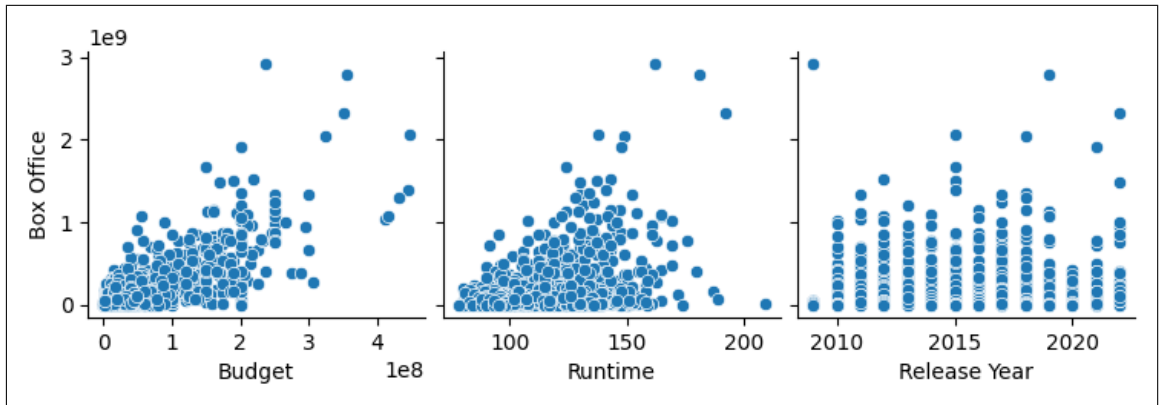


Figure 4.1: Scatter Plot of Budget, Runtime, Release Year and Box Office

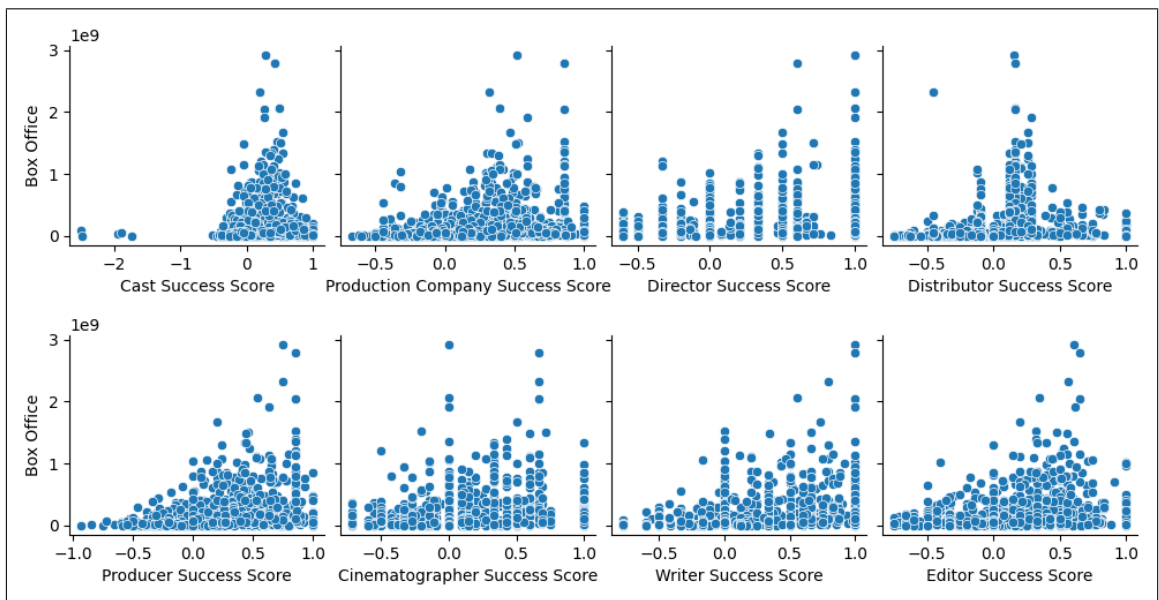


Figure 4.2: Scatter Plot of Success Scores and Box Office

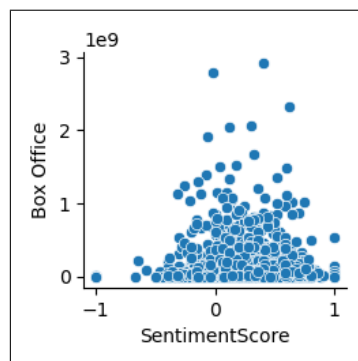


Figure 4.3: Scatter Plot of Success Scores and Box Office

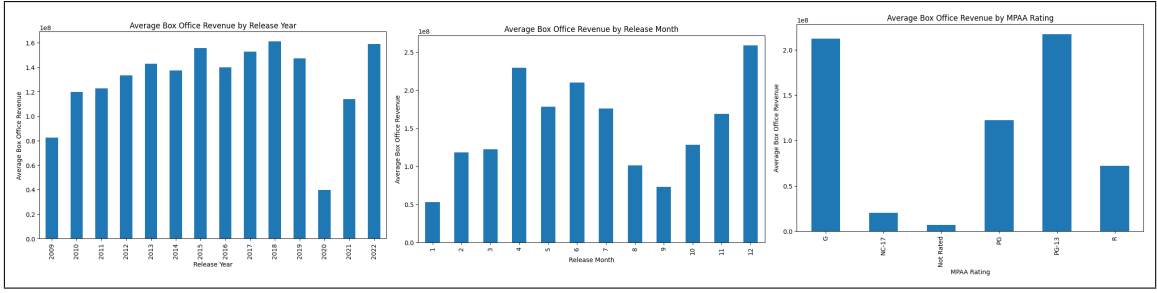


Figure 4.4: Average Box Office Revenue by Release Year, Release Month, and MPAA Rating

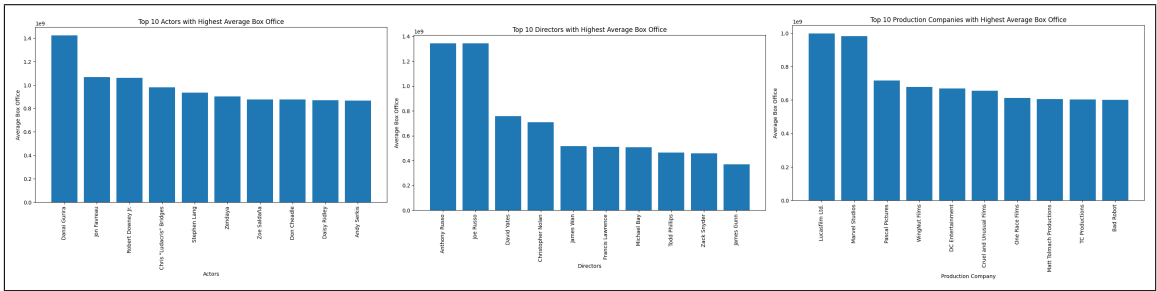


Figure 4.5: Top 10 Actors, Directors, and Production Companies with Highest Average Box Office Revenue

4.2 Average Box Office Revenue

Figure 4.4 shows the average movie box office revenue for different years, months, and MPAA ratings. It can be observed that 2018 and 2022 had the highest average box office revenues, for movies released between the years 2009 and 2022, with an average value of \$160 million. The year 2020 saw a massive fall in average box office revenue, which can be attributed to fewer people going to movie theaters due to the restrictions put in place during the 2020 COVID-19 pandemic. For release months, the month of December seems to be the time when most people visit movie theaters, as indicated by an average monthly box office value of over \$250 million. The months of April and June also show respectable box office numbers with both having a monthly average box office revenue of over \$200 million. On the other hand, the month of January seems to be the least popular time for people to see movies, with an average monthly box office of only \$50 million. Movies with an MPAA rating of G and PG-13 have performed the best within the last decade. This is not surprising as most franchise films have PG-13 ratings and G films are usually family films watched by many families during the Summer and Christmas holidays. Likewise, films with NC-17 and R ratings performed less well because of the niche nature of their audiences. Films that have no rating performed the worst.

Figure 4.5 shows the top 10 actors, directors, and production Companies with the highest average box office revenue. Danai Gurira, Jon Favreau, and Robert Downey Jr are the top 3 highest-average box office grossing actors. Anthony and Joe Russo, along with David Yates are the top 3 highest average grossing directors. The top 3 highest average grossing production studios are Lucasfilm Ltd, Marvel Studios, and Pascal Pictures.

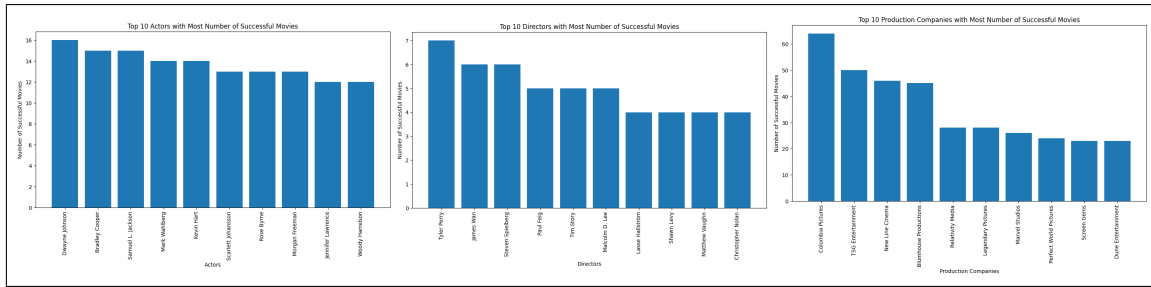


Figure 4.6: Top 10 Actors, Directors, and Production Companies with Most Number of Successful Movies

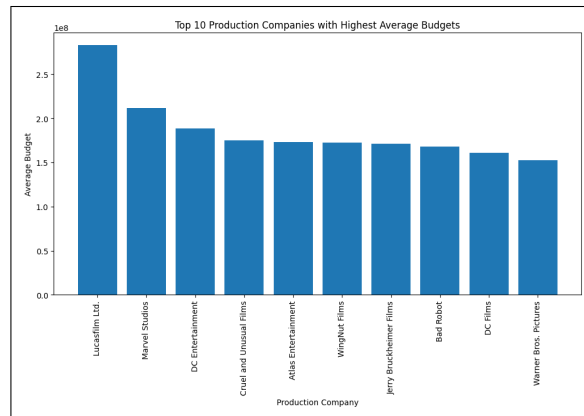


Figure 4.7: Top 10 Production Companies with Highest Average Budgets

4.3 Most Number of Successful Movies

Figure 4.6 shows the top 10 actors, directors, and production Companies with the most number of successful movies. Dwayne Johnson, Bradley Cooper, and Samuel L Jackson are the top 3 actors with most number of successful movies with 16, 15, and 15 successful movies respectively. Tyler Perry, James Wan, and Steven Spielberg are the top 3 directors with the most number of successful movies. The top 3 production studios with the most successful movies are Columbia Pictures, TSG Entertainment, and New Line Cinema.

4.4 Highest Average Budgets

Figure 4.7 depicts the top 10 production companies with the highest average budgets. Lucasfilm Ltd, Marvel Studios, and DC Entertainment are the top 3 production companies with the highest average budgets. This can be attributed to the fact all three studios produce large-scale blockbuster franchises such as Star Wars, Marvel, and DC, which have large crews and huge production budgets.

4.5 Most Number of Movies Released

Figure 4.8 shows the number of movies released in different years, months, and for different MPAA ratings. It can be observed that 2011 and 2010 had seen the most numbers of movies released, between the years 2009 and 2022, with more than 150

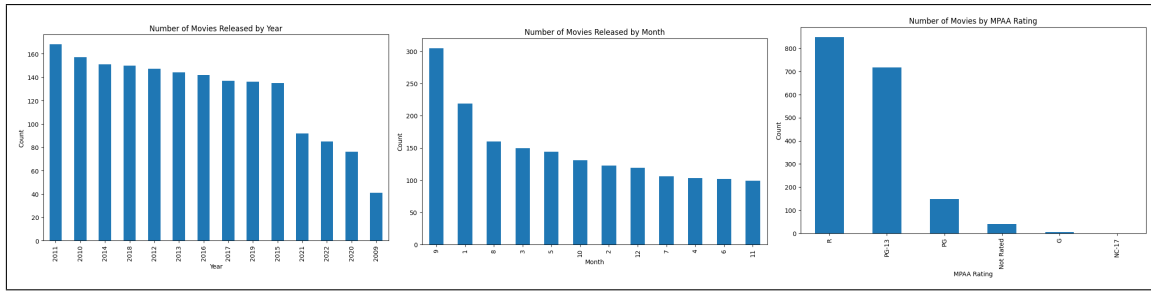


Figure 4.8: Number of Movies Released by Release Year, Release Month, and MPAA Rating

major films being released. The year 2020 saw a massive fall in the number of movies released, which can be attributed to movie theaters being closed and film production coming to a halt due to the restrictions put in place during the 2020 COVID-19 pandemic. For release months, the month of September seems to be the time when most movies are released in movie theaters, as indicated by a release count of over 300 movies, between the years 2009 and 2022. In terms of MPAA ratings, the vast majority of films released in the last decade have been rated either R or PG-13.

Chapter 5

Data Preprocessing

Before any machine learning algorithms could be implemented on the data, it needed to be preprocessed in order to remove irrelevant features, address null values, and to make it suitable for machine learning implementation. This was done in four steps. Firstly, the data was cleaned in order to remove irrelevant features and address null values. Sentiment analysis was then performed on the YouTube trailer comments. The success scores for the cast, director, producers, etc were calculated. And finally, the MPAA rating column was encoded using both Label Encoding and One Hot Encoding.

5.1 Data Cleaning

During the data cleaning process, all irrelevant features were removed from the dataset. These features included: **Country**, **Story by**, **Music by**, **Based on**, **Trailer ID**, **Trailer Title**, **Language**, and **Title**. **Trailer ID** and **Trailer Title** were removed as they were only used to collect YouTube comments data using the YouTube Data API. **Country** and **Language** were removed because 99% of the movies in the dataset were produced in the United States and 99% of the languages spoken in the movies were English, which gave no significant diversity to the features. Part of the data cleaning process was also to address null values within the features and samples. All null were filled in manually by the authors by looking up the missing information corresponding to the movie's title on IMDB. For features such as **Cast**, **Producers**, **Directors**, etc., the features were converted to lists of strings from strings. This was done to make sure that the names of people or companies within those lists could be evaluated individually. The **Release Date** feature was split into three distinct **Release Year**, **Release Month**, and **Release Day** features.

5.2 Success Score Calculation

The success score is a measure that was formulated by the authors of this study to calculate the likelihood of an individual actor, director, producer, etc.'s movie being successful. Each individual was assigned a success score which was calculated by subtracting the number box office failures that they appeared in from the number of box office successes that they appeared in, divided by the total number of movies

that they appeared in. This is demonstrated in Equation 5.1. After the individual success score were calculated, the scores were assigned to all groups of individuals involved in each movie. The success score for the Cast feature was assigned weights, with the lead actor having the most weight. The lead actor contributed to 20% of the total cast score, while the other actors jointly contributed to the other 80%. This was done because not all roles contribute to the movie’s success equally, with lead actors traditionally contributing the most. A group success score was then calculated by summing up the success score of each individual in the group and dividing that sum by the number of individuals in the group. This is demonstrated in Equation 5.2.

$$\text{Individual Success Score} = \frac{\text{No. of Successful Movies} - \text{No. of Failed Movies}}{\text{Total No. of Movies}} \quad (5.1)$$

$$\text{Group Success Score} = \frac{\sum \text{Individual Success Scores}}{\text{No. of Individuals in the Group}} \quad (5.2)$$

5.3 Data Encoding

The **MPAA Rating** feature was composed of categorical values. To make it suitable for machine learning implementation, encoding techniques were used to convert the categorical data into numerical figures on which machine learning algorithms can be trained. Label Encoding and One Hot Encoding are both techniques which help to convert categorical data into numerical data. Figure 5.1 illustrates the difference between the two encoding techniques. Both techniques were used in order to evaluate which one gave better results. The target column **Box Office Status** was also converted into numerical values through binary encoding, with *Success* being 1, and *Failure* being 0.

5.4 Data Splitting

Initially, the dataset was split into features and target variables. The features were then normalized using a MinMaxScaler which transforms the features by scaling each feature to a given range. The normalized values and the target variable was then split into three parts for training, validating, and testing, with a random state of 42. This was done in the ratio 0.7: 0.06: 0.24, which in numerical terms was 1232: 106: 423 respectively. The training set of 1,232 movies was used to train the five machine-learning algorithms. The validation set of 106 films was used for hyper-parameter tuning. And, the test set of 423 movies was used to test the prediction capabilities of the algorithms.

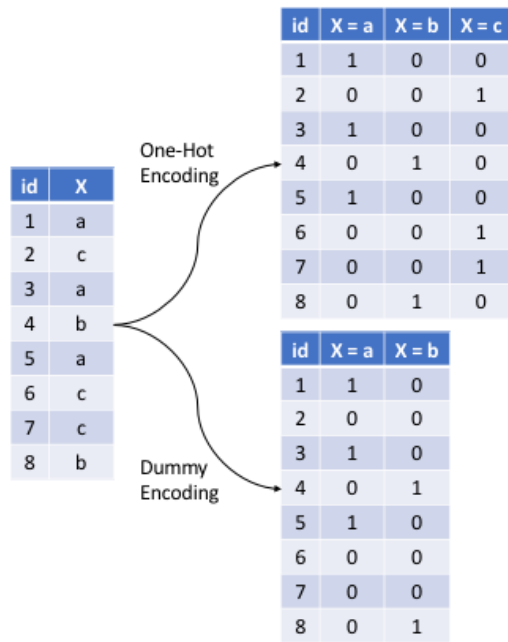


Figure 5.1: Label Encoding vs. One Hot Encoding

Feature	Data Type
Cast Success Score	float
Director Success Score	float
Producer Success Score	float
Cinematographer Success Score	float
Runtime	int
Distributor Success Score	float
Writer Success Score	float
Editor Success Score	float
Production Company Success Score	float
Release Day	int
Release Month	int
Release Year	int
Age Rating	int
SentimentScore	float
Budget	float
Box Office Status	int

Table 5.1: Dataset After Preprocessing

Chapter 6

Methodology

6.1 Performance Metrics

The results of the machine learning models for both the validation set and test set were evaluated using five performance metrics and visualized using a Confusion Matrix and ROC Curve.

- **Accuracy:** A measure of the proportion of correctly classified instances (true positives + true negatives) out of the total number of instances [1].

$$\text{Accuracy} = \frac{\text{True Positives} + \text{True Negatives}}{\text{Total Instances}} \quad (6.1)$$

- **Precision:** A measure of the proportion of true positive instances out of the total number of instances [1].

$$\text{Precision} = \frac{\text{True Positives}}{\text{True Positives} + \text{False Positives}} \quad (6.2)$$

- **Recall:** A measure of the proportion of true positive instances out of the total actual positive instances [1].

$$\text{Recall} = \frac{\text{True Positives}}{\text{True Positives} + \text{False Negatives}} \quad (6.3)$$

- **F1 Score:** The harmonic mean of precision and recall. Provides a balance between the two metrics and is useful when dealing with imbalanced data such as the one used in this study [1].

$$\text{F1 Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (6.4)$$

- **ROC AUC Score:** A representation of the degree or measure of separability. It tells how much the model is capable of distinguishing between classes [1].
- **Confusion Matrix:** A representation used to evaluate the performance of a classification model. It provides counts, for true positives, true negatives, false positives, and false negatives allowing assessment of the models' accuracy and error rates [1].

- **ROC Curve:** Illustrates the diagnostic ability of a classifier system by plotting the True Positive Rate against the False Positive Rate at various threshold settings [1].

6.2 Model Description

6.2.1 Sentiment Analysis Model

The model used for Sentiment Analysis is a pre-trained sentiment analysis model from Hugging Face’s Transformers library. Python’s TextBlob library was also used to perform sentiment analysis on the data, but upon manual inspection, it proved to be inaccurate. Transformers is a state-of-the-art NLP library that provides access to thousands for pre-trained models to performs various text based tasks. It was developed by Hugging Face using PyTorch and TensorFlow. The pipeline function was used to load a pre-trained sentiment analysis model. The model was then applied to the comments that were gathered using the YouTube Data API, which were then labeled as either POSITIVE or NEGATIVE. The accuracy of the sentiments was judged through manual inspection by the authors by looking at random samples of comments and their corresponding sentiments. A Sentiment Score was then formulated by subtracting the number of negative comments for a movie from the number of positive comments for that movie and then dividing that value by a total number of comments for that movie. This Sentiment Score was then added to the Wikipedia dataset as a feature.

6.2.2 Machine Learning Models

Five sentiment analysis models were trained and evaluated on the validation and test sets:

- **Random Forest** is an ensemble machine learning algorithm that uses the combined results from multiple decision trees to make classifications. It builds decision trees by taking random subsets of the data and then averages the results of those trees to classify the data [1]. This model was chosen because of its ability to handle categorical variables such as rating.
- **Decision Tree** is a supervised machine learning algorithm that breaks down the data into smaller subsets to develop a tree-like structure to make classification decisions [1]. A decision tree is effective in handling categorical variables and provides an estimate of which features are most important in making the classifications.
- **K Nearest Neighbors** is a supervised machine learning algorithm where classifications are made based on the k closest training examples [1]. KNN was chosen because of its ability to handle multi-dimensional data effectively.
- **Logistic Regression** classifies data by using a logistic function to find the find probability of each class [1]. Logistic Regression was chosen because it is suitable for binary classification tasks, such as predicting whether a movie will be a box office success or failure.

- **Support Vector Classifier** is a supervised machine learning algorithm that classifies data by finding the optimal hyperplane that separates the different classes in a dataset [1]t. SVC was chosen for its ability to handle complex relationships between features.

6.3 Model Training

The machine learning models were trained on the training set, and hyper-parameter tuning was done by testing the results on the validation set. All of the models were trained on both label encoded data and one hot encoded data to determine the best encoding method for each model. Different hyper-parameters were tested for each model to find the most suitable one.

6.3.1 Random Forest Classifier Training

The Random Forest classifier was trained using 10, 25, 50, 100, and 200 trees. All variations were tested on the validation set with both label encoded data and one hot encoded data to identify the optimum number of trees. The results are compared in Table 6.1.

Label Encoded Data					
No. Of Trees	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
10 Trees	0.9151	0.9214	0.9151	0.9141	0.9064
25 Trees	0.9622	0.9647	0.9622	0.9620	0.9574
50 Trees	0.9434	0.9486	0.9434	0.9428	0.9362
100 Trees	0.9434	0.9486	0.9434	0.9428	0.9362
200 Trees	0.9434	0.9486	0.9434	0.9428	0.9362
One Hot Encoded Data					
No. Of Trees	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
10 Trees	0.9528	0.9565	0.9528	0.9525	0.9468
25 Trees	0.9716	0.9730	0.9717	0.9715	0.9681
50 Trees	0.9622	0.9647	0.9622	0.9620	0.9574
100 Trees	0.9528	0.9565	0.9528	0.9525	0.9468
200 Trees	0.9622	0.9647	0.9622	0.9620	0.9574

Table 6.1: Results for Different Number of Trees for Random Forest Classifier On Validation Set

As shown in Table 6.1, the Random Forest classifier model achieved scores of over 0.9 on the validation set across all performance metrics for both label encoded and

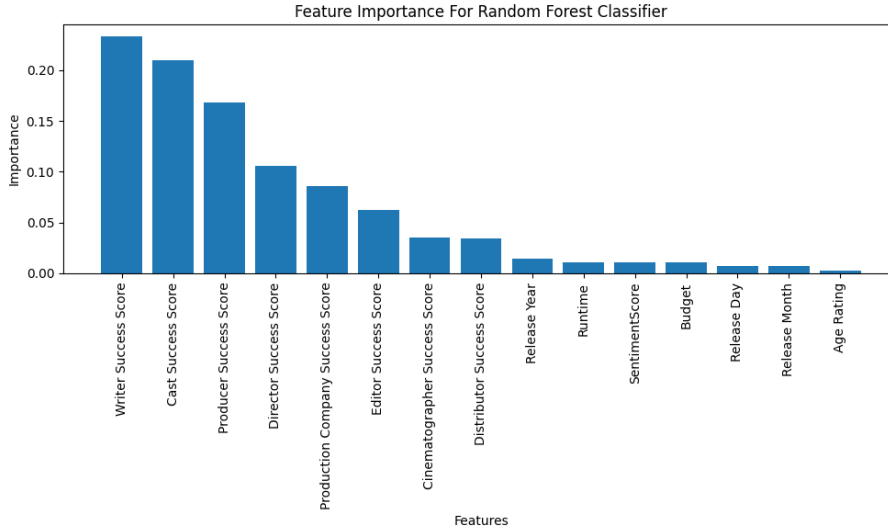


Figure 6.1: Feature Importance for Random Forest Classifier

one hot encoded data. On label encoded data, it can be observed that maximum performance is achieved at 25 trees, with an accuracy of 0.9622 and an F1 score of 0.9620, and increasing the number of trees over 25 brings down the model's performance. The performance stays the same for every subsequent increase in the number of trees. On one hot encoded data, peak performance is also achieved at 25 trees, with an accuracy of 0.9716 and an F1 score of 0.9715, with the performance fluctuating below the peak with each subsequent increase in the number of trees. To evaluate the Random Forest classifier model on the test data; the combination of 25 trees with One Hot Encoding was chosen as the optimum method for training the model, as it gives the best balance between accuracy and training time.

Figure 6.1 shows the feature importance for the Random Forest classifier. It can be observed that Writer Success Score, Cast Success Score, and Producer Success Score are the three features that hold the most weight when classifying box office success using the Random Forest model.

6.3.2 Decision Tree Classifier Training

The Decision Tree classifier model was trained on 2 criterions: Gini Impurity and Entropy. The two criterions are described in Equations 6.5 and 6.6. All variations were tested on the validation set with both label encoded data and one hot encoded data to identify the optimum criterion. The results are compared in Table 6.2.

- **Gini Impurity:** A measure that quantifies the likelihood of an incorrect classification of a randomly chosen element from the dataset. Used in decision tree algorithms to evaluate the quality of a split in a dataset.

$$\text{Gini}(D) = 1 - \sum_{i=1}^c (p_i)^2 \quad (6.5)$$

- **Entropy:** A measure that quantifies the amount of uncertainty or randomness in the data. Used in decision tree algorithms to assess the information gain from splitting a dataset based on a particular attribute.

$$\text{Entropy}(D) = - \sum_{i=1}^c p_i \log_2(p_i) \quad (6.6)$$

Label Encoded Data					
Criterion	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
Gini Impurity	0.9528	0.9565	0.9528	0.9524	0.9468
Entropy	0.9151	0.9159	0.9151	0.9147	0.9107
One Hot Encoded Data					
Criterion	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
Gini Impurity	0.9339	0.9372	0.9339	0.9334	0.9276
Entropy	0.9150	0.9159	0.9150	0.9147	0.9107

Table 6.2: Results from Different Criteria for Decision Tree On Validation Set

As shown in Table 6.2, the Decision Tree classifier model achieved scores of over 0.9 on the validation set across all performance metrics for both label encoded and one hot encoded data. On label encoded data, it can be observed that maximum performance is achieved with criterion Gini Impurity, with an accuracy of 0.9528 and an F1 score of 0.9524. On one hot encoded data, peak performance is also achieved with criterion Gini Impurity, with an accuracy of 0.9339 and an F1 score of 0.9334. To evaluate the Decision Tree classifier model on the test data; the combination of Gini Impurity with Label Encoding was chosen as the optimum method for training the model, as it is seen to give the best performance.

Figure 6.2 shows the feature importance for Decision Tree classifier. It can be observed that the Writer Success Score holds the most weight when classifying box office success using the Decision Tree model. Cast Success Score also carries a great deal of importance, although not to the level of Writer Success Score.

6.3.3 K Nearest Neighbors Classifier Training

The KNN classifier was trained using 5, 10, 15, 20, 30 neighbors. All variations were tested on the validation set with both label encoded data and one hot encoded data to identify the optimum number of neighbors. The results are compared in Table 6.3.

As shown in Table 6.3, the KNN classifier model achieved scores of over 0.88 on the validation set across all performance metrics for both label encoded and one hot encoded data. On label encoded data, it can be observed that maximum performance

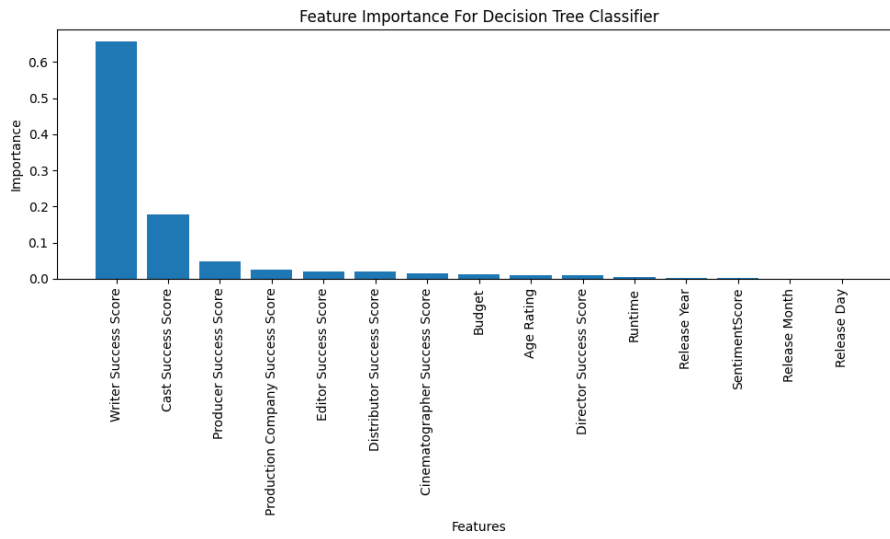


Figure 6.2: Feature Importance for Decision Tree Classifier

Label Encoded Data					
No. Of Neighbors	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
5 Neighbors	0.9245	0.9335	0.9245	0.9234	0.9148
10 Neighbors	0.8962	0.9125	0.8962	0.8937	0.8829
15 Neighbors	0.9056	0.9193	0.9056	0.9037	0.8936
20 Neighbors	0.9056	0.9193	0.9056	0.9037	0.8936
30 Neighbors	0.8867	0.9059	0.8867	0.8837	0.8723
One Hot Encoded Data					
No. Of Neighbors	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
5 Neighbors	0.9056	0.9193	0.9056	0.9037	0.8936
10 Neighbors	0.8962	0.9125	0.8962	0.8937	0.8829
15 Neighbors	0.9150	0.9263	0.9150	0.9136	0.9042
20 Neighbors	0.8867	0.9059	0.8867	0.8837	0.8723
30 Neighbors	0.8773	0.8995	0.8773	0.8736	0.8617

Table 6.3: Results from Different Number of Neighbors for KNN Classifier On Validation Set

is achieved at 5 neighbors, with an accuracy of 0.9245 and an F1 score of 0.9234, and increasing the number of neighbors over 5 brings down the model's performance. On one hot encoded data, peak performance is also achieved at 5 neighbors, with an accuracy of 0.9056 and an F1 score of 0.9037, with the performance fluctuating below the peak with each subsequent increase in the number of neighbors. To evaluate the KNN classifier model on the test data; the combination of 5 neighbors with Label Encoding was chosen as the optimum method for training the model, as it gives the best balance between accuracy and training time.

6.3.4 Logistic Regression Training

The Logistic Regression model was trained on both label encoded data and one hot encoded data to identify the best encoding method for the model. The results are compared in Table 6.4.

Label Encoded Data				
Accuracy	Precision	Recall	F1 Score	ROC AUC Score
0.9150	0.9214	0.9150	0.9140	0.9064
One Hot Encoded Data				
Accuracy	Precision	Recall	F1 Score	ROC AUC Score
0.9150	0.9214	0.9150	0.9140	0.9064

Table 6.4: Results from Different Encodings for Logistic Regression On Validation Set

As shown in Table 6.4, the Logistic Regression model achieved scores of over 0.91 on the validation set across all performance metrics for both label encoded and one hot encoded data. On label encoded data, the model achieved an accuracy of 0.9150 and an F1 score of 0.9140. On one hot encoded data, the model achieved an accuracy of 0.9150 and an F1 score of 0.9140. Both encoding methods had similar performances. To evaluate the Logistic Regression model on the test data; Label Encoding was chosen as the optimum method for training the model, as it is seen to give the best performance.

6.3.5 Support Vector Classifier Training

The Support Vector classifier was trained using linear, polynomial, RBF, and sigmoid kernels. All variations were tested on the validation set with both label encoded data and one hot encoded data to identify the optimum kernel. The results are compared in Table 6.5.

As shown in Table 6.5, the Support Vector classifier model achieved scores of over 0.92 on the validation set across all performance metrics for both label encoded and one hot encoded data for all kernels except Sigmoid. On label encoded data, it

Label Encoded Data					
Kernel	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
Linear	0.9245	0.9292	0.9245	0.9237	0.9170
Polynomial	0.9339	0.9409	0.9339	0.9331	0.9255
RBF	0.9433	0.9455	0.9433	0.9430	0.9383
Sigmoid	0.2358	0.2031	0.2358	0.2177	0.2140
One Hot Encoded Data					
Kernel	Accuracy	Precision	Recall	F1 Score	ROC AUC Score
Linear	0.9245	0.9292	0.9245	0.9237	0.9170
Polynomial	0.9433	0.9455	0.9433	0.9430	0.9383
RBF	0.9339	0.9372	0.9339	0.9334	0.9276
Sigmoid	0.2075	0.2075	0.2075	0.2075	0.1972

Table 6.5: Results from Different Kernels for SVCs On Validation Set

can be observed that maximum performance is achieved with the RBF kernel, with an accuracy of 0.9433 and an F1 score of 0.9430. On one hot encoded data, peak performance is achieved with the Polynomial kernel, with an accuracy of 0.9433 and an F1 score of 0.9430. To evaluate the Support Vector classifier model on the test data; the combination of Polynomial kernel with One Hot Encoding was chosen as the optimum method for training the model, as it gives the best accuracy and has a shorter training time.

Chapter 7

Results

All models were trained with the most suitable hyper-parameters that were determined using the validation test set. The results of the machine learning models were analyzed using 5 performance metrics and visualized using a Confusion Matrix and ROC Curve.

7.1 Individual Results Analysis

7.1.1 Random Forest Classifier Results

Model	F1 Score	Accuracy	Precision	Recall	ROC AUC Score
Random Forest	0.9598	0.9598	0.9598	0.9598	0.9598

Table 7.1: Random Forest Classifier Performance Metrics on Test Data

As shown in Table 7.1 the Random Forest model was able to accurately classify 95.98% of movies, suggesting that the model can make correct predictions effectively. The accuracy of the model is demonstrated in greater detail with the confusion matrix in Figure 7.1. The model also achieved a precision of 0.9598, which indicates that the model is able to avoid making false positive errors 95.98% of the time, which can help prevent movie studios from making incorrect decisions and thus save them from a potential failure. The model achieved a recall of 0.9598, which indicates that it correctly identified 95.98% of all actual box office success. An F1 Score of 0.9598, suggests that the model is well balanced in making accurate positive predictions and identifying all positive cases. A ROC AUC score of 0.9598 implies that the model does well in distinguishing a box office success from a box office failure.

7.1.2 Decision Tree Classifier Results

As shown in Table 7.2 the Decision Tree model was able to accurately classify 94.09% of movies, suggesting that the model can make correct predictions effectively. The accuracy of the model is demonstrated in greater detail with the confusion matrix

Model	F1 Score	Accuracy	Precision	Recall	ROC AUC Score
Decision Tree	0.9409	0.9409	0.9420	0.9409	0.9406

Table 7.2: Random Forest Classifier Performance Metrics on Test Data

in Figure 7.2. The model also achieved a precision of 0.9409, which indicates that the model is able to avoid making false positive errors 94.09% of the time, which can help prevent movie studios from making incorrect decisions and thus save them from a potential failure. The model achieved a recall of 0.9598, which indicates that it correctly identified 94.20% of all actual box office success. An F1 Score of 0.9409, suggests that the model is well balanced in making accurate positive predictions and identifying all positive cases. A ROC AUC score of 0.9406 implies that the model does well in distinguishing a box office success from a box office failure.

7.1.3 K Nearest Neighbors Classifier Results

Model	F1 Score	Accuracy	Precision	Recall	ROC AUC Score
K Nearest Neighbors	0.9166	0.9173	0.9274	0.9173	0.9152

Table 7.3: K Nearest Neighbors Classifier Performance Metrics on Test Data

As shown in Table 7.3 the KNN model was able to accurately classify 91.66% of movies, suggesting that the model can make correct predictions effectively. The accuracy of the model is demonstrated in greater detail with the confusion matrix in Figure 7.3. The model also achieved a precision of 0.9274, which indicates that the model is able to avoid making false positive errors 92.74% of the time, which can help prevent movie studios from making incorrect decisions and thus save them from a potential failure. The model achieved a recall of 0.9173, which indicates that it correctly identified 91.73% of all actual box office success. An F1 Score of 0.9166, suggests that the model is well balanced in making accurate positive predictions and identifying all positive cases. A ROC AUC score of 0.9152 implies that the model does well in distinguishing a box office success from a box office failure.

7.1.4 Logistic Regression Results

As shown in Table 7.4 the Logistic Regression model was able to accurately classify 93.83% of movies, suggesting that the model can make correct predictions effectively. The accuracy of the model is demonstrated in greater detail with the confusion matrix in Figure 7.4. The model also achieved a precision of 0.9416, which indicates

Model	F1 Score	Accuracy	Precision	Recall	ROC AUC Score
Logistic Regression	0.9383	0.9385	0.9416	0.9385	0.9374

Table 7.4: Logistic Regression Performance Metrics on Test Data

that the model is able to avoid making false positive errors 94.16% of the time, which can help prevent movie studios from making incorrect decisions and thus save them from a potential failure. The model achieved a recall of 0.9385, which indicates that it correctly identified 93.85% of all actual box office success. An F1 Score of 0.9383, suggests that the model is well balanced in making accurate positive predictions and identifying all positive cases. A ROC AUC score of 0.9374 implies that the model does well in distinguishing a box office success from a box office failure.

7.1.5 Support Vector Classifier Results

Model	F1 Score	Accuracy	Precision	Recall	ROC AUC Score
Support Vector Classifier	0.9574	0.9574	0.9594	0.9574	0.9566

Table 7.5: Support Vector Classifier Performance Metrics on Test Data

As shown in Table 7.5 the SVC model was able to accurately classify 95.74% of movies, suggesting that the model can make correct predictions effectively. The accuracy of the model is demonstrated in greater detail with the confusion matrix in Figure 7.5. The model also achieved a precision of 0.9594, which indicates that the model is able to avoid making false positive errors 95.94% of the time, which can help prevent movie studios from making incorrect decisions and thus save them from a potential failure. The model achieved a recall of 0.9574, which indicates that it correctly identified 95.74% of all actual box office success. An F1 Score of 0.9574, suggests that the model is well balanced in making accurate positive predictions and identifying all positive cases. A ROC AUC score of 0.9566 implies that the model does well in distinguishing a box office success from a box office failure.

7.2 Results Visualization

Continued on Next Page

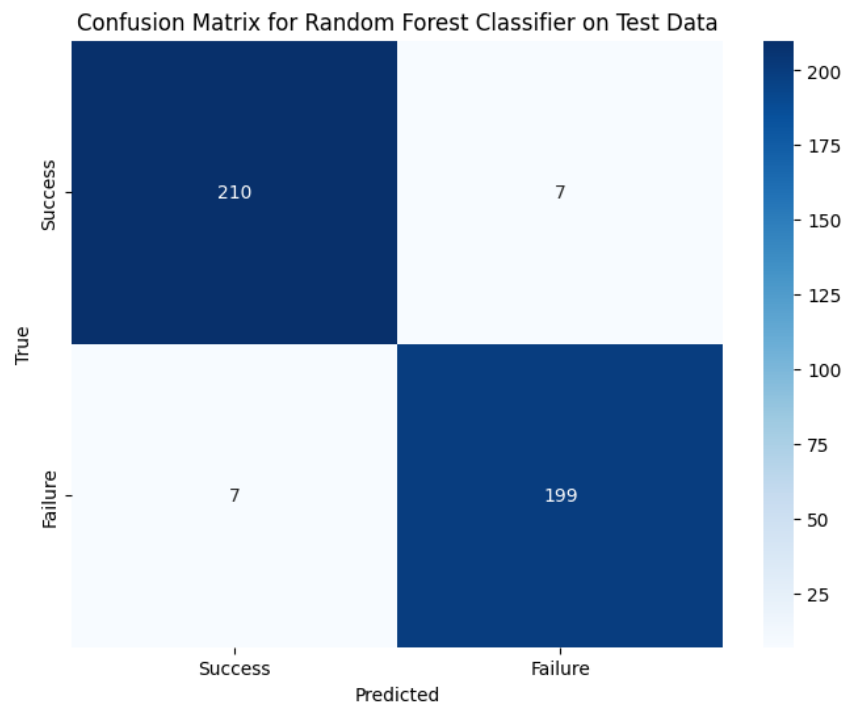


Figure 7.1: Confusion Matrix for Random Forest Classifier on Test Data

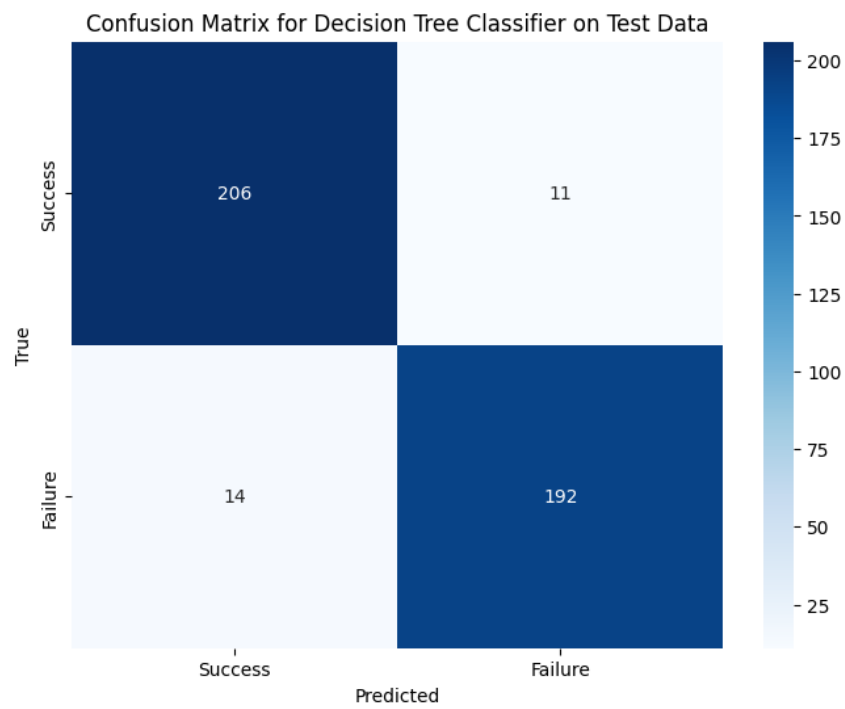


Figure 7.2: Confusion Matrix for Decision Tree Classifier on Test Data

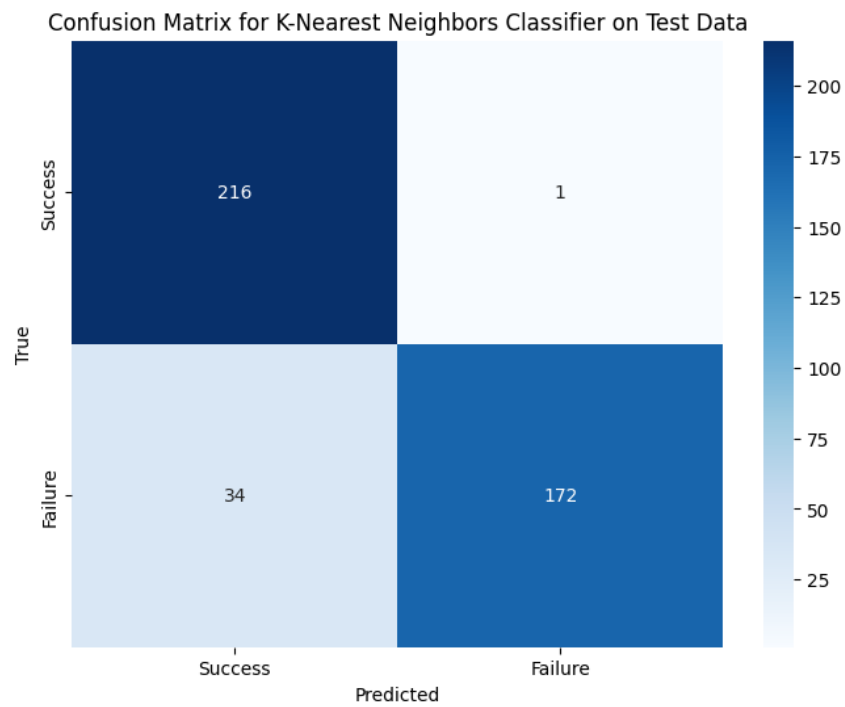


Figure 7.3: Confusion Matrix for K Nearest Neighbors Classifier on Test Data

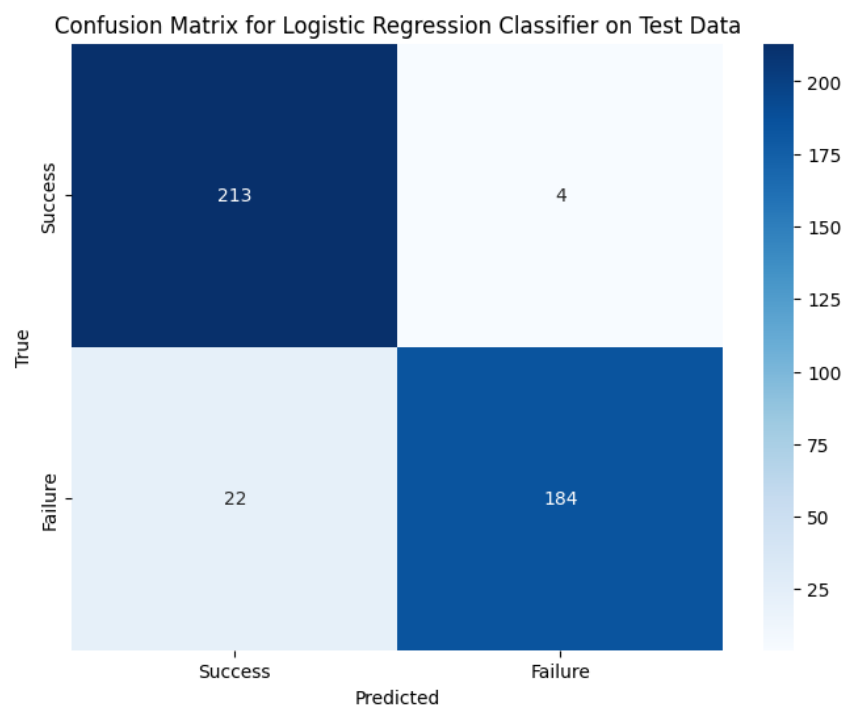


Figure 7.4: Confusion Matrix for Logistic Regression on Test Data

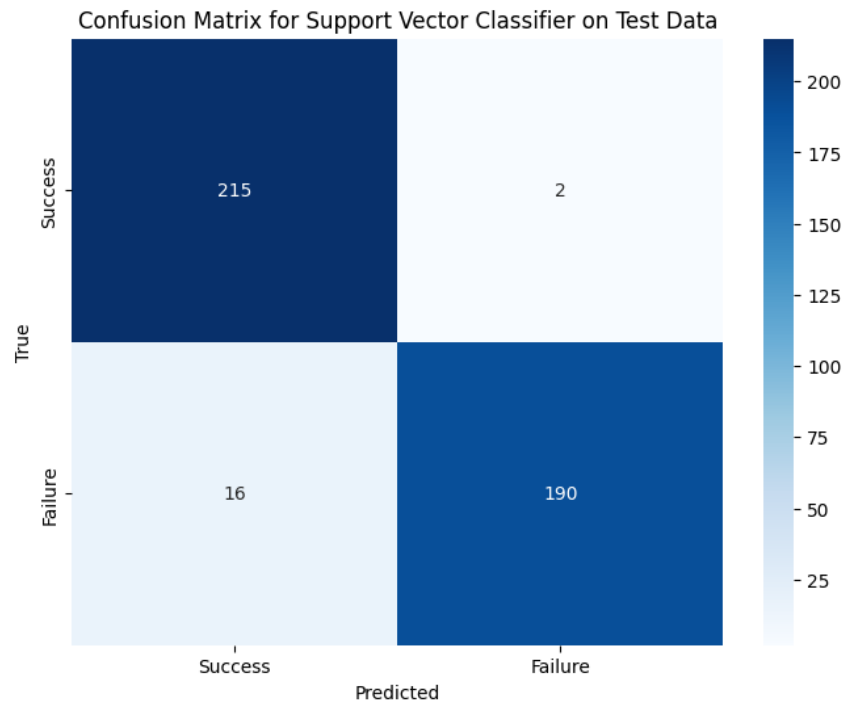


Figure 7.5: Confusion Matrix for Support Vector Classifier on Test Data

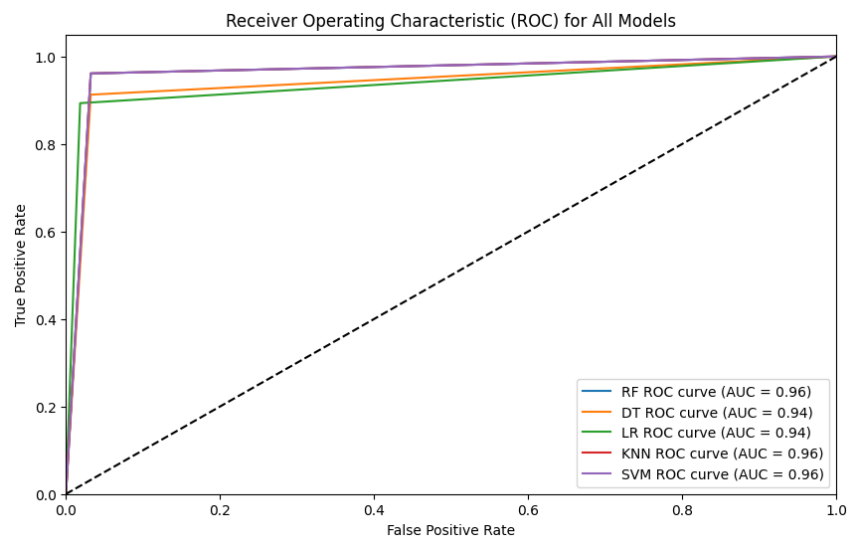


Figure 7.6: Receiver Operating Characteristic (ROC) for All Models

7.3 Combined Results Analysis

Model	F1 Score	Accuracy	Precision	Recall	ROC AUC Score
Random Forest	0.9598	0.9598	0.9598	0.9598	0.9598
Decision Tree	0.9409	0.9409	0.9420	0.9409	0.9406
K-Nearest Neighbor	0.9166	0.9173	0.9274	0.9173	0.9152
Logistic Regression	0.9383	0.9385	0.9416	0.9385	0.9374
Support Vector Classifier	0.9574	0.9574	0.9594	0.9574	0.9566

Table 7.6: Performance Metrics for all Classifiers on Test Data

Table 7.6 shows the combined results for all machine learning models on the test set in descending order of accuracy. All the models tested in this study were able to classify the box office status of movies with an accuracy of over 90%. The Random Forest classifier is seen to perform the best across all models with an accuracy and F1 score of 0.9598. The Support Vector classifier is a close second with an accuracy and F1 score of 0.9574. A ROC AUC score of 0.9598 and 0.9566 respectively, also shows that these models can effectively distinguish between success and failure. Based on these performance scores, it can be argued that these models can provide a reliable basis for movie studios to make informed decisions about which movies they should make and expect a net profit on.

7.4 Comparative Results Analysis

Model	Accuracy
Our Random Forest Model	95.98%
M. S. Neethu and R. Rajasree [4]	90%
M. D. Athira and K. S. Lakshmi [11]	90%
N. Darapaneni et al. [12]	90%
Verma, Garima and Verma, Hemraj [10]	80%
Sivakumar et al. [13]	70%

Table 7.7: Accuracy Comparison with Multiple Papers

Table 7.7 shows the comparison of the accuracy of our Random Forest model with other research papers that have also attempted to predict the box office success of movies. Our model was able to outperform other models in this area of research by a significant margin.

Chapter 8

Conclusion

This study aims to contribute to movie box office prediction by integrating YouTube trailer comments and Wikipedia data. Our approach offers an enhanced understanding of audience sentiments and expectations, enabling us to predict movie box office success accurately. The research outcomes provide valuable insights into the film industry, assisting stakeholders in making informed decisions to maximize box office performance and overall success. However, market trends and unforeseen events can significantly influence predictions' accuracy. Future research can build upon these findings to make further advancements in this field.

8.1 Limitations

While continuing the process of our analysis, we noted multiple limitations regarding the data collection process and strategies. Due to these limitations, the ultimate findings from this analysis may be impacted compromising their stability and robustness.

The limitations that we could identify are as follows:

- **Compact Dataset:** The model's potential for scaling to a wider spectrum of movies may be limited due to assembling only 1780 samples. Such a small sample size can impose limits on the model's ability to provide accurate prediction while tested on fresh unseen data resulting in overfitting and mediocre outcomes. Data augmentation techniques along with opting for transfer learning using pre-trained models might help us deal with this issue.
- **Credibility of Wikipedia Data:** Since Wikipedia is an open-source platform, there is a huge risk of having inaccurate and outdated data. Solely relying on Wikipedia data imposes a huge risk on the validity of the model's prediction compromising its credibility. Cross-referencing of key information with other reliable sources like official datasets, websites, and reports could be a great solution to this problem.
- **Biased Data:** Having considered only American theatrical releases, the dataset used in the analysis is inflicted with several unfortunate biases. Since the selected dataset may not represent a wide portion of international or other

non-theatrical movies, it may limit the generalizability of the prediction analysis. Moreover, movies without sufficient info and trailers were excluded from this research which certainly raises a question of its impartiality towards less prominent films.

- **Confined Sentiment Analysis:** Since only 100 comments were chosen from each movie trailer it may not represent the sentiment of the broader audience. Furthermore, relying only on YouTube comments may not provide a comprehensive overview of public opinion since a huge part of the audience uses social media rather than the YouTube comment section to express themselves. Combining other social media platform comments and reactions could provide a more generic overview while adding an extra dimension to this analysis.

8.2 Future Work

Our goal was to provide a better prediction model to the film industry and with our predictive research, we have been able to conclude a movie's box office performance successfully. Still, there is ample scope for improvement in the future that will make the model capable of ensuring a much more robust, unbiased, and reliable outcome.

The future enhancements that we anticipate include:

- **Incorporating Deep Learning Models:** Integrating deep learning models is a promising strategy to enhance prediction accuracy. In processing sequential and textual data, deep learning architectures for instance transformer models like BERT or recurrent neural networks (RNNs) have demonstrated impressive expertise. Utilizing these models, we can capture more detailed insights from sources such as social media discussions and YouTube comments, identifying slight variances in audience emotions and interaction that may go overlooked by traditional methods.
- **Addition of Diverse Features:** Considering a wider range of significant features like genre, popularity of the cast, release timing, marketing & promotional strategies, audience preferences can certainly enhance the accuracy level of the prediction. The more detailed and diverse information can be collected, the more there will be a chance of capturing a better comprehensive understanding of the factors impacting a movie's success.
- **Integrating Social Media Insights:** For implementing sentiment analysis, surveying versatile indicators for example reactions, comments, tweets, hashtags, and mentions can enhance the performance of the prediction system. Counting in public opinions from various platforms can ensure a holistic overview of audience engagement and track any minor factors influencing their attitudes. Combined with deep learning techniques, this elevated sentiment analysis will enable a more precise and in-depth assessment of a movie's box office potential.

Bibliography

- [1] C. M. Bishop, *Pattern Recognition and Machine Learning*. New York, NY: Springer, 2006.
- [2] K. R. Apala, M. Jose, S. Motnam, C.-C. Chan, K. Liszka, and F. Gregorio, “Prediction of movies box office performance using social media,” Aug. 2013, pp. 1209–1214. DOI: 10.1145/2492517.2500232.
- [3] V. Jain, “Prediction of movie success using sentiment analysis of tweets,” *The International Journal of Soft Computing and Software Engineering [JSCSE]*, vol. 3, no. 3, 2013, ISSN: 2251-7545.
- [4] M. S. Neethu and R. Rajasree, “Sentiment analysis in twitter using machine learning techniques,” Jul. 2013, pp. 1–5. DOI: 10.1109/ICCCNT.2013.6726818.
- [5] M. S. Usha and M. I. Devi, “Analysis of sentiments using unsupervised learning techniques,” Feb. 2013, pp. 241–245. DOI: 10.1109/ICICES.2013.6508203.
- [6] V. Nithin, M. Pranav, P. Sarathbabu, and A. Lijiya, “Predicting movie success based on imdb data,” *International Journal of Data Mining Techniques and Applications*, vol. 3, pp. 365–368, Jun. 2014. DOI: 10.20894/IJBI.105.003.002.004.
- [7] A. Bhawe, H. Kulkarni, V. Biramane, and P. Kosamkar, “Role of different factors in predicting movie success,” Jan. 2015, pp. 1–4. DOI: 10.1109/PERVASIVE.2015.7087152.
- [8] R. Dhir and A. Raj, “Movie success prediction using machine learning algorithms and their comparison,” Dec. 2018, pp. 385–390. DOI: 10.1109/ICSCCC.2018.8703320.
- [9] W. Lu, “Research on prediction of movie box office based on internet comments,” Dec. 2019, pp. 11–14. DOI: 10.1109/ISCID.2019.00010.
- [10] G. Verma and H. Verma, “Predicting bollywood movies success using machine learning technique,” Feb. 2019, pp. 102–105. DOI: 10.1109/AICAI.2019.8701239.
- [11] M. D. Athira and K. S. Lakshmi, “Movie success prediction using ensemble classifier,” Jan. 2020, pp. 1–5. DOI: 10.1109/ICCCI48352.2020.9104183.
- [12] N. Darapaneni, C. Bellarmine, A. R. Paduri, *et al.*, “Movie success prediction using ml,” Oct. 2020. DOI: 10.1109/UEMCON51285.2020.9298145.
- [13] P. Sivakumar, V. P. Rajeswaren, K. Abishankar, E. Ekanayake, and Y. Mehendran, “Movie success and rating prediction using data mining algorithms,” *Journal of Information Systems & Information Technology (JISIT)*, vol. 5, no. 2, pp. 72–80, Feb. 2021, ISSN: 2478-0677. DOI: 10.13140/RG.2.2.18052.86402.

- [14] D. M. Qaseem, N. Ali, W. Akram, A. Ullah, and K. Polat, “Movie success-rate prediction system through optimal sentiment analysis,” *Journal of the Institute of Electronics and Computer*, vol. 4, pp. 15–33, 2022. DOI: 10.33969/JIEC.2022.41002.