

UNDERSTANDING THE POLITICAL OPINION OF SRI LANKANS THROUGH DEEP LEARNING BASED SOCIAL MEDIA DATA ANALYSIS

A.A. Ameen
(209309F)

MSc in Computer Science Specializing in
Data Science Engineering and Analytics

Department of Computer Science and Engineering

University of Moratuwa

Sri Lanka

May 2022

UNDERSTANDING THE POLITICAL OPINION OF SRI LANKANS THROUGH DEEP LEARNING BASED SOCIAL MEDIA DATA ANALYSIS

A.A. Ameen
(209309F)

This dissertation submitted in partial fulfillment of the requirements for the Degree of
MSc in Computer Science specializing in Data Science Engineering and Analytics

Department of Computer Science and Engineering

University of Moratuwa

Sri Lanka

May 2022

Declaration

I declare that this is my own work and this dissertation does not incorporate without acknowledgment any material previously submitted for degree or Diploma in any other University or institute of higher learning and to the best of my knowledge and belief, it does not contain any material previously published or written by another person except where the acknowledgment is made in the text.

Also, I hereby grant to the University of Moratuwa the non-exclusive right to reproduce and distribute my dissertation, in whole or in part in print, electronic or other media. I retain the right to use this content in whole or part in future works (such as articles or books).

Signature:

Date:

Name: A.A. Ameen

The supervisor/s should certify the thesis/dissertation with the following declaration.

I certify that the declaration above by the candidate is true to the best of my knowledge and that this report is acceptable for evaluation for the CS5999 PG Diploma Project.

Signature of the supervisor:

Date:

Name: Dr Surangika Ranathunga

Acknowledgment

I would like to express profound gratitude to my advisor, Dr. Surangika Ranathunga, for her invaluable support in guiding the selection, scoping, and completion of this research study.

Further, I would like to thank all my colleagues for their help in finding relevant research material, for sharing knowledge and experience, and for their encouragement. Big gratitude goes to University of Moratuwa Codestars team and dear friend Sanjeth Kathir for supporting me to complete the thesis. I am as ever, especially indebted to my family for their love and support throughout my life. Finally, I wish to express my gratitude to all my colleagues at work, for the support given to me to manage my MSc research work.

Abstract

With the rising popularity of social media usage, the data generated through it has significantly increased. By focusing critically on these data, certain patterns, traits and opinions can be obtained which can be used for the betterment of the society. This research focuses on finding out the possibility of understanding the political opinion of the country based on these social media data. Understanding the impact of these data will prompt the public to use these platforms to a greater effect thus guiding the decision-makers to make better and informed decisions.

To achieve the above objective English and Sinhala comments from 60 posts from six prominent politicians in Sri Lanka were collected from November 2019 to December 2021. These data was then annotated as positive, negative or neutral sentiments.

6591 annotated comments were used to fine-tune the XLM-RoBERTa (XLM-R) pre-trained model for a text classification task. XLM-R is the new state-of-the-art multilingual masked language model which performs exceptionally well in cross-lingual understanding. To improve the performance of the baseline model a novel approach of adding the context of the comment as a feature in the comment is proposed in the thesis. The XLM-R baseline model achieved an F1 score of 79% while the model using politicians' representation in parliament as a context obtained an F1 score of 91%. The models performed exceptionally well for unseen data as well, when tested with data related to politicians not considered in the training data, the model reported an F1 score of 86%.

Predicting the sentiments using the best model for the latest posts of the six main politicians in the study, the current opinion of the people was derived. Based on it, the Government representatives President Gotabaya Rajapaksha, Prime Minister Mahinda Rajapaksha, and Former President Maithripala Sirisena obtained negative sentiments while Opposition MPs Sajith Premadasa and Anura Kumara obtained positive sentiments.

TABLE OF CONTENTS

Declaration	i
Acknowledgment	ii
Abstract	iii
Chapter 1 INTRODUCTION.....	1
1.1 Research Problem	2
1.2 Research Objectives	2
1.3 Organization of the report	3
Chapter 2 LITERATURE REVIEW.....	4
2.1 Understanding public opinion using Sentiment Analysis Techniques	4
2.1.1 Sentiment Analysis.....	4
2.1.2 Sentiment Analysis to understand political opinion.....	7
2.2 Understanding public opinion using Structural Feature Analysis	10
2.3 Hybrid Approach: Using both Sentiment Analysis and Structural Analysis.....	11
2.4 Evaluating Performance of Text classification tasks.....	12
Chapter 3 METHODOLOGY	14
3.1 Data Collection	15
3.1.1 Data Collection Period	15
3.2 Data Cleaning	16
3.3 Data Annotation.....	16
3.4 Model Building.....	17
3.4.1 Using XLM-R to model the sentiments of the people	18
3.4.2 Model Performance and Evaluation.....	21
3.5 Predicting sentiments for New Politicians.....	21
3.6 Predicting the sentiment for current data.....	22
Chapter 4 ANALYSIS AND RESULTS	23
4.1 Performance evaluation of the models built	25

4.2	Model Validation with new data	26
4.2.1	Results of the Validation data	26
4.3	Predicting the sentiment for the current data.....	27
4.3.1	Understanding the political opinion of people toward the politicians	28
4.3.2	Understanding the overall political opinion of the people	29
Chapter 5	CONCLUSIONS	30
5.1	Conclusions	30
5.2	Limitations and future work	31
Chapter 6	REFERENCES.....	32

LIST OF FIGURES

Figure	Description	Page
Figure 2.1	Sentiment analysis process	5
Figure 2.2	Sentiment classification techniques.....	5
Figure 2.3	Lexicon based sentiment analysis framework	8
Figure 2.4	Machine Learning framework	9
Figure 3.1	Summary of the methodology undertaken in the study	14
Figure 3.2	Summarized view of the data gathering period	15
Figure 3.3	Text classification architecture with XLM-R.	17
Figure 3.4	Illustration of the models built in the study	18
Figure 4.1	Distribution of the sentiments.....	23
Figure 4.2	Sentiment score of the politicians over the period based on human annotation	24
Figure 4.3	Predicted sentiment score for the current data.....	27
Figure 4.4	Sentiment score of the politicians.....	28
Figure 4.5	Overall sentiment score of the Government and Opposition over the years ..	29

LIST OF TABLES

Table	Description	Page
Table 2.1	Fleiss' Kappa value interpretation	13
Table 3.1	Sample comments and Labels used for the baseline model	18
Table 3.2	Sample comments and Labels used for the Language Predictor.....	19
Table 3.3	Sample comments and Labels used for the language tag model.....	19
Table 3.4	Page tag corresponding to the politician	20
Table 3.5	Sample comments and Labels used for the page name as the tag model.....	20
Table 3.6	Tag corresponding to the politician.....	20
Table 3.7	Sample comments and Labels used for the Government tag model	21
Table 4.1	Distribution of comments obtained for the politicians.....	24
Table 4.2	Summary of the sentiment analysis model performance	25
Table 4.3	Summary of the Language predictor model performance.....	25
Table 4.4	The distribution of comments in the validation dataset	26
Table 4.5	Representation in the parliament.....	26
Table 4.6	Performance of the models for the new data.....	26
Table 4.7	Predicted sentiments for the current data	27

Chapter 1 INTRODUCTION

Public opinion represents people's collective preference on matters of interest. These opinions will be varying among people based on their expression of beliefs and ideologies. In the political context, public opinion is the collective sentiment of people about the rulers of the country/region and the actions and decisions, they make [1]. The concept of public opinion arose with the rise of democracy, premised on the notion that governments should rule with the consent of the governed.

Understanding public opinion is important as it guides governments, influences public policies and gives feedback to politicians. Assessing the opinion of the people is a difficult task as it varies from person to person and changes over time too. Governments conduct various kinds of sentiment analyses of people's opinions towards them. Intelligence agencies use different types of techniques to access the opinion of the people. At the time of elections, universities and non-government organizations also conduct surveys and opinion polls to understand these sentiments and opinions of people and try to predict how they'll impact the elections. Using this information, governments are supposed to take action and course-correct to win back the support of the people. The parties and politicians who understand these sentiments of the people better tend to perform well and win people's support in elections. Sri Lanka too followed the traditional methods to understand public opinion but the inception of social media in Sri Lanka has created a medium for the people to communicate and give feedback while also creating a platform for the governing bodies to influence people, which has changed the whole political game in Sri Lanka.

Currently, there are plenty of social media platforms that have over 4.2 billion users around the world, equating to more than 53 percent of the total global population [2]. With the increase of internet users in Sri Lanka, in 2020 there were 6.4 million social media users [2]. Social media penetration stands at 30% in 2020 and is expected to grow significantly in the future [2].

Social media platforms such as Facebook and Twitter have given Sri Lankans the ability to share their thoughts, reactions, and feedback to all sorts of things occurring in Sri Lanka and around the world. In the Facebook context, people can share their views in terms of posts, comments, shares, and post reactions through emojis (eg: Like, Wow, Angry, Haha). These platforms have created an opportunity and a medium of communication also in the political context between ordinary people and governing bodies. People can communicate to views on the updates given by politicians and news items also post grievances/issues directly to the politicians. For politicians, these platforms help to update people on what they are doing, influence the public to support them, and most importantly understand the sentiment of people towards them and their political parties.

1.1 Research Problem

All these interactions in social media are valuable data that can be analyzed and used for greater effect. Using these data, studies in some countries have shown that the political opinion of people can be mined and has been used as a tool to win elections [3], [4], [5]. Even though there have been studies conducted in finding the political opinion in Sri Lanka [6], [7] none has used social media data. Despite Sri Lankan politicians and the public often believe and use the trends and sentiments on social media platforms as public opinion there have been no formal studies that have been conducted to investigate their relationship. Thus, a formal study on this area would provide an affirmation of the relationship.

In the Sri Lankan context, gauging public opinion through social media data is challenging with the use of different means of communication and multiple language usage on Sri Lankan social media platforms. Despite the challenges, there are some studies conducted in the Sri Lankan context to analyze the sentiments from the text in social media data in the field of sports [8], and transport [9] but no study has been conducted to analyze the sentiments in the field of politics. Also, in the Sri Lankan context, no studies have used English, Sinhala, and English and Sinhala mixed social media to extract the sentiments.

Research on this area will enable to identify whether the interactions happening on social media can be used as a proxy for people's sentiments and opinions. It will help to identify the impact of these data and will urge people to use them to a greater effect to shape, guide, and direct rulers to make better decisions and steer us to a better Sri Lanka.

1.2 Research Objectives

The objectives of this research are:

- To build a process to capture data from Facebook posts of politicians and feed it into the analysis.
- To build a model to assess the sentiment of the public using the Sinhala and English comments on Facebook Posts of Sri Lankan politicians and political parties.
- To identify public opinion over different politicians using the built models and identify the overall political opinion of the public at a given time.

1.3 Organization of the report

This dissertation consists of six chapters inclusive of the Introduction chapter. The rest of the thesis is as follows, with a detailed introduction which comprises the background, the research objectives, and the significance of the research were explained in Chapter 1. Chapter 2 inaugurates the research by providing the literature on the area of study. Moreover, in Chapter 3 the methodology undertaken in the study is explained. Chapter 4 provides the results of the analysis and the insight derived from the study. To wrap up things, chapter 5 concludes the study by highlighting the general discussion, summarizing key findings, limitations, and future work along with the conclusion of the study. The reference used in the thesis is stated in chapter 6.

Chapter 2 LITERATURE REVIEW

Understanding public opinion is vital as it gives governing bodies an indication of their success of how they are governing the country. Even though formal ways of studying public opinion are relatively new but practical study of it has been there as long as there have been governments [10]. Early rulers tend to measure public opinion through rebels and tax receipts. Higher the rebels and lower tax collection was an indication that public opinion was turning against them [10]. In the modern democratic environment, researchers have developed many methods to understand public opinion using formal academic studies.

The most used method to understand public opinion is through sample surveys. In sample surveys, public opinion is obtained through analyzing the answers given to a few questions related to a particular issue of interest by a certain sample of the public. Survey studies differ based on the types of questions asked the order of the questions, sample size, sampling methods and timing of polls in relation to an election or an event, the approach to dealing with non-response, and the context surrounding the poll [11].

The emergence of new technologies such as smart devices and social media platforms has opened new avenues through which researchers access and analyze public opinion. As these technologies grow so does the access to public opinion; researchers find a way to the thoughts, feelings, and actions of the public. This has in turn led to a transformation of methods in accessing and sharing opinions, attitudes, and behaviours and the trends keep evolving with time. Thus, the abundance of social media and the opinion of the people for easy access has enabled the development of new data collection tools and new sources of qualitative and quantitative data to enhance or provide alternatives to the traditional data collection methods [12]

Many of the recent studies in this space use a variety of social media data to identify the political opinion of the public. These studies can be categorized as studies using the sentiment-based approach, studies using structural feature analysis of social networks, or studies with a combination of both. On-demand interest in gauging political opinion is to predict the results of elections as it is a validation of the accuracy of the analysis.

2.1 Understanding public opinion using Sentiment Analysis Techniques

2.1.1 Sentiment Analysis

The most used method in gauging public opinion using social media is to use sentiment analysis of social media data content. Sentiment Analysis (SA), which is also known as opinion mining or contextual mining, is used in Natural Language Processing (NLP), computational linguistics which helps in identifying, systematically extracting and

quantifying the subjective information in the text to understand the opinions and sentiments stated in a particular text [13].

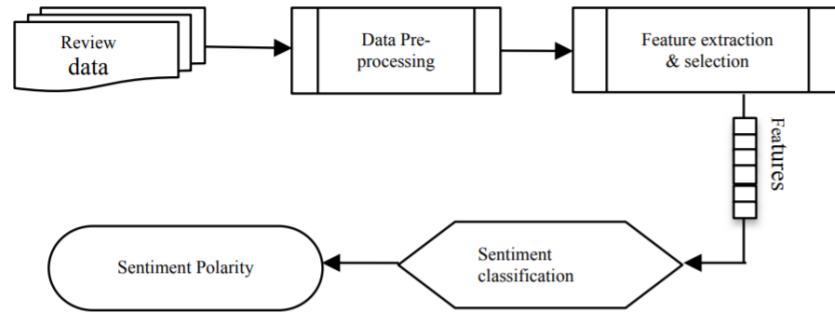


Figure 2.1 Sentiment analysis process

Source: Adapted from [13]

Figure 2.1, shows the Sentiment analysis process which needs to be carried out to get the sentiments out of text contents [13]. To classify sentiment requires data pre-processing on the collected data, feature extraction, sentiment classification, and evaluation. Data pre-processing includes tokenization, stop word removal, stemming, lemmatization, etc. Using appropriate text pre-processing methods including data transformation and filtering can significantly enhance the performance of the sentiment classifier [14]. In the feature extraction phase, the text is converted into a feature vector using different methods. Term Presence versus Term Frequency, N-gram Features, Parts of Speech and Term Position are the most commonly used methods in Sentiment Analysis to obtain features [13].

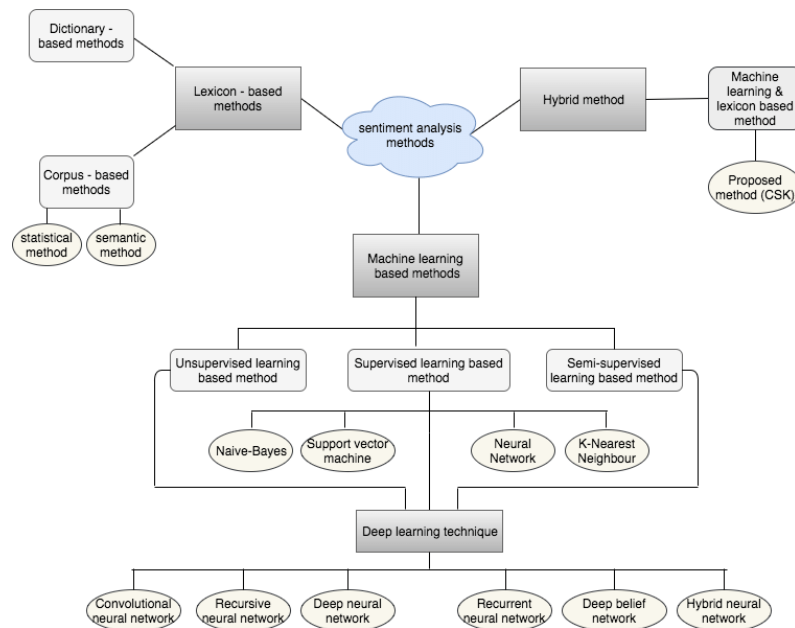


Figure 2.2 Sentiment classification techniques.

Source: Adapted from [15]

Once the features are extracted, the sentiment of the text is obtained through either using a Lexicon-based approach, a Machine Learning approach, or using a hybrid approach. Figure 2.2 summarizes the techniques used in classifying the polarity in sentiment analysis.

Lexicon-based techniques were the first to be used for sentiment analysis tasks and they are divided into two approaches: dictionary-based and corpus-based approaches. [16]. The dictionary-based approach looks for opinion seed words in a text and it searches its synonyms and antonyms from a dictionary [13]. The other approach is using a corpus that starts with a set of seed lists of opinion words and using opinion words in a large corpus helps find context-specific orientations of opinion words. This could be done by using statistical or semantic methods [13].

Machine Learning techniques are the techniques which classify sentiments by using Machine Learning algorithms on text data. Machine Learning techniques can be categorized mainly into supervised and unsupervised methods of learning [16]. Supervised methods of learning rely on labeled training manuals. Using an annotated/label dataset as a training dataset these models are developed. Unsupervised learning techniques do not use pre-labeled data to train the classifier. They find associations in the data to cluster the text into opinions [17]. Machine Learning techniques use traditional models and Deep Learning models to obtain the polarity of a text. Traditional models refer to classical Machine Learning techniques like Support Vector Machine (SVM), K nearest neighbors (KNN), and Decision Trees [18]. These models use features extracted from the text as inputs to build the models. The more commonly used unsupervised Machine Learning algorithms are K means and Apriori Algorithms [17]. The accuracy of the Machine Learning models depends on which features are chosen and the size of the training dataset available [18].

With the increase in the use of neural networks, text classification tasks such as sentiment analysis have seen an increase in the use of Deep Learning-based models as opposed to the traditional Machine Learning-based approaches [19]. Neural approaches have been explored to address the limitations due to the separate process required in feature extracting in Machine Learning models. In Deep Learning models, the core component is to use a machine-learned embedding model that maps text into a low-dimensional continuous feature vector, thus not requiring to extract features manually using certain logics [19]. The simplest Deep Learning model used in this task is the Feed-forward network which view a text as a bag of words. RNN-based models view a text as a sequence of words and are intended to capture word dependencies and text structures for text Classification tasks. Among many variants to RNNs, Long Short-Term Memory (LSTM) is the most popular architecture, which is designed to better capture long-term dependencies and is commonly used in sentiment analysis [19]. In a sentiments analysis classification on the Stanford Sentiment Treebank (SST) dataset [20], Deep Learning models fitted using FFN, LSTM models, and other Deep Learning models such as CNN and BERT outweighed Machine Learning-based Naïve Bayes model on accuracy [19].

Comparing the Deep Learning models, models which use transformer-based architecture have improved the performance of text classification tasks vastly as they capture the context from whole sentences rather than individual words when encoding text to word embeddings [21]. BERT [22] developed by Google is a transformer-based model and uses bidirectional training using MLM (Masked Language Model) and NSP (Next Sentence Prediction) techniques to learn the text representation. To classify text in multiple languages including low resource languages, mono lingual model like BERT [22], RoBERTa [23], multilingual models like mBERT which is an extension of BERT, XLM [24], and XLM-R [25] were introduced. Using transfer learning these models can be used for text classification tasks of any domain by fine-tuning with specific data and can achieve higher performance even with less labeled data [21]. XLM-R model which was proposed by Facebook [25] is the new state-of-the-art multilingual masked language model trained on 2.5 TB of data created from Common Crawl data in 100 languages. XLM-R has shown vast improvement over multilingual models such as mBERT and XLM on classification, sequence labeling, and question answering downstream tasks [25]. It is the first multilingual model to outperform traditional monolingual baselines as it performs well on low-resource languages. It also achieves state-of-the-art performance in cross-lingual understanding, where the model is trained in one language and then used with other languages without additional training data [25].

2.1.2 Sentiment Analysis to understand political opinion.

The public gives their opinion on various subjects through social media platforms using tweets, stories, status updates, and comments on social media platforms. By analyzing this data, it is possible to gain an understanding into the public opinion on any subject of interest. By performing sentiment analysis of a given domain, it can identify the effect of domain information on sentiment classification [26]. In the political domain, studies extract positive and negative opinions as a proxy for voting behavior.

Meduru et al. [26] used a lexicon-based approach to gauge public opinion with respect to government issues and political reforms. The researchers used tweets from Twitter which includes specific hashtags and further analyzing the data, the opinion of the public was gauged.

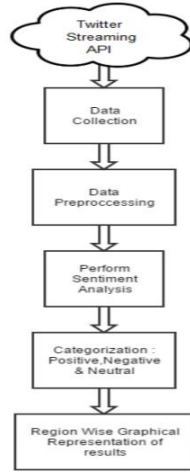


Figure 2.3 Lexicon based sentiment analysis framework

Source: Adapted from [20]

Figure 2.3 shows the steps undertaken by the researchers in their analysis. In the preprocessing step, data is cleaned by removing stop words, slang, URLs, hashtags, emojis, and special characters. Using a predefined lexicon called Vader, the sentiment polarity of the tweets is measured and categorized as positive, negative, or neutral. Using the geographical data obtained from profiles of people who tweeted the tweet, geographical region-wise analysis is further conducted to see how the views are different between the regions.

A similar approach was followed by Tumasjan et. al [27], using a pre-existing lexicon to identify the sentiments. Researchers used a LIWC2007 [28] tool for extracting sentiment from the tweets to predict the German Federal election in 2009. LIWC is a text analysis software that revealed people's thoughts, emotions, cognition, and personalities using text samples. They analyzed over 100,000 tweets identifying either a politician or a political party and concluded that the sentiments obtained are highly correlated with the chances of winning the election. Burnap [29] used sentiment analysis on a set of tweets related to elections to predict the United Kingdom General Election in 2015. In their research, they applied sentiment analysis using software developed by Thelwall [30] which allocates a string of text a positive and negative score ranging from -5 (extreme negative) to +5 (extremely positive), where each score is produced based on words in the string that are known to carry such emotive meaning (eg 'love'=5, hate='-4'). They calculated sentiment scores for each tweet and applied a rationale that tweets containing party or leader names can be treated as vote intentions, they calculated an overall sentiment score for the political party to predict the winners of the election.

Most studies in finding political opinion using social media content have used Machine Learning approaches to classify the sentiments [1]. Marko [1] pointed out the difficulty in using pre-existing lexica to identify the sentiments as “(1) incorrect classification of the word in the lexicon, (2) the mismatch of parts-of-speech (POS) in tweets due to their poor

grammar, (3) lack of word disambiguation, and (4) their reliance on word polarity rather than context inference”. Their investigation also showed that sentiment analysis using the Machine Learning approach has obtained better results compared to analysis that used existing lexica.

Using the Machine Learning approach suggested by Elghazaly [31], tweets were classified supporting the contestant of the Presidential Election of Egypt in 2012. Figure 2.4 outlines the framework of the study undertaken by the researchers.

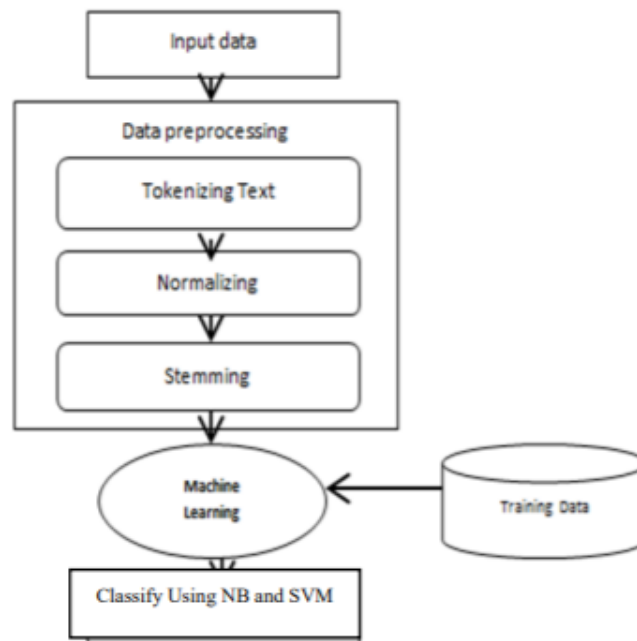


Figure 2.4 Machine Learning framework

Source: Adapted from [30]

Researchers have collected over 18,000 tweets and each tweet is then annotated to positive or negative to create the labels for training data. Preprocessing steps involve tokenization, normalization, and stop word removal. Finally using the term frequency-inverse document frequency (TF-IDF) each tweet is represented as a weighted vector of the terms found in the super vector. Using this matrix as features and annotated polarity as labels, a Machine Learning model is built using Support Vector Machine (SVM) and Naïve Bayes (NB) algorithms. Based on the result, Naïve Bayes (NB) model had a better result compared to the SVM model but both models had more than 80% recall and precision. Researchers believe that using this model future tweets can be classified as supports of the candidates and can be used as a gauge to predict the public support to candidates in the election.

Felipe et al. [32] used a similar Machine Learning approach to compute the sentiment polarity of tweets and then used a time-series approach to forecast public opinion in the

future. The polarity of over 250,000 tweets for each candidate in the Presidential Election of the United States in 2008 was computed and modeled as time series data based on the occurrence of the tweet. ARIMA, ARMA, and GARCH models were fitted to these data and forecasted the polarity of people's sentiments up to the next six months. Based on these forecasts, how opinion dynamics move can be assessed and used in election campaigns.

2.2 Understanding public opinion using Structural Feature Analysis

Using structural features of social media is another way that researchers have used to predict public opinion in the political context. Researchers have used structural features like the number of followers, the number of likes, number of friends, mentions, comments, shares, etc. to understand the popularity in the social media space and use it as a proxy for the public opinion and election result prediction.

Cameron et al. [33] measured the popularity of politicians in the general election of New Zealand in 2011 by using the number of friends on Facebook and the number of followers on Twitter of the candidates. The number of friends/followers in Facebook and Twitter, 2 months prior, 1 month prior, 1 week prior and 1 day prior of the election of 453 candidates were used as explanatory variables and actual results was used as response variable to build a Machine Learning model to understand the relationship. An ordinary least square regression model was built to understand the factors influencing the voter share of the election while a classification model using logistic regression was built to understand the capability of structural data in social media platforms to correctly predict the winners and losers of the election. Results of this research found out that there is a significant relationship between the size of the social media networks and election voting and election results and will be a crucial variable to look at in a closely conducted election.

Reima et al. [5] used Facebook likes to predict the results of the 2015 Finnish Parliamentary elections. It attempted to predict the Finnish Parliamentary election which took place in April 2015. Data from official Facebook walls of Politicians who ran for the election were gathered in this study. Over 750,000 Facebook likes from 1131 candidates were used in the analysis. Then the data was used to create an electoral forecast. The results compared with the actual election result indicated that the forecast of the election only based on Facebook "likes" was not accurate, but the research discovered a clear statistical relationship between "likes" and votes. A similar kind of research was done in the Indian context to predict the Lok Sabha elections in 2014. The number of 'likes' in different time frames from Politicians' verified Facebook pages and Pages of Political Parties were used as input data for the study. A strong positive correlation between voter share obtained in the election and the number of 'likes' a political party or its party leader had on their official Facebook page was observed in this study. It also concluded that the no. of likes month before the voting period was the best variable to predict the vote share using 'likes'—with 86.6% accuracy [3].

2.3 Hybrid Approach: Using both Sentiment Analysis and Structural Analysis

Some studies use both structural data in social media platforms as well as sentiments analysis of text to understand the public opinion of the country. Kalampokis et al. [4] exploit both volume and sentiment-related variables to create a predictive model to predict the votes of three major parties in the General Election of the United Kingdom in 2010. They define two predictor variables, Relative Frequency (RF) and Positive Negative Ratio (PNR) of tweets in the analysis. The relative frequency of a political party is the tweet volume per day for a party while the PNR ratio is used to quantify the sentiments expressed in tweets about the party. PNR is the ratio between positive and negative sentiment tweets per day. This study used a Machine Learning approach to compute the sentiments of tweets and using DynamicMLClassifier the tweets were classified as positive or negative. DynamicMLClassifier is a language model classifier that accepts training events of categorized character sequences. Training is based on a multivariate estimator for the category distribution and dynamic language models for the per-category character sequence estimators. Using these two variables and polling data a regression model was built to predict the voter share of the General election.

With the introduction of emotions of reactions in 2016, the like button in posts on Facebook was redesigned to have additional reactions "love", "haha", "sad", "wow", and "angry" along with the reaction "Like". With this development, investigators used these reactions as a sentiment analysis tool in research. Tian et al. [34] suggested that "Facebook reactions are a good data source for identifying people's feelings". They claimed that Facebook reactions ("love", "haha", "sad", "wow", and "angry") express the overall sentiment of Facebook users. In this study, the researchers analyzed both pro and anti-comments and observed the reactions/emojis used for these comments types. They analyzed over 21 thousand Facebook posts, 72 million reactions, and 8 million comments on Facebook across 4 countries in this study.

In a study done in Mexico in 2017 to analyze if the Facebook reactions of users reflected the outcome of the elections and the user emotion towards the candidates, reactions from 4128 Facebook posts were analyzed [35]. Researchers used the "like", "love", "haha", "wow", "sad", and "angry" reactions in this analysis. In this study, they categorized the reaction 'Like' in a particular post as a positive expression towards the content of the post while the reaction 'Love' was categorized as very positive. Similarly, the reaction 'Sad' was categorized as a negative expression, and the reaction 'Angry' as a very negative expression towards the content of the post. The expressions 'Haha' and 'Wow' were expressed as neutral as they had mixed expressions and they defined their polarity based on positive or negative expressions of the post based on other reactions. They defined the below formulas to measure the sentiment polarity for each candidate in terms of the reactions obtained in their Facebook Posts.

$$\text{Total Positive Sentiment} = \text{Wow} + \text{Likes} + 2.\text{Love}$$

$$\text{Total Negative Sentiment} = \text{Haha} + \text{Sad} + 2.\text{Angry}$$

Using these sentiments along with analyzing the impacts of comments and shares, Sandoval [35] were able to measure the impact of posts by political parties and the sentiment of Facebook users toward each political party during the political campaign. Their findings showed that despite not being the party with the largest positive sentiment, the winning political party won the election with a lesser number of posts and fewer comments on candidates in the Facebook comments section.

2.4 Evaluating Performance of Text classification tasks

In text classification, it is required to evaluate the data set used and the models built to validate and improve the performance of the outcome. To check the validation of the dataset, inter-agreement scores of the judges who annotated the data are measured. It measures how uniformly the judges understood the text and its sentiments, and the reproducibility of the annotation task in the consideration. To measure this agreement score, the most used approaches are to use Cohen's Kappa score and Fleiss' Kappa Score [36].

Fleiss' Kappa Coefficient (K) is an extension of Cohen's Kappa Coefficient and is a statistic that is used to measure inter-annotator agreement for qualitative data. Cohen's Kappa measures how often two annotators agree with each other while Fleiss' Kappa works for any number of annotators, after taking into account the possibility of they would agree by chance [37]. The equation for Fleiss' Kappa measure is as follows

$$K = \frac{P - P_e}{1 - P_e}$$

Where,

$P - P_e$ = Degree of the agreement actually achieved above chance

$1 - P_e$ = Degree of agreement that is attainable above chance

Fleiss' Kappa's score ranges from -1 to 1. A score closer to 1 means there is a full agreement between the annotators while a score closer to zero or less means that the annotators have no agreement between them [37]. The below table summarizes the Kappa score's interpretation for different ranges.

Table 2.1 Fleiss' Kappa value interpretation

Value Range	Fleiss's Interpretation
<0	Poor agreement
0.01-0.20	Slight Agreement
0.21-0.40	Fair Agreement
0.41-0.60	Moderate Agreement
0.61-0.80	Substantial Agreement
Above 0.80	Almost perfect Agreement

Source: Adapted from [37]

Accuracy, Precision, Recall, and F1 score are the popular measures to check the correctness of the classification in a text classification task [38]. Accuracy measures the percentage of text which are correctly classified, Precision computes the percentage of text the classifier got right out of the total number of examples that is predicted for a given tag, Recall measures the percentage of examples the classifier predicted for a given tag out of the total number of examples it should have predicted for that given tag. F1 score is the harmonic mean of recall and precision. Based on the task in consideration, each of these measures is computed and compared to obtain a better model [38].

Chapter 3 METHODOLOGY

This chapter discusses research methodology of the dissertation. It also discusses the fundamental theoretical background used to achieve the objectives.

Figure 3.1 illustrates the summary of the methodology undertaken in the study. The approach and techniques used under each step are comprehensively explained below.

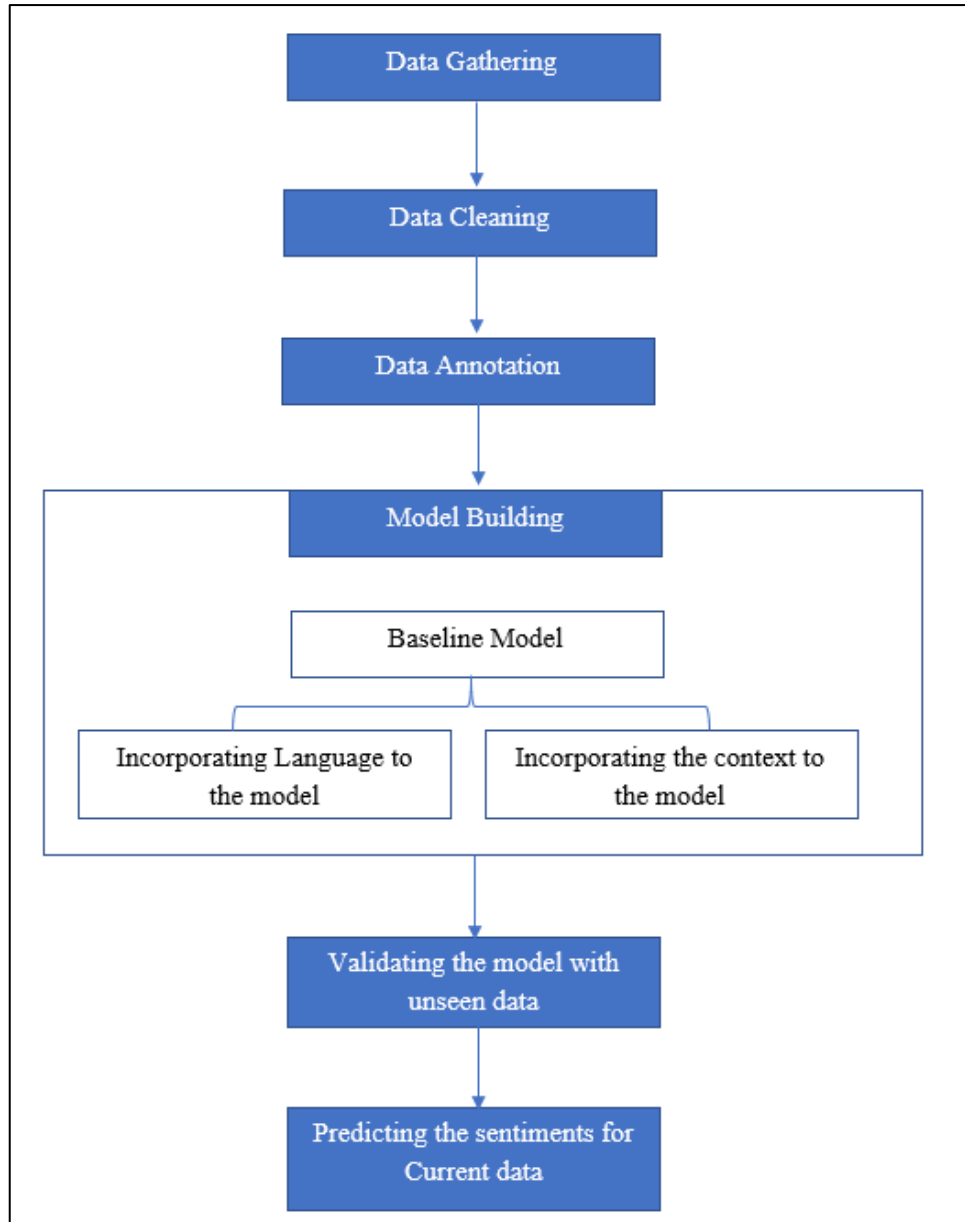


Figure 3.1 Summary of the methodology undertaken in the study

3.1 Data Collection

Data required for the analysis were gathered from the Facebook pages of 6 prominent politicians of Sri Lanka. The basis of selection was based on seniority and the level of influence they have in the Sri Lankan political framework.

Following are the six politicians from whom the data was gathered.

1. Gotabaya Rajapaksha – The President of Sri Lanka
2. Mahinda Rajapaksha – The Prime minister of Sri Lanka
3. Maithripala Sirisena – The former President of Sri Lanka
4. Ranil Wickramasinghe – The former Prime Minister of Sri Lanka
5. Sajith Premadasa – The Opposition Leader of the Parliament of Sri Lanka
6. Anura Kumara Dissanayake – Party leader of National People’s Power

3.1.1 Data Collection Period

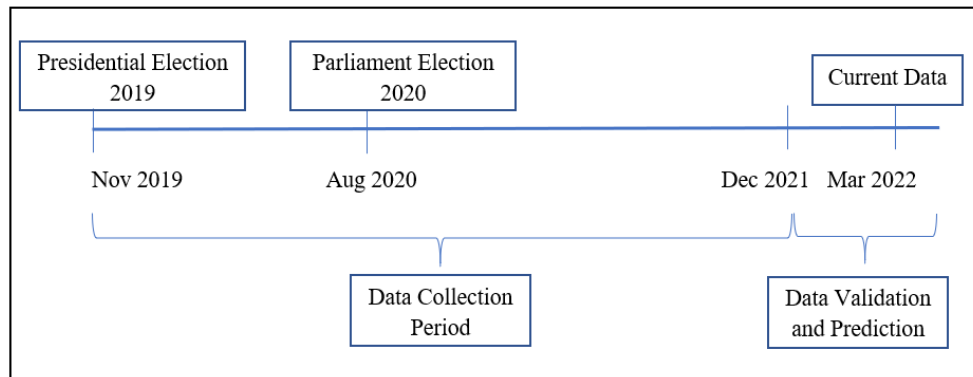


Figure 3.2 Summarized view of the data gathering period

As shown in Figure 3.2, Data for the preliminary analysis were collected from November 2019 to December 2021. In this period, there were two main elections held in Sri Lanka namely, the Presidential election of 2019 (November 2019) and the General Election of 2020 (August 2020). The basis for the selection of the above period was that with the elections being held, the public actively used social media platforms to express their opinion, and gauging sentiment out of it is important. Using the election results as a proxy for people's sentiments, it could be checked whether there is a correlation between the sentiments from data and election results.

For each of the six politicians in the study, 10 Facebook posts from this period were selected at random while ensuring at least one post is captured for 2019, 2020, and 2021. Comments from these 60 posts were analyzed and modeled to bring out the sentiment of the people of Sri Lanka.

For the validation of the model to unseen data, comments from the three latest posts from the politicians Namal Rajapaksha, Harin Fernando, and Wimal Weerawansa's Facebook pages were also used.

Finally, for the prediction of the current sentiment of people, the model built was used to predict the sentiment of the three latest posts of the six main politicians in the study.

3.2 Data Cleaning

The scraped comments from the Facebook posts of the politicians were then cleaned to bring out meaningful patterns in the data. Below data processing steps were conducted to cleanse the data.

1. Removed picture comments
2. Removed mentions/tags
3. Removed links
4. Removed emoji only comments
5. Removed transliteration comments
6. Removed comments that were not in Sinhala or English

Since mentions/tags and links do not give out any sentiments, they were removed from the study. Assessing the sentiments from picture comments and emotes require image processing techniques and hence was opted out. Since the objective was to build a model using Sinhala and English data, the non-Sinhala and non-English data along with transliterating data was removed from the study.

3.3 Data Annotation

Data annotation plays a key role in text analysis as labeled data are necessary for machine and Deep Learning studies so that machines can interpret the input sequence accurately and clearly. The comments obtained after cleaning were annotated under three labels "Positive", "Negative" and "Neutral".

- Positive – The comment agrees with the post or the comment is in support of the politician.
- Negative - The comment disagrees with the post or the comment is not in support of the politician.
- Neutral – It is not clear if the comment agrees or disagrees with the post or it is not clear if the comment is in support of the politician or not.

The dataset was annotated by three judges and the inter agreement score was measured for the data using the Fleiss' Kappa score to further proceed with the analysis.

3.4 Model Building

In order to model the data to bring out the sentiments of the people, XLM-R pre-trained model was fine-tuned in this study.

XLM-R is a multilingual version of RoBERTa. It was trained with the Masked language modeling (MLM) approach. In the MLM approach, 15% of the words in a sentence are masked randomly and then the model is trained to predict the masked words. This way, the model learns an inner representation of the 100 languages which can be used to extract features for downstream tasks.

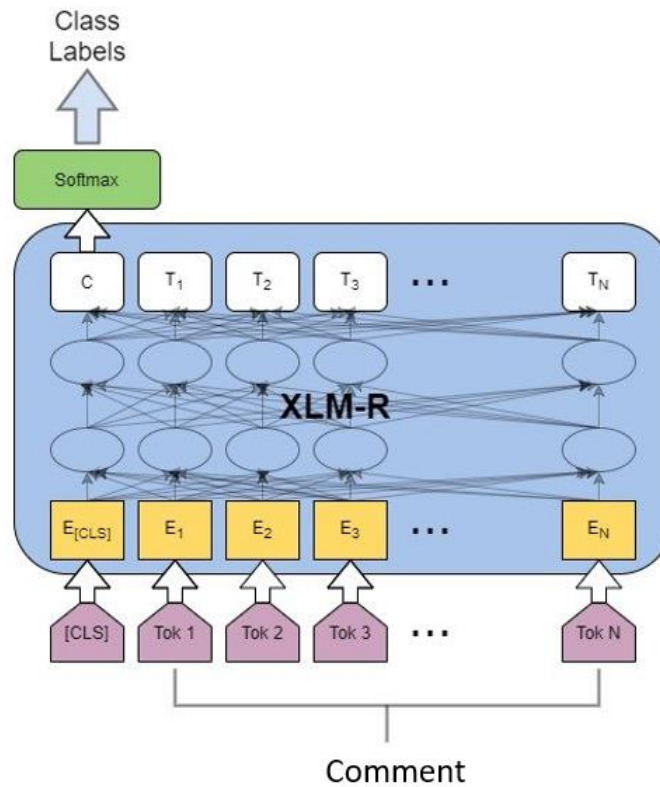


Figure 3.3 Text classification architecture with XLM-R.

Source: Adapted from [39]

XLM-R model contains a total of 12 layers which different semantic information [40]. It contains 768 hidden states, 3072 feed-forward hidden states, and 8 heads [39]. It takes the input of a set of tokens no more than 512 tokens and outputs the representation of the sequence. When classifying text, XLM-R takes the final hidden state of the first token [CLS] as the representation of the whole sequence. Softmax classifier is added to predict the labels from the model [39].

3.4.1 Using XLM-R to model the sentiments of the people

As shown in figure 3.4, five models were developed by fine-tuning XLM-R models to predict the sentiments from cleaned and annotated comments obtained from the Facebook posts of the 6 main politicians.

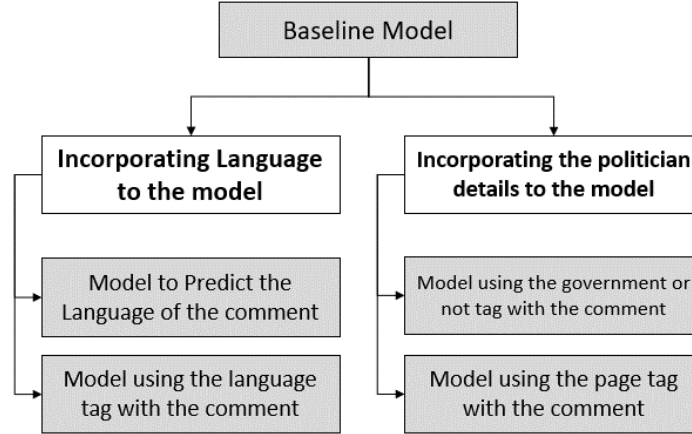


Figure 3.4 Illustration of the models built in the study

1. Baseline Model

Using the extracted comment as sentence sequence and annotated sentiments as labels the Baseline model was created by fine-tuning the XLM-R model for all English, Sinhala, and English – Sinhala mixed comments.

Table 3.1 Sample comments and Labels used for the baseline model

Comment	Label
අනිවාර්යයෙන් මෙවර අපි ජවිපෙ දිනවනවා.	Positive
Superb leadership! Salute you sir...!!!	Positive
අපින් ගොන්නු උනා... ගෝටා වෙනසක් කරයි කියලා ඔය කොරෙ	Negative
Waiting for "අපි නමයි හොඳම කලේ.	Negative
බැන්කෝ හදාගන්න නේ මැට්ටෝ	Neutral

In order to improve the performance of the sentiment analysis task, a novel approach is introduced in this study by applying certain context as features to the comment and adding it as a tag in the comment when fine-tuning XLM-R models.

2. Incorporating the language of the model as Context

In this study, comments from the English language, Sinhala Language, and Sinhala-English code-mixed are considered. According to Fishman [41], the choice of language of people is based on the situation and carries certain sentiments. To check if the language of the comment in the Sri Lankan context can be used as a feature to improve the

performance of the sentimental analysis task, the language of the comment was further processed in this study. Following two models were built.

1. Predicting the language – XLM-R model was fine-tuned to predict the language of a comment using the dataset.

Table 3.2 Sample comments and Labels used for the Language Predictor

Comment	Language
අනිවාර්යයෙන් මෙවර අපි ජවිපෙ දිනවනවා.	Sinhala
Superb leadership! Salute you sir..!!!	English
අපිත් ගොන්නු උනා... ගෝටා වෙනසක් කරයි කියලා ඔය කොරේ	Sinhala
Waiting for "අපි තමයි හොදටම කලේ.	Mix
බැන්නේ හදාගන්න නේ මැට්ටෝ	Sinhala

2. The predicted language was appended as a tag at the beginning of the sentence and the XLM-R model was fine-tuned to predict the sentiments.

This approach was used by Google [42] and showed promising results in multilingual Neural Machine Translation system. It improved the translation quality of all the language pairs while keeping other model parameters unchanged.

Table 3.3 Sample comments and Labels used for the language tag model

Language	Tag	Updated Comment	Label
Sinhala	<sn>	<sn>අනිවාර්යයෙන් මෙවර අපි ජවිපෙ දිනවනවා.	Positive
English	<en>	<en>Superb leadership! Salute you sir..!!!	Positive
Mix	<mx>	<mx>Waiting for "අපි තමයි හොදටම කලේ.	Negative
Sinhala	<sn>	<sn>බැන්නේ හදාගන්න නේ මැට්ටෝ	Neutral

3. Incorporating details of politicians as context

When analyzing the labeled dataset, it was found that the same comment on two different politicians' pages, could be interpreted as different sentiments and be labeled differently. Following are two examples of the above issue.

Page	Comment	Label
Gotabaya Rajapaksha	"ස" තමා හොදටම කලේ	Negative
Sajith Premadasa	"ස" තමා හොදටම කලේ	Positive

Page	Comment	Label
------	---------	-------

Gotabaya Rajapaksha	We still trust our president	Positive
Anura Kumara	We still trust our president	Negative

To reduce the above ambiguity, it was necessary to give some context to the comment so that model can understand the difference. Two approaches were taken to give some details as context to the comment.

1. Page name as a tag – Using the page name as a feature and adding it as a tag in the comment.

Table 3.4 Page tag corresponding to the politician

Page	Tag
Gotabaya Rajapaksha	<gr>
Mahinda Rajapaksha	<mr>
Maithripala Sirisena	<ms>
Ranil Wickramasinghe	<rw>
Sajith Premadasa	<sp>
Anura Kumara Disanayake	<ak>

Table 3.5 Sample comments and Labels used for the page name as the tag model

Updated Comment	Label
<gr>"ස" නමා භොද්ධම කලේ	Negative
<sp>"ස" නමා භොද්ධම කලේ	Positive
<gr>We still trust our president	Positive
<ak>We still trust our president	Negative

2. Representation side in the parliament as a tag (Government representation) – Adding whether the politician in consideration represents the government or not as a tag and adding it as a feature as above.

Table 3.6 Tag corresponding to the politician

Page	Representation	Tag
Gotabaya Rajapaksha	Government	<gt>
Mahinda Rajapaksha	Government	<gt>
Maithripala Sirisena	Government	<gt>
Ranil Wickramasinghe	Opposition	<ngt>
Sajith Premadasa	Opposition	<ngt>
Anura Kumara Disanayake	Opposition	<ngt>

Table 3.7 Sample comments and Labels used for the Government tag model

Updated Comment	Label
<gt;"ස" නමා හොදටම කලේ	Negative
<ngt>"ස" නමා හොදටම කලේ	Positive
<gt>We still trust our president	Positive
<ngt>We still trust our president	Negative

3.4.2 Model Performance and Evaluation

Model evaluation is an important part of the model development process. It helps to find the best model that brings out the sentiment of the people through the data. In this study, Accuracy and F1 score was compared to find the best model.

Based on the predicted labels to get the sentiment score following approach was used.

Total Sentiment of the post = No. of Positive Comments – No. of Negative Comments

$$\text{Sentiment score} = \text{Total sentiment as \% comments} = \frac{\text{Total Sentiment of the post}}{\text{Total no. of Comments}}$$

Since the number of comments for each post varies, to make it relative, the sentiment score of a post is defined in this study as the total sentiment of a post as a percentage of comments obtained for the post.

3.5 Predicting sentiments for New Politicians

To validate the model performance for unseen data, data was gathered from three influential politicians who have a large following and are quite active on social media platforms. Comments on the three latest posts (March 2022) of the politicians were collected to validate. Comments of these three politicians are not in the training set.

1. Namal Rajapaksha (1.34 million Likes)
2. Harin Fernando (0.25 million Likes)
3. Wimal Weerawansa (0.30 million Likes)

Using the models built, sentiments were predicted for the comments of the validation dataset, and accuracy and F1 score was measured. By comparing model performance to the validation data set, how well the model can be extrapolated to bring out the sentiment of other politicians in Sri Lanka can be evaluated.

3.6 Predicting the sentiment for current data

For the six main politicians in the study, sentiment from the comments from their latest three posts (March 2022) was predicted using the built models. For each politician, the sentiment score was measured and compared with his scores in the past and against other politicians. By combining the sentiments score of politicians with the same party, the overall sentiment score was calculated. This way public opinion over different politicians, different parties, and the overall political opinion of the public at a given time was measured.

Chapter 4 ANALYSIS AND RESULTS

Based on the methodology proposed in the previous section, 6591 comments were pre-processed and cleaned and were annotated by three annotators.

The inter-annotator score obtained using Fleiss' Kappa statistics were 0.92. Comparing the results against Fleiss' Kappa interpretation table in Table 4.1, the agreements between the annotators are almost perfect thus the dataset can be used for further analysis.

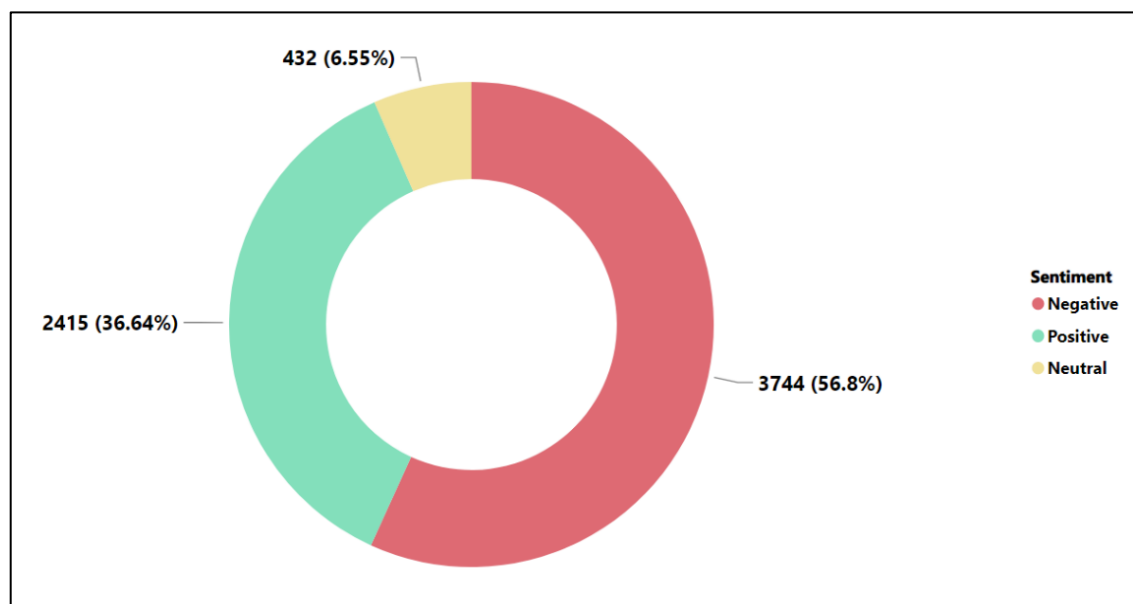


Figure 4.1 Distribution of the sentiments

Figure 4.1 shows the proportion of the comments obtained in the study based on the sentiments. A significant number of comments are negative while 37% of the comments are positive. Out of the total comments, only 432 comments are labeled as neutral.

Table 4.1 Distribution of comments obtained for the politicians

Page Name	No. of Comments	As % of total comments	Positive Comments	Negative Comments	Neutral Comments
Anura Kumara	645	9.8%	261	313	71
Gotabaya Rajapaksha	3562	54%	1262	2057	243
Mahinda Rajapaksha	961	14.6%	174	748	39
Maithripala Sirisena	659	10%	386	227	46
Ranil Wickramasinghe	198	3%	114	79	5
Sajith Premadasa	566	8.6%	218	320	28

Table 4.2 shows the distribution of the comments obtained for each of the politicians in consideration. Based on the result, out of the sampled posts majority of the comments are obtained by President Gotabaya Rajapaksha while the least number of comments are obtained for Former prime minister Ranil Wickramasinghe. Comparing the positive and negative comments apart from Former president Maithripala Sirisena and Ranil Wickramasinghe all others have obtained more negative comments than positive comments.

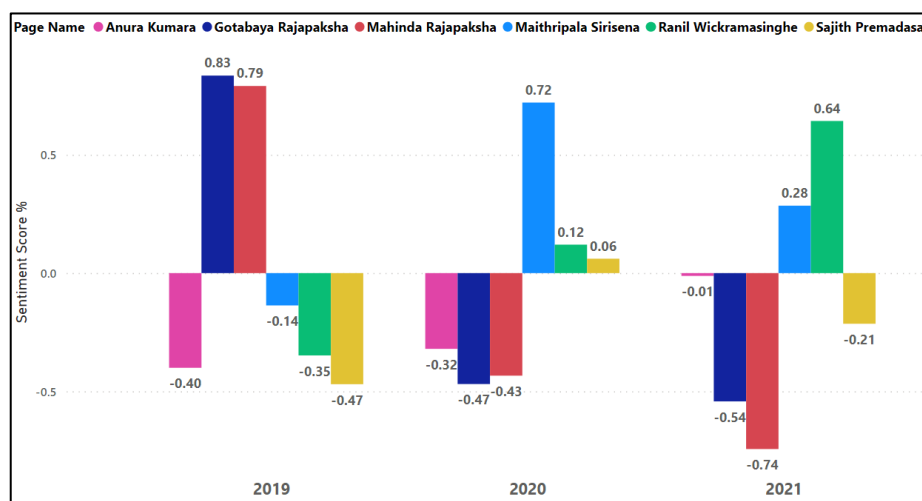


Figure 4.2 Sentiment score of the politicians over the period based on human annotation

Figure 4.2 plots the sentiment score obtained for each of the politicians for the sampled posts based on the year of posting. Based on the figure, for the posts in 2019, the sentiment score obtained was highest for Gotabaya Rajapaksha and followed by Mahinda Rajapaksha. The sentiment score is negative for all other politicians. This result is in line with the election results in 2019 as there was a regime change as President Gotabaya Rajapaksha got elected as president while Mahinda Rajapaksha was appointed as Prime Minister. For 2020, the highest sentiment score was obtained for former President Maithripala Sirisena while in 2021, Ranil Wickramasinghe topped the chart. Comparing scores over the period for the sampled posts the sentiment score has declined for Gotabaya Rajapaksha and Mahinda Rajapaksha while there's an increment for other politicians.

4.1 Performance evaluation of the models built

As proposed in the methodology, five models were built by fine-tuning the XLM-R pre-trained model. Four out of them were used to predict the sentiments of the comments while the other model was used to predict the language of the comment.

XLM-R models were fine-tuned using Adam as the optimizer, learning rate of 5e-5, and CrossEntropy loss. The dataset was split to 90:10 as training and test data when modeling the data.

Table 4.2 and 4.3 summarizes the model performance obtained for the models built in the study.

Table 4.2 Summary of the sentiment analysis model performance

Model	Task	Accuracy	F1 Score
Baseline Model	Predict the Sentiment	0.82	0.79
Language as a tag model	Predict the Sentiment	0.74	0.73
Page Name as a tag model	Predict the Sentiment	0.82	0.8
Government representation as a tag model	Predict the Sentiment	0.92	0.92

Table 4.3 Summary of the Language predictor model performance

Model	Task	Accuracy	F1 Score
Language Predictor	Predict the language	0.98	0.98

Based on the results summarized in Table 4.2, for the baseline model the sentiments of the comments can be predicted with an accuracy of 82%. Looking at the prediction for each label, the model has failed to predict any of the neutral categories.

One of the major tasks of the study is to model the language of the comments published in posts by politicians. Using the cleaned data and the approach in Table 3.3, the XLM-R model was fine-tuned to predict the language. Based on the results for the Language Predictor model in Table 4.3, the ability to predict the language using the XLM-R model is significantly high as the model obtained an accuracy of 98%.

By adding language as a tag in the comment as in Table 3.4, an XLM-R model was fine-tuned to check if it improves the baseline model accuracy. The Language as a tag model achieves an accuracy of 74%. Compared with the baseline model the accuracy has dropped, suggesting that using the language of the comment as a feature doesn't impact sentiments prediction greatly.

To remove the ambiguity of having two same comments with different labels, page names and Government representation were added as a tag as shown in Table 3.6 and Table 3.8.

The accuracy of the Page as a tag model and Government representation as a tag model is 82% and 92% respectively as shown in Table 4.3. Based on the two model performances, using the politicians' information as a tag has improved the performance compared to the baseline model. Comparing the two approaches, using Government representation as a tag has improved performance significantly as the accuracy obtained is 92%, which is a significant result in a sentimental analysis task. The F1 score of the model is the highest comparing all the sentimental analysis models built in this study. In comparison to the baseline model, the model which used Government representation as a tag has been able to improve the prediction capability of neutral comments as well. Based on the results, the novel approach of using some context in predicting sentiments has improved the performance.

4.2 Model Validation with new data

Comments were collected from three influential politicians for their latest three posts in March 2022. The distribution of the comments obtained is as follows.

Table 4.4 The distribution of comments in the validation dataset

Page Name	No. of Comments	Positive Comments	Negative Comments	Neutral Comments
Namal Rajapaksha	874	410 (47%)	447 (51%)	17 (2%)
Harin Fernando	297	289 (98%)	4 (1%)	4 (1%)
Wimal Weerawansa	441	89 (21%)	346 (78%)	6 (1%)
Total	1612	788 (49%)	797 (49%)	27 (2%)

For the validation dataset, the baseline and the model which uses Government representation as a tag were used to predict the sentiments of the comments. Government representation for the three politicians was tagged using the following manner.

Table 4.5 Representation in the parliament

Page	Representation
Namal Rajapaksha	Government
Wimal Weerawansa	Government
Harin Fernando	Opposition

4.2.1 Results of the Validation data

Table 4.6 Performance of the models for the new data

Model	Accuracy	F1 Score
Baseline Model	0.87	0.86
Government representation as a tag model	0.85	0.84

Based on the results obtained in Table 4.6, the Baseline and the Government tag model perform significantly better having an accuracy rate of more than 85%. This means using the models in the consideration, the approach suggested in the thesis can be generalized for other politicians and predicting the sentiments of the people can be done with greater accuracy.

4.3 Predicting the sentiment for the current data.

Comments from the latest three posts (March 2022) were collected and using the models built their sentiments were predicted. 2999 comments were collected, and the outcome of their prediction is given below.

Table 4.7 Predicted sentiments for the current data

Page Name	Comments	Predicted Positive	Predicted Negative	Predicted Neutral
Anura Kumara	439	285 (65%)	142 (32%)	12 (3%)
Gotabaya Rajapaksha	1816	359 (20%)	1411 (78%)	46(2%)
Mahinda Rajapaksha	274	73 (27%)	190 (69%)	11 (4%)
Maithripala Sirisena	290	96 (33%)	191 (66%)	3 (1%)
Sajith Premadasa	180	131 (72%)	46 (26%)	3 (2%)
Total	2999	944 (32%)	1980 (66%)	75 (2%)

For the 2999 comments, 66% of the comments are predicted to be a negative sentiment, while 944 data gives positive and 75 comments were predicted to be neutral.

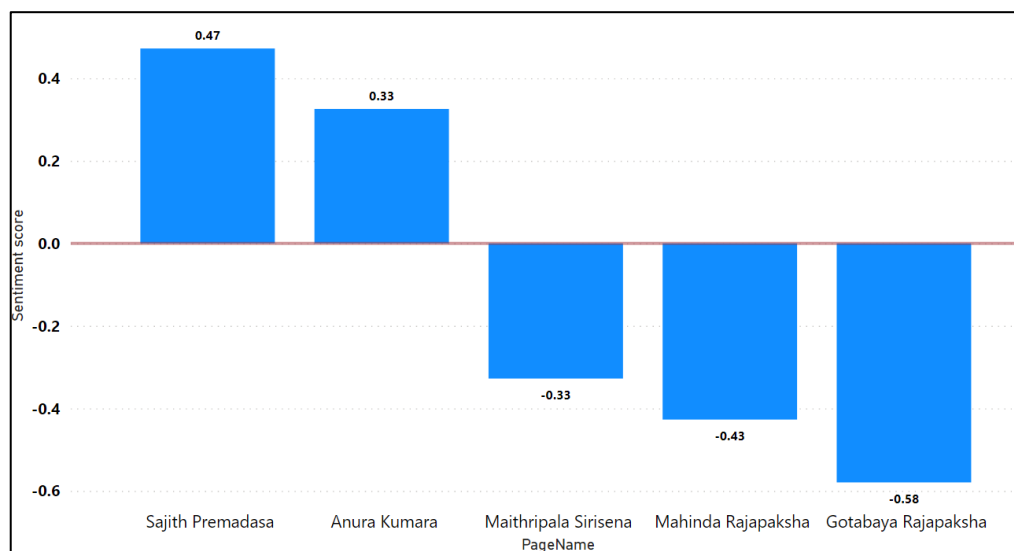


Figure 4.3 Predicted sentiment score for the current data

Figure 4.3, plots the sentiment score derived out of the predicted sentiments. Based on it, the sentiment score is positive for Sajith Premadasa and Anura Kumara Disanayake while the score is negative for others. Ranil Wickramasinghe hasn't posted a post in the given period.

4.3.1 Understanding the political opinion of people toward the politicians

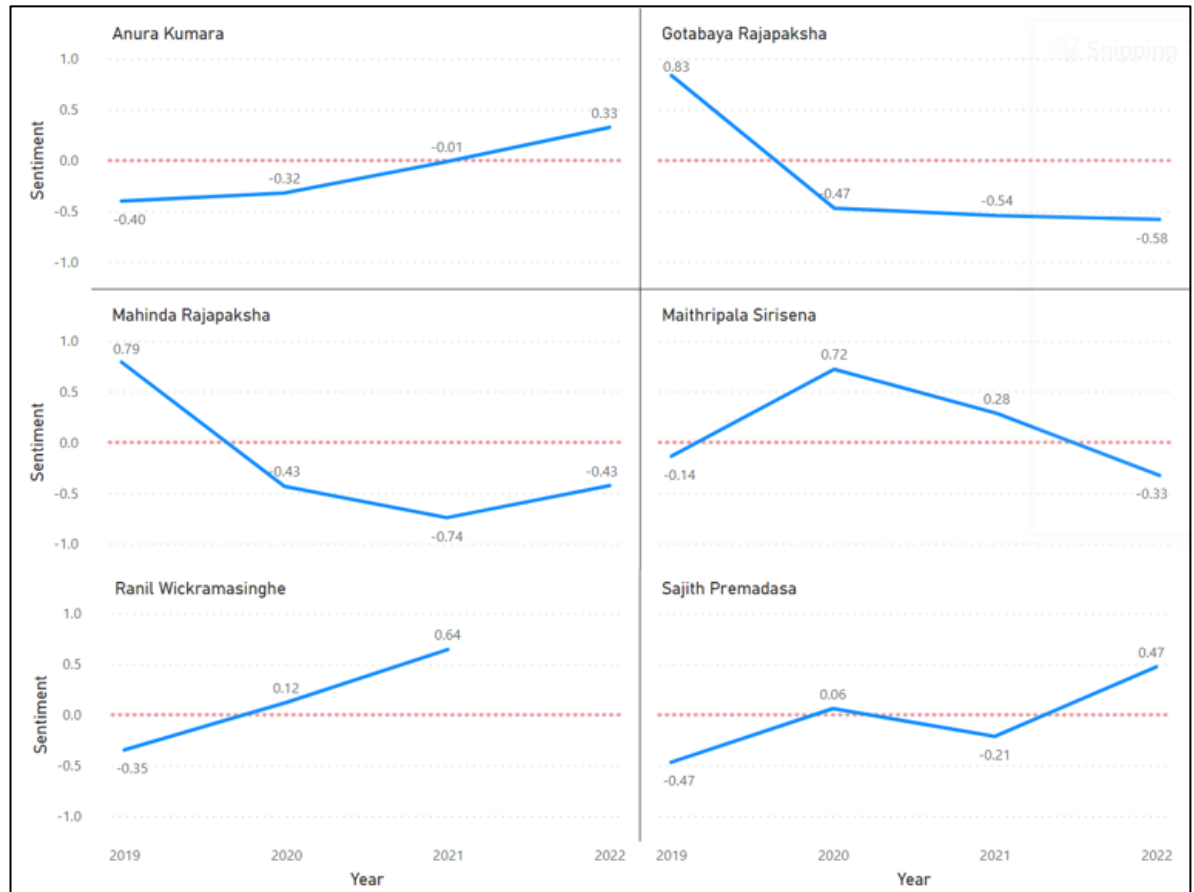


Figure 4.4 Sentiment score of the politicians

Figure 4.4 plots the individual sentiment score for each of the main politicians in the study for the collected and predicted data based on the year. The popularity or the positive sentiments Gotabaya Rajapaksha and Mahinda Rajapaksha had in 2019 has significantly deteriorated while the sentiments for Anura Kumara, Sajith Premadasa, and Ranil Wickramasinghe has moved from negative to positive in the span of three years.

4.3.2 Understanding the overall political opinion of the people

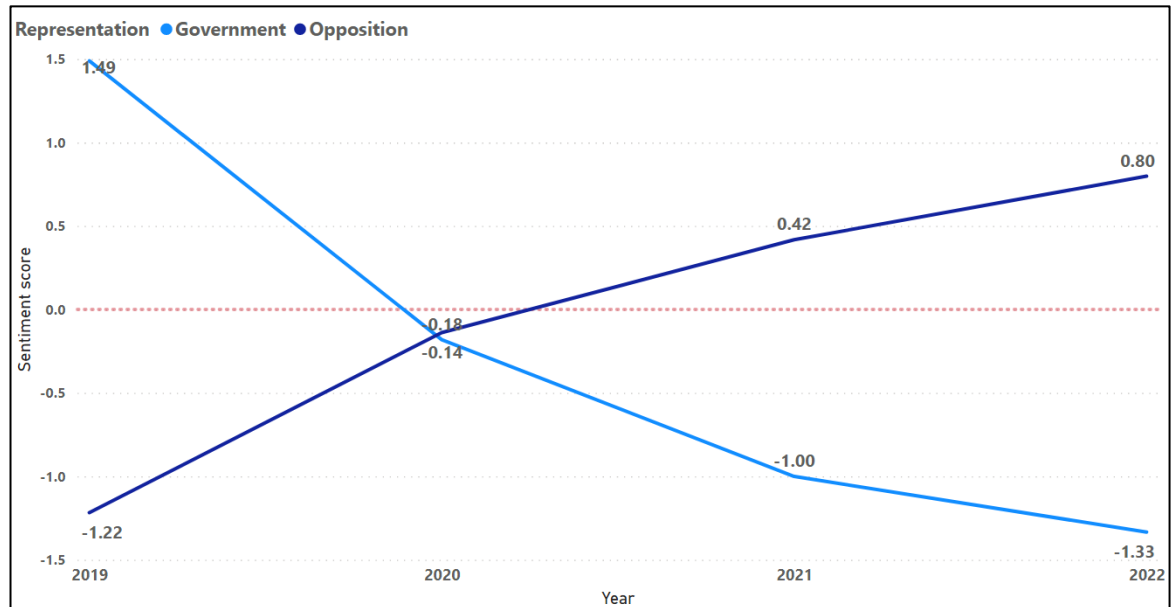


Figure 4.5 Overall sentiment score of the Government and Opposition over the years

Figure 4.5 illustrates the overall sentiment score for the government and the opposition. By combining the sentiment scores of the government representatives and opposition representatives the above line chart is plotted. Based on the chart, the government has lost its popularity over the period while the opposition has gained out of it. Looking at the predicted sentiments, the overall sentiment of the people is negative towards the government and positive toward the opposition.

According to the results, the true sentiment of the people can be predicted with an accuracy as high as 92%. This means the data plotted in Figure 4.12 represents the actual opinion of the people with greater accuracy. The government which had a significant positive score and to have lost it with just three years in power means that people's sentiments have not been monitored and not taken into account when making decisions.

Chapter 5 CONCLUSIONS

This chapter summarizes the results obtained through the analysis in the previous chapter to highlight the important findings and the conclusions derived as a result. The chapter also discusses the limitations of the study and concludes by proposing further areas of research.

5.1 Conclusions

The public opinion of a country is the collective preference of people on a matter of interest. Opinions of people will be based on the education and the experience they possess, the culture they leave in, and the beliefs they have. Gauging the opinion of people is a difficult task as many variables impact it and human behaviors change from time to time as well. In this study, the political sentiment of the Sri Lankan people is attempted to be gauged through analyzing data. Understanding political opinion is important as it guides governments, influences public policies, and gives feedback to the decision-makers. Assessing the right sentiment at right time will influence the decision-makers to make better decisions and work for the betterment of the country.

The baseline model achieving an accuracy of 82% means, that predicting the sentiment of Sri Lankan people in the political context by fine-tuning the XLM-R pre-trained model using English and Sinhala comments is a success and gives reliable results. Also, the XLM-R model can be used as a tool for Language detector tasks as it has shown greater results in this study.

The novel approach of adding the context of the comment as a tag when training a model has an impact on the performance of the model. When language as a tag is used, the model performance reduces while when the Government representation tag is used, the model performance significantly improved. Thus, the results can be concluded that applying the novel approach and adding the right context will improve sentiment analysis task performance.

Concluding the study, at present, President Gotabaya Rajapaksha, Prime Minister Mahinda Rajapaksha, and Former President Maithripala Sirisena have negative sentiments from people while Sajith Premadasa and Anura Kumara have positive sentiments. The overall sentiment of people towards the government is negative while the opposition has more support from the people.

Thus, this thesis provides a framework to understand people's sentiments through social media data. More importantly, this dissertation proposes an application of using comments of the public and analyzing through a Deep Learning model to bring out meaningful insights that the decision-makers can use. The novel approach introduced in this thesis of

adding context as a feature in sentiment analysis can be used in other domains and used as an approach to improve performance. Therefore, this thesis can inspire others to pursue future research in this area.

5.2 Limitations and future work

The main limitation of the research was the availability of data. Facebook's privacy policy and data protection laws prevent data gathering in bulk which makes data extraction an extremely difficult task.

This thesis used comments from English and Sinhala only which meant comments from Tamil language, Sinhala comments written in English, emojis, and picture comments had to be removed from the dataset. These preprocessing steps removed a lot of data that could have been used to bring out more opinions from people. This model can only be used to predict the sentiments of politicians whose majority of followers comment either in English or Sinhala. To capture sentiments from politicians who represent the Tamil-speaking community requires modeling Tamil language data as well.

By analyzing the comments, it was found that some comments had been put after a long time since the original date of the post. These comments might give a wrong interpretation when the sentiment trend is measured. But these comments which are posted on older posts also need to be modeled as it captures the current sentiments. It was also reported that there are paid campaigns to post certain comments on politicians' posts. These artificial comments will impact the outcome of the research.

Even though the proposed approach in the study is capable enough to successfully predict the sentiments with high accuracy, whether the comments represent the whole population of Sri Lanka is a question. Since the social media penetration stands around 30%, can social media data be used as a proxy for people's opinion is a question to be addressed. Even though election results and social media sentiments in 2019 had similar patterns, further analysis is required to conclude that social media data can be correlated to election results. Addressing these limitations and improving the model using structural features of social media data is recommended for future work.

Chapter 6 REFERENCES

- [1] M. Skoric, L. Jing and J. Kokil, "Electoral and Public Opinion Forecasts with Social Media Data: A Meta-Analysis," *Information*, vol. 11, p. 187, 03 2020.
- [2] "Datareportal," 18 02 2020. [Online]. Available: <https://datareportal.com/reports/digital-2020-sri-lanka#:~:text=There%20were%206.40%20million%20social,at%2030%25%20in%20January%202020..>
- [3] F. Barclay, C. Pichandy, A. Venkat and S. Sudhakaran, "India 2014: Facebook 'Like' as a Predictor of Election Outcomes," *Asian Journal of Political Science*, vol. 23, pp. 1-27, 2015.
- [4] E. Kalampokis, A. Karamanou, E. Tambouris and K. Tarabanis, "On predicting election results using twitter and linked open data: The case of the UK 2010 election," *Journal of Universal Computer Science*, vol. 23, pp. 280-303, 2017.
- [5] V. T. Reima, H. Li and Suomi, "Facebook likes and public opinion: Predicting the 2015 Finnish parliamentary elections," *Government Information Quarterly*, 05 2017.
- [6] T. Krishnamohan and S. Kandasamy, "The Ninth Parliamentary Election of Sri Lanka in 2020: An Analysis of the Outcomes," *International Journal of Scientific and Research Publications (IJSRP)*, vol. 10, p. 434, 2020.
- [7] K. tokke and P. Peiris, "Public Opinion on Peace as a Reflection of Social Differentiation and Politicisation of Identity in Sri Lanka," *PCD Journal*, vol. 2, 2010.
- [8] P. Jayasuriya, B. Kumarasinghe, S. Ekanayake, R. Munasinghe, S. Thelijjagoda and I. Weerasinghe, "Sentiment classification of Sinhala content in social media," 2020.
- [9] B. Sharounthan, D. P. Nawinna and R. De Silva, "Singlish Sentiment Analysis Based Rating For Public Transportation," in *2021 International Conference on Computer Communication and Informatics (ICCCI)*, 2021.
- [10] R. G. Brooker, "Methods of Measuring Public," 2003.
]
- [11] R. Macreadie, "Public Opinion Polls," 07 2011.
]

- [12 J. Murphy and M. W. Link, "Social Media in Public Opinion Research," American Association for Public Opinion Research, 2014.
- [13 V. Kumar, "Sentiment Analysis Techniques for Social Media Data: A Review," 2019.
- [14 E. Haddi, X. Liu and Y. Shi, "The Role of Text Pre-processing in Sentiment Analysis," in *Procedia Computer Science*, 2013.
- [15 S. Choudhary, "Choudhary, Seema. (2018). Comparative study of deep learning based sentiment analysis with other existence techniques and challenges.," 2018.
- [16 Hemamalini and D. Perumal, "International Journal of Scientific and Technology Research," *Literature-Review-On-Sentiment-Analysis*, vol. 9, no. 04, 2020.
- [17 A. Shathik and K. Prasad, "A Literature Review on Application of Sentiment Analysis Using Machine Learning Techniques," *International Journal of Applied Engineering and Management*, vol. 04, 2020.
- [18 N. C. Dang, M. N. Moreno-García and F. D. I. Prieta, "Sentiment Analysis Based on Deep Learning: A Comparative Study," *Electronics*, vol. 9, no. 3, 2020.
- [19 S. Minaee, N. Kalchbrenner, E. Cambria, N. Nikzad Khasmakhi, M. Asgari-Chenaghlu and J. Gao, "Deep Learning Based Text Classification: A Comprehensive Review," 2020.
- [20 R. Socher, A. Perelygin, J. Wu, J. Chuang, C. D. Manning, A. Ng and C. Potts, "Recursive Deep Models for Semantic Compositionality Over a Sentiment Treebank," in *Empirical Methods in Natural Language Processing*, 2013.
- [21 Y. K. Wiciaputra, J. C. Young and A. Rusli, "Bilingual Text Classification in English and Indonesian via Transfer Learning using XLM-RoBERTa," *Advance Soft Compu*, vol. 13, no. 3, 2021.
- [22 J. Devlin, W. Chang, K. Lee and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformer for Language Understanding," *Google AI Language*, pp. 1-16, 2019.
- [23 Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer and V. Stoyanov, "RoBERTa: A Robustly Optimized BERT Pretraining Approach," *Facebook AI*, pp. 1-13, 2019.
- [24 G. Lample and A. Conneau, "Cross-lingual Language Model Pretraining," *Facebook AI*, pp. 1-10, 2019.

- [25 K. K. A. Conneau, N. Goyal, V. Chaudhary, G. Wenzek, F. Guzmán, M. O. E. Grave,
] L. Zettlemoyer and V. Stoyanov, "Unsupervised cross-lingual representation learning
at scale," *arXiv*, 2019.
- [26 M. Meduru, A. Mahimkar, K. Subramanian, P. Y. Padiya and P. & N. Gunjgur,
] "Opinion Mining Using Twitter Feeds for Political Analysis," *International Journal
of Computer (IJC)*, pp. 116-123, 2017.
- [27 Tumasjan, T. Sprenger, P. Sandner and I. Welp, "Election forecasts with Twitter:
] How 140 characters reflect the political landscape," *Social Science Computer Review*,
vol. 29, pp. 402-418, 2011.
- [28 J. Pennebaker, R. Booth and M. Francis, Linguistic Inquiry and Word Count
] (LIWC2007), 2007.
- [29 Burnap, R. Gibson, Sloan, R. Southern and M. Williams, "140 characters to victory?:
] Using Twitter to predict," *Electoral Studies*, vol. 41, pp. 230-233, 2016.
- [30 Thelwall, Buckley, Paltogou, Cai and Kappas, "Sentiment strength detection in short
] informal text," *Sci Technol*, p. 61, 2010.
- [31 T. Elghazaly, A. Mahmoud and H. A. Hefny, "Political Sentiment Analysis Using
] Twitter Data," 2016.
- [32 B.-M. Felipe, D. Gayo-Avello, M. Mendoza and B. Poblete, "Opinion Dynamics of
] Elections in Twitter," 2016.
- [33 M. Cameron, P. Barrett and B. Stewardson, "Can Social Media Predict Election
] Results? Evidence From New Zealand," *Journal of Political Marketing*, 2014.
- [34 Y. Tian, T. Galery, Dulcinati, Giulio, E. Molimpakis and C. Sun, "Facebook
] sentiment: Reactions and emojis," in *5th International Workshop on Natural
Language Processing for Social Media*, 2017.
- [35 R. Sandoval-Almazan and D. Valle-Cruz, "Sentiment Analysis of Facebook Users
] Reacting to Political Campaign," 2020.
- [36 McHugh, "Interrater reliability: the kappa statistic.," *Biochemia medica*, vol. 22(3),
] p. 276–282, 2012.
- [37 J. L. Fleiss, "Measuring nominal scale agreement among many raters," *Psychological
] Bulletin*, vol. 76, p. 378–382, 1971.
- [38 M. Sokolova and G. Lapalme, "A systematic analysis of performance measures for
] classification tasks," *Information Processing & Management*, vol. 45, no. 9, pp. 427-
437, 2009.

- [39 T. Ranasinghe and M. Zampieri, "Multilingual Offensive Language Identification
] with Cross-lingual Embeddings," 2020.
- [40 L. Xu, J. Zeng and S. Chen, " Fine-tune XML-RoBERTa for Hate Speech
] Identification," in *yasuo at HASOC2020*, 2020.
- [41 J. A. Fishman, "Who speaks what language to whom and when?," in *The Bilingualism
] Reader*, 2007, p. 16.
- [42 M. Johnson, M. Schuster, Q. Le, M. Krikun, Y. Wu, Z. Chen, N. Thorat, F. Viégas,
] M. Wattenberg, G. Corrado, M. Hughes and J. Dean, "Google's Multilingual Neural
Machine Translation System: Enabling Zero-Shot Translation," *Transactions of the
Association for Computational Linguistics*, 2016.
- [43 K. Jaidka, S. Ahmed, M. Skoric and M. Hilbert, "Predicting elections from social
] media: A three-country,," *Asian J. Commun*, vol. 23, pp. 252-273, 2019.
- [44 M. Farhadloo and E. Rolland, *Fundamentals of Sentiment Analysis and Its
] Applications*, 2016.
- [45 "Napoleoncat," [Online]. Available: [https://napoleoncat.com/stats/facebook-users-
\] in-sri_lanka/2021/01/](https://napoleoncat.com/stats/facebook-users-in-sri_lanka/2021/01/).
- [46 M. Johnson, M. Schuster, Q. V. Le, M. Krikun, Y. Wu, Z. Chen, N. Thorat, F. Viégas,
] M. Wattenberg, G. Corrado, M. Hughes and J. Dean, "Google's Multilingual Neural
Machine Translation System: Enabling Zero-Shot Translation," *arXiv*, 2016.
- [47 M. M. L., "Interrater reliability: the kappa statistic,," *Biochemia medica*, vol. 22(3),
] p. 276–282, 2012.