

Aspect Based Sentiment Analysis Model for Hotel Services in Amharic Language Using Machine Learning Techniques



By

Yordanos Athinafe Bethiha

A Thesis Submitted to the Department of Computer Science and Engineering
School of Electrical Engineering and Computing
Presented in Partial Fulfillment for the Degree of Master of Science in
Computer Science and Engineering

Office of Graduate Studies
Adama Science and Technology University

August 2021
Adama, Ethiopia

**Aspect Based Sentiment Analysis Model for Hotel Services in Amharic
Language Using Machine Learning Techniques**

By

Yordanos Athinafe Bethiha

Advisor

Bahiru Shifaw Yimer (Ph.D.)

A Thesis Submitted to the Department of Computer Science and Engineering
School of Electrical Engineering and Computing
Presented in Partial Fulfillment for the Degree of Master of Science in
Computer Science and Engineering

Office of Graduate Studies
Adama Science and Technology University

August, 2021
Adama, Ethiopia

Declaration

I hereby declare that this Master Thesis entitled “Aspect Based Sentiment Analysis Model for Hotel Services in Amharic Language Using Machine Learning Techniques” is my original work. That is, it has not been submitted for the award of any academic degree, diploma or certificate in any other university. All sources of materials that are used for this thesis have been duly acknowledged through citation.

Yordanos Athinafe Bethiha

Name of the Student

Signature

Date

Recommendation

I, the advisor of this thesis, hereby certify that I have read the revised version of the thesis entitled “Aspect Based Sentiment Analysis Model for Hotel Services in Amharic Language Using Machine Learning Techniques” prepared under my guidance by Yordanos Athinafe submitted in partial fulfillment of the requirements for the degree of Master of Science in Computer Science and Engineering.

Therefore, I recommend the submission of revised version of the thesis to the department following the applicable procedures.

Bahiru Shifaw Yimer (Ph.D.)

Advisor

Signature

Date

Approval Page

I, the advisors of the thesis entitled “Aspect Based Sentiment Analysis Model for Hotel Services in Amharic Language Using Machine Learning Techniques” and developed by Yordanos Athinafe, hereby certify that the recommendation and suggestions made by the board of examiners are appropriately incorporated into the final version of the thesis.

Bahiru Shifaw Yimer (P.H.D)

Advisor

Signature

Date

We, the undersigned, members of the Board of Examiners of the thesis by Yordanos Athinafe (Name of student) have read and evaluated the thesis entitled “Aspect Based Sentiment Analysis Model for Hotel Services in Amharic Language Using Machine Learning Techniques” and examined the candidate during open defense. This is, therefore, to certify that the thesis is accepted for partial fulfillment of the requirement of the degree of Master of Science in Computer Science and Engineering.

Chairperson

Signature

Date

Internal Examiner

Signature

Date

External Examiner

Signature

Date

Finally, approval and acceptance of the thesis are contingent upon submission of its final copy to the Office of Postgraduate Studies (OPGS) through the Department Graduate Council (DGC) and School Graduate Committee (SGC).

Department Head

Signature

Date

School Dean

Signature

Date

Office of Postgraduate Studies, Dean

Signature

Date

ACKNOWLEDGMENT

I would like to begin with say a huge thank you to my advisor Dr. Bahiru Shifaw for all the support and encouragement he gave me during this thesis work, and I also would like to thank my friend Fikadu Eshetu for your valuable support and motivation.

TABLE OF CONTENTS

ACKNOWLEDGMENT	iv
TABLE OF CONTENTS	v
LIST OF TABLES	ix
LIST OF FIGURES	x
LIST OF ACRONYMS	xi
ABSTRACT	xii
CHAPTER ONE	1
INTRODUCTION	1
1.1. Background of the study	1
1.2. Background of Hotel Industry	3
1.3. Motivation of the study	3
1.4. Statement of the Problem	3
1.5. Research Questions	4
1.6. Objectives of the Study	4
1.6.1. General Objective	4
1.6.2. Specific Objectives	4
1.7. Scope and Limitations.....	5
1.7.1. Scope.....	5
1.7.2. Limitations.....	5
1.8. Application of Results.....	5
1.9. Organization of the Study.....	5
CHAPTER TWO	7
LITERATURE REVIEW AND RELATED WORKS	7
2.1. Sentiment Analysis	7
2.2. Basic Components of Sentiment Analysis	7
2.3. Levels of Sentimental Analysis	8
2.3.1. Sentence Level Sentimental Analysis	8
2.3.2. Document Level Sentimental Analysis.....	8
2.3.3. Feature Level Sentimental Analysis	9
2.4. Common Approaches to Sentiment Analysis.....	9

2.4.1.	Rule-based Approach.....	9
2.4.2.	Machine Learning Approach.....	10
2.4.3.	Hybrid Approach	11
2.5.	Background of Amharic Language	11
2.6.	The Amharic Writing System.....	11
2.7.	Punctuation	12
2.8.	Numbers	12
2.9.	Characteristics of the Amharic Writing System	12
2.10.	Contextual Shifter (Negation).....	13
2.11.	Grammatical Rules of Amharic	13
2.12.	Amharic Phrases.....	14
2.13.	Related Works	15
CHAPTER THREE.....		22
METHODOLOGY.....		22
3.1.	Research Methodologies	22
3.2.	Building a Corpus	22
3.2.1.	Data Collection	22
3.2.2.	Method of Data Collection.....	23
3.2.3.	Hotel Selecting Criteria.....	23
3.3.	Dataset Preparation	29
3.4.	Dataset Annotation.....	29
3.4.1.	Annotation Guideline.....	29
3.4.2.	Annotation Procedure	29
3.4.3.	Inter Annotator Agreement	30
3.4.4.	Dataset Annotation Result.....	31
3.5.	Aspect Based Sentimental Analysis Modeling.....	32
3.5.1.	Data pre-processing in the Amharic language.....	32
3.5.2.	Normalization	33
3.5.3.	Tokenization	33
3.5.4.	Stop-words removing.....	33
3.6.	Machine Learning Classification Algorithm.....	34

3.7. Cosine Similarity	36
3.8. Evaluation.....	37
3.8.1. Evaluation Techniques (Model Evaluation).....	37
CHAPTER FOUR.....	40
PROPOSED SOLUTION FOR ASPECT BASED SA	40
4.1. The proposed solution for Aspect Based SA for the Amharic Language	40
4.2. The proposed Aspect Based SA Architecture	40
4.3. Proposed Amharic Text Preprocessing	42
4.3.1. Data Cleaning	42
4.3.2. Amharic Normalization Character.....	43
4.3.3. Tokenization	43
4.3.4. Stop-Words Removal.....	44
4.3.5. Part of Speech (POS) tagging:.....	44
4.4. Proposed Feature Extraction Methods	45
CHAPTER FIVE	47
IMPLEMENTATION AND EXPERIMENTATION.....	47
5.1. Implementation of the proposed solution.....	47
5.2. Implementation Environment.....	47
5.3. Data Preprocessing Implementation	49
5.3.1. Implementation of Removing Irrelevant Characters.....	49
5.3.2. Implementation of Normalization of Amharic Character	49
5.3.3. Implementation of Tokenization.....	50
5.3.4. Implementation Stop-Words Removal.....	50
5.4. Feature Extraction Implementation.....	50
5.4.1. Implementation of Count Vectorize and N-gram	51
5.5. Machine Learning Models Implementations.....	51
5.6. Cosine Similarity	53
5.7. Model Testing and Evaluation.....	53
5.8. User Interface Design.....	54
5.8.1. Flask Frameworks.....	54
5.8.2. Bootstrap	55

CHAPTER SIX.....	57
RESULTS AND DISCUSSIONS.....	57
6.1. Discussion of Results	57
6.2. Models Evaluation Results	57
6.2.1. Ternary Classification Models Evaluation Results.....	57
6.3. Normalized Confusion matrix for SVM, NB, LR, and GB.....	65
6.4. Discussions	67
CHAPTER SEVEN.....	68
CONCLUSION AND FUTURE WORKS.....	68
7.1. Conclusion.....	68
7.2. Contributions of the Thesis	68
7.3. Future Works	69
Reference	70
Appendix A: Sample letter Permission	74
Appendix B: Approval Letter for Linguistic Expertise.....	75
Appendix C: ከዚህ በታች የተዘረዘሩት ቃላት አዎንታን ያሳያሉ::	76
Appendix D: ከዚህ በታች የተዘረዘሩት ቃላት አፍራሽነትን ያሳያሉ::	77
Appendix E: Amharic alphabets and their pronunciation.	78
Appendix F: Amharic Punctuation	79
Appendix G: Some of Amharic Number.....	79
Appendix H: Some of Amharic stop-word list.....	80
Appendix I: Amharic ABSA Hotel Reviews Annotation Guidelines.....	81
Appendix J: Implementation of Removing Irrelevant Characters	84
Appendix K: Implementation of Normalization of Amharic Character	85
Appendix L: Implementation of Normalizing of Confusion Matrix	86

LIST OF TABLES

Table 1.1: The Previous Research Thesis for Aspect Level SA.....	19
Table 1.2: Distributing Comment to Annotators	30
Table 1.3: Sample inter-annotator agreement.....	31
Table 1.4: Confusion matrix of Ternary class	38
Table 1.5: Some of the Tools and Python Package.....	47
Table 1.6: SVM Models Average Accuracy for Each Feature	57
Table 1.7: NB Models Average Accuracy Scores for Each Feature.....	59
Table 1.8: LR Models Average Accuracy Scores for Each Feature	60
Table 1.9: GB Models Average Accuracy Scores for Each Feature.....	61
Table 1.10: Classification Performance Result.....	64

LIST OF FIGURES

Figure 1.1: Process in Training and prediction.....	34
Figure 1.2: Amharic Aspect Based SA Architecture	41
Figure 1.3: Figure Feature Extraction Flow Diagram	45
Figure 1.4: Python Code for Load the Dataset	49
Figure 1.5: Sample Code for Tokenization.....	50
Figure 1.6: Sample Code for Stop word	50
Figure 1.7: Sample Code for Count vectorizer and Ngram	51
Figure 1.8: Importing the necessary Package	51
Figure 1.9: Instantiating SVM and Fit the Model.....	52
Figure 1.10: Instantiating Naïve Bayes and Fit the Model.....	52
Figure 1.11: Instantiating Logistic Regression and Fit the Model	52
Figure 1.12: Instantiating Gradient Boosting and Fit the Model.....	53
Figure 1.13: Sample code for Cosine Similarity in aspect Extraction	53
Figure 1.14: Performing 10-fold CV on GB Models	54
Figure 1.15: User Interface Design	55
Figure 1.16: Positive Request and Response	55
Figure 1.17: Negative Request and Response	56
Figure 1.18: Neutral Request and Response.....	56
Figure 1.19: Confusion Matrix of GB based on Extracted Features.....	65
Figure 1.20: Confusion Matrix of SVM, LR, and NB Based on Extracted Feature.....	66

LIST OF ACRONYMS

NB	Bayesian Network
CV	Cross-Validation
DM	Data Mining
F1	F1- Score
FN	False Negative
FP	False Positive
GB	Gradient Boosting
ML	Machine Learning
MoCT	Ministry of Culture and Tourism of Ethiopia
NLP	Natural Language Processing
NLTK	Natural Language Toolkit
OM	Opinion Mining
P	Precision
POS	Part of Speech Tag
R	Recall
SA	Sentiment Analysis
SVM	Support Vector Machine
TF-IDF	Term Frequency-Inverse Document Frequency
TN	True Negative
TP	True Positive
WWW	World Wide Web

ABSTRACT

In the past years, the World Wide Web (www) has come to be a large source of user-generated content and opinionated data. used by social media, such as YouTube, Facebook, organizational websites, etc. Opinion mining focuses on the sentiment which may be positive, negative, and neutral. Aspect-based sentiment analysis where particular aspects are extracted, sentiment polarity of the aspects is determined. Sentimental analysis carries through in one of three different levels namely, sentence, document, and aspect /feature. Among the three levels, aspect level sentimental analysis is detail and complex but has a better advantage to meet customers and the organization's needs. Comments written by customers have a huge advantage on the success of the hotel. contrarily, it would be difficult for a hotel to manually analyze a numerous amount of commented data to know whether a customer is satisfied or not. In this study, to alleviate this problem, Aspect Based Sentimental Analysis Model through implementing machine learning techniques is proposed. In this research, four Machine learning classification approaches are used. these include Naive Bayesian, Logistic Regression, Support Vector Machine, and Gradient Boosting which were experimented for building and evaluating the sentiment (polarity) model with the extracted features based on the 2124 datasets collected from 10 hotels' social media pages as well as their websites and annotated by Amharic linguistic expert. the cosine similarity technique for extracting aspects (features) of the hotels is also used. The cross-validation method is known for its skill in the less biased or less optimistic estimation over the previous simple train/split approach. Dataset is partitioned into k equal-sized groups or folds, and each fold is treated as a validation set, while the rest k-1 is used as a training set to fit the model. The 10-fold cross-validation of the Gradient boosting model based on the countvectorizer+unigram+Trigram feature with an accuracy of 92.8% and an f1-score of 95% has shown the best performance as compared to SVM, NB, and LR models. Aspect based SA has evolved as an active research area that dominates all the sciences in the world, so this research proposes aspect-based sentimental analysis to predict the aspect (like food, ambiance, drink, price) and polarity (+ve, -ve or neutral) by gathering data from Facebook pages and websites. Due to a shortage of time stemmer is not apply, so we recommend for the others researchers to apply stemmer.

Keywords: Machine learning, Hotel Service, social media, Sentiment Analysis, Feature Extraction, Aspect

CHAPTER ONE

INTRODUCTION

1.1. Background of the study

Social communication is an internet-based application that creates the ideological and innovative basis for Web 2.0 and permits user-created content to be delivered and exchanged (Kaplan & Haenlein, 2010). Sentimental analysis is a modern and evolving field of research that uses Data Mining and Natural Language Processing techniques to extract information and discover text knowledge, so the goal of OM is to make the machine able to recognize and express feelings. A feeling, opinion, or attitude based on emotion is called sentiment (Khan et al., 2009).

Business companies spend too much fee on customer feelings and views about their goods through consultants and surveys. Similarly, other views regarding goods, services, issues, and events are of interest to individuals to find the best choices (Khan et al., 2009). Many people can quickly interact and reflect thoughts and feeling on topics placed on social media, but it is difficult to extract well-structured knowledge from a higher number of social network feelings. We should design and deploy automated sentiment analysis to tackle this type of problem.

Sentiment analysis is generally applied for analyzing reviews like movie reviews, customer product reviews on social media for the desire of determining the attitude and sentiments of the writer/opinion holder concerning some topics (Mebratu Tadesse, 2018). It determines the subjectivity whether the opinion is subjective or objective and uses automated tools to figure out the subjective expressions of opinion holders (Mebratu Tadesse, 2018).

People can post their opinions in different ways; some give general information while others provide detail. And also some bounce from one product feature to another with only a brief description while others elaborate on certain features (TULU TILAHUN, 2013). According to (TULU TILAHUN, 2013), These factors have particular importance during the classification of the text orientation as, positive or negative, Two or more sentences may refer to the same or different features. Similarly, depending on the user's interest one sentence can express many opinions within it. Accordingly, there are three stages of opinion mining: sentence level, document level, and feature level opinion mining.

According to(TULU TILAHUN, 2013), Sentence level sentimental analysis polarity determination is at the sentence level towards the target object, there are two main tasks in sentence-level opinion mining. The first is subjectivity classification which classifies the given sentence either as subjective or objective and the second is sentiment determination which determines the sentiment of the subjective (positive or negative). For example, “ይህ ጥሩ ወርቅ ነው / **This is a beautiful gold**”, this sentence is classified as subjective, and its sentiment is determined as positive.

According to(TULU TILAHUN, 2013), The author focused on classifying the document based on the overall opinion expressed by the user or customer. All the user’s review sentences are considered. For instant, “ቆንጆ መኪና ነው። ነገር ግን ፍጥነት የለውም /**It is a beautiful car. However, its speed is poor**”. The primary sentence is positive whereas the second one is negative while the resultant polarity is neutral. According to(TULU TILAHUN, 2013), the mining opinion at the aspect level first defines the detailed objectives on which sentiment has been expressed in a sentence level and then decides whether the views are positive, negative, or neutral. Objects, and their elements, attributes, and pieces, are the targets.

According to(TULU TILAHUN, 2013), An item, service, individual, organization, event, subject, etc. may be an entity. An item review sentence, for example, defines product features that the reviewer has written on and chooses whether the comments are positive or negative. For example, ጥሩ አገልግሎት ያለዉ ሆቴል ነዉ /The hotel has good service. For example, ጥሩ አገልግሎት ያለዉ ሆቴል ነዉ /**The hotel has good service**. The feature in this sentence is አገልግሎት (Service) of the ሆቴል (Hotel) entity and the opinion is positive, ጥሩ (Good). To make product changes, many real-life applications need thorough analysis. Consumers need to know what components of the product are liked and disliked. Both sentence and document-level opinion classification do not discover detailed descriptions.

1.2. Background of Hotel Industry

Hotel is the segment of the service industry that deals with visitor convenience or lodgings. there are a variety of accommodation types used to describe services related to guest accommodation like food, drink, price, ambiance. MoCT is getting prepared to dole out star evaluations to 200 hotels. According to MoCT hotels with at slightest 11 room service are qualified for grading. The hotels moreover must pass checks for wellbeing, safety, and security, environmental services prerequisites for squander administration, and have certified narrative prove of compliance. Ethiopia has more than 600 hotels but most of them are placed in Addis Ababa, Hawassa, Bahir Dar, Adama, and Mekele.

1.3. Motivation of the study

The main motivations to investigate in this area came from two things: Firstly, Satisfying the needs and wants of the customers manually are the main challenges in operating a hotel industry, Secondly, Hotels in Ethiopia do not properly automate customer reviews which is crucial to managing their services efficiently. Automatic sentiment mining can help any user of the web including hotels to get access throughout the reviews by reducing time lost in reading information and analyzing it.

1.4. Statement of the Problem

The recognition and classification of people's opinions in plain content and from a variety of social media are of key significance in various industries, government organizations, and services(Pang & Lee, 2008; Schrauwen, 2010). For instance, in product improvement systems, producers of specific products need consumers' opinions and feelings towards the product they want to improve their quality(Liu, 2012). Hence people's opinions convey the positivity, negativity, and neutrality of some products, and neutral, so Sentimental Analysis can be used for analyzing people's opinions and feelings about the product the producers want to improve its quality as per the need of the users(Ravendra Ratan and Singh Jandail, 2014). The increased acceptance of these social sites resulted in a huge collection of people's opinions on the web in an unstructured manner(Liu, 2012). Extracting useful content from these opinion sources becomes a challenging task and created a new field of research called sentiment analysis (SA)(Ravendra Ratan and Singh Jandail, 2014). Several sentiment analysis (SA) systems have been creating in a variety of languages using different approaches such as English(Balijepalli, 2007; Esuli & Sebastiani, 2006; Kennedy &

Inkpen, 2006), Chinese(Xu et al., 2009), French(Maurel et al., 2008), Amharic(Selama Gebremeskel, 2010), etc. But, as far as the knowledge of the researchers is concerned, there is no Automatic sentimental analysis system was developed for the Amharic Language using Machine learning at the Aspect (Feature) level.

Satisfying the feeling of the customers are the most challenges in working hotels. comment written by the customers has a great advantage on the success of a hotel. It would be difficult for the hotels to manually analyze a large amount of data to get it whether the customers are satisfied or not.

According to(Anderson et al., 1994; Yeung et al., 2002), the researcher shows that client satisfaction has an impact on commerce results, and they concluded that client satisfaction positively influences hotels benefit. most of the studies have explored the relationship between client behavior designs. The satisfaction of client increase client loyalty impacts repurchase intentions and lead to certain word-of-mouth(Kandampully & Suhartanto, 2000; Olorunniwo et al., 2006; Söderlund, 1998).

1.5. Research Questions

This research set the following research question to examine the problem in the Aspect Based SA Model.

RQ1: Which machine learning approach produces a better classification model for aspect-based sentimental analysis?

RQ2: How to extract aspects/features from comments?

RQ3: Which Feature Extraction method produces a good result?

1.6. Objectives of the Study

1.6.1. General Objective

The general objective of this research is Aspect Based Sentimental Analysis Model for Hotel services in Amharic Language using Machine Learning Techniques.

1.6.2. Specific Objectives

The specific objectives (building block) of this research work are:

- To Review different kinds of literature on sentiment analysis, techniques of sentiment analysis, and the linguistic behaviors of Amharic languages.

- To analyze the structure of Amharic statements.
- To apply necessary ML algorithms to realize the proposed model in developing an Amharic aspect-based sentiment analysis model.
- To test the proposed model.

1.7. Scope and Limitations

1.7.1. Scope

The scope of this research is to predicting aspect-based sentimental analysis for hotel services in Amharic language using machine learning techniques. The research aimed to train the dataset on SVM, NB, LR, and GB. From the cross-validation method, cross-validation using 10-fold is used for model performance evaluation with different metrics. As there is no available corpus for the language considered in this research, creating a new corpus is required, and data is gathered from Facebook pages and the website from hotels.

1.7.2. Limitations

This research arises from two major limitations. The first limitation is due to the lack of standard Amharic language stemmer tools, so this study did not apply these preprocessing methods. The second limitation of this study is a shortage of time to prepare more datasets.

1.8. Application of Results

Amharic Opinion Mining model can be used for different Application purposes. Some of them are:

- The study has a huge value for the Hotels by taking user's reviews on the service they gave for the organization to make efficient and effective decisions.
- It also provides input for predicting aspect-based Sentimental analysis for the Amharic Language.
- The most beneficiaries of this study are users of the hotel.

1.9. Organization of the Study

This study is arranged into seven chapters in the following ways:

Chapter one discusses the background of the study, the motivation of the study, and the statement of problems, the research questions that would be answered in the proposed solutions, the objective of the study, the scope and limitation, and the application results of the study.

Chapter two discusses literature reviews and related works. Under this section, we tried to address an introduction of Sentimental analysis, a component of the sentimental analysis, level of sentimental analysis, the approach of sentimental analysis, the background of the Amharic language, and related works.

Chapter three deals with the research methodology used in this research work, which are methods used to build the corpus, data collection method, hotel selecting criteria, dataset preparation, dataset annotation, data preprocessing, machine learning classification technique, and Evaluation technique.

Chapter four deals with the proposed solution for Aspect Based SA, proposed Architecture, proposed Amharic text preprocessing algorithm and proposed feature extraction method.

Chapter five discusses the implementation and experimentation of the proposed solution. Including data preprocessing implementation, feature extraction implementation, machine learning model implementation, model testing/evaluation, and user interface

Chapter six discusses the result and discussion of the proposed solution. Including discussion of ternary result, the model evaluation result, and normalization confusion matrix.

Chapter seven presents the conclusion and future work. Including the conclusion, contribution of the research, and future work.

CHAPTER TWO

LITERATURE REVIEW AND RELATED WORKS

2.1. Sentiment Analysis

In NLP, extracting relevant information from comprehensive text is a great idea. Factual information extraction is part of an information retrieval process (IR). opinion analysis in the field of computer science that explores techniques for evaluating people's sentiments, thoughts, appraisals, and emotions of individuals. we are now focused on the advancement of the web and opinion mining, since in Internet forums, discussion groups, blogs, videos, and social networking sites, people can reflect their feeling. All these systems have a huge amount of useful data that we are involved in collecting and analyzing.

2.2. Basic Components of Sentiment Analysis

Opinion holders, sentiment, and objects(entities) are basic components of opinion mining. An opinion holder is either a person or organization that use a specific sentiment on a specific object whereas a sentiment is a view, feeling of an object by an opinion holder as sentiment and an attribute is an entity that may be a commodity, a person, an event, an organization, a subject, or even an opinion that is evaluated by the opinion holder. Generally, there are two ways of expressing a sentiment: direct opinions and comparative. Typically, direct opinions describe one object and contain certain adjectives that refer to it. For instance, **the touch-screen interface of my mobile phone is good**. This sentence consists of an entity “mobile phone” and an adjective “good” which modifies the **touch-screen** interface of the mobile phone. More than one item is listed in the comparative statements and some kind of relation is defined. For example, **the touch-screen interface of my mobile phone is much better than Fikadu mobile phone**. The superlative statements, on the contrary, address more than two items and explain an item's connection to the remaining items. For example, **the touch-screen interface of the fikadu phone is best when compared to ours**.

2.3. Levels of Sentimental Analysis

Different researchers have focused on sentiment classifications using supervised learning and unsupervised learning methods at different levels (document-level, sentence-level, and feature). As presented in (Bhuiyan et al., 2009) Classifying opinions at the document or sentence level does not tell precisely what the holder of the opinion likes and dislikes, But The feature-level, however, explores ways of classifying each feature. The classification of feedback at the feature level, however, is harder to do.

2.3.1. Sentence Level Sentimental Analysis

Sentence level Sentimental analysis is a decision of polarity against the target object at the sentence level. There are two tasks in sentence-level Sentimental analysis: the first is the classification of subjectivity that classifies the given sentence either as subjective or objective, and the second is the determination of sentiment that decides the sentiment of the subjective sentence whether it is +ve or -ve or neutral sentiment.

The researcher conducted (Jagtap & Pawar, 2013) machine learning approaches at sentence-level polarity determination. This ML approach treats the problem of opinion classification as a problem of topic-based classification. Any topic-based classification technique can be used, e.g. naïve Bayes, SVM, MEntropy, etc. The authors use the [Naïve Bayes, MEntropy, SVM] approach used to classify movie reviews into two classes: positive and negative. The study compared naïve Bayes, Maximum Entropy, and SVM. It uses 700 positive reviews and 700 negative reviews were used. The highest classification accuracy (82.9%) was achieved using SVM with cross-validation using 3-fold.

2.3.2. Document Level Sentimental Analysis

The aim here is to consider the overall opinion of an entire document. For instance, given a hotel service review, the task is to determine whether sentiment is positive or negative or neutral opinion ns about the hotel service. This level looks at the document as a distinct entity, thus it is not extensible to multiple documents.

2.3.3. Feature Level Sentimental Analysis

Sentiment classification is common at both document and sentence levels; the most comprehensive research is aspect-level opinion mining. Commented aspects are identified, extracted and the sentiment toward this feature is determined. The goal of aspect-level classification aims to generate a compilation of multiple reviews based on aspect-based opinions. It primarily has three-phases. The first phase is to recognize and extract object aspects that have been commented on by the user. The second phase is to determine the polarity of sentiment on aspect classes: +ve, -ve, and neutral, and finally the third task is related to the group feature synonyms.

2.4. Common Approaches to Sentiment Analysis

Many research studies on sentiment mining have been achieved using various approaches. The most successful approaches to opinion mining, depending on how rules are learned, are machine learning, Rule-based, and hybrid approaches. Every approach is outlined in subsequent parts.

2.4.1. Rule-based Approach

To distinguish subjectivity, sentiment, or the subject of an opinion, a rule-based framework typically uses a collection of human-crafted rules. Rule-based approaches incorporate different natural processing like Stemming, tokenization, speech tagging, parsing, and Dictionaries (i.e., records of words and expressions). The rule-based method has been used to examine the text sentence by sentence and extract the relationships that express feelings by using linguistic awareness of sentimental terms and heuristic rules relevant to opinion as a hint for opinion analysis (Wiebe et al., 2004). The interpretation of the text in the development of the sentiment lexicon is based on the words in the lexicon that have been given unique characteristics that mark positive or negative or neutral feelings. Most of the words are verbs and adjectives but also some common nouns and adverbs.

Hu and Liu (Hu & Liu, 2004) authors develop a lexicon-based technique to use the opinion mining term in the opinion mining task. They dealt with product reviews from Websites to create a report of +ve, -ve, or neutral statements made about product features. They categorized aspects like camera size, image quality, etc. by selecting the frequent term, assuming that people often use the same words to determine features. Then they identified opinion sentences and their orientation. According to them, an opinion sentence is a sentence that contains both features.

Here's a simple example of how a rule-based method works:

- i. Defines two polarized word lists (e.g., negative terms such as መጥፎ, በጣም መጥፎ, አስቀያሚ, etc and +ve words such as ጥሩ, ምርጥ, ቆንጆ, etc).
- ii. Compute the number of terms that are +ve and -ve in a given text that appears
- iii. If the number of positive term occurrences is greater than the number of negative term occurrences, the proposed system returns a positive opinion and vice versa. If the numbers are even, the proposed system will return a neutral opinion.

2.4.2. Machine Learning Approach

Automatic approaches do not rely on manually designed rules, however on machine learning techniques, compared to rule-based methods. A technique of machine learning using unsupervised or/and supervised learning to construct a model from a huge dataset. Classification is a method in which a model learns to identify and define significant groups of data and training a data label set to predict discrete class labels on new samples. According to (Anderson et al., 1994; Mowlaei et al., 2020) commonly used algorithms for machine learning sentiment classification are Naïve Bayes (NB) classification, maximum entropy (ME) classification, and SVM classification. The experimental results produced through the techniques of machine learning are very good. NB tends to work the worst in terms of relative efficiency, and SVM tends to work the best, even though the differences are not very significant. Although techniques of machine learning have been found to produce good results. The collection of such a training set is challenging, as the collection and human classification of the vast number of different features are involved. Typically, supervised machine learning algorithms are corpus-based methods of classification to find patterns of co-occurrence (e.g., frequency) of terms in the corpus to assess the feelings of words or phrases.

Work in (Pang et al., 2002) also tested multiple supervised machine learning methods for movie review sentiment classification and concluded that machine learning techniques outperform the system based on human-tagged characteristics, although none of the current methods could manage the classification of sentiment with acceptable accuracy. In this research, our emphasis is on the classification of the aspect-level sentiment using a supervised machine learning approach for Amharic opinionated text in the area of sentiments in social media content.

2.4.3. Hybrid Approach

Hybrid systems incorporate into one method the desirable elements of rule-based and automated techniques. The huge benefit of these techniques is that the result is always more precise. For example, According to (Hovy, 2006) both rule-based and statistical methods were integrated into the work They proposed a methodology that could classify opinion holders and opinion topics. They performed two tasks to achieve their goal: first, using FrameNet data, they collected opinion words and related frames. Then, to mark semantic parts in a sentence, they applied a statistical approach. Their experimental outcome reveals that their system performs much better than the baseline.

2.5. Background of Amharic Language

Ethiopia is a linguistically diverse nation above 80 languages are used. In Ethiopia, numerous languages are spoken, Amharic language is one of the dominant because a significant part of the population speaks it as a mother tongue and is the second language most commonly learned throughout the country. The Language is the working language of the country's federal government. From the Ge'ez writing script, the present Amharic writing system has been adopted. Ge'ez, which belongs to the Semitic language class.

The repetition in the number of symbols with the same accent is a distinct outcome of Ge'ez's growth. For example, the three different ሀ, ሐ, and ሓ symbols are used interchangeably in text written in Amharic although they gave different meanings to words in the Ge'ez language. Likewise, both ሠ and ሡ have the same accent (/s/), both ፀ and ፂ (/s'/), and both ኦ and ዐ (/a/). This redundancy has been recognized in(Haile, 1967) as a problem of the language.

2.6. The Amharic Writing System

The Ethiopic writing system, which the Amharic language uses, consists of a core of thirty-three characters (ፈገገ, Fidel) each of which occurs in one basic form and six other forms all known as orders. The seven orders (the first basic order and the other six orders) of the Ethiopic script represent the different sounds of a consonant-vowel combination (a characterization known as syllabic). The 33 core characters then yield 231 distinct symbols. In addition to the 231 characters, others include special features usually describing labialization, the script of the Amharic language

is phonetic (የድምፅ አወጣጥ), and it has 32 consonants and 7 vowels. Each symbol represents a consonant together with its vowel. The vowels are combined to the consonant shape within the frame of diacritic markings. The diacritic markings are strokes attached to the base characters to change their order (e.g. the first order ‘le’ ለ is transformed into the second-order symbol lu “ሊ” by attaching “.” to it; the same first order is transformed into the fifth-order symbol lai(as inlaid) “ሊ” by attaching “.” to it, and so on).

2.7. Punctuation

In the Amharic language, there are about 17 punctuation marks like a colon (: Hulet netib) are word delimiters. The equivalent of a full stop is four dots arranged in a square pattern (:# arat netib). Some others include equivalents for the comma (፣ netela serez) and the semi-colon (፤ dirib serez), as well as a number of the borrowed symbols (?!, “,” ‘, /, \, etc.). The word delimiter (two dots) is mostly used in handwritten text. Space is being utilized as a word separator instead. The sentence delimiter (full stop) however, continues to be used. Other punctuations as comma, colon, and semi-colon are used where appropriate.

2.8. Numbers

Numbers in Amharic consist of single characters for one to ten, for multiples of ten (twenty to ninety), hundred, and thousand. The number in the Amharic language is derived from Greek letters. **Example** አራት መቶ ሰላሳ አንደኛ።

1	2	3	4	5	6	7	8	9	10
፩	፪	፫	፬	፭	፮	፯	፰	፱	፲

2.9. Characteristics of the Amharic Writing System

The different symbols with similar accents also pose a problem in making words appear different (not in meaning, but spelling.) Although in the Ge’ez language, these different symbols give each word different meanings, in the Amharic language they have been used interchangeably(Bender, M. L., Sydney W. Head, 1976; Haile, 1967). The class of symbols with the same sound falls into two:

1. the same sound for the first and fourth-order alphabets.
2. different alphabets that share the same sound.

The following belong to the class of the first type. Letters in the same row share the same sound.

<u>1st order</u>	<u>4th order</u>
ሀ	ሃ
ሐ	ከ
ኀ	ከ
አ	አ
ዐ	ዓ

እሳት and እሃት, the symbols (letters) “ሳ” and “ሃ” have a similar sound and do not change the meaning of the word, which is “fire”. Likewise, the two words “ጸሐይ” and “ፅሐይ”, “ጸሀይ” and “ፅሀይ” all mean the same, “sun” although they are written differently (indicating alternate use of ጸ and ፀ, and ሀ and ሐ respectively).

2.10. Contextual Shifter (Negation)

Handling negation in the study of sentiment and emotion can be a significant concert. Negation is key and must be treated carefully. A single negation can shift the sentiment of the entire sentence, such as the expression "I like this car" and "I don't like this car," the two phrases considered by the most largely used similarity measures to be quite similar, and only the negation word is the difference. Negation, however, does not force the second sentence into the opposite polarity, so there is a list of negations that can change the term's sentiment polarity from +ve to -ve and vice versa that occurred before the word of sentiment. In Amharic language negation like ነገር ግን(But), አይደለም(Not), ቢሆንም (Even so), አይደለም(Not).

2.11. Grammatical Rules of Amharic

Word Series: The Amharic language follows the grammatical pattern of the subject-object-verb as in opposition to, for example, the English language that has an SVO word sequence. For example, the Amharic equivalent of "yordanos ate bread" is written as "የሮዳኖስ (yordanos)

ዳቦ(dabo/bread) በላ(bäla/ate)." The speech tags of individual words are used here as inputs to verify grammatical validity.

Sentences in Amharic

An Amharic sentence is collected of NP and VP. NP comes first, then VP follows, about their order in a given sentence. A sentence is several words organized formally. In Amharic, a noun sentence is a phrase with a noun as its starting word. An example of NP in Amharic is አንበሳ ሁለት ላሞችን ገደለ because the starting term አንበሳ/Lion is a noun. In a sentence, a verb means action or the state of being of something. In ካሳ መስኮቱን ሰበረ, for instance, the word " መስኮቱን "means that the window has been shattered by ካሳ in ካሳ ነጋዴ ሆኑ, the word " ነጋዴ ሆኑ " indicates the state of being, i.e., the individual called ካሳ has become a merchant.

2.12. Amharic Phrases

A phrase is a structure in a language that is constructed from one or more words in the language. In Amharic, phrases are categorized into five categories, namely noun phrase, verb phrase, adjectival phrase, adverbial phrase, and prepositional phrase(Baye Ymam, 2002).

Noun Phrases (NP)

An NP is a syntactic unit where a noun or a pronoun is the head (H). It can be easy or intricate. A single noun or pronoun is the easy NP, such as እሱ (he), እሷ (she), እነሱ(they), etc. A composite NP may incorporate a noun (called the head) and other components (such as specifiers, adverbial and adjectival modifiers) that change the head from various components(Baye Ymam, 2002). For instance, ለገና አንድ ቆንጆ ቡችላ አፈልጋለሁ።

Note that, an NP can be also constructed from another different possible combination.

Adjectival Phrases (AP)

The construction of an Amharic adjectival phrase is similar to that of a noun and verb phrase. It can be composed of an adjective (head), and other constituents such as complements, modifiers, and specifiers(Baye Ymam, 2002). For instance, የመኪናው ዋጋ በጣም ውድ ነበር።

The AP can be obtained from nouns and verbs (Baye Ymam, 2002). Adjectival noun phrases are phrases that are constructed from nouns and called adjectival phrases. The AP can also be derived from verbs called adjectival verb phrases. አስቴር እንደ ካሳ ሽማግሌ አታላይ. In this example, እንደ ካሳ

ሸማግሌ ኢታላይ is an adjectival phrase, its headword is the adjective ኢታላይ/etalay derived from the verb ኢታለለ/etalele.

2.13. Related Works

This chapter briefly reviews some of the Aspect level Sentimental Analysis for both English and Local languages. In this review, we are going to see the employed approaches, domain, target language, corpus source, procedures, experimental results, and performance are the main emphasize points to be given focused when going through the different works.

According to (Mei et al., 2007), the key points of the mining feature level opinion were drawn up by the researcher. To the analysis of the feature level opinion, according to the researcher, first determines the goals for which sentiment has been expressed in a sentence and then decides whether the sentiment is +ve, -ve, or neutral. The researchers' goals were: objects, components, features, and characteristics on which the reviewer commented and determined whether the comments were positive or negative. This paperwork is on the latest issue of simultaneously forming subtopics and sentiments in Weblogs. It describes the problem of Topic-opinion mining (TOM) and proposes a probabilistic mixture model called Topic opinion Mixture (TOM) in a series of blog articles to extract multiple subtopics and opinions.

According to (Tripathi & S, 2015), sentiment analysis has arisen as a common and efficient technique for information extraction and data analysis. The quick development of user-created content has opened up modern skylines for investigation within the field of sentiment analysis this paper uses a combination of NLP and machine learning approaches to propose a model for sentiment analysis (including 2000 film reviews) of film reviews. First, on the dataset, various data pre-processing schemes are implemented. Secondly, in combination with various variable selection schemes, the behavior of two classifiers techniques, NB and SVM, is examen to get the results for sentiment analysis. Thirdly, the proposed demonstration for opinion analysis is expanded to get the comes about for higher-order n-grams. The result shows that the TF-IDF scheme obtains maximum accuracy for linear SVM. The term, occurrence gives maximum accuracy for the Naïve Bayes classifier. 84.75% accuracy acquire from unigrams and Linear SVM with the TF-IDF scheme.

According to(Kumar et al., 2016), sentiment mining is an encouraging research area of data mining used to extract precious knowledge from the huge volume of users' remarks, reviews, and feedback on any topic or product and event. According to this research using ontologies for evaluating customer feedback about the performance of the product and the extraction of the product's performance feature is a significant task of review analysis and summarization. The proposed approach provides a new Score Calculation Algorithm (SCA) based method for score calculation computed on various features from the customer opinions. It is based on three special stages: (i) feature identification, (ii) a new SCA-based method for score calculation computed on various features in the customer opinions (iii) compile and display the features along with appropriate score values in the order of document-based. Firstly, identify the features was based on the ontology-based aspect identification method. Secondly calculated the TF-IDF values of each feature and evaluate, score values for each feature. Finally, this system ranked the features on the origin of their score rates. The methodology was implemented and tested in the canon digital camera feedback dataset. The performance exhibit result shows that the average accuracy of the correctly classified feature was found to be 81.11%.

According to(Varghese & Jayasree, 2013), Opinion analysis involves the method of distinguishing the polarity of opinionated texts. Lots of social network sites are being used for expressing thoughts and opinions by users to rate products. This user opinionated text is highly unstructured and thus involves the application of various natural language processing techniques. In feature-based sentiment analysis, the various aspect of a product is identified through the training process. The quantitative reasoning of each feature is done using an SVM classifier. In most of the earlier works, a product review is analyzed as an entire rather than considering each aspect of it. feature-based sentimental analysis is challenging because the extraction of individual features is in itself a tedious task. The system which also uses coreference resolution produces results with an average accuracy rate of 77.98%.

According to(Saritha & Nayak, 2019), opinion mining is extremely helpful in social network monitoring as it enables us to pick up an outline of the more extensive popular sentiment behind certain topics. People share their opinions on products and services through social network applications such as blogs, Facebook, Twitter, and Instagram, etc. These online reviews provide new customers with a sense that the particular business is genuine and provides a real product or service. Sentiment analysis focuses on the opinions which may be positive or negative. Sentiment

analysis based on feature-based, particular aspects are extracted and opinion polarity of the aspects is determined. The proposed model is analyzed using a reference dataset of mobile comments collected from Amazon. Authors use different machine learning techniques such as Support vector machine, Naïve Bayes for the classification of human opinion, and Evaluation parameters such as precision, recall, F-score, and accuracy are used to assess the performance of a classifier. The Accuracy using SVM is 88.43 and NB is 84.61.

According to(Tamrakar et al., 2020), authors propose feature-based sentiment analysis to help in understanding the opinion of the related entities helping for a better quality of service. A model is developed to detect the aspect-based sentiment in Nepali text using Machine Learning (ML) classifier algorithms namely Support Vector Machine and Naïve Bayes. The system collects Nepali text data from various websites and Part of Speech tagging is applied to extract the desired features of aspect and sentiment. Manual labeling is done for each sentence to identify the sentiment of the sentence. TF-IDF is applied to compute the importance of the words. The feature vectors thus produced are then applied to the Classifier algorithms to predict and classify the sentence. The accuracy obtained by the SVM classifier is 76.8% whereas Bernoulli NB is 77.5%

According to(Devi et al., 2016), In this modern era of globalization, e-commerce has become one of the most convenient ways to shop. Every day people buy many products online and post their reviews about the product which they have used. Even so, it would be a difficult task to manually analyze. This problem can be alleviated by an automated system called 'Sentiment Analysis and Opinion Mining' that can analyze and extract the users' perceptions in the whole review. In this work, they developed a comprehensive process of 'Aspect based opinion mining' by using Support Vector Machine in a novel approach. It is proved to be one of the ultimate good ways to analyze and extract the overall users' view about the particular feature and whole product as well. The accuracy by using SVM is 90.85%.

According to(Mowlaei et al., 2020), Reviews posted on commerce websites have been a significant source of information for customers as well as companies. Text sentiment analysis approaches aim to detect the sentiments of written reviews to obtain a deeper understanding of public opinion towards entities. opinion mining at aspect-based deals with an occupying opinion expressed about each feature of entities. A known approach in opinion mining problems is to take advantage of

lexicons to produce a feature for the classification of reviews. aspect-based techniques fail to adapt properly common dictionaries to the conditions of aspect-based corpus which comes about in reduced performance. To address this issue this paper proposes expansions of two vocabulary generation strategies for feature-based issues; one utilizing statistical strategies, and another employing a genetic algorithm displayed in our previous works. To classify the features in reviews, the previous lexicons are then combined with important static lexicons; this best our previous work based on t-test results (p-value of < 0.001). Experimental findings show that the proposed methodology outperforms baseline approaches on Bing Liu's consumer review datasets in aspect-based sentiment classification and increases recall, precision, and F-score by 1.0, 6.0, and 7.4 percentage respectively.

According to (TULU TILAHUN, 2013), The author proposes a sentiment analysis at the aspect-level model for the Amharic language by employing manually crafted rules and lexicon. The model proposed includes five major parts that can extract features, determine opinion words regarding identified features with their semantic orientation, aggregate multiple opinions, and generate a structured summary. Two experiments have been conducted for feature extraction and opinion word determination by using 484 reviews from three different domains. The first experiment indicated that an average precision of 95.2% and recall of 26.1% were achieved in the feature's extraction and an average precision of 78.1% and recall of 66.8% were achieved in the determination of opinion words. The precision of the second experiment in features extraction gets lower by 15.4% because the precision of opinion word determination gets higher by 1.9% and the recall of both features extraction and opinion word determination gets higher by 7.8% and 25.9% respectively when compared to the first experiment. Thus, our experimental results demonstrate the effectiveness of the techniques we have applied.

According to (MELESE MIHRET, 2017), The researcher works on feature a level opinion mining model for the Awnji language using Machine learning. the author has used three classification techniques namely, Maximum Entropy, Naive Bayesian, and SVM for building the model and they used precision, recall, f-score, and accuracy for evaluating the model experiments, the author used unigram words as features, because no more bigram words as features and customized manually crafted rules as a feature selector and a simple bag of the word as a feature extractor achieved the best result. The outcome produced via the machine learning approach is promising. As the result

of the experiment shows, the NB classifier algorithm achieved an average accuracy of 75%, 70% accuracy for MxEn, and outperforms 79 % of SVM performances were obtained. Hence, the SVM algorithm outperforms those of NB and MxEn. The author further recommended research to be conducted on sentiment analysis by provided the availability of standardized text corpus and POS Tagger.

The related works that can be examined and discussed above were summarized based on the following features, such as the name of the author, objective, techniques, and algorithms used, Corpus, and Target languages to conduct the study.

Table 1.1: The Previous Research Thesis for Aspect Level SA

Researcher Name	Objective	Approach	Algorithm	Features	Corpus	Target Language
Qiaozhu Mei Xu Ling Matthew Wondra Hang Su ChengXiang Zhai	Topic Sentiment Mixture: Modeling Facets and Opinions in weblogs	Unified probabilistic approach	▪ Expectation-Maximization algorithm. ▪ Baum-Welch algorithm. ▪ Viterbi algorithm.	Unigram, Bag of words	Weblog datasets	English
Gautami Tripathi Naganna S	Feature Selection and Classification Approach For SA	NLP + ML approaches	NB & SVM,	TF-IDF, N-grams	(Polarity dataset V2.0)	English
Kumar, A Nayaki, A P Devi, M	Customer Feedback Evaluation System Using Feature-Based OM	score calculation	Score Calculation Algorithm	TF-IDF	canon digital camera	English

Raisa Varghese Jayasree M	Aspect Based Sentiment Analysis using SVM	Machine Learning	SVM	Part-Of-Speech tagging	Amazon dataset (digital camera)	English
B Saritha1 Janmenjoy Nayak	Aspect Based Sentiment Analysis using NB & SVM	Machine Learning	SVM and NB	POS Tagging, TF-IDF	Amazon (movie review)	English
Sujan Tamrakar, Bal Krishna Bal, Rajendra Bahadur Thapa	Aspect Based Sentiment Analysis Of Nepali Text Using SVM and NB	Machine Learning	SVM and NB	Part of Speech (POS)	Nepali websites	Nepali
D V Nagarjuna Devi, Chinta Kishore Kumar, Siriki Prasad	A Feature-Based Approach for Sentiment Analysis by Using SVM	Machine Learning	SVM	POS Tagging	Collected from companies like HP, APPLE, etc	English
Mohammad Erfan Mowlaei, Mohammad Saniee Abadeh	Aspect-based sentiment analysis using adaptive aspect-based lexicons	Machine Learning	Genetic Algorithm	PoS-tags	Dataset of restaurant reviews	English
Tulu Tilahun Hailu	Opinion Mining from Amharic Blog	Rule-Based Approach	Lexicon	PoS-tags N-grams	Manually collected from hospital, hotel.	Amharic

Melese Mihret	Sentiment Analysis Model for Opinionated Awngi Text	Machine Learning	SVM, NB, MxEn	POS Tagging, TF-IDF	Manually collected from Awngine wspaper	Awngi
My Research	Develop feature/aspect level Sentiment Analysis, models	Machine Learning+ Cosine Similarity	NB, LR, SVM, and GB [for opinion analysis] And Cosine similarity [for aspect extraction]	POS Tagging, Ngram, and TF-IDF, Countvecctor izer	Hotel from their own social media site, and manual collection	Amharic

CHAPTER THREE

METHODOLOGY

3.1. Research Methodologies

In this chapter, we have gone through a research methodology for creating a dataset as well as strategies for achieving research goals and answering the research question. The methodology used in conducting the study on Aspect Based Sentimental Analysis for the Amharic language is explained and justified in the following section.

3.2. Building a Corpus

This research aims to detect Amharic Aspect-based SA. As a result, it would need to build a new Amharic Aspect-based Sentimental analysis dataset. There are three key phases in the process of creating the dataset for Amharic Aspect-based Sentimental analysis.

1. Collecting textual data (comment) from public social media.
2. Collected data is prepared, filtered, or consolidated into a single file dataset.
3. Annotating the results.

3.2.1. Data Collection

Social network users usually access services using web-based innovations on desktops, mobile, and laptops. There are numerous types of social media like Facebook, Twitch, YouTube, Instagram which may be a popular social networking site that permits registered users to set up profiles, release photos, and audio, send a message and keep in touch with companions, family, and colleagues. Facebook has been chosen from other social network platforms because of the ubiquity of the platform, The build and look of the site are very attractive which gives the user a more prevalent interface to work on. And according to Internet Users Statistics in 2020 for Africa, Facebook had 6,745,000 clients in Ethiopia in January 2020, which accounted for 5.3% of its whole populace(*Africa Internet Users, 2021 Population and Facebook Statistics*, n.d.). The study collects reviews from a Facebook public page using Face pager.

This research aimed to develop text-based classification methods for hotel reviews. People read reviews to see if previous guests were pleased with their hotel reservations in the digital age that we live in. Before making a hotel room reservation, users go online to read the business's reviews to learn more about the hotel's accommodations, to get a sense of how clean the rooms are, to learn about the amenities, security, and food options, and ultimately, people read reviews to see if previous guests were satisfied with their hotel reservations.

To collect the comments, a list of Hotels in Ethiopia is considered, based on the content of service that gives to the customer and social media (Facebook pages followers and website). Because the study aims to have a diverse source of the hotel and have a representative dataset. Social media (The hotel should have a social media site like Facebook, website), and the number of followers and likes (Each hotel's social media has a high follower and like rate) is considered. Ethiopia has more than 600 hotels, however, most of them are placed in Addis Ababa, Hawassa, Bahir Dar, Adama, and Mekele. most of the hotels in Ethiopia have their own social media site to get customer reviews, but analyzing customer reviews is a tedious task because of manually analyzing, so this research focuses on analyzing customer feedback in an automotive way.

3.2.2. Method of Data Collection

Data is the foremost profitable asset today's businesses have. The more data regarding your customers, the superior you'll be able to get their wants and needs. data collection methods are categorized into primary and secondary data gathering. The data accumulated by gathering primary data methods are particular to the research's rationale. Primary data collection strategies can be isolated into two-stage: quantitative method (which is numeric and quantifiable) and qualitative methods (which is expressive, instead of numeric). Besides, data regarding customer opinions about hotels are collected by using Face-pager tools which are software useful to collect primary data.

3.2.3. Hotel Selecting Criteria

To select hotels and to determine features there is an inclusion criterion like service provided by hotels (Value for money, Location, Local experience, Cleanliness, etc.), and the hotel should have social media site and the number of followers has significant value for selecting a hotel (≥ 12000 site follower selected).

The hotels that are selected for this research based on the above inclusion criteria are the following.

1. Capital Hotel and Spa, Addis Ababa

Capital Lodging and Spa is a 5-Star Lodgings within the provincial capital Addis Ababa-Ethiopia. the hotel has rich conference rooms; luxury housing; premium feasting rooms; a dedicated sports Bar; fine restaurants; an authentic Ethiopian Cultural Restaurant nearby social music and dance performance; a luxurious Spa, an Exercise center, and a swimming pool; and one kind Hotel and Administration College.

Capital Hotel and Spa		
Website Address	https://capitalhotelandspa.com/	
Social Media (Facebook Address)	https://www.facebook.com/CapitalHotelandSpa	
	Number of Followers	Number of Like
	27,676 people follow this	27,396 people like

2. Radisson Blu Hotel, Addis Ababa

Radisson Blu Hotel in Addis Ababa is centrally found inside Ethiopia's capital city-adjacent to the United Nations area and 6 kilometers from Bole International Airport. The Radisson Blu Lodging, Addis Ababa highlights 212 rooms for trade or relaxation travelers, total with free Wi-Fi and in-room coffee and tea facilities.

Radisson Blu Hotel		
Website Address	https://www.radissonhotels.com/	
Social Media (Facebook Address)	https://www.facebook.com/RadissonHotels/	
	Number of Followers	Number of Like
	87,185 people follow this	85,290 people like

3. Hilton Hotel, Addis Ababa

Hilton hotel is located in Addis Ababa within the heart of Africa. Select from an assortment of restaurants and bars in Hilton Addis Ababa, advertising remarkable menus, live amusement, and romantic settings. Have an assembly in one of our advanced occasion spaces or catch up with work within the trade center.

Hilton Hotel		
Website Address	http://addisababa.hilton.com/	
Social Media (Facebook Address)	https://www.facebook.com/HiltonAddisAbaba/	
	Number of Followers	Number of Like
	103,667 people followers	103,288 people like

4. Sheraton Addis, a Luxury Collection Hotel, Addis Ababa

Found within the heart of the Ethiopian capital, Sheraton Addis, a Luxury Collection Hotel, Addis Ababa sits on a peak overlooking the city, securely settled between the National Palace and the prime minister's residence.

Sheraton Addis Hotel		
Website Address	http://www.sheratonaddis.com/	
Social Media (Facebook Address)	https://www.facebook.com/AddisSheraton/	
	Number of Followers	Number of Like
	21,599 people followers	21,128 people like

5. Ethiopian Skylight Hotel

Ethiopian Skylight Hotel, claimed by Ethiopian Airlines, is the foremost lavish and the biggest hotel in Ethiopia found at the heart of Africa's diplomatic center Addis Ababa, just six minutes away from Bole International Airport.

Ethiopian Skylight Hotel		
Website Address	https://www.ethiopianskylighthotel.com/	
Social Media (Facebook Address)	https://www.facebook.com/ethiopianskylighthotel/	
	Number of Followers	Number of Like
	24,368 people followers	23,593 people like

6. ELGEL Hotel and Spa

Elgel Lodging and Spa offers the preeminent amenity and installment services to our visitors, so they can first-hand encounter supreme splendor all through their remain. room sees are the foremost noteworthy particularly at characteristic timberland side, it has a mind-blowing see of the evergreen timberland with best arrive include. they have 52 rooms that are all prepared with modern furniture and verandas.

ELGEL Hotel and Spa		
Website Address	https://elgelhotelandspa.com/	
Social Media (Facebook Address)	https://www.facebook.com/elgelhotelandspa/?ref=page_internal	
	Number of Followers	Number of Like
	16,133 people followers	15,897 people like

7. HAILE HOTELS AND RESORTS

Haile Inns and Resorts are titled behind the Ethiopian myth long-distance runner Major Haile Gebreselassie. His eminent state “it is Possible!” has made Haile Inns & Resorts a reality and keeps boosting the neighborliness industry over Ethiopia & East Africa. We yearn to be the benchmark inborn inn chain engineer and administrator in East Africa by 2025.

HAILE HOTELS AND RESORTS		
Website Address	https://www.hailehotelsandresorts.com	
Social Media (Facebook Address)	https://www.facebook.com/Hailehotelsandresorts/?ref=page_internal	
	Number of Followers	Number of Like
	123,999people followers	121,676 people like

8. Harmony Hotel

Placed in Addis Ababa, Harmony Hotel is in the city center and near the airport. Edna Shopping center and Selam City Shopping center are worth checking out on the off chance that shopping is on the plan, whereas those wishing to involve the area's common excellence can investigate Solidarity. Park and Sheger Park. Besides 2 restaurants, this smoke-free inn includes a full-service spa and an indoor pool, Free buffet breakfast is given, as well as free WIFI in open ranges, free valet stopping, and a free air terminal shuttle. Additionally, a nightclub, 3 bars/, and a health club are onsite.

Harmony Hotel		
Website Address	http://www.harmonyhotelethiopia.com/	
Social Media (Facebook Address)	https://www.facebook.com/HarmonyAddis/	
	Number of Followers	Number of Like
	12,060 people followers	11,802 people like

9. Intercontinental Hotel, Addis Abeba

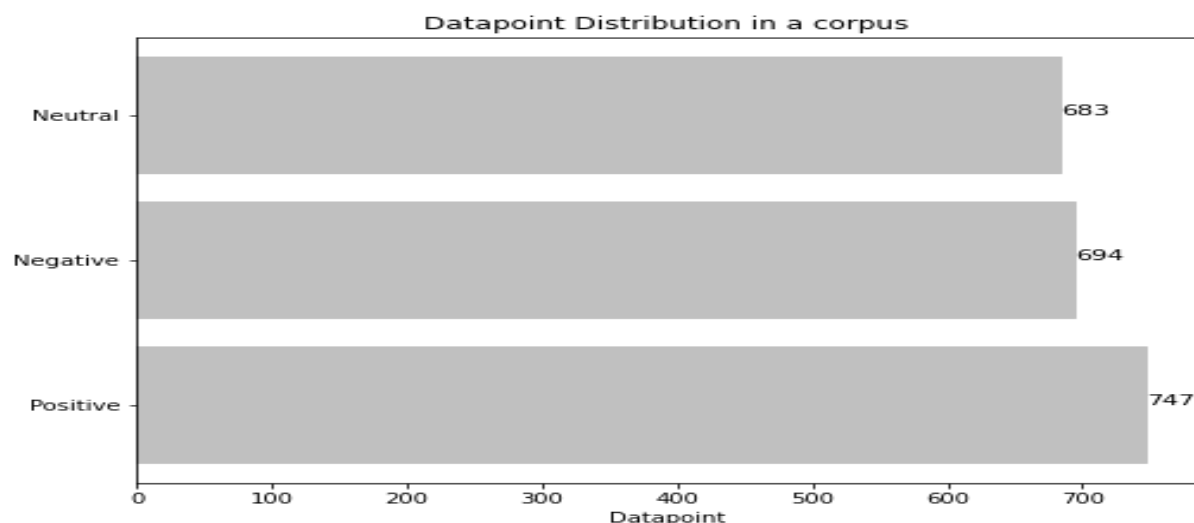
Intercontinental hotel located in Addis Ababa. specifically, in Kirkos; one of the city's most prevalent districts. The lodging lies 0.5 Km from the city center and gives openness to critical town offices.

10. Planet International HOTEL Mekelle

Planet Hotel places you within the heart of Mek'ele, a 4-minute drive from Memorial for Martyred Freedom Warriors and Loyalists and 7 minutes from Mekele College. Highlighted comforts incorporate a 24-hour business center, complimentary daily papers within the campaign, and dry cleaning/laundry administrations, and have a social network to connect with clients.

Total 2124 Data was collected using its Facebook and website from 10 selected hotels. the below table shows the number of data collected from selected hotels.

Hotel Name	Number of Data collected
Capital Hotel	220
Radisson Blu Hotel	200
Hilton Hotel	215
Skylight Hotel	209
Intercontinental Hotel	315
Haile hotel and Resort	250
Elgel Hotel and Spa	200
Harmony Hotel	233
Planet International HOTEL	108
Sheraton Addis	174
Total	2124



3.3. Dataset Preparation

Following the data collection, a data processing process is carried out, which includes data collection, cleaning, filtering, and consolidation into a single file or data table. Preparation methods such as spreadsheet software (Microsoft Excel) were used in this procedure. Filtering and cleaning the data is primarily used for the next step, which is the annotation of the comments in the dataset, and then for predict an Amharic Aspect-based SA training model.

The following tasks are done to prepare the dataset for annotation:

- All non-Amharic and non-textual comments have been removed.
- Getting rid of all null, blank values, and whitespace.
- Combining data into a single dataset

3.4. Dataset Annotation

The annotation process in this case must label a comment and feature to create Aspect based SA datasets. Three annotators participated in the labeling process. The annotator was selected from the Social Sciences Faculty and Humanities of Amharic Language and Literature in woldia university and Adiss Abeba university.

3.4.1. Annotation Guideline

A guideline was created and provided in Appendix I to make the annotation process transparent. The annotation guidelines are also prepared by consulting Linguistic experts. Using this guideline, the annotators label each comment into three different classes (+ve, -ve, and neutral, and also label its feature/aspect).

3.4.2. Annotation Procedure

The dataset annotation procedure, which is primarily handled by the researchers, also includes three additional annotators. Because of a lack of resources, primarily budget and time, the number of annotators participating in the study was limited. The annotators were chosen based on their ability to complete the mission and their knowledge of the Amharic language. All the annotators have master's degrees in linguistic. The annotators were given a random subset of an equal number of comments from the dataset to annotate. Three annotators were given the task of annotating the same number of comments to determine the inter-rater agreement, which indicates how closely their annotation decisions matched.

Due to the challenging nature of the manual annotation process of Aspect Based SA datasets, the annotators were allowed to annotate in their own time freely. However, for the annotation task to be easier, we give brief insights into the annotation guidelines provided for labeling comments.

3.4.3. Inter Annotator Agreement

Inter annotator agreement assessment calculates the agreement of the dataset's human annotators by measuring how similar the annotators labeled the dataset's comments and features. Cohen's kappa was used to compute the intensity of agreement in the analysis. We told them to annotate the same number of comments and aspects in the dataset, as well as a similar subset of them. To compute the kappa using a dataset that has been annotated by all annotators. Cohen's Kappa value ranges between 0 and 1. The following is the standard interpretation of Kappa values(“*Cohen’s Kappa Statistic - Statistics How To.*”).

Cohen’s Kappa

According to (“*Cohen’s Kappa Statistic - Statistics How To.*”), Cohen's kappa is an analytical compute of rater similarity (sometimes called interobserver agreement). Interrater reliability, or exactness, happens when your information raters (or collectors) allow similar esteem to the same information item. In this research, three annotators participate and reliably use the criterion to make the same assessment on the same targets, then their agreement going to be very high and provide the evidence we have reliable ratings.

Table 1.2: Distributing Comment to Annotators

Annotator Name	Annotated Data
Annotator 1	2124
Annotator 2	2124
Annotator 3	2124

The three annotators annotated the same data and analyzed it by using the Kappa statics tool, and we get 0.78, which is substantial agreement.

The Kappa statistic varies from 0 to 1, where.

- 0 = agreement equivalent to chance.
- 0.1 – 0.20 = slight agreement.
- 0.21 – 0.40 = fair agreement.
- 0.41 – 0.60 = moderate agreement.
- 0.61 – 0.80 = substantial agreement.
- 0.81 – 0.99 = near perfect agreement
- 1 = perfect agreement.

3.4.4. Dataset Annotation Result

Each annotator annotates 2124 data, and we decided on the final class using the majority votes method. The annotator rater agreement for 2124 comments that receive annotation from three annotators results in 0.78 Kappa. The kappa value interpretation indicates that a substantial agreement between annotators.

Table 1.3: Sample inter-annotator agreement

No	Comments	1 st annotator		2 nd annotator		3 rd annotator	
		polarity	aspect	polarity	aspect	polarity	aspect
1	ጥሩ የደንበኞች አገልግሎት በመስጠታቸው ለሁሉም ሰው በጣም ትኩረት ይሰጣል።	pos	አገልግሎት	pos	አገልግሎት	pos	አገልግሎት
2	የጠየቅነውን ሰላጣ በጭራሽ አላመጡም።	neg	ምግብ	neg	ምግብ	neg	ምግብ
3	ምግቡ አስገራሚ ነበር።	pos	ምግብ	pos	ምግብ	pos	ምግብ
4	አስተዳደሩ ጨዋነት የጎደለው ነው።	neg	አስተዳደር	neg	ማናጀር	neg	ማናጀር
5	ሰራተኞቹ ተግባቢ ናቸው።	Pos	ሰራተኛ	pos	ሰራተኛ	pos	ሰራተኛ
6	ሆቴሉ ጥሩ ድባብ አለው።	Pos	ድባብ	pos	ድባብ	pos	ሆቴል
7	የሆቴሉ ዋጋዎቹ ድንቅ ነበሩ።	Pos	ክፍያ	pos	ክፍያ	pos	ክፍያ
8	የዶሮ ወጥ አሰራሩ ልዩ ነው።	Pos	ምግብ	pos	ምግብ	pos	ምግብ

9	በግ ስጋው አይመችም።	Neg	ምግብ	neg	ምግብ	neg	ምግብ
10	የምግቡ አሰራር	neu	ምግብ	neu	ምግብ	neu	ምግብ
11	በዚህ ሆቴል ለአምስት ቀናት ቆየን።	neu	ሆቴል	neu	ሆቴል	neu	ሆቴል

According to (Hernik & Grīnberga-Zālīte, 2017), There are different features that hotels give like Value for money, Location, Local experience: Cleanliness, SafeFree internet access: Service, but For this research features that selected by questionnaires like ሆቴል(Hotel), ምግብ(food), መጠጥ(drink), አገልግሎት(service), አካባቢ(Ambience), ቦታ(Location), ክፍያ(Price), ጥራት(Quality).

Throughout this study, the word feature refers to the aspect or part of the Hotel. We manually gathered features as POS taggers to see how many of them the algorithm can recognize. Identifying a word as a feature may be difficult. The cosine Similarity Algorithm is used to evaluate the aspect because of the use of cosine similarity, which is characterized by the comparison of similarity among two vectors of an inner product space features in this analysis, but the frequency of features is used to decide the best feature using the proposed algorithms. Then, to make the data suitable for python within NLTK, feature transliteration was used to compare the applicability of machine learning classification algorithms (LR, SVM, NB, and GB) are compared with the same data set and categories.

3.5. Aspect Based Sentimental Analysis Modeling

3.5.1. Data pre-processing in the Amharic language

Filtering out non-Amharic content was the first pre-processing phase performed on the collected data. "This is often particularly important when dealing with "web" or "social media" information. It is obtained in raw format whenever the data is elicited from various sources, which is not reasonable for study. Therefore, some steps are then taken to transfer the data into a small, clean data set.

First, all the sentences in the data are tokenized to their words, and all the stop words are detached from the training data as the next step. Stopwords are the widely used words in the language that do not add any meaning that helps the process of text classification. The main preprocessing measures that are used during the prediction of our model are removing numbers, accent marks, whitespaces, and other diacritics. In the preprocessing step, the normalization of Amharic

characters is also achieved. The normalization issue that we gave attention is the replacement of homophonic characters with their common forms; ሐ, ኸ and ኹ are replaced with ሀ, ሠ with ሰ, ፀ with አ, and ፀ with ጸ, normalization of punctuation marks, different styles of punctuation marks exist in texts from social media. For instance, double quotation marks, “:”, “ « or » are available.

3.5.2. Normalization

About punctuation marks, the normalization part of the preprocessing step handles the issue of Amharic text writing. The main function of the component is to replace the alphabets that have the same pronunciation and use one of the alphabets. For instance, the letters ሀ:ሐ:ኀ:ኸ are represented by letter ሀ and its variants with all their respective variants.

3.5.3. Tokenization

Tokenization may be used to break the sentence of the query into an instance of the sequence of characters that can be grouped for processing as an important semantic unit. During tokenization, we chop on whitespaces, throw away punctuation characters and choose the correct token to use. For instance, in the statement “ሙስና ከመፈጸሙ በፊት ለመከላከል እየሰራ መሆኑን ጸረ ሙስና ኮሚሽን አስታወቀ:”, unlike human beings, a computer cannot spontaneously understand that the sentence has 10 words; rather it understands the statement only as it has a sequence of 49 characters. Therefore, to convert them into meaningful tokens, they have to be tokenized using certain Amharic language tokenization criteria as:

ሙስና	ከመፈጸሙ	በፊት	ለመከላከል	እየሰራ	መሆኑን	ጸረ	ሙስና	ኮሚሽን	አስታወቀ
-----	-------	-----	--------	------	------	----	-----	------	-------

3.5.4. Stop-words removing

Stop words are words that do not use important meaning and are usually discarded from texts. It is not possible to remove Amharic stop words using Natural Language Toolkit (NLTK), we designed our own stop words because NLTK doesn't recognize Amharic stop words. Stop word in Amharic like ነው:ሆነ:ወደ:ናቸው:ውስጥ: ነበር :ብቻ :እኔ :እሷ :እሱ :እነርሱ. Etc.

3.6. Machine Learning Classification Algorithm

This research aims to use a machine learning approach for Aspect-based SA is categorized and labeled datasets. For comparison and detection accuracy, we use a multiple supervised machine learning algorithm. The algorithms were chosen for their high classification accuracy. However, for this study, a supervised ML approach was used, which included developing a collection of labeled training and test data sets for analysis tasks, which is required when using supervised learning and having prior knowledge of the data.

Supervised learning

The availability of annotated training datasets is the distinguishing feature of supervised learning. The name suggests a "supervision" that directs the learning system on which labels to associate with which training examples (Guarino, 1998). We choose an ML algorithm to train the function mapping from the dependent to the independent variable using supervised learning as an input variable (x) and an output variable (Y).

$$Y = f(X)$$

The aim is to approximate the function mapping to the point that we can predict the dependent variable (Y) for new text input data (x).

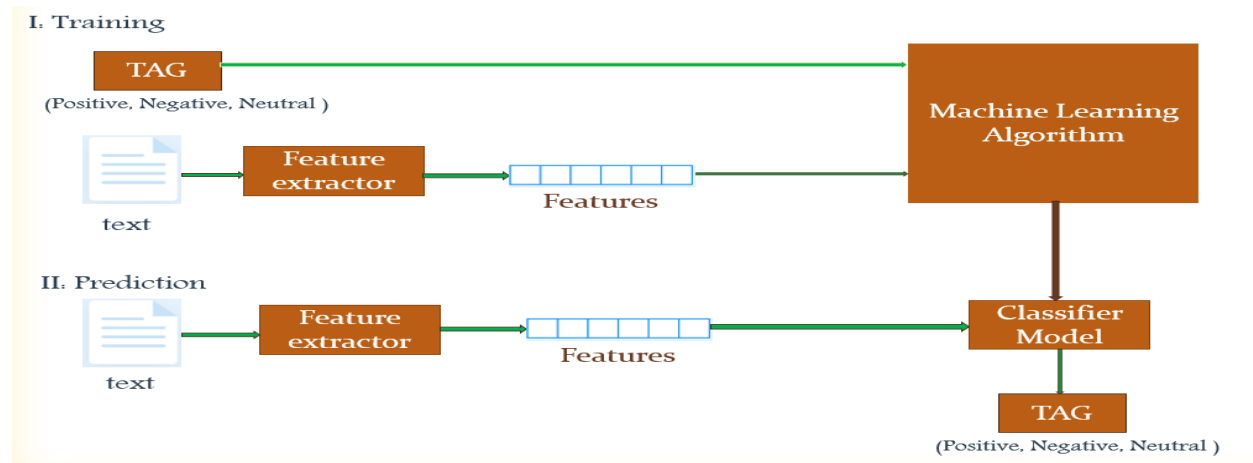


Figure 1.1: Process in Training and prediction

The following four machine learning algorithms are used to build Aspect based SA model.

Support Vector Machine (SVM)

According to (*Maximum Margin Hyperplane - an Overview | ScienceDirect Topics*, n.d.), SVM is a supervised machine learning technique. The SVM algorithm distinguishes groups or data points in feature space by finding a hyperplane in n-dimensional space. Hyperplanes are decision-making boundaries that aid in data classification. The equation below depicts the form of a hyperplane for a set of input points X in general.

$$w * x + b = 0 -$$

Where w is the hyperplane support vectors, b is the learning algorithm's intercept, and x is a fixed input data point. For predicting a new input, the SVM classifier can be written as,

$$fw, b(x) = g(wx + b)$$

Here, then $f(x) > 0$ for points on one side of the hyperplane, and $f(x) < 0$ for points on the other side. Training data is used to create a high-size space that can be divided into positive and negative instances using a hyperplane. The location of new instances in space concerning the hyperplane is used to classify them. Support vector machines find the maximum line that divides the training space into classes as far apart from each other as possible.

Naïve Bayes (NB)

Multinomial Bayesian classifier technique that assumes predictor independence and is based on Bayes' Theorem. To put it another way, the NB classifier assumes that the existence of one feature in a class has no bearing on the presence of any other feature. The Bayes' Theorem equation (Harrington, 2012).

$$P(c|x) = P(x|c) / P(c)P(x)$$

Logistic Regression Algorithm

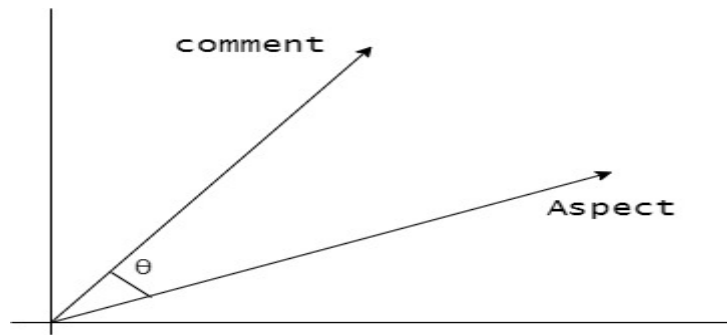
Logistic regression has a place in the family of classifiers known as exponential or log-linear classifiers. It is utilized to decide the yield or result when there is one or more than one independent variable. The yield esteem can be in the shape of 0 or 1 i.e., in binary shape.

Gradient Boosting Algorithm

One can subjectively indicate both the loss function and the base-learner models on demand. In practice, given a few particular loss functions $\Psi(y, f)$ and/or a custom base-learner $h(x, \theta)$, the solution to the parameter estimates can be difficult to get. To deal with this, it was proposed to select a new function $h(x, \theta)$ to be the most parallel to the negative angle $\{g_t(x_i)\}_{i=1}^N$ along with the observed data (Natekin & Knoll, 2013).

3.7. Cosine Similarity

According to (Bhattacharjee et al., 2015), Cosine similarity could be a metric utilized to decide how similar the records are independent of their size. Mathematically, it computes the cosine of the point between two vectors anticipated in a multi-dimensional space. The cosine similarity is profitable since indeed on the off chance that the two similar documents are distance separated by the Euclidean distance (due to the estimate of the document), the chance is they may still be arranged closer together. The smaller the point, the higher the cosine similarity. In this research, we applied a cosine similarity approach for supervised learning in aspect extraction. In the below diagram, each Aspect and comment are represented as vectors, where each vector has the word frequencies, and then, the cosine formula is applied as follows.



$$\text{similarity}(A, B) = \frac{A \cdot B}{\|A\| \times \|B\|} = \frac{\sum_{i=1}^n A_i \times B_i}{\sqrt{\sum_{i=1}^n A_i^2} \times \sqrt{\sum_{i=1}^n B_i^2}}$$

Where the Aspect is represented by A, and the comment is represented by B.

3.8. Evaluation

This research has been evaluated at various stages of the implementation of the Aspect based SA model. After all, the evaluation addressed the findings and recommended a model for detecting Amharic Aspect-based SA.

3.8.1. Evaluation Techniques (Model Evaluation)

Each system should be tested with criteria that compare the number of reviews that are categorized correctly or incorrectly classified. Typically, the comparison is done between the reviews categorized by the proposed system and that of the manually labeled (categorized) reviews. System evaluation is the method of deciding how well a system meets the information needs of its users

According to (*Confusion Matrix for Machine Learning*, n.d.), There are several available performance metrics to measure the performance of a classifier.

Confusion matrix: It is a $N \times N$ matrix that is used to describe a classification model's performance on a collection of the test set for which the actual values are known. The number of occurrences associated with the actual class is represented in each row and each column shows the number of occurrences associated with the predicted class or vice versa. So, this study uses a normalized confusion matrix for three-class models (positive, negative, neutral).

The following four-parameter describe the overall activities to study to determine actual and predicted results:

True Positives (TP) - This is one of the parameters used to correctly estimate positive values, which indicates that yes is the value of the true class and yes is also the value of the expected class.

True Negatives (TN) - used to correctly estimate negative values, which indicate that the value of the true class is no and value of the predicted class is also no (negative comments not classified under positive for positive class and vice versa).

False Positives (FP) – are negative comments that are incorrectly identified as positive for positive class and positive comments that we incorrectly classified as negatives for negative class. (When true class is no and predicted class is yes).

False Negatives (FN) – positive comments that are incorrectly not classified under positive for positive class and negative comments that are incorrectly not classified under negative for negative class.

Generally, evaluating the performance of the classifier is also made in terms of the different confusion matrices (True Positive (TP), False Positive (FP), True Negative (TN), and False Negative (FN)).

Table 1.4: Confusion matrix of Ternary class

		Predictive Value		
		Positive	negative	neutral
Actual Value	Positive	True positive	False Negative	False neutral
	Negative	False Positive	True negative	False neutral
	neutral	False Positive	False Negative	True neutral

In table 1.4, True positive, True negative, and True neutral are the observations that are calculated correctly. We want to decrease false positives, false negatives, and false neutral. Generally, using the standard measurement metrics of precision, recall, F-score, and accuracy, the total positive and negative, and neutral content is usually measured.

Accuracy: is a percentage of correctly predicted observations to the total observations.

$$\text{Accuracy} = \frac{TP+TN}{TP+FP+FN+TN}$$

Recall: is the rate of correctly estimated positive observations for all observations in true class(yes).

$$\text{Recall} = \frac{TP}{TP+FN}$$

Precision: is the ratio of correctly calculated +ve observations to the total predicted +ve observations.

$$\text{Precision} = \frac{TP}{TP+FP}$$

F-Measure: F1 Score is the weighted average of Precision and Recall.

$$\text{F1 Score} = 2 * (\text{Recall} * \text{Precision}) / (\text{Recall} + \text{Precision})$$

To test the efficiency of the classifiers, these evaluation methods are usually used for this research. Because the accurate estimation of the generalization error on the trained model and larger K make low bias and highly accurate, this study used proposed cross-validation using 10-fold.

CHAPTER FOUR

PROPOSED SOLUTION FOR ASPECT BASED SA

4.1. The proposed solution for Aspect Based SA for the Amharic Language

In this chapter, we briefly discuss the most important component and proposed a model for sentiment analysis with opinionated Amharic text. The suggested model has the following components: pre-processing, identifying features, polarity identifications, determining opinions regarding identified frequent and infrequent features, extracting an opinion word. Also, the proposed tools and algorithms are furthermore explored in this chapter.

4.2. The proposed Aspect Based SA Architecture

The proposed architecture is used to detect Aspect-based SA, and the procedure starts by taking Amharic datasets as input and preprocesses them based on the language's existence, such as punctuation removal, normalization, tokenization, stop word removal. After preprocessing is finished, feature extraction is performed using the TF-IDF, Ngram, and POS tagger to extract features. By using the Cosine Similarity Algorithm, we can detect the aspect/feature. Then the supervised learning algorithms are applied to polarity detection to build a sentiment model from the labeled training set. The succeeding task is the evaluation of the built model and the performance status of the model. Any of the unique features which be chosen or created to be applied to the matrix when constructing this model. The ML classifier is then trained with some data and then checked with the remaining data to see how well it performs.

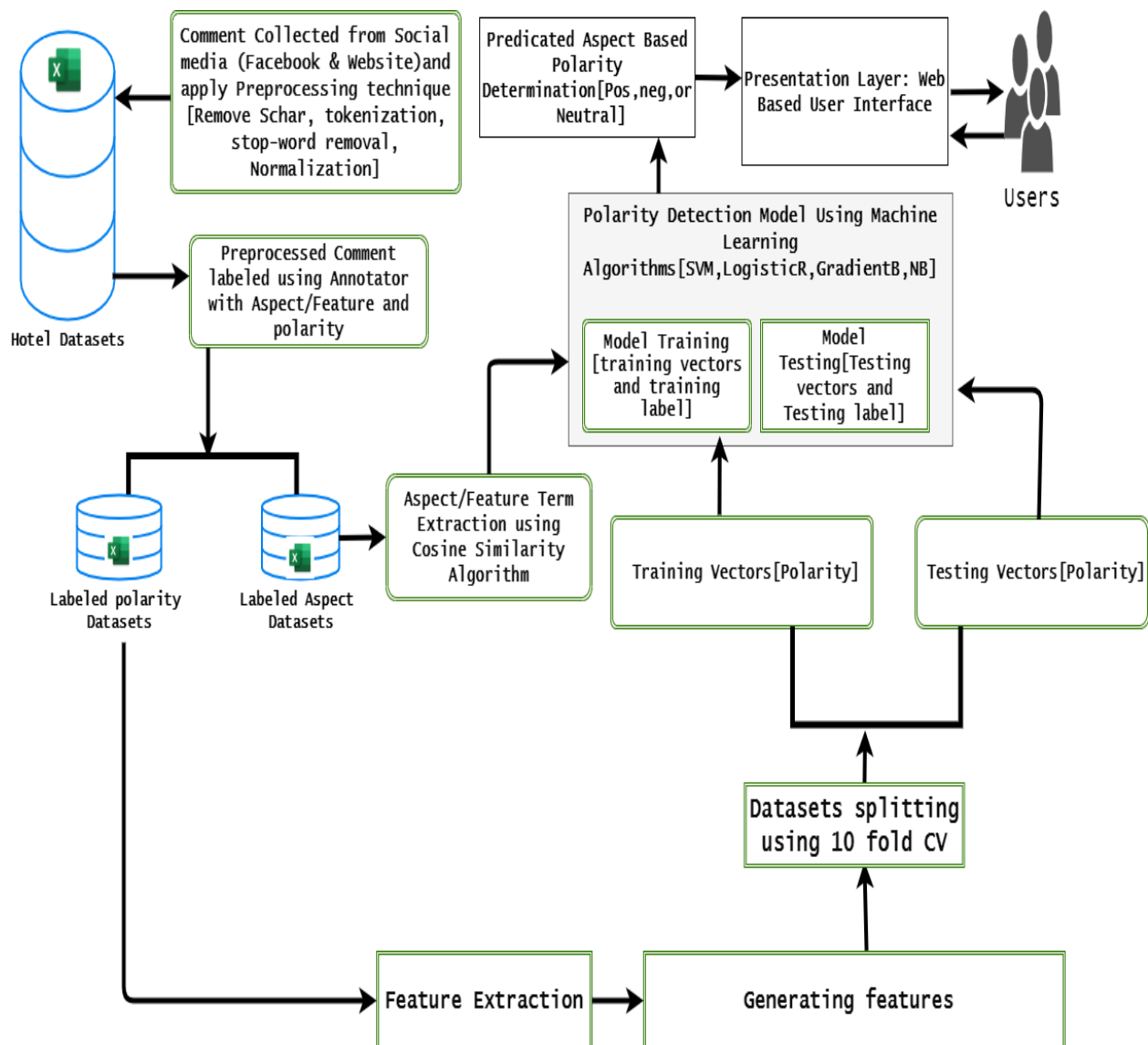


Figure 1.2: Amharic Aspect Based SA Architecture

4.3. Proposed Amharic Text Preprocessing

To preprocessing the Amharic comment for the model's training and testing. This preprocessing is achieved using the Amharic language and simple text preprocessing techniques including punctuation elimination (cleaning), normalization, tokenization, stop word removal.

4.3.1. Data Cleaning

Cleaning data is the main stage in any machine learning model, but it's particularly important in NLP. Without the cleaning process, the dataset is often a jumble of words that the machine can't decipher. Such kinds of data cleaning Irrelevant Character, Punctuations Symbol, and Emojis.

4.3.1.1. Algorithm for Removing Irrelevant Characters

Input: Collected Amharic comments in a Corpus.

Output: Clean text

Start:

- 1.** Read Collected Comments in a Corpus.
- 2.** While (! end of the text in a corpus):
 - 2.1.** If the collected comment contains stopwords like [‘ራሱ’, ’እነዚህ’, ’እና’] then
Remove stopwords
 - 2.2.** If the collected comment contains Symbols like [[<>«»≡:-`~_ /] then
Remove Symbols
 - 2.3.** If the collected comment contains Punctuation like [*::: ፣ ፤ ፥ :- ። ::] then
Remove Punctuation
 - 2.4.** If the collected comment contains special characters like [, '*^@#\$\$%#! ~] then
Remove special characters
 - 2.5.** If the collected comment contains English alphanumeric characters like[^a-zA-Z0-9]
then Remove English alphanumeric characters
 - 2.6.** If the collected comment contains geez_number like [፩ ፪ ፫ ፬ ፭ ፮ ፯ ፰ ፱ ፲ ፳ ፴]
 - then Remove geez_number
 - 2.7.** If the collected comment contains emojis like [😊🙏...] then
Remove emojis
 - 2.8.** If the collected comment contains Unwanted space, then
Remove unwanted space
- 3.** Return clean_text;

Stop

4.3.2. Amharic Normalization Character

In the previous chapter from the Amharic language perspective, normalization was used to address the redundancy issue in the Amharic language, where the same sound character had different character forms. It's a way of changing words into a single type by using a character substitution with a sound that's identical to a common character. The transformation of a character into a single common representation does not affect its meaning.

4.3.2.1. Algorithm for Normalization

Input: Collected Amharic comments in a Corpus.

Output: Normalized comment

Start:

1. Read Collected Comments in a Corpus.
2. While (! end of the text in a corpus):
 - 2.1. If the collected comment contains Characters like ['ገ', 'ኀ', 'ኒ', 'ቃ', 'ኔ', 'ኅ', 'ኖ'] then
Replace Respectively with ['ሀ', 'ሁ', 'ሂ', 'ሀ', 'ሄ', 'ህ', 'ሆ']
 - 2.2. If the collected comment contains Characters like ['ሐ', 'ሑ', 'ሒ', 'ሓ', 'ሔ', 'ሐ', 'ሐ'] then
Replace Respectively with ['ሀ', 'ሁ', 'ሂ', 'ሀ', 'ሄ', 'ህ', 'ሆ']
 - 2.3. If the collected comment contains Characters like ['ጸ', 'ጹ', 'ጺ', 'ጻ', 'ጼ', 'ጸ', 'ጸ'] then
Replace Respectively with ['ፀ', 'ፀ', 'ፉ', 'ፉ', 'ፊ', 'ፀ', 'ፀ']
 - 2.4. If the collected comment contains Characters like ['ሠ', 'ሡ', 'ሢ', 'ሣ', 'ሤ', 'ሥ', 'ሦ'] then
Replace Respectively with ['ሰ', 'ሰ', 'ሱ', 'ሱ', 'ሲ', 'ሱ', 'ሰ']
 - 2.5. If the collected comment contains Characters like ['ዐ', 'ዑ', 'ዒ', 'ዓ', 'ዤ', 'ዐ', 'ዐ'] then
Replace Respectively with ['አ', 'አ', 'አ', 'አ', 'አ', 'አ', 'አ']
3. Return replaced text;

Stop

4.3.3. Tokenization

After the cleaning and normalization steps, the Tokenization process is used to break down feedback text into individual words or tokens using spaces between words or punctuation marks. This is relevant since the meaning of the text is often determined by the relationships between

words in that text, and it helps feature extraction methods in extracting suitable features from the dataset.

4.3.4. Stop-Words Removal

Stop words are terms that don't add much to a sentence's meaning. Without jeopardizing the sentence's meaning, they can be safely ignored.

Input: list of words in the corpus

Output: list of words doesn't have a stop word

Variable: ListOfStopwords[], ListOfRemovedStopWord[]

Start:

1. Read Collected Comments in a Corpus.
2. WHILE (list of words in the corpus):
 - 2.1. IF list of words in the corpus not in ListOfStopwords: THEN
 - 2.2. Append the word in ListOfRemovedStopWord
3. Return replaced ListOfRemovedStopWord
4. **Stop**

4.3.5. Part of Speech (POS) tagging:

Because there is no publicly available Amharic POS tagger for use, we use manual tagging with the assistance of language professionals for each opinion to identify and extract aspects and opinions from Amharic text. Using this tagger, however, we were able to classify and label text. To check morphological structure, we consulted a language professional.

POS tagged words are commonly used as features for supervised learning in sentiment analysis (Prabowo & Thelwall, 2013). Adjectives (JJ), adverbs (RB), verbs (VB), and nouns (N) were among the features chosen (NN). When treated as features in feature spaces, the combination of adjectives, adverbs, and nouns performed better than individuals, as discussed in (Meng, 2012). Because there is no publically (automatic) POS tagger for the Amharic language, we use nouns or noun phrases for feature extraction as aspects of manually identifying nouns from opinion in this study.

4.4. Proposed Feature Extraction Methods

Techniques for minimizing redundant features and dimensionality were applied to feature extraction methods based on state-of-the-art text mining. The methods include taking a subset of specific features that can be utilized in the modeling of detection problems to aid in detecting from a dataset. This method takes an input of data preprocessed and tokenized dataset words and performs feature extraction. At that point, extracted features were utilized for training the models and also to predict the class of comments as positive, negative, and neutral.

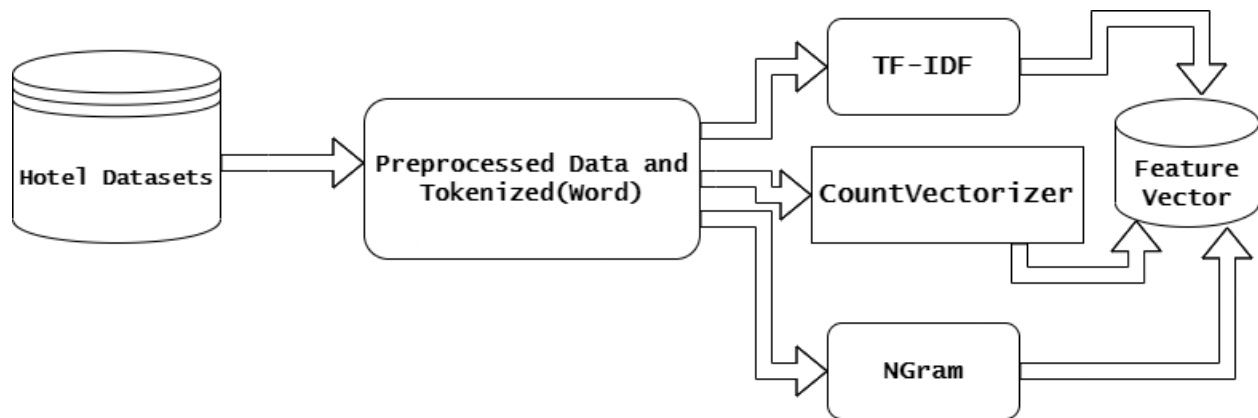


Figure 1.3: Figure Feature Extraction Flow Diagram

Since N-gram, TF-IDF, and POS tagger feature extraction are better and more common feature extraction methods for text-based classification problems, they were used in this research. For modeling Aspect based SA for Amharic, the following feature extraction methods are used.

TF-IDF

measure of a word's significance in a dataset, and it rises in proportion to the number of times the word appears in the document. It computes weight for ranking, the features using term presence and term frequency, and Term Frequency and Inverse document frequency.

N-Gram Feature Extraction

This research proposed a word N-gram feature extraction strategy tested on a distinctive value of N that ranges from one to three. Where N is the number of words utilized within the likelihood sequences. An n-gram of two words is called bigram (2-gram). They include extraction performed by utilizing unigram (1), bigram (2), trigram (3), and combination n-grams. The execution of the Ngram needs an appropriate choice of the 'n' value. Moreover, it gives a distinctive highlight demonstrate to prepare and to comparing n-gram highlights to one another.

Count Vectorizer

is utilized to change a group of text-based documents to a vector token count. It moreover enables the pre-processing of text data before producing the vector representation. This usefulness makes it a highly flexible highlight representation module for text.

Part of Speech (POS) tagging

POS tagged terms are often used as supervised learning features in sentiment analysis. Adjectives (JJ), adverbs (RB), verbs (VB), and nouns (N) were among the features chosen (NN). because there is no automated speech tagger for the Amharic language, we use nouns or noun phrases for feature extraction as aspects of manually distinguishing nouns from sentiments.

CHAPTER FIVE

IMPLEMENTATION AND EXPERIMENTATION

5.1. Implementation of the proposed solution

This Research follows an experimental Based approach to decide the finest combination of the supervised machine learning algorithms and features extraction for the last models' comparison.

5.2. Implementation Environment

This study utilized several development tools and bundles to implement the proposed solution and the python programming language used for implementing and experimenting with each proposed solution from the data preprocessing to the model building stage. Table 1.5 shows the list of tools and python packages with their version and description which is used in this study.

Table 1.5: Some of the Tools and Python Package

Development Tools		
Tools	Version	Description
Anaconda Navigator	4.9.2	A desktop user interface included in Anaconda® conveyance permits you to dispatch applications and effectively oversee conda bundles, environments, and channels without utilizing command-line commands.
Colaboratory(Colab)		Permits anyone to write and execute arbitrary python code and utilized in particularly well-suited to machine learning, information examination, and instruction
Jupyter notebooks	6.0.1	The Jupyter Note pad is a non-proprietary software that permits you to make and share reports that include live code, conditions, visualizations, and story content.
Python	3.7.4	The Python programming language best fits machine learning due to its independent platform's, vast features, applicability, simplicity. and its popularity in the programming community.

Facepager	4.2.14	Enables you to collect public data from platforms on the social web (such as Facebook, Twitter, or YouTube) which these platforms make available through program interfaces (APIs).
diagrams.net	14.6.5	Building diagramming applications, and the world's most broadly utilized browser-based end-user diagramming program
Microsoft Office	2016	helps simplify basic office tasks and improve work productivity. Each application is designed to address specific tasks, such as word accessing, data management, making presentations, and organizing emails.
Phyton Package		
Package Name	Version	Description
Scikit-learn	0.22.2	Most important library for machine learning in Python. The sklearn library include different effective tools for machine learning approach and factual modeling including regression, classification.
pandas	0.25.1	Which is used for data cleaning and analysis. It has features that are utilized for investigating, cleaning, changing, and visualizing from information
NumPy	1.16.5	contains a multi-sized array and matrix data structures. It can be utilized to perform several mathematical operations on arrays
NLTK	3.4.5	Python programs that work with human dialect information for applying in factual characteristic dialect preparing

Flask	1.1.1	a micro web framework is written in Python? It can create a REST API that allows you to send data, and receive a prediction as a response.
--------------	--------------	--

5.3. Data Preprocessing Implementation

To implement the Amharic Data preprocessing method, the study utilized python programming language and modules. To perform the entire implementation of the methods, we perform the following common activities, which are importing important libraries bundles and load the dataset file to utilizing. The loaded dataset using panda is used throughout the whole implementation of this study. The Amharic Data preprocessing method uses Python RegEx (regular expression) to create a search pattern and NLTK modules.

Loading dataset and show some datapoint in dataset(Read the Dataframe)

```
df=pd.read_excel('C:\\Users\\yewe\\Final Msc Thesis Implementasion ABSA\\new dataset\\polarity.xlsx')
# import POS tag
data = pd.read_excel('C:\\Users\\yewe\\Final Msc Thesis Implementasion ABSA\\new dataset\\postag.xlsx')
df.head()
my_labels = ['positive', 'negative', 'neutral']
```

Figure 1.4: Python Code for Load the Dataset

5.3.1. Implementation of Removing Irrelevant Characters

Before applying machine learning algorithms, data must be cleaned properly. we implement a `clean_text()` method that performs removal and replaces the punctuation mark, special character, alphanumeric, symbol, emojis, etc. a method is written using the Python `re` module. The code for cleaning irrelevant characters is shown in Appendix J.

5.3.2. Implementation of Normalization of Amharic Character

Normalization is a way of changing words into a single type by using a character substitution with a sound that's identical to a common character, we implement this method using Python with the `re` module. The code for normalizing the Amharic characters is shown in Appendix K.

5.3.3. Implementation of Tokenization

The tokenization process is used to break down feedback text into individual words or tokens using spaces between words or punctuation marks. We used a python module NLTK.

Tokenized process

```
▶ def clean_text(x):  
    words=[]  
    for sent in nltk.sent_tokenize(x):  
        for word in nltk.word_tokenize(sent):  
            if len(word)<2:  
                continue  
            words.append(word)  
    return words
```

Figure 1.5: Sample Code for Tokenization

5.3.4. Implementation Stop-Words Removal

Without jeopardizing the sentence's meaning, they can be safely ignored.

Implementing of StopWord

```
▶ STOPWORDS = stopwords.words(r'C:\Users\yewe\AppData\Roaming\nltk_data\corpora\stopwords\stopwordyo.txt')  
def clean_text(x):  
    x=" ".join([word for word in str(x).split() if word not in STOPWORDS])  
df['comment'] = df['comment'].apply(lambda x: clean_text(x))  
df.head()
```

Figure 1.6: Sample Code for Stop word

5.4. Feature Extraction Implementation

This method takes an input of data preprocessed and tokenized dataset words and performs feature extraction. Because machine learning models work on numbers (vectors) instead of words to train models.

5.4.1. Implementation of Count Vectorize and N-gram

To work on count vectorizer this study used a Count Vectorizer class of sci-kit learns package. This class is used to convert a given text into a vectorizer based on the frequency (count) of each word that occurs in the entire text.

Word embedding and Transformation

```
▶ # converting in to ML format (machine understandable using TFIDF)

tfidf = CountVectorizer(max_features=5000, ngram_range=(1,3),tokenizer=clean_text,analyzer='word')
X = df['comment']
y = df['polarity']

X = tfidf.fit_transform(X)
X
```

Figure 1.7: Sample Code for Count vectorizer and Ngram

5.5. Machine Learning Models Implementations

Python sci-kit learning library package is used to build machine learning models. The following are the python package that is used in model building.

Importing nessary packages

```
▶ import re
import math
from collections import Counter
import pickle
import random
import numpy as np
import pandas as pd
import nltk
from nltk.corpus import stopwords
from sklearn.feature_extraction.text import TfidfVectorizer,CountVectorizer
from nltk import word_tokenize
from nltk.tokenize import word_tokenize
from sklearn.metrics import accuracy_score
import matplotlib.pyplot as plt
import seaborn as sns
import string
%matplotlib inline
from sklearn.model_selection import train_test_split
from sklearn.metrics import classification_report,confusion_matrix
import warnings
warnings.filterwarnings("ignore")
from sklearn.feature_extraction.text import TfidfTransformer
from sklearn.naive_bayes import MultinomialNB
from sklearn.linear_model import LogisticRegression
from sklearn.svm import LinearSVC
from sklearn.ensemble import GradientBoostingClassifier
from sklearn.model_selection import cross_val_score
from sklearn.metrics import plot_confusion_matrix
```

Figure 1.8: Importing the necessary Package

The study utilized the LinearSVC (Support Vector Classifier) classifier of the sklearn bundle to build the SVM model. LinearSVC is used to fit the data you give, returning a "best fit" hyperplane

Train using Support Vector Machine Learning Algorithm

```
▶ #using Supervised ML[SVM]
sm = LinearSVC(multi_class='ovr',penalty='l2',loss='squared_hinge',dual=True,C=0.01,class_weight='balanced',verbose=0,
               max_iter=10000)
sm.fit(X_train, y_train)
y_pred = sm.predict(X_test)
filename = 'svm.sav'
pickle.dump(sm, open(filename, 'wb'))
```

Figure 1.9: Instantiating SVM and Fit the Model

The multinomial Naive Bayes classifier is appropriate for classification with discrete, so this study uses a sklearn MultinomialNB() classifier.

Train using Naive bayese Machine Learning Algorithm

```
▶ # using Supervised ML[Naive bayese]
nb = MultinomialNB(alpha=1.0,fit_prior=True,class_prior=None)#instantiate naive bayes classifier
nb.fit(X_train, y_train)#to train NB
y_pred = nb.predict(X_test)
filename = 'NaiveBayi.sav'
pickle.dump(nb, open(filename, 'wb'))
```

Figure 1.10: Instantiating Naïve Bayes and Fit the Model

The study utilized the LogisticRegression() is utilized to describe data and to clarify the relationship between one dependent binary variable

Train using Logistic Regration Machine Learning Algorithm

```
▶ # Use Supervised ML algo [Logistic Regration]
lr = LogisticRegression(penalty='l2',dual=False,C=0.01,class_weight='balanced',max_iter=10000,multi_class='ovr',
                        verbose=0,solver='lbfgs')
lr.fit(X_train, y_train)
y_pred = lr.predict(X_test)
filename = 'LogisticRegr.sav'
pickle.dump(lr, open(filename, 'wb'))
```

Figure 1.11: Instantiating Logistic Regression and Fit the Model

The study utilized the GradientBoostingClassifier() since it minimizes the overall prediction blunder and advantage from regularization strategies that penalize diverse parts of the algorithm.

```
❏ # Use Supervised ML algo [Gradient boosting]
gb = GradientBoostingClassifier(loss='deviance',n_estimators=100,max_features='auto',verbose=0)
gb.fit(X_train, y_train)
y_pred = gb.predict(X_test)
filename = 'GradientBo.sav'
pickle.dump(gb, open(filename, 'wb'))
```

Figure 1.12: Instantiating Gradient Boosting and Fit the Model

5.6. Cosine Similarity

For the good extracting of feature/aspect, the study utilized the cosine similarity. After extracting aspect/feature using cosine similarity then the machine learning takes and predicts aspects-based polarity determination.

```
❏ def GetInput():
    x=input("please enter your comment")
    data['input']=x
    # preprocessed unseen text and input
    x=clean_text(x)

    vec = tfidf.transform([x])
    vec.shape
    y=gb.predict(vec)
    #readData()
    data['vector1']=data['comment'].apply(lambda x: text_to_vector(x))
    data['vector2']=data['input'].apply(lambda x: text_to_vector(x))
    data['Distance']=data.apply(lambda x: get_cosine(x['vector1'],x['vector2']),axis=1)
    column = data["Distance"]
    max_index = column.idxmax()
    print("The aspect is "+ data['aspect'].values[max_index]+ " The polarity is "+y)
    return y
```

Figure 1.13: Sample code for Cosine Similarity in aspect Extraction

5.7. Model Testing and Evaluation

When we evaluate the model, performance is based on an error metric to decide the accuracy of the model. This method, however, isn't very reliable as the exactness obtained for one test set can be very different from the accuracy gotten for a different test set, K-fold Cross-Validation (CV) gives a solution to this issue by partitioning the data into 10 folds and guaranteeing that each fold is utilized as testing set at a few points.

```
► #Kfold validation
scores = cross_val_score(gb, X, y, cv=10, scoring='accuracy')
print(scores)
kgb=scores.mean()
print("The Average Kfold Score",kgb)
```

Figure 1.14: Performing 10-fold CV on GB Models

The study utilized `classification_report()`, `accuracy_score()`, and `confusion_matrix()` methods, which are utilized to assess the performance of the built models utilizing predict the value of 10-fold CV for the complete dataset. These strategies allow a result of evaluation metrics such as accuracy, precision, recall, and F1-score value of the model under testing. At long last, A model for recognizing Aspect Based SA is implemented using the chosen best model then deployed using flask web services for it to be able to acknowledge new input comments and then return prediction the results.

5.8. User Interface Design

The proposed system was mainly developed to make computers communicate with users through the Amharic language. The technique is used to create a web application that is integrated with the model. To do this we used the flask framework. And there is a technique called a bootstrap to make the app easily compatible with different window size devices.

5.8.1. Flask Frameworks

Flask is a python web framework that's utilized to creating web applications. The proposed framework was created by utilizing this framework. Flask has numerous preferences. From those advantages, a few are; flexibility and simplicity. For this reason, we utilized the Flask framework.

When we use the flask framework, we used two directories, which are known as a template and static. In the static folder, we put the files that are used to make attractive the system. those files are an image, CSS files, JavaScript files. After we create those, all files the next step is loading the trained model file that is saved in our hard disk and make it to predict the users' request and finally give an appropriate response.

5.8.2. Bootstrap

Bootstrap is a Cascading Style sheet (CSS) system that's utilized for creating responsive web applications. It is an open-source framework utilized at the front end of application development. It contains diverse HTML layouts that are utilized to create a graphical user interface.

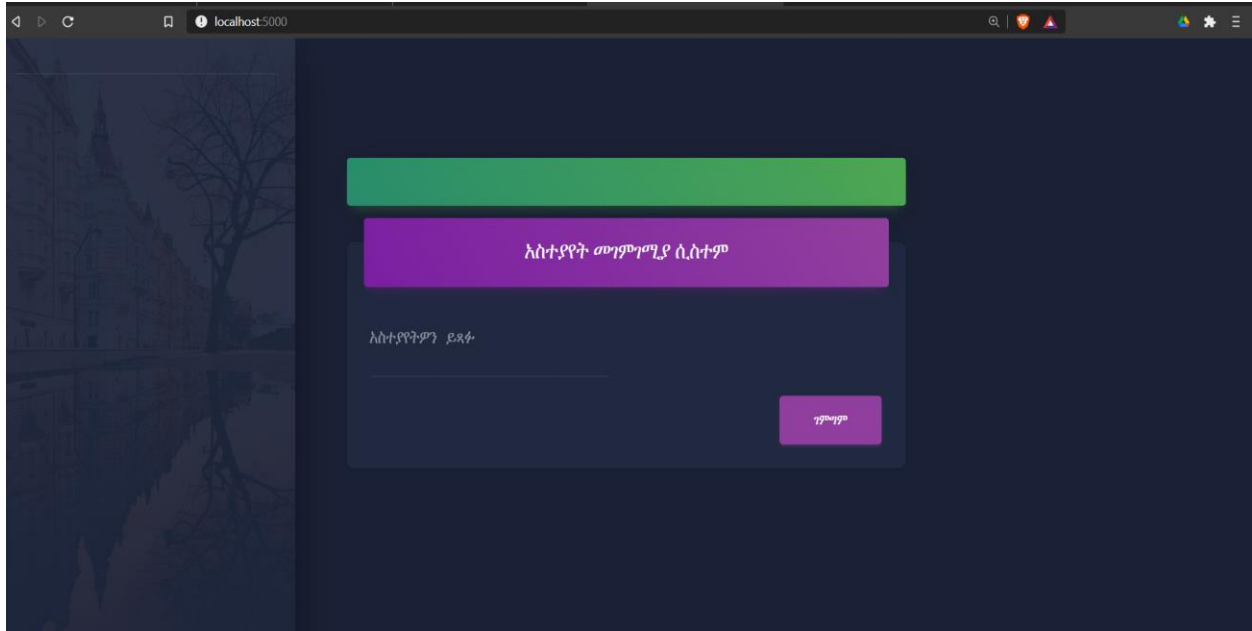


Figure 1.15: User Interface Design

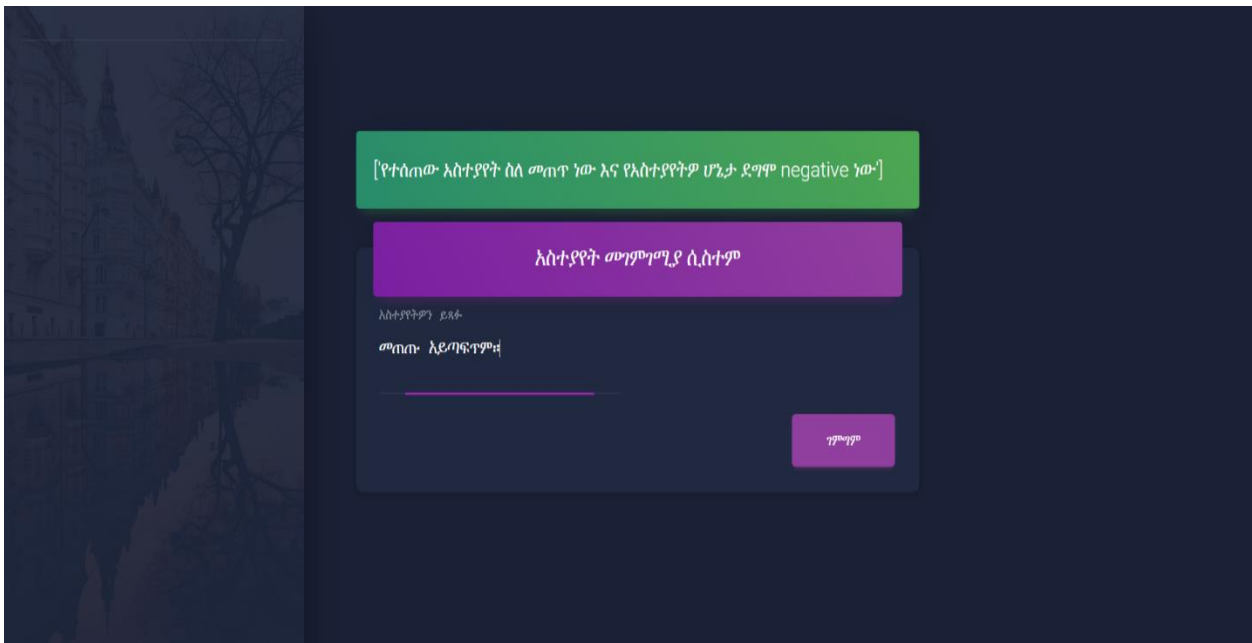


Figure 1.16: Positive Request and Response

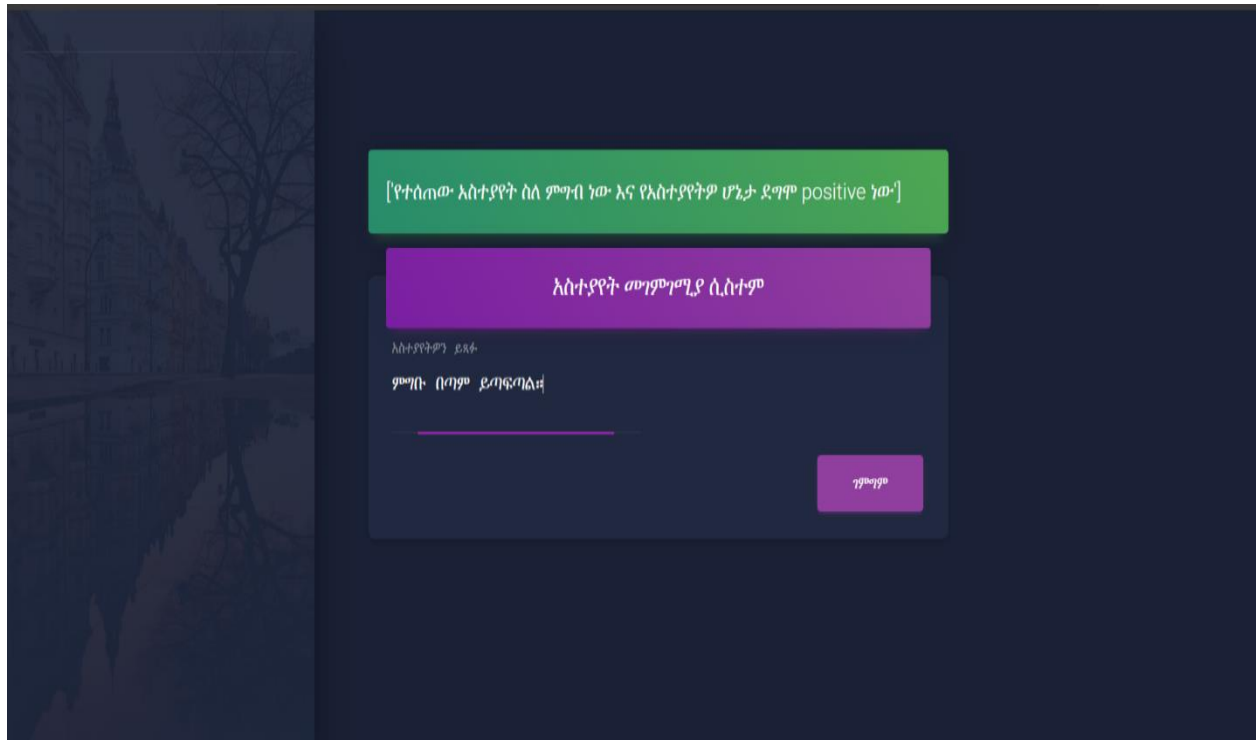


Figure 1.17: Negative Request and Response

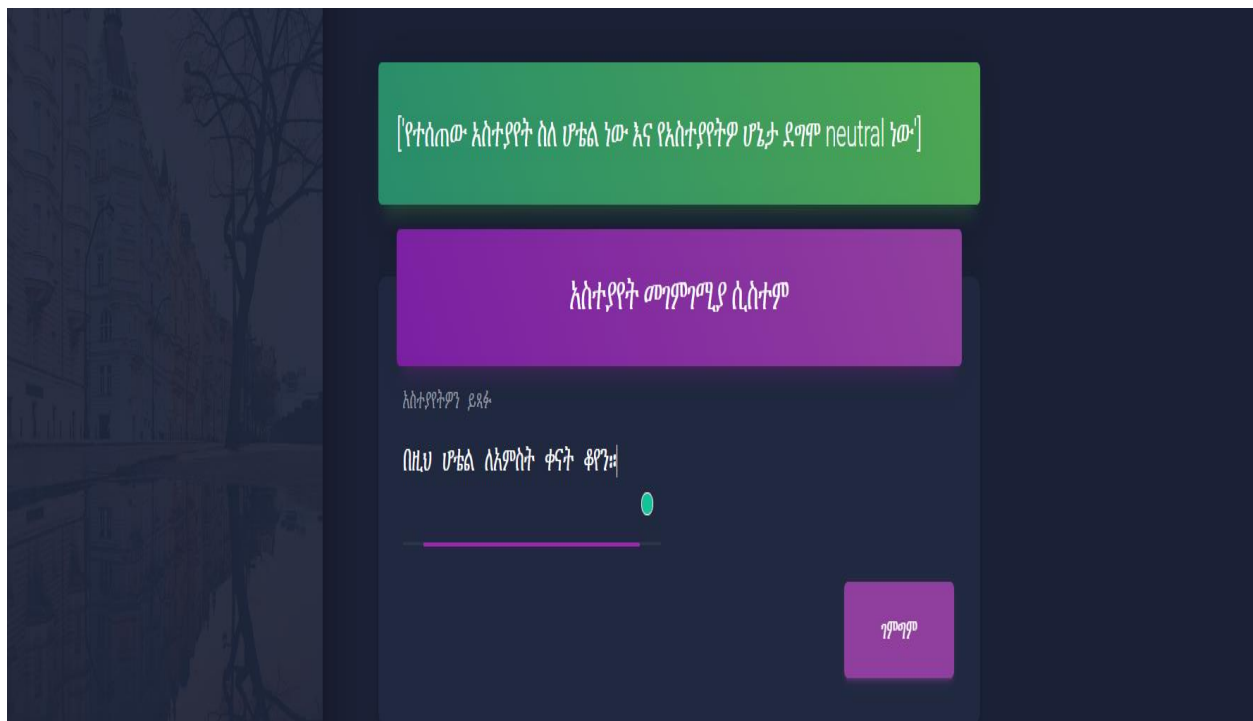


Figure 1.18: Neutral Request and Response

CHAPTER SIX

RESULTS AND DISCUSSIONS

6.1. Discussion of Results

Within the case of extracting important features and deciding the polarity of opinion words. The result of the experiments of the proposed solution for Aspect Based Sentimental Analysis using machine learning approaches is examined.

6.2. Models Evaluation Results

The researchers are comparing the machine learning algorithms (that are discussed in chapter 3) to choose a more fitting model for an Aspect Based SA with the same data set and categories for feature selection and polarity determination. Testing these trained models is performed using 10-fold cross-validation. This method randomly split the dataset into ten equal size sets.

6.2.1. Ternary Classification Models Evaluation Results

Ternary Classification models are built using the three-class dataset that is converted from the annotated dataset that means the dependent variables are positive, negative, or neutral. The trained models are tested using the 10-fold CV. The result of each test accuracy score for SVM, NB, LR, and GB models based on extracted feature vectors are presented in tables separately below.

Table 1.6: SVM Models Average Accuracy for Each Feature

Feature Extraction	10 Folds CV Accuracy (%) for SVM model										Average Accuracy (%)
	1 st	2 nd	3 rd	4 th	5 th	6 th	7 th	8 th	9 th	10 th	
TF-IDF	53.5	67.6	79.8	77.9	87.3	75.9	67.9	78.7	76.4	65.1	73.3
Count Vectorizer	54.9	74.6	82.6	83.1	90.6	87.3	82.5	85.4	85.4	71.7	79.8
TF-IDF + Unigram	53.5	67.6	79.8	77.9	87.3	75.9	67.9	78.8	76.4	65.1	73.0
TF-IDF + Bigram	62.4	81.2	89.2	85.4	93.9	94.8	82.1	92.9	87.3	80.7	84.9

TF-IDF + Trigram	61.9	78.9	86.9	87.8	91.9	94.8	86.8	95.8	90	87.7	86.3
TF-IDF + Unigram+Bigram	60	78.4	85.9	86.4	93.9	91.9	81.1	90.6	84.9	76.4	82.9
TF-IDF + Bigram+Trigram	62.9	80.8	88.3	88.3	94.3	95.8	84.9	94.3	87.7	87.3	86.5
TF-IDF + Unigram+Trigram	62.9	80.3	88.3	87.3	94.3	94.8	83.9	93.4	86.8	82.5	85.5
Countvectorizer + unigram	54.9	74.6	82.6	83.1	90.6	87.3	82.5	85.4	85.4	71.7	79.8
Countvectorizer + Bigram	69.0	89.2	92.0	91.5	93.7	96.7	95.8	94.8	97.2	87.7	90.4
Countvectorizer + Trigram	68.5	86.9	92.5	91.5	91.9	96.2	93.9	97.2	96.7	91.5	90.7
Countvectorizer + Unigram+Bigram	69.9	88.7	90.1	91.1	94.3	97.2	95.3	95.8	97.6	90.6	91.1
Countvectorizer + Bigram+Trigram	71.4	89.7	93.9	92.9	92.9	96.2	94.8	96.7	97.6	92.5	91.9
Countvectorizer + Unigram+Trigram	72.3	90.1	92.0	92.9	94.3	98.1	95.3	96.7	98.1	92.5	92.2

The outcome in Table 1.6 shows the cross-validation accuracy scores for the Support vector machine model based on features with corresponding each fold test. The lower accuracy of 73.0% (0.730) outcome was recorded on the TF-IDF+unigram feature model and the higher accuracy of 92.2% (0.922) resulted using countvectorizer+Unigram+Trigram feature model.

Table 1.7: NB Models Average Accuracy Scores for Each Feature

Feature Extraction	10 Folds CV Accuracy (%) for NB model										Average Accuracy (%)
	1 st	2 nd	3 rd	4 th	5 th	6 th	7 th	8 th	9 th	10 th	
TF-IDF	44.6	60.6	65.7	71.4	85.4	60.8	61.3	73.6	66.9	56.6	64.7
Count Vectorizer	44.1	61.9	67.1	71.8	86.3	64.6	58	73.1	68.4	59.4	65.5
TF-IDF + unigram	44.6	60.6	65.7	71.4	85.4	60.8	61.3	73.6	66.9	56.6	64.7
TF-IDF + bigram	64.3	84.0	86.4	84.0	89.2	81.6	71.7	88.2	81.1	73.1	80.4
TF-IDF + Trigram	67.6	82.1	84.5	83.6	90.0	83.5	72.6	91.0	83.5	75.9	81.5
TF-IDF + Unigram+Bigram	61.9	79.8	84.5	83.1	91.0	80.2	70.8	87.7	81.1	70.8	79.1
TF-IDF + Bigram+Trigram	66.2	83.6	85.9	84.0	89.6	83.0	72.2	90.6	82.5	75.9	81.4
TF-IDF + Unigram+Trigram	64.3	81.2	84.9	83.1	90.0	83.0	71.2	89.2	82.5	75.5	80.5
Countvectorizer + unigram	44.1	61.9	67.1	71.8	86.3	64.6	58.0	73.1	68.4	59.4	65.5
Countvectorizer + Bigram	61.9	84.5	86.4	82.2	89.6	81.6	71.7	86.8	82.5	74.5	80.2
Countvectorizer + Trigram	66.2	84.0	84.9	81.7	89.2	83.0	71.7	89.6	85.4	76.9	81.3
Countvectorizer + Unigram+Bigram	58.2	78.9	83.6	81.2	89.6	80.6	71.7	86.3	81.1	73.1	78.4
Countvectorizer + Bigram+Trigram	63.4	84.0	85.9	83.6	88.7	81.6	71.7	88.2	84.4	75.5	80.7

Countvectorizer + Unigram+Trigram	62.9	81.2	85.4	83.1	89.6	80.1	71.2	87.3	83.0	76.4	80.0
---	------	------	------	------	------	------	------	------	------	------	------

The outcome in Table 1.7 shows the cross-validation accuracy scores for the Naive Bayes classifier based on features with corresponding each fold test. The lower accuracy of 64.7% (0.647) outcome was recorded on the TFIDF and TFIDF+unigram feature model, and the higher accuracy 81.5% (0.815) resulted using the TFIDF+Trigram feature model.

Table 1.8: LR Models Average Accuracy Scores for Each Feature

Feature Extraction	10 Folds CV Accuracy (%) for LR model										Average Accuracy (%)
	1 st	2 nd	3 rd	4 th	5 th	6 th	7 th	8 th	9 th	10 th	
TF-IDF	49.3	62.4	76	75.6	86.8	70.3	61.3	76.9	69.3	64.2	69.2
Count Vectorizer	51.6	65.3	78.9	77.9	87.3	74.5	70.8	80.1	75.5	60.8	72.3
TFIDF + Unigram	49.3	62.4	76.1	75.6	86.8	70.3	61.3	76.9	69.3	64.2	69.2
TFIDF + Bigram	61.0	77.5	87.3	84.5	90.6	89.2	79.7	90.0	83.9	78.3	82.2
TFIDF + Trigram	58.7	76.1	85.4	84.5	91.9	91.9	79.2	92.9	84.4	85.8	83.1
TFIDF+ Unigram+Bigram	56.3	74.2	83.6	84.0	91.0	86.8	75.9	88.2	81.6	72.3	79.4
TFIDF + Bigram+Trigram	61.0	76.1	86.4	84.9	92.9	89.6	79.7	91.0	84.9	84.9	83.2
TFIDF+ Unigram+Trigram	58.2	75.6	85.9	86.4	92.9	89.2	78.8	91.0	83.9	79.7	82.2
Countvectorizer+ unigram	51.6	65.3	78.9	77.9	87.3	74.5	70.8	80.2	75.5	60.1	72.3
Countvectorizer+ Bigram	63.4	79.8	84.0	86.4	91.9	96.2	86.3	89.2	91.0	79.2	84.7
Countvectorizer+ Trigram	62.9	79.3	86.4	84.5	90.6	94.8	89.2	94.3	90.0	86.8	85.9

Countvectorizer+ Unigram+Bigram	62.9	78.4	87.3	85.9	94.3	93.9	86.3	90.6	90.0	77.4	84.7
Countvectorizer+ Bigram+Trigram	67.1	82.2	88.7	88.3	93.4	97.2	91.5	94.3	92.9	83.9	87.9
Countvectorizer+ Unigram+Trigram	64.8	83.1	87.3	89.7	93.3	96.7	89.6	91.9	92.9	83.0	87.3

The outcome in Table 1.8 shows the cross-validation accuracy scores for the Logistic Regression classifier based on features with corresponding each fold test. The lower accuracy of 69.2% (0.692) outcome was recorded on the TF-IDF and TF-IDF+unigram feature model and the higher accuracy of 87.9% (0.879) resulted using the Countvectorizer+Bigram+Trigram feature model.

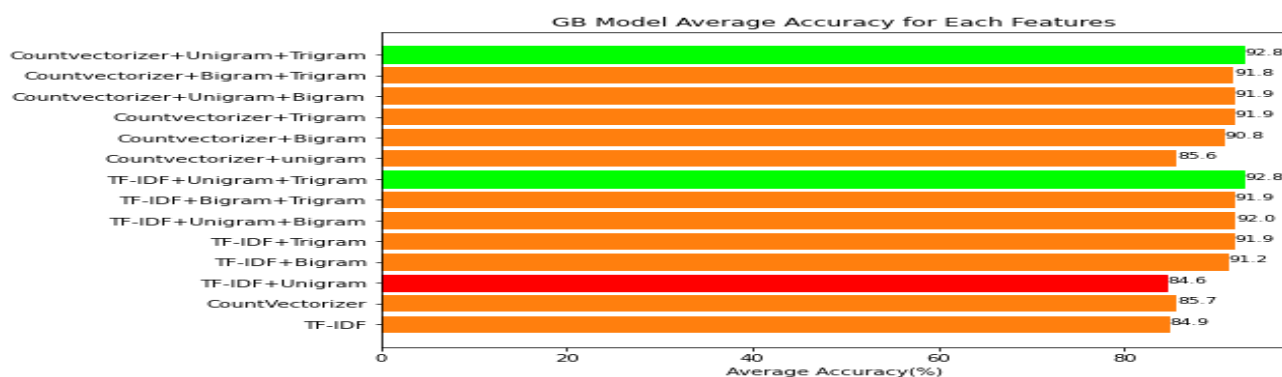
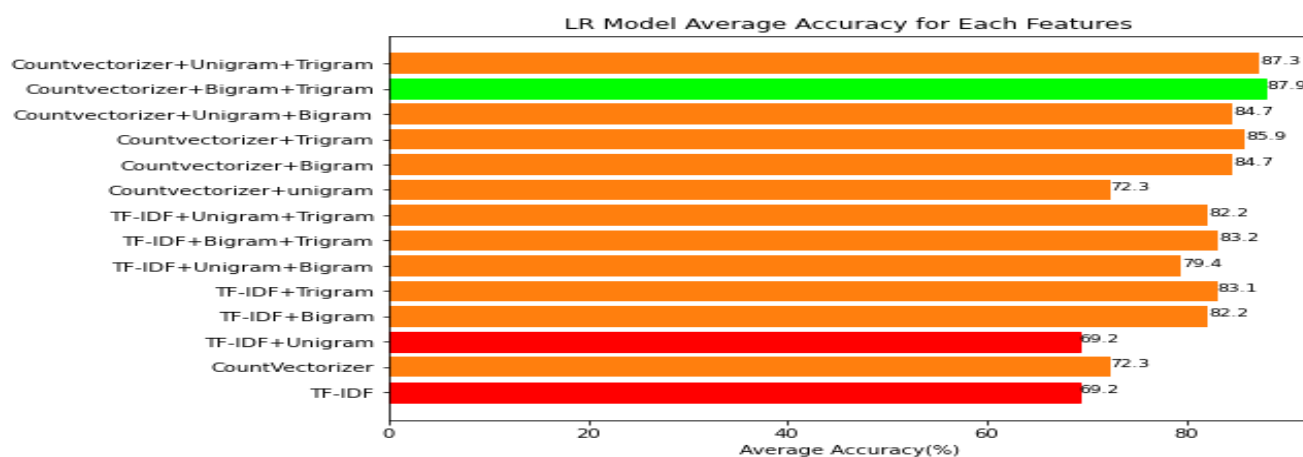
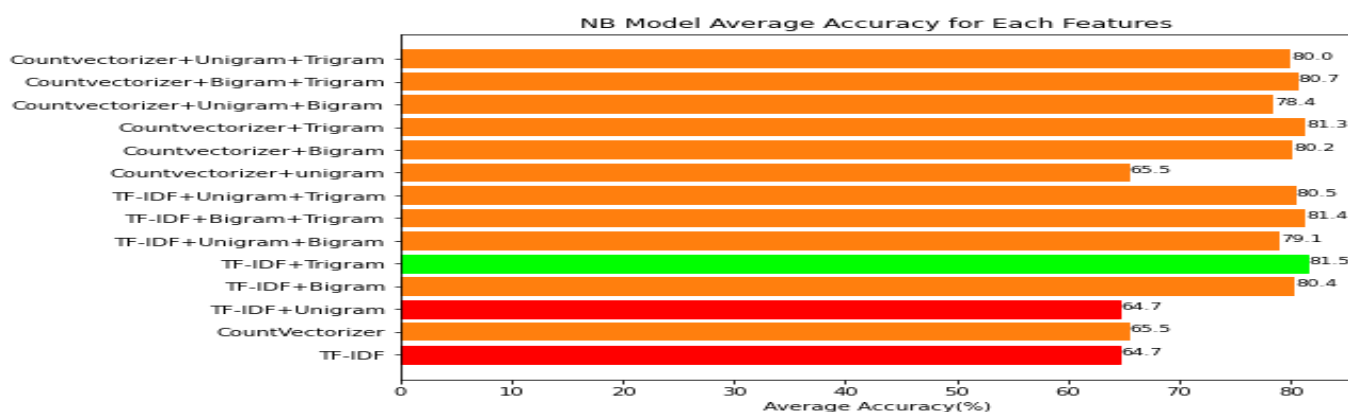
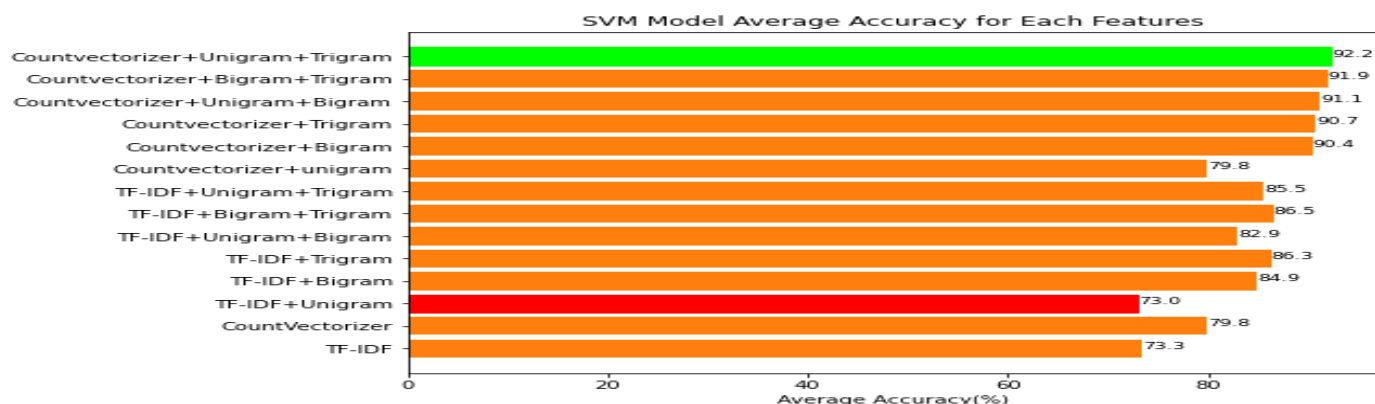
Table 1.9: GB Models Average Accuracy Scores for Each Feature

Feature Extraction	10 Folds CV Accuracy (%) for GB model										Average Accuracy (%)
	1 st	2 nd	3 rd	4 th	5 th	6 th	7 th	8 th	9 th	10 th	
TF-IDF	59.6	79.3	88.7	83.6	92.4	96.2	90.1	90.1	92.9	75	84.9
Count Vectorizer	61.0	80.8	88.7	83.1	93.4	94.3	89.2	89.6	94.8	81.6	85.7
TFIDF + Unigram	59.6	79.3	89.2	83.6	92.5	96.2	89.2	90.6	92.9	73.1	84.6
TFIDF + Bigram	70.9	89.7	93.	91.1	94.3	97.2	93.4	96.2	94.3	91.0	91.2
TFIDF + Trigram	72.3	88.7	96.7	92.5	91.9	97.6	91.5	97.2	98.1	92.9	91.9
TFIDF+ Unigram+Bigram	69.5	89.2	92.9	91.5	94.3	98.6	94.3	97.2	97.6	94.8	92.0
TFIDF + Bigram+Trigram	71.4	90.1	95.3	91.5	94.8	97.2	93.3	96.2	95.8	92.9	91.9
TFIDF+ Unigram+Trigram	72.3	90.6	95.3	92.5	94.8	98.6	93.9	96.7	97.6	95.8	92.8
Countvectorizer+ Unigram	61.0	80.7	89.2	84.9	93.4	94.3	89.2	89.6	94.8	78.8	85.6

Countvectorizer+ Bigram	69.9	87.3	94.8	91.1	93.9	97.2	93.9	95.3	96.2	88.2	90.8
Countvectorizer+ Trigram	71.3	87.8	96.2	93.9	92.9	97.6	93.9	96.7	97.6	91.5	91.9
Countvectorizer+ Unigram+Bigram	70.4	88.7	95.8	91.5	94.8	98.6	94.8	95.8	98.6	90.0	91.9
Countvectorizer+ Bigram+Trigram	70.9	89.2	95.8	92.3	91.0	97.6	93.4	96.7	97.6	92.9	91.8
Countvectorizer+ Unigram+Trigram	72.3	88.3	95.8	93.4	95.8	98.6	94.8	96.7	98.6	93.3	92.8

The outcome in Table 1.9 shows the cross-validation accuracy scores for the Gradient Boosting classifier based on features with corresponding each fold test. The lower accuracy of 84.6% (0.846) outcome was recorded on the TFIDF +Unigram feature model and the higher accuracy of 92.8% (0.928) resulted using the TFIDF+Unigram+Trigram and Countvectorizer +Unigram+Trigram feature model.

When we generalized the above four tables, 92.8% of the accuracy of the Gradient Boosting model based on the TFIDF+Unigram+Trigram and Countvectorizer+ Unigram+Trigram features greater than that of the SVM, NB, and LR models. The accuracy of the SVM models based on the Countvectorizer+Unigram+Trigram features is greater than that of the LR and NB models. The below bar chart shows four Machine learning Models Average Accuracy Scores for Each Feature, and GB based on Countvectorizer+ Unigram+Trigram and TFIDF+Unigram+Trigram makes good accuracy than others.



The study also uses other model performance evaluation metrics, as discussed in chapter three. These metrics are Precision(P), Recall(R), and F-Measure(F1). Table 1.10 appears the results of the evaluation metric for each model based on features extracted.

Table 1.10: Classification Performance Result

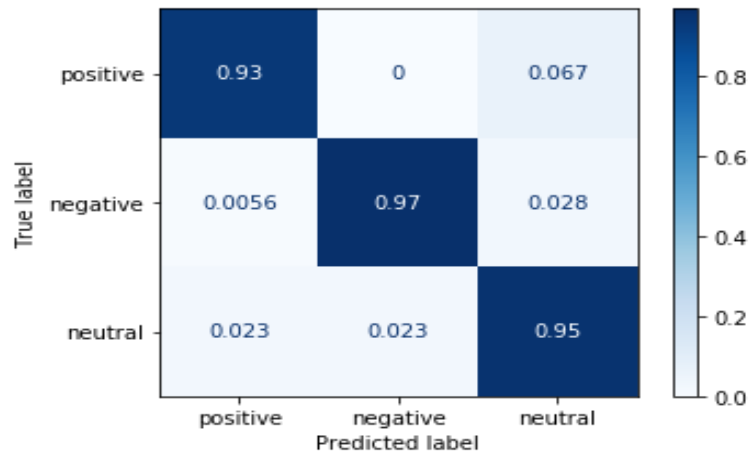
Feature	SVM			NB			LR			GB		
	P	R	F1	P	R	F1	P	R	F1	P	R	F1
TF-IDF	77	77	76	73	74	73	74	74	74	87	87	87
Count Vectorizer	86	85	85	74	74	74	76	76	76	91	91	91
TFIDF +Unigram	79	79	78	76	76	76	74	74	74	89	89	89
TFIDF + Bigram	89	89	89	85	85	85	87	86	86	92	92	92
TFIDF + Trigram	89	87	87	89	89	89	87	85	85	94	93	93
TFIDF+ Unigram+Bigram	87	86	86	86	86	86	85	83	84	93	93	93
TFIDF + Bigram+Trigram	92	92	92	88	88	88	88	86	86	89	89	89
TFIDF+ Unigram+Trigram	89	88	88	87	87	87	87	85	85	94	93	93
Countvectorizer+ unigram	84	84	84	76	76	76	76	76	76	95	95	95
Countvectorizer+ Bigram	93	93	93	86	86	86	89	88	88	90	90	90
Countvectorizer+ Trigram	92	91	91	88	88	88	91	90	90	92	91	91
Countvectorizer+ Unigram+Bigram	92	92	92	86	86	86	86	86	86	95	95	95
Countvectorizer+ Bigram+Trigram	93	93	93	89	89	89	92	90	90	94	94	94
Countvectorizer+ Unigram+Trigram	95	95	95	89	89	89	89	89	89	95	95	95

From Table 1.10 F1-score determine the weighted average of Precision and Recall, and F1-score metrics selected to compare the models based on the features because it's more useful than accuracy when the dataset contains uneven class distribution, so based on F1-score SVM and GB model based on countvectorizer+unigram+bigram obtain a higher score of 95% than NB, LR, and GB models.

6.3. Normalized Confusion matrix for SVM, NB, LR, and GB

A normalized confusion matrix makes it easier to visualize and interpret how the labels are being predicted.

normalized confusion matrix of GB Model Based on countvectorizer+Unigram+Trigram



normalized confusion matrix of GB Model Based on TF-IDF+Unigram+Trigram

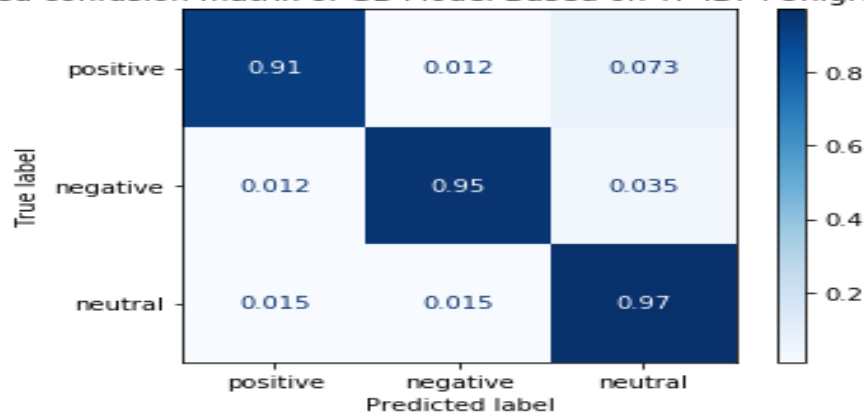


Figure 1.19: Confusion Matrix of GB based on Extracted Features

Figure 1.19 illustrates the sampled confusion matrix of GB models. Based on Countvectorizer+Unigram+Trigram, GB classifies 93% of True positive, 97% of True Negative, and 95% of True neutral class correctly, and 0.56% and 2.3% of FP, 2.3% of FN, 6.7% and 2.8% of False neutral are misclassified. Based on TFIDF+Unigram+Trigram, GB classifies 91% of True positive, 95% of True Negative, and 97% of True neutral class correctly, and 1.2% and 1.5% of FP, 1.2% and 1.5% of FN, 7.3% and 3.5% of False neutral are misclassified.

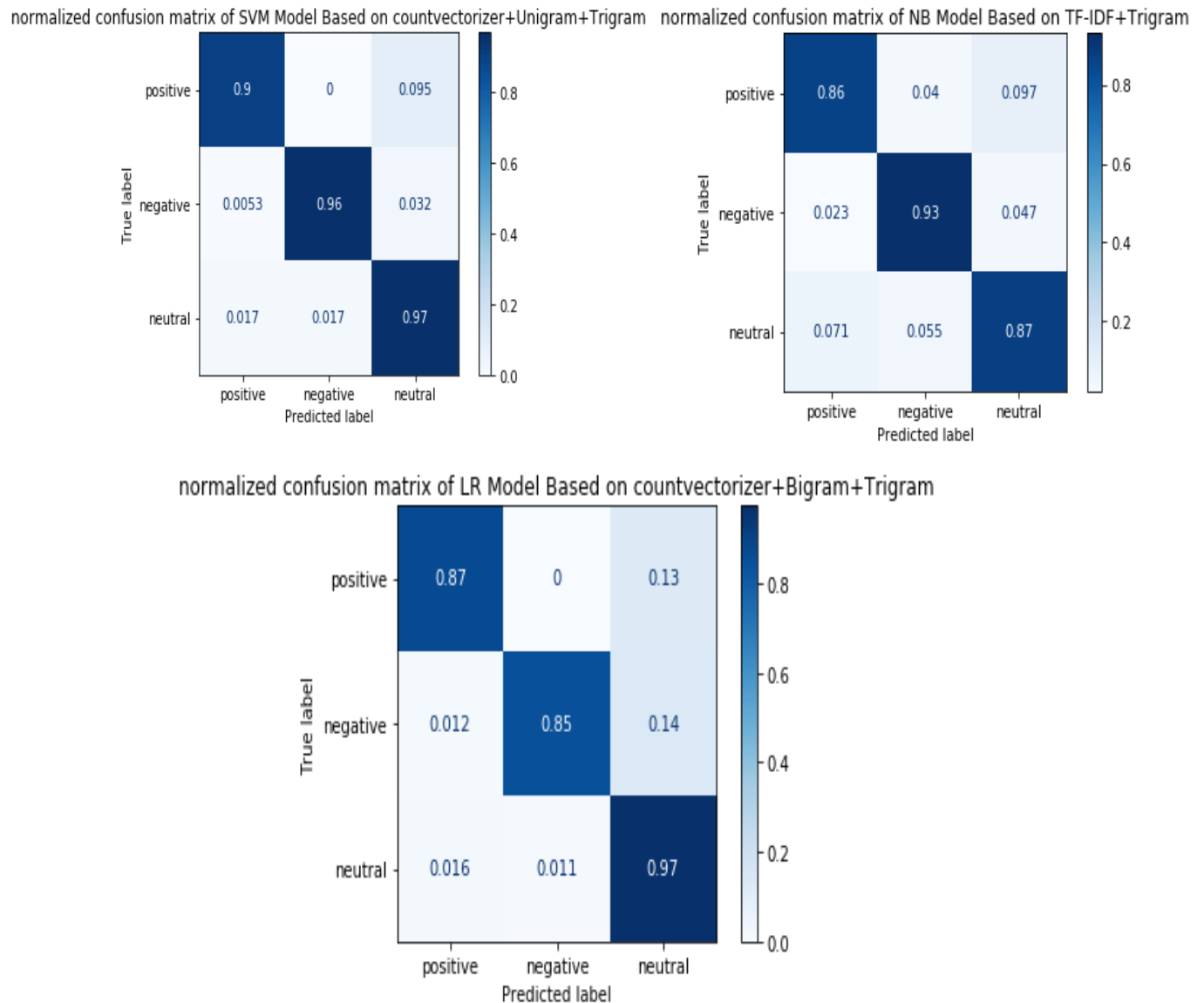


Figure 1.20: Confusion Matrix of SVM, LR, and NB Based on Extracted Feature

Figure 1.20 illustrates the sampled confusion matrix for SVM, NB, LR, and GB models. Based on Countvectorizer+ Unigram+Trigram, SVM classifies 90% of True positive, 96% of True Negative, and 97% of True neutral class correctly, and 0.53% and 1.7% of FP, 1.7% of FN, 9.5% and 3.2% of False neutral are misclassified. Based on TFIDF+Trigram, NB classifies 86% of True positive, 93% of True Negative, and 87% of True neutral class correctly, and 2.3% and 7.1% of FP, 4% and 5.5% of FN, 9.7% and 4.7% of False neutral are misclassified. Based on Countvectorizer+ Bigram+Trigram, LR classifies 87% of True positive, 85% of True Negative, and 97% of True neutral class correctly, and 1.2% and 1.6% of FP, 1.1% of FN, 13% and 14% of False neutral are misclassified.

6.4. Discussions

This study experimented with different feature extraction and machine learning algorithms like SVM, NB, LR, and GB to model Aspect/feature Based Sentimental Analysis evaluated by 10-fold CV. In the beginning, the dataset builds from scratch by collecting data from the hotel's social media (Facebook page and website) and after collecting data we apply different data preprocessing techniques to the collected data in the dataset, and label (its polarity and aspect/feature) the data by using 3 annotators. This study shows SVM makes the accuracy of 92.2% (0.922) resulted in using the countvetorizer+unigram+Trigram feature model, NB makes the accuracy of 81.5% (0.815) resulted using TFIDF+trigram feature model, LR makes the accuracy of 87.9% (0.879) resulted using the Countvectorizer+bigram+trigram feature model, and GB makes the accuracy of 92.8% (0.928) resulted using the TF-IDF+unigram+trigram and countvetorizer+unigram+Trigram feature model. Finally, GB achieves high accuracy than other algorithms.

CHAPTER SEVEN

CONCLUSION AND FUTURE WORKS

7.1. Conclusion

These days, the application of opinion analysis technology has been applied in different commerce issues, services, movie, and Hotel businesses, products, and a few specific subject surveys. But it would be difficult for the hotel management to manually analyze a large amount of data to understand whether the customer is satisfied or not. In this research Aspect based sentimental Analysis is proposed in the Hotel review domain. we use data preprocessing tasks to extract relevant data and applying a machine learning approach to learn the data.

This research applied four supervised machine learning classification techniques, Naive Bayesian, Logistic Regression, Support Vector Machine, and Gradient boosting were experimented with for building and evaluating the sentiment (polarity) model, and we use the cosine similarity technique for extracting aspects (features) of the hotels. There is Different performance evaluation metrics are going to be used like accuracy, precision, recall, and F- Measure to help us in evaluating the performance of the proposed prototype and 10-fold CV used for the effectiveness of the performance.

The 10-fold accuracy of the GB model makes of 92.8% based on the TF-IDF+N-grams and Count vectorizer+N-gram feature greater than that of the SVM, NB, and LR models. The accuracy of the SVM models makes of 92.2% based on the Countvectorizer+unigram+trigram features greater than that of the LR and NB models.

7.2. Contributions of the Thesis

Some of the main contributions of this research work are given below.

- We proposed an ML model for predicting Aspect Based sentiment of Amharic text with good performance metrics.
- We Build Hotel Corpus by collecting 2124 Amharic hotel reviews.
- This research can be used as a baseline for Aspect Based sentiment mining-related research works for opinionated Amharic text.

7.3. Future Works

The study has shown that Aspect Based sentiment analysis can be done automatically for the Amharic language by using supervised machine learning. However, many issues need further consideration for future research. In this section we recommend further research:

- Opinion spam detection or fake review detection can also be a future research direction.
- Opinion holder detection or the reason for positive, negative, and neutral opinions can also be a future research direction.
- Due to no standard available Amharic language stemmer tool. The study did not apply these preprocessing methods.

Reference

- Africa Internet Users, 2021 Population and Facebook Statistics*. (n.d.). Retrieved June 2, 2021, from <https://www.internetworldstats.com/stats1.htm>
- Anderson, E., Fornell, C., & Lehmann, D. (1994). Customer Satisfaction, Market Share, and Profitability: Findings from Sweden. *Journal of Marketing*, 58, 53–66.
- Balijepalli, S. (2007). *Blogvox2: A Modular Domain Independent Sentiment Analysis System*.
- Baye Ymam. (2002). *ye amargha sewasew* (2nd editio). Mega printing enterprise.
- Bender, M. L., Sydney W. Head, R. C. (1976). The Ethiopian Writing System. In *Bender et Al (Eds.) Language in Ethiopia*. London: Oxford University Press. <https://doi.org/10.1016/j.cosrev.2017.10.002>
- Bhattacharjee, S., Das, A., Bhattacharya, U., Parui, S., & Roy, S. (2015). *Sentiment analysis using cosine similarity measure*. <https://doi.org/10.1109/ReTIS.2015.7232847>
- Bhuiyan, T., Xu, Y., & Josang, A. (2009). State-of-the-art review on opinion mining from online customers' feedback. In Y. Aruka & A. Namatame (Eds.), *Proceedings of the 9th Asia-Pacific Complex Systems Conference* (pp. 385–390). Chuo University. <https://eprints.qut.edu.au/29301/>
- “Cohen’s Kappa Statistic - Statistics How To.” <https://doi.org/10.15520/ajcsit.v4i8.9>
- Confusion Matrix for Machine Learning*. (n.d.). Retrieved June 3, 2021, from <https://www.analyticsvidhya.com/blog/2020/04/confusion-matrix-machine-learning/>
- Devi, D. V. N., Kumar, C. K., & Prasad, S. (2016). *A Feature Based Approach for Sentiment Analysis by Using Support Vector Machine*. <https://doi.org/10.1109/IACC.2016.11>
- Esuli, A., & Sebastiani, F. (2006). *SentiWordNet: A Publicly Available Lexical Resource for Opinion Mining*.
- Guarino, N. (1998). *Formal Ontologies and Information Systems*.
- Haile, G. (1967). *The Problems of Amharic Writing System*. unpublished. <https://doi.org/10.5281/zenodo.3977576>
- Harrington, P. (2012). *Machine Learning in Action*. <https://doi.org/10.4018/978-1-4666-0059-1.ch008>
- Hernik, J., & Grīnberga-Zālīte, G. (2017). Criteria of hotel services quality evaluation in the opinion of tourists on the objects in Świnoujście and Jurmala examples. *Marketing i Zarządzanie*, 47, 209–219. <https://doi.org/10.18276/miz.2017.47-19>

- Hovy, M. K. and E. (2006). Extracting Opinions Expressed in Online News Media Text with Opinion Holders and Topics. In *COLING-ACL06*. <https://doi.org/10.1109/WIAT.2008.388>
- Hu, M., & Liu, B. (2004). Mining and Summarizing Customer Reviews. *Proceedings of the Tenth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 168–177. <https://doi.org/10.1145/1014052.1014073>
- Jagtap, V., & Pawar, K. (2013). Analysis of different approaches to Sentence-Level Sentiment Classification. *International Journal of Scientific Engineering and Technology*, 2, 164–170.
- Kandampully, J., & Suhartanto, D. (2000). Customer loyalty in the hotel industry: the role of customer satisfaction and image. *International Journal of Contemporary Hospitality Management*, 12(6), 346–351. <https://doi.org/10.1108/09596110010342559>
- Kaplan, A. M., & Haenlein, M. (2010). Users of the world, unite! The challenges and opportunities of Social Media. *Business Horizons*, 53(1), 59–68. <https://doi.org/10.1016/j.bushor.2009.09.003>
- Kennedy, A., & Inkpen, D. (2006). Sentiment Classification of Movie Reviews Using Contextual Valence Shifters. *Computational Intelligence*, 22, 6–7. <https://doi.org/10.1111/j.1467-8640.2006.00277.x>
- Khan, K., Baharudin, B., Khan, A., & E-Malik, F. (2009). *Mining opinion from text documents: A survey*. <https://doi.org/10.1109/DEST.2009.5276756>
- Kumar, A., Nayaki, A. P., & Devi, M. (2016). *Customer Feedback Evaluation System Using Feature Based Opinion Mining*.
- Liu, B. (2012). Sentiment Analysis and Opinion Mining. In *Synthesis Lectures on Human Language Technologies* (Vol. 5). <https://doi.org/10.2200/S00416ED1V01Y201204HLT016>
- Maurel, S., Curtoni, P., & Dini, L. (2008). A Hybrid Method for Sentiment Analysis. In *Statistical Analysis Software (SAS) press, Grenoble, France*,. <http://editions-rnti.fr/?inprocid=1000881>
- Maximum Margin Hyperplane - an overview | ScienceDirect Topics*. (n.d.). Retrieved June 2, 2021, from <https://www.sciencedirect.com/topics/computer-science/maximum-margin-hyperplane>
- Mebrahtu Tadesse. (2018). Trilingual Sentiment Analysis on Social Media. *Unpublished Masters Thesis, Addis Ababa University*.
- Mei, Q., Ling, X., Wondra, M., Su, H., & Zhai, C. (2007). Topic sentiment mixture: Modeling facets and opinions in weblogs. *16th International World Wide Web Conference, WWW2007*,

- 171–180. <https://doi.org/10.1145/1242572.1242596>
- MELESE MIHRET. (2017). SENTIMENT ANALYSIS MODEL FOR OPINIONATED AWNGI TEXT. *Master Thesis, Unpublished, Gonder University*.
- Meng, Y. (2012). *Sentiment Analysis: A Study on Product Features*.
- Mowlaei, M. E., Abadeh, M. S., & Keshavarz, H. (2020). Aspect-based sentiment analysis using adaptive aspect-based lexicons. *Expert Systems With Applications*, 148, 113234. <https://doi.org/10.1016/j.eswa.2020.113234>
- Natekin, A., & Knoll, A. (2013). Gradient Boosting Machines, A Tutorial. *Frontiers in Neurorobotics*, 7, 21. <https://doi.org/10.3389/fnbot.2013.00021>
- Olorunniwo, F., Hsu, M. K., & Udo, G. J. (2006). Service quality, customer satisfaction, and behavioral intentions in the service factory. *Journal of Services Marketing*, 20(1), 59–72. <https://doi.org/10.1108/08876040610646581>
- Pang, B., & Lee, L. (2008). Opinion Mining and Sentiment Analysis. *Foundations and Trends® in Information Retrieval*, 2(1–2), 1–135. <https://doi.org/10.1561/15000000011>
- Pang, B., Lee, L., & Vaithyanathan, S. (2002). Thumbs up? Sentiment Classification using Machine Learning Techniques. *Proceedings of the 2002 Conference on Empirical Methods in Natural Language Processing ({EMNLP} 2002)*, 79–86. <https://doi.org/10.3115/1118693.1118704>
- Prabowo, R., & Thelwall, M. (2013). Sentiment analysis: A combined approach. *Journal of Informetrics*, 143–157. <https://doi.org/10.1016/j.joi.2009.01.003>
- Ravendra Ratan and Singh Jandail. (2014). A proposed Novel Approach for Sentiment Analysis and Opinion Mining. *International Journal of UbiComp*, 5, 1–2. <https://doi.org/10.5121/iju.2014.5201>
- Saritha, B., & Nayak, J. (2019). *Aspect Based Sentiment Analysis using Naïve Bayes and Support Vector Classifiers*. XI(336), 336–344.
- Schrauwen, S. (2010). *MACHINE LEARNING APPROACHES TO SENTIMENT ANALYSIS USING THE DUTCH NETLOG CORPUS*.
- Selama Gebremeskel. (2010). SENTIMENT MINING MODEL FOR OPINIONATED AMHARIC TEXTS. *Unpublished Masters Thesis, Addis Ababa University*.
- Söderlund, M. (1998). Customer satisfaction and its consequences on customer behaviour revisited. *International Journal of Service Industry Management*, 9, 169–188.

- Tamrakar, S., Bal, B. K., & Thapa, R. B. (2020). *ASPECT BASED SENTIMENT ANALYSIS OF NEPALI TEXT*. 2(1), 22–29.
- Tripathi, G., & S, N. (2015). Feature Selection and Classification Approach for Sentiment Analysis. *Machine Learning and Applications: An International Journal*, 2, 1–16. <https://doi.org/10.5121/mlaij.2015.2201>
- TULU TILAHUN. (2013). Opinion Mining From Amharic Blog. *Unpublished Masters Thesis, Addis Ababa University*.
- Varghese, R., & Jayasree, M. (2013). Aspect based Sentiment Analysis using support vector machine classifier. *Proceedings of the 2013 International Conference on Advances in Computing, Communications and Informatics, ICACCI 2013*, 1581–1586. <https://doi.org/10.1109/ICACCI.2013.6637416>
- Wiebe, J., Bruce, R., Martin, M., Wilson, T., & Bell, M. (2004). Learning subjective language. *Computational Linguistics*, 30(3), 277–308. <https://doi.org/10.1162/0891201041850885>
- Xu, X., Li, Y., Hu, L., & Tao, J. (2009). *Categorizing terms' subjectivity and polarity manually for opinion mining in Chinese*. <https://doi.org/10.1109/ACII.2009.5349566>
- Yeung, M. C. H., Ging, L. C., & Ennew, C. (2002). Customer satisfaction and profitability: A reappraisal of the nature of the relationship. *Journal of Targeting, Measurement and Analysis for Marketing*, 11, 24–33.

Appendix A: Sample letter Permission



አዳማ ሳይንስና ቴክኖሎጂ ዩኒቨርሲቲ
ADAMA SCIENCE & TECHNOLOGY UNIVERSITY

ASTU, P.O. Box 1888 ☎ 022-110-0073 Fax: 0251-022-112-01-50 E-mail: soecdean@astu.edu.et

ሰላም ሆኖል
 ለኤስኤስኤስ ሆኖል
 ለአዳማ ጋርመንት ሆኖል
 አዳማ
 ለአዲስ አበባ ዩኒቨርሲቲ
 አዲስ አበባ

ቀን: **07 መጋቢት 2013**
 ቁጥር: **7/2ተ/7/30/7606/13**
 ኤሌክትሮኒክስ ምህንድስና እና ኮምፒውተር
 ት/ቤት ዲን
 ዶ/ር ደረጀ ተክሉ አበፋ

ጉዳዩ:- ትብብር ስለመጠየቅ ይሆናል

በአዳማ ሳይንስና ቴክኖሎጂ ዩኒቨርሲቲ በኤሌክትሮኒክስ ምህንድስና እና ኮምፒውተር ት/ቤት ስር የሚገኘው የኮምፒውተር ሳይንስና ምህንድስና ት/ክፍል የድህረ ምረቃ ተማሪ የሆኑት የርዳኛስ አጽናፍ የመመረቂያ ዕረፋቸውን “Aspect Based Sentimental Analysis for the Amharic Language Using Machine Learning Approach” በሚል ርዕስ እየሰሩ ይገኛሉ።

በመሆኑም የመመረቂያ ዕረፋቸውን በተመለከተ ለሚያደርጉት የመረጃ አሰባሰብ አስፈላጊው ትብብር እንዲደረግላቸው እየጠየቅን ለሚደረግላቸው ትብብር በቅድሚያ እናመሰግናለን።

ከሰላምታ ጋር



ዶ/ር ደረጀ ተክሉ አበፋ
 ኤሌክትሮኒክስ ምህንድስና እና
 ኮምፒውተር ት/ቤት ዲን



ግልግጭ:

- ለት/ቤቱ ዲን
- ለት/ቤቱ አካ/ጉ/ተ/ዲን
- ለኮምፒውተር ሳይንስና ምህንድስና ት/ክፍል

አ.ሳ.ቴ.ዩ
 አዳማ

Appendix B: Approval Letter for Linguistic Expertise

አዳማ ሳይንስና ቴክኖሎጂ ዩኒቨርሲቲ
ኮምፒውተር ሳይንስ እና ኢንጅነሪንግ ትምህርት ክፍል

የመጠየቁ አላማ፡

በአዳማ ሳይንስና ቴክኖሎጂ ዩኒቨርሲቲ በኮምፒውተር ሳይንስ እና ኢንጅነሪንግ ትምህርት ክፍል ለሁለተኛ ዲግሪ (ማስተርስ) ምርምር ማግኒያ የሚሆን በአማርኛ ቋንቋ ላይ መሰረት ያደረገ ኮምፒውተራዊዝድ የሰዎች እስተዋት መለያ ሲስተም ለመስራት ግብአት የሚሆኑ 418 አዎንታ እና 577 አሉታ ቃላት እና Annotation Guideline ተዘጋጅቶል ::

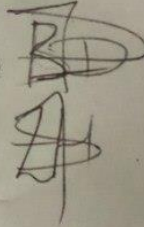
በቅድሚያ ለሚያደርጉልን ትብብር ክልብ እያመሰገንን፤ የእነዚህን ቃላት ትክክለኛነት እና Annotation Guideline እንዲያረጋግጡልን በትህትና እንጠይቃለን፡፡

የመጠየቁ አዘጋጅ፡

ዮርዳኖስ አጽናፍ በኮምፒውተር ሳይንስ እና ኢንጅነሪንግ ትምህርት ክፍል የሁለተኛ ዲግሪ(ማስተርስ) ተማሪ።

ያረጋጠው፡

1. በርህን ደጅኝ የ አማርኛ ትምህርት ክፍል ሀላፊ(Msc)
2. ዘነበ ሀይሉ ከ አማርኛ ትምህርት ክፍል (Msc)



Appendix C: ከዚህ በታች የተዘረዘሩት ቃላት አዎንታን ያሳያሉ፡፡

[ገለልተኛ፣ ቆንጆዎች፣ ጠንካራ፣ የሚያበረታታ፣ ሀይማኖተኛ፣ በረከት፣ ባለሙያ፣ የሚገርም፣ እድገት፣ ፍቃድ፣ ስልት፣ ነቃ፣ ተሀድሶ፣ ልማት፣ ተፈላጊ፣ ጤነኛ፣ እድለኛ፣ ያለው፣ ያለው፣ እሺ፣ ከብረት፣ የበለጠ፣ ግርማ፣ አንጠልጣይ፣ ፈውስ፣ የሚያዝናና፣ ፀጋ፣ ፀፀት፣ እምቅ፣ አርአያ፣ ማስወገድ፣ ቅልጥ፣ ሳቂታ፣ ሰናይ፣ ገነት፣ ነፃ፣ ተነሳሽነት፣ መልካም፣ ይቅርታ፣ ሙደኛ፣ የማይካድ፣ ቃል፣ ምርጥ፣ ስንዱ፣ ዋስትና፣ የተካነ፣ ሰፋፊ፣ እሸሩሩ፣ ድምቀት፣ ትህትና፣ የደመቀ፣ ደህና፣ እርካታ፣ እሙን፣ የተወደ፣ ሎጋ፣ ንፁህ፣ አንበሳ፣ መልእክት፣ ትጋት፣ ገቢራዊ፣ ጠንቃቃ፣ እንቁ፣ ተግሳፅ፣ ባለውለታ፣ ጋባዥ፣ ሀይለኛ፣ አወንታዊ፣ መሳካት፣ እድል፣ ህጋዊ፣ ርህራሄ፣ ወርቃማ፣ እጣ፣ ጭብት፣ እፁብ፣ ብልህተኛ፣ አእምሮ፣ ጀግንነት፣ ተመራጭ፣ የሚበረታታ፣ ቅልጥፍና፣ ደስታ፣ ማለፍያ፣ ደፋር፣ ሞገስ፣ ንቃት፣ ቀልጣፋ፣ ጥበብ፣ ጥበበኛ፣ ሰፊ፣ አላቸው፣ ፍላጎት፣ ንዋይ፣ ስጦታ፣ ጀግና፣ አለቃ፣ ብልጥ፣ ወርቃማ፣ ጣፋጭ፣ ቆንጆ፣ ጠቀሜታ፣ ታታሪ፣ ጥብቅ፣ ሀብት፣ መተዛዘን፣ ህሊና፣ ፅናት፣ ተድላ፣ ጠቃሚ፣ ሞያ፣ ታዋቂ፣ ጣእም፣ ቀኝ፣ ጀብድ፣ ቀሽት፣ እውቅ፣ አስቂኝ፣ ስልጡን፣ አሜን፣ አስተማማኝ፣ ልምላሜ፣ ምሩቅ፣ ምክር፣ ተአምራዊ፣ ፍቃደኛ፣ ተአምር፣ የላቀ፣ ሞቅ፣ ጥቅም፣ ትእግስት፣ ፅኑ፣ ለጋሽ፣ ተመጣጣኝ፣ ብርቱ፣ አዝናኝ፣ ራእይ፣ ዘንካታ፣ አስተዋይ፣ ፕ፣ ብልህ፣ ያማረ፣ እውነተኛ፣ ጌጣጌጥ፣ ተሰጥኦ፣ አስደናቂ፣ አሳሳቢ፣ ደንታ፣ ተልእኮ፣ ጤና፣ ጎበዝ፣ ጠባ፣ ጥሩ፣ ልዩ፣ ምቹት፣ ቁጥብ፣ ተጫዋች፣ ብልሆች፣ የምያምር፣ የሚል፣ ውል፣ መረጋጋት፣ እመርታ፣ አሸናፊ፣ ታጋሽ፣ ተቀባይነት፣ ስምምነት፣ ውብ፣ መግባባት፣ ቸር፣ ልክ፣ ዘላቂነት፣ ቁርጠኝነት፣ ፈታኝ፣ ባህላዊ፣ ተአማኒነት፣ በቂ፣ ድህነት፣ ሽልማት፣ አወንታስምምነት፣ ፅዱ፣ ዝምተኛ፣ ልብ፣ አስፈላጊ፣ ሲሳይ፣ የተስተካከለ፣ ቆራጥ፣ ምስጋና፣ ጎሽ፣ አላማ፣ የተመረጠ፣ ሽጋ፣ ገናናነት፣ ፀባይ፣ አብነት፣ ደህንነት፣ ወሳኝ፣ አስደማሚ፣ ማስተዋል፣ የበላይ፣ ዋነኛ፣ ስብእና፣ ቅንጦት፣ ውድ፣ ስልጣን፣ ጥዱ፣ መታዘዝ፣ ዘለአለማዊ፣ ትልቅ፣ ገናና፣ በጎፍቃድ፣ ደስተኛ፣ ሩሁሩ፣ ሽታ፣ ማራኪ፣ ፍቅር፣ የተረጋገጠ፣ መረዳዳት፣ ተስፋ፣ በጎ፣ ባልንጀራ፣ የማይጎዳ፣ ዘላቂ፣ ይመቻል፣ ንቁ፣ እርቅ፣ ሀላፊነት፣ መንፈሳዊ፣ ሀቅ፣ ታጋች፣ ምሳሌ፣ ጌታ፣ ዘመናዊነት፣ ልሙጥ፣ ደንብ፣ ምቹ፣ ድንግል፣ ጨዋ፣ እመቤት፣ ትጉ፣ ትጋተኛ፣ አሳማኝ፣ ግብ፣ ህግ፣ ብልህት፣ ድሎት፣ አርበኛ፣ አልማዝ፣ ሰላም፣ አራፍ፣ ወሮታ፣ ልምድ፣ የማይረሳ፣ ታላቅ፣ ረጋ፣ ዋጋ፣ እርግጠኛ፣ ውጤታማ፣ የሚያስቅ፣ ትግል፣ ነጋሲ፣ በፍፁም፣ አቅል፣ ተጨባጭ፣ ቅዱስ፣ ጋደኛ፣ ጎበዞች፣ መርዘኛ፣ ጉጉ፣ ብሩክ፣ እቁብ፣ መተጋገዝ፣ ታድያ፣ ስ፣ አመኔታ፣ አሸበረቀ፣ ጥጋብ፣ ሀዋሪያ፣ ተከላካይ፣ ቁጥብ፣ ይሁንታ፣ ባለፀጋ፣ ምህረት፣ ቀላል፣ ዋና፣ ሰብአዊ፣ ሀመልማል፣ ትክክለኛ፣ ግርማዊ፣ እውነት፣ ወግ፣ ሚዛን፣ ደግነት፣ ቁጥብነት፣ ኩራት፣ አጋዥ፣ ዋዉ፣ ሹም፣ አድናቆት፣ ዝነኛ፣ መድሀኒት፣ ጤንነት፣ ፍሬ፣ ሀያል፣ ድል፣ ደንበኛ፣ ንብረት፣ ፊደዳ፣ ወሰክ፣ ቸሎታ፣ ፈንጠዝያ፣ ጉልበታም፣ ክብር፣ ትጉህ፣ ወዳጅ፣ ለጋስ፣ ቁርጠኛ፣ ተጠቃሚ፣ ጠቃሚነት፣ ሙድ፣ አህዛብ፣ የማይበገር፣ የዋህ፣ ቅን፣ መልከመልካም፣ ስኬታማ፣ ውይይት፣ አቅም፣ ዘለቄታ፣ ረዣዥም፣ ትብብር፣ አግባብ፣ ፀዳል፣ ሀቀኛ፣ ተወዳጅ፣ ውጤት፣ ህክምና፣ ህዝባዊ፣ ግልፅ፣ የምያረካ፣ የሚያረካ፣ ፍሰህ፣ አጥጋቢ፣ ዝግጁ፣ አሳመረ፣ ቀጥታ፣ ታዛዥ፣ ልባዊ፣ ፅድቅ፣ ትሁት፣ ፋና፣ ማአረግ፣ ለዛ፣ አዘኔታ፣ አይነተኛ፣ ሙያዊ፣ አመርቂ፣ ጭብጥ፣ ሰላማዊ፣ አንድነት፣ የተረጋጋ፣ ዘመናዊ፣ ፈጣን፣ ቸርነት፣ አዛኝ፣ ብርቅ፣ ረዥም፣ ትንግርት፣ ወዘና፣ የሚያረካ፣ የምመች፣ የሚመች፣ ብራሽ፣ ከበሬታ፣ ልእልና፣ ስመጥር፣ ብልፅግና፣ ምሁራዊ፣ ብርሀን፣ ሀብታም፣ አይብ፣ ሳያመነታ፣ ወዳጅነት፣ ክብር፣ ቻይ፣ አስተካካይ፣ ደማቅ፣ ህያው፣ ሚና፣ ደመቀ፣ ታማኝ፣ ፍሬአማ፣ የተማረ፣ ሻካራ፣ ቅጥ፣ ፍትሀዊ፣ ቀያይ፣ መሰረት፣ ተአምረኛ፣ ሀርነት፣ ወሽን፣ ወሽኔ፣ ድንግልና፣ ፈገግታ፣ ዝና፣ ሩህሩህ፣ ከፍተኛ፣ ጥረት፣ ድንቅ፣ ተግባቢ፣ ደግ፣ ጠቢብ፣ ሰላምተኛ፣ እምነት፣ ግብረገብ፣ ይሰራሉ፣ አዳኛ፣ ምሁር፣ ጥርት]

Appendix D: ከዚህ በታች የተዘረዘሩት ቃላት አፍራሽነትን ያሳያሉ፡፡

[ቂመኛ፣የማይመች፣ነሁላላ፣አክሳሪ፣ትእቢት፣እቡይ፣ጮሌ፣ኢፍትህዊ፣ዘራፊ፣ደነዝ፣ያልታወቀ፣ወገናዊ፣ህፀፅ፣ኢፍትህዊነት፣እልህኛ፣ተቃዋሚ፣መሸወድ፣የምያሳጣ፣ተቃራኒ፣ጥገኝነት፣ብልጣብልጥ፣ጀልጋጋ፣አድካሚ፣መቅጫ፣ግፈኛ፣ጨለምተኛ፣እነካኪ፣እዳ፣እምባ፣ባርነት፣ነዳይ፣ፈተና፣ወልጋዳ፣ገዳይ፣ቱባ፣እብለት፣መዘዝ፣ወቀሳ፣እሪታ፣መልቲ፣በዘፈቀደ፣ጭቅጭቅ፣እንከን፣አለቅጥ፣አስደንጋጭ፣ኮተታም፣ዱብዳ፣የማያስደስት፣የሚያበሳጭ፣ለቅሶ፣ፈሳም፣ስሜታዊ፣ጉርምርምታ፣አለመሆን፣አለመሆኑ፣ፍርሀት፣ደረቅ፣ቋጣሪ፣መርዝ፣ጠባቦች፣የሚያዳግ፣እድል፣ወሬኛ፣ጨፍጫፊ፣አልሆነም፣አልተመቸኝም፣ጠላት፣ዝንጉ፣መንዛዛት፣መአት፣ጨቅጫቃ፣ጦርነት፣ተፃራሪ፣ቅሬታ፣ኡኡታ፣ቀውስ፣አረንቋ፣ወድ፣አፍራሽ፣ወለፌንድ፣ሞልጣፋ፣አንገብጋቢ፣የሞተ፣ጨሀት፣ግድፈት፣ስርቅታ፣የሚያስጠላ፣ዝቃጭ፣መና፣ዝሙት፣አጣዳፊ፣ሀሜት፣ግፊት፣ውርጅብኝ፣ፍጥጫ፣በጎሪጥ፣አይደለም፣ሌባ፣አስከፊ፣ዘግናኝ፣አረመኔ፣ዱርየ፣አስቸካይ፣አላግባብ፣ብስጭት፣አለመተማመን፣ቀጣፊ፣የሚያስቀይም፣እኩይ፣የባሰ፣ፌዝ፣ቀጣት፣ከባድ፣ያለመሳካት፣ምላሽ፣በሽታ፣ህቅታ፣የተጋለጠ፣ወንጀለኛ፣ተራ፣ሌሊት፣በሽተኛ፣እርካሽ፣ንጭንጭ፣ቸልተኛ፣ወረርሽኝ፣ያልነቃ፣አጭር፣አዋረድ፣ባዶ፣ከረከረ፣አስጠያፊ፣ረብሻ፣እንቅፋት፣ምፀት፣የረጅ፣የማይታለም፣የማይል፣ቀማኛ፣ለፍላፊ፣የማይገኝ፣አለመስማማት፣ገደብ፣ድንጋጤ፣ተመዳቂ፣እከካም፣ዝርክርክ፣የሚቃወም፣ግፍ፣ከንቱ፣አንጎል፣ቅማላም፣አክራሪ፣መጋኛ፣አርነት፣ፈት፣እለታዊ፣ሲአል፣እብድ፣ደካማ፣ንፉግ፣በቀል፣አይደለም፣ሽፍታ፣ወላዋይ፣ቂል፣አሉታዊ፣ሻገተ፣ቂም፣ጥፋት፣ያልነቃ፣የተዘበራረቀ፣ምቀኛ፣ጥፋ፣ጥድፍያ፣ተናደደ፣ፈዛዛ፣ሰነፍ፣ቅጥፈት፣ሽቦት፣አይታዘዙም፣መከራ፣የማይስማማ፣ሽብርተኛ፣ባእድ፣አስጠሊ፣ጉረኛ፣ከሀዲ፣ነቀፋ፣ተጠራጣሪ፣ዱርዬዎች፣እርቃን፣መቃወም፣ሰለባ፣ስድ፣ዝሙት፣ተንኮል፣ቁስል፣እጦት፣አባሳ፣ድባቅ፣አጠራጣሪ፣ሽንፈት፣ውስን፣ጅል፣ትንሽ፣ብኩን፣አጥፊ፣መጫር፣ኢምንት፣ሸባ፣ኋላቀርነት፣ትርምስ፣ገውጋዋ፣ልል፣የማይረባ፣ጉደኛ፣ሙጭጭ፣ሀሰት፣አሳሳች፣ወራዳ፣ደመኛ፣ቸግር፣ክስ፣የሌለዉ፣እርጉም፣አማፅያን፣ሞልቃቃ፣ተፅእነኖ፣ማታለል፣ጠብ፣ስንኩል፣ጉዳት፣ኮርማታ፣አማፂ፣ንፍግ፣ብልሹ፣ውሻ፣ሳፋሪ፣ጥል፣ቅናት፣ሀጢአት፣አልባሌ፣ነፃነት፣ቅናታም፣ሰይጣን፣ስጋት፣ያልተገደበ፣ሰነፍ፣ሽክም፣ህመም፣አንዛራጭ፣መቅዘፍት፣ካልቶ፣ፈጣጣ፣አድላዊ፣ምስኪን፣አምባገነን፣ዘረኛ፣ይጎላል፣ንዝንዝ፣መግረፍያ፣ቆራጣ፣ክፉ፣እስከነአካቴው፣ግርፍያ፣ቅር፣ልቅ፣ያለአግባብ፣አይልም፣ገብጋባ፣ጥለኛ፣ኮተት፣ማግለል፣ጭንቅ፣አይዴሉም፣ውዥምብር፣ሴራ፣ዜሮ፣ወንጀል፣ወረኛ፣ኋላቀር፣ጠባሳ፣አይመችም፣ግትር፣ከፋፋይ፣አዳሚ፣ዳተኛ፣ጥቁር፣ጠሳ፣ጉስቁል፣እርኩስ፣ሀሴት፣አያዉቁም፣አላስፈሊጊ፣ማደናቀፍ፣ወረተኛ፣ይረብሻል፣መደናገር፣አጭበርባሪ፣ንጭጭ፣ማጭበርበር፣ሽባ፣ስድብ፣አሸባሪ፣ሞልፋግ፣ቆፎ፣ወከባ፣በዘኔ፣የተጋነነ፣እንቆቅልሻዊ፣አደገኛ፣ጨኸት፣ጥርሳም፣አደናጋሪ፣ፍጋሪ፣ቅፅበታዊ፣ያልተማሩ፣አፋኝ፣ክልክል፣ሸፋፋ፣ቂላቂል፣አድማ፣ጠባብ፣ወሬ፣ማንባት፣ድንገተኛ፣ጥፋተኛ፣ልፍያ፣ጥንታዊ፣ሴሶኛ፣አስቀያሚ፣ነገረኛ፣ቀልቃላ፣ጎርባጣ፣መቃብር፣ትምክህተኛ፣ጥገኛ፣አመፅ፣ደዌ፣አስቃቂ፣አውደልዳይ፣አደናቃፊ፣ሹጣም፣መጥፎ፣ተልካሽ፣እርግጫ፣ትችት፣ግራ፣ኩራተኛ፣ያልተረጋገጠ፣ማስገደድ፣ጥማት፣ተመፅዋች፣ደባሪ፣ገገማ፣ያልተገራ፣የሚያሳፍር፣አስመሳይ፣ተመናመነ፣ጋጠወጥ፣ሸራፋ፣መርዛም፣ጥርጣሬ፣ስርዝ፣ስህተት፣ቂርጠት፣ሸረኛ፣ነውጠኛ፣ንደት፣ብቸኛ፣ወለብለባ፣ማሸንክ፣ጌጃ፣ይሸታሉ፣ጠማማ፣ይጎላቸዋል፣ጅንን፣ይደብራል፣ቅሌት፣ጤናማጣት፣ውስብስብ፣አለመግባባት፣የተጨናነቀ፣ቡካን፣ግርፋት፣ፍዳ፣ቀለጤ፣ዬለዉም፣ያበደ፣የተሳሳተ፣የማያረካ፣ውሸት፣ባለጌ፣ብላሽ፣የማየሻሻል፣ያስጠላል፣ባይሆንም፣ሀሰተኛ፣ሆዳም፣አውዳሚ፣ያልታጠበ፣አስፈሪ፣አፈንጋጭ፣ፅንፈኛ፣የሚያወላውል፣ነፍጠኛ፣ሞኝ፣ዘዋሪ፣ግጭት፣ጣልቃ፣የማይለወጥ፣የሚያሳዝን፣የሚያመነታ፣መጮህ፣ውርደት፣ህገወጥ፣አሰልቺ፣ጋጋታ፣ግም፣ውሸታም፣ማስጨነቅ፣የሚረብሽ፣ናፋቂ፣ማጣት፣መላምታዊ፣አመል፣መቅሰፍት፣ጭፍጨፋ፣ጅሎች፣ፈሪ፣ያልተዛባ፣ፀፀት፣አሉባልታ፣መካን፣ለምፍ፣ምቅኝነት፣ጭቀን፣ኮስታራ፣ጨላማ፣ፋራ፣ዝቅተኛ፣ጎስቃላ፣ልሞሽ፣ሀዘን፣ቀዥቃዣ፣ባርያ፣ምላሰኛ፣ቸኮላ፣ቅልጣን፣አታላይ፣ቆሻሻ፣ከይሲ፣ክህደት፣ቸኩላ፣ቅፅበት፣ፀያፍ፣ግድያ፣ሸካካ፣ሳይሳካ፣ግሸበት፣ምስጢር፣ፍጭት፣ትምክህት፣አድመኛ፣ጥፊ፣ከዳተኛ፣ልግመኛ፣ንትርክ፣ሁከት፣እጥረት፣ህፍረት፣መማቀቅ፣እክል፣መደብደብያ፣ማስጠንቀቅያ፣ቸስታ፣ሙጥኝ፣ብቻ፣አዛዥ፣ሽብር፣እምቢተኛ፣ተፅእኖ፣ፈርሳም፣የሚያጎድል፣ቁጣ፣ነጭናጫ፣ሎሌ፣ጨካኝ፣አፀያፊ፣ቃጠሎ፣ስስታም፣አዋራጅ፣የሚቃረን፣ቀበጥ፣ጭንቀት፣ድሀ፣ብድር፣ድርቅ፣ቅራኔ፣ስደተኛ፣ሙሰኛ፣አስቸጋሪ፣በደለኛ፣ነውጥ፣የማያምር፣ቀሳፊ፣ቀናተኛ፣ትረባ፣ዋልጌ፣ድካም፣ቀሽም፣ኪሳራ፣ተጠቂ፣ሴረኛ፣ረባሽ፣እልቂት፣ከሀድ፣ሸርሙጣ፣የማይታገስ፣ጉድለት፣ድንክ፣ዘልዛላ፣ውድመት፣አቤቱታ፣ያልሰመረ፣ቀውላላ፣ሚስጢር፣ደብዛዛ፣ሀራም፣ተብታባ፣ሀካይ፣ዘገምተኛ፣ድልዝ፣ከውካዋ፣ቅንጣት፣ዩላቸዉም፣የማይታመን፣በደል፣አስጨናቂ፣ጋኔን፣ጥላሽት፣በሊታ፣ባይሆን፣ገሀነም፣ይሉኝታ፣ነውር፣ያልተጠበቀ፣ዘረኝነት፣ትእቢተኛ፣ነፍናፍ፣መሀይምነት፣ድቃላ፣ከሲታ፣የማያስመሰግን፣ስቃይ፣ወሸታም፣መራራ፣ለዘብተኛ፣መራራ፣ጥንብ፣ብክለት፣ፋንጋ፣የማይሰጥ፣ጎሰኛ፣መንድስ፣ጭንቅንቅ፣ነዝነዛ፣ድክመት፣ጫና፣ሸሌ፣የተወናበደ፣ሂስ፣ብቸኝነት፣ርካሽ፣ኩምትር፣አይደለም]

Appendix E: Amharic alphabets and their pronunciation.

[illegible]

ᱠ	ᱡ	ᱢ	ᱣ	ᱤ	ᱥ	ᱦ	ᱧ	ᱨ
[lʰa]	[hʰa]	[mʰa]	[sʰa]	[rʰa]	[sʰa]	[ʃʰa]	[kʰo]	[kʰi]
ᱩ	ᱪ	ᱫ	ᱬ	ᱭ	ᱮ	ᱯ	ᱰ	ᱱ
[kʰa]	[kʰe]	[kʰiʰu]	[nʰa]	[vʰa]	[tʰa]	[ʈʰa]	[hʰo]	[hʰi]
ᱲ	ᱳ	ᱴ	ᱵ	ᱶ	ᱷ	ᱸ	ᱹ	ᱺ
[hʰa]	[hʰe]	[hʰiʰu]	[nʰa]	[nʰa]	[ʔa]	[kʰo]	[kʰi]	[kʰa]
ᱻ	ᱼ	ᱽ	᱾	᱿	ᱺ	ᱻ	ᱼ	ᱽ
[kʰe]	[kʰiʰu]	[hʰo]	[hʰi]	[hʰa]	[hʰe]	[hʰiʰu]	[zʰa]	[ʒʰa]
᱾	᱿	ᱺ	ᱻ	ᱼ	ᱽ	᱾	᱿	ᱺ
[tsʰa]	[tʰa]	[gʰo]	[gʰi]	[gʰa]	[gʰe]	[gʰiʰu]	[tʰa]	[ʈʰa]
᱾	᱿	ᱺ	ᱻ	ᱼ	ᱽ			
[pʰa]	[tsʰa]	[tʰa]	[pʰa]					

Appendix F: Amharic Punctuation

✱	section mark	⋮	colon
⋮	word separator	⋮-	preface colon
⋮⋮	full stop (period)	⋮	question mark
⋮	comma	⋮⋮	paragraph separator
⋮	semicolon		

Appendix G: Some of Amharic Number

Western Numeral	Greek Numeral	Coptic Numeral	Ethiopic Numeral	Ethiopic Syllable
1	α'	ⲁ	፩	አ
2	β'	Ⲃ	፪	በ
3	γ'	Ⲃ	፫	ገ
4	δ'	Ⲉ	፬	ደ
5	ε'	Ⲉ	፭	ሀ
6	ς'	Ⲉ	፮	ወ
7	ζ'	Ⲉ	፯	ዘ
8	η'	Ⲉ	፰	ሐ
9	θ'	Ⲉ	፱	ጠ
10	ι'	Ⲉ	፲	የ
20	κ'	Ⲉ	፳	ከ
30	λ'	Ⲉ	፴	ለ
40	μ'	Ⲉ	፵	መ
50	ν'	Ⲉ	፶	ነ
60	ξ'	Ⲉ	፷	ሠ
70	ο'	Ⲉ	፸	ዐ
80	π'	Ⲉ	፹	ፈ
90	ρ'	Ⲉ	፺	ጸ
100	ρ'	Ⲉ	፻	ቀ

Appendix H: Some of Amharic stop-word list

['አኔ'፣'የአኔ'፣'አኔራሴ'፣'እኛ'፣'የእኛ'፣'እኛራሳችን'፣'አንቺ'፣'ከህ'፣'አላችሁ'፣'እርስዎ'፣'ትፈልጋለህ'፣'ያንተ'፣'ራስህን'፣'አራሳችሁ'፣'እሱ'፣'የእሱ'፣'ራሱ'፣'እሷ'፣'እሷናት'፣'የእሷ'፣'እራሷ'፣'ነው'፣'እነሱ'፣'እነሱን'፣'የእነሱ'፣'ራሳቸው'፣'ምንድን'፣'የትኛው'፣'ማን'፣'ይህ'፣'የሚልነው'፣'ያ'፣'እነዚህ'፣'እነዚያ'፣'ነኝ'፣'ናቸው'፣'ነበር'፣'ነበሩ'፣'ሁን'፣'ቆይቷል'፣'መሆን'፣'አላቸው'፣'አለው'፣'ነበረው'፣'ያለው'፣'መስራት'፣'ያደርጋል'፣'አደረገ'፣'ማድረግ'፣'ሀ'፣'አንድ'፣'የ'፣'እና'፣'ከሆነ'፣'ወይም'፣'ምክንያቱም'፣'አንደ'፣'እስከ'፣'እያለ'፣'በ'፣'ለ'፣'ጋር'፣'ስለ'፣'ላይ'፣'መካከል'፣'ወደ'፣'በኩል'፣'ወቅት'፣'ከዚህበፊት'፣'በኋላ'፣'ከላይ'፣'ከታች'፣'ከ'፣'ወደላይ'፣'ታች'፣'ውስጥ'፣'ውጭ'፣'በላይ'፣'እንደገና'፣'ተጨማሪ'፣'ከዚያ'፣'አንድጊዜ'፣'እዚህ'፣'እዚያ'፣'መቼ'፣'የት'፣'ለምን'፣'እንዴት'፣'ሁሉም'፣'ማንኛውም'፣'ሁለቱም'፣'እያንዳንዳቸው'፣'ጥቂቶች'፣'በጣም'፣'ሌላ'፣'አንዳንድ'፣'እንደዚህ'፣'ብቻ'፣'የራሱ'፣'ተመሳሳይ'፣'ስለዚህ'፣'ይልቅ'፣'እንዲሁ'፣'ት'፣'ይችላል'፣'ይገባል'፣'ይገባኛል'፣'አሁን'፣'መ'፣'ም'፣'አ'፣'ዳግም'፣'መሆን'፣'ሁለ'፣'ሁለም'፣'ሕዝብ'፣'ሀሙስ'፣'ለመሆኑ'፣'ለምንድን'፣'ሌሎች'፣'መጽሀፍ'፣'ማክሰኞ'፣'ምን'፣'ሰኞ'፣'ሰው'፣'ሲሆን'፣'ስንት'፣'ረቡዕ'፣'ቅዳሜ'፣'በዚህ'፣'በላ'፣'ነገር'፣'አለ'፣'አርብ'፣'አንተ'፣'አንዳንድ'፣'ኢትዮጵያ'፣'እሁድ'፣'እናንተ'፣'እንኳን'፣'እግር'፣'ከመሆን'፣'ወይንም'፣'ዋና'፣'ዘንድ'፣'የሚከተለው'፣'ያኔ'፣'ይኼው'፣'ገጽ'፣'እነርሱ'፣'ን'፣'ና'፣'ዎች'፣'ይጠበቃል'፣'በለዋል'፣'ሆ'፣'ሁሉ'፣'አንቀጽ'፣'እንደሆነ'፣'በማይበልጥ'፣'መሰረት'፣'ሁኔታ'፣'ይሆናል'፣'ሆኖ'፣'ከአንድ'፣'በማናቸውም'፣'ወር'፣'ከአምስት'፣'በሆነ'፣'ከዚህ'፣'የሆነ'፣'ሀያ'፣'ሆነ'፣'በኩላ'፣'በአንድ'፣'የሆኑ'፣'ከአስራ'፣'የሆነውን'፣'መሆኑ'፣'ሌላውን'፣'ከሰባት'፣'ለሌላ'፣'አለበት'፣'ሲል'፣'ይሆናሉ'፣'በሙሉ'፣'አስራ'፣'ቢሆንም'፣'አንዱ'፣'የሌላውን'፣'ከሁለት'፣'የሆኑትን'፣'በሆኑ'፣'ጀምሮ'፣'በመሆን'፣'ባለ'፣'ይህንን'፣'እንዲቆይ'፣'ሌላው'፣'የሚሆነው'፣'በአንዱ'፣'ሲባል'፣'ሳለ'፣'የሆነው'፣'መሆናቸው'፣'በዋና'፣'በማቀድ'፣'ጊዜና'፣'ለዚህ'፣'ሶስተኛ'፣'የነገሩ'፣'ስድስት'፣'በሆነው'፣'ይሁን'፣'ከዚሁ'፣'በእነዚህ'፣'ከማናቸውም'፣'ከነበረው'፣'በአንዳንድ'፣'በእያንዳንዱ'፣'ጊዜም'፣'እስከ'፣'የሌሎች'፣'የሚሆኑት'፣'ከሆነው'፣'የነበረውን'፣'ያሉ'፣'ከሌሎች'፣'እንዲት'፣'ለሌሎች'፣'ለሆነው'፣'ስላት'፣'በሎ'፣'ከሰላሳ'፣'የሚሆኑ'፣'ላይም'፣'የሆናል'፣'ከእነዚህ'፣'ያህል'፣'ከሆነና'፣'ለሆኑት'፣'እነዚሁ'፣'እንደሆኑ'፣'ስለማናቸውም'፣'ስለዚሁ'፣'ከአንዳንድ'፣'በእነዚሁ'፣'በአምስት'፣'የሆኑበታል'፣'ለነዚህ'፣'ለማንኛውም'፣'አንደኛ'፣'ይኸኛው'፣'ከርሱ'፣'መሆኑን'፣'ለዚያው'፣'ለዚሁ'፣'ለእነርሱም'፣'እዚሁ'፣'ሐ'፣'ረ'፣'ሸ'፣'አምስት'፣'ከሶስት'፣'በተለይም'፣'በሌላ'፣'ሺህ'፣'ማናቸውንም'፣'ከአስር'፣'የማይበልጥ'፣'እንዲሁም'፣'ይህን'፣'የዚህ'፣'ማናቸውም'፣'ከስድስት'፣'መቶ'፣'ያለ'፣'አንድን'፣'ያላቸውን'፣'ሊሆን'፣'ሶስት'፣'ካልሆነ'፣'ቢያንስ'፣'ቢሆን'፣'እነዚህን'፣'አንዱን'፣'ሁለት'፣'ወይዘሮ'፣'ተብሎ'፣'ሳይሆን'፣'እንደሆነና'፣'ከብር'፣'ሆኖም'፣'የሌላ'፣'ያላቸው'፣'ይህንን'፣'ሆነው'፣'በስተቀር'፣'ስም'፣'እንደገና'፣'የማያንስ'፣'እጅግ'፣'እንዲሆን'፣'እንኳ'፣'ከሀያ'፣'ከሀምሳ'፣'ይኸው'፣'ለአንድ'፣'የሚችለውን'፣'በሚገባ'፣'ይህም'፣'እንዲሆኑ'፣'ከሌላ'፣'ለሆነ'፣'በሌሎች'፣'አንደሆነ'፣'እንዲህ'፣'በነዚሁ'፣'በእንደዚህ'፣'ስምንት'፣'ሲሆንና'፣'ምንጊዜም'፣'ለማናቸውም'፣'የአንድ'፣'እነዚህ'፣'ሲሆኑ'፣'በሁለቱም'፣'እንደዚህ'፣'የሆኑት'፣'የማናቸውም'፣'ይህንንም'፣'የአንድን'፣'በሙሉም'፣'በነዚህ'፣'የዚሁ'፣'ለእያንዳንዱ'፣'ስለሆነ'፣'መሆናቸውን'፣'ማንኛውንም'፣'ሁለቱ'፣'እንጂ'፣'ከስምንት'፣'ሁለቱንም'፣'በሁለት'፣'በአስር'፣'በሚል'፣'ቁጥር'፣'ባሉ'፣'ከመቶ'፣'እነዚህም'፣'ሲኖር'፣'ሰላሳ'፣'ለሆኑ'፣'ሰባት'፣'እንደሆነ'፣'ይህቸው'፣'ከእነዚህ'፣'ከእነዚሁ'፣'የአንቀጹ'፣'ወይ'፣'የሆነችን']

Appendix I: Amharic ABSA Hotel Reviews Annotation Guidelines

Aspect Based Sentiment Analysis (ABSA) Hotel Reviews Annotation Guidelines

Introduction

The objective of this annotation(labeling) guideline task is to recognize opinions expressed inside of the Hotel review towards particular their feature/Aspect. This labeling strategy is performed by using an Amharic linguistic expert. A feature/Aspect to be assessed can be (a) the restaurant as an entire, (b) it's ambiance (climate), (c) its area(location), (d) the service is given, (e) the food, (f) the drinks that are advertised, and (g) the price. In specific, given a Hotel review, the task of the annotator is to recognize the following sorts of information:

For the Hotel Review data, two types of information need to be annotated by the three annotators:

1. Aspect/Feature terms

Single or multiword words naming particular aspects of the target aspect. For instance, in “አገልግሎቱን እና ሰራተኞችን ወደጄ ነበር ፣ ግን ምግቡን አልወደድኩትም።”, the aspect terms are “አገልግሎት”, “ሰራተኞች” and “ምግብ”; in “ምግቡ በጣም ይጣፋጣል” the only aspect term is “ምግብ”. The task of the annotator is to identify the aspect discussed in a sentence given the following aspect categories: ሆቴል(Hotel)

-  አገልግሎት(service)
-  ክፍያ(price)
-  ምግብ(food)
-  መጠጥ(drink)
-  አካባቢ(Location)
-  ድባብ(ambiance) (sentences referring to the atmosphere and the environment of a hotel)
-  anecdotes/miscellaneous (sentences that do not belong to the above four categories)

A sentence may be classified into one or more aspects based on its overall meaning. For example, the sentence “ሆቴሉ ውድ ነው ፣ ግን ምግቡ በጣም ጥሩ ነበር” discusses the aspect categories “ሆቴል” and “ምግብ”, But this research only focused on single polarity(+ or -) determination for one or more aspect

2. Aspect/ Feature term polarity Determination

Each Feature/aspect term has to commit one of the following polarities based on the sentiment that is described in the sentence about it:

- o አዎንታዊ(positive)
- o አሉታዊ(negative)
- o Neutral

For example, in “ምግቡ በጣም ያስጠላል።”, the aspect term “ምግብ” has negative(ያስጠላል), and “መጠጥ ደስ ይላል።”, the aspect term “መጠጥ” has positive (ደስ ይላል).

Annotation guidelines for aspect terms

❖ **ምግብ(FOOD)** for opinions focusing on the food in general or in terms of specific dishes, dining options, etc. Below are some examples:

1. ምግባቸው በጣም ጥሩ ነበር። ☐ ☐ {ምግብ}
2. ምግቡ በጣም ባህላዊ ነበር። ☐ ☐ {ምግብ}
3. ፒዛው ጥሩ ነበር ። ☐ ☐ {ምግብ}

❖ **መጠጥ(DRINKS)** for opinions focusing on the drinks in general or in terms of specific drinks, drinking options, etc. Below are some examples:

1. ምግቡ ጥሩ ነው ፣ በተለይም መጠጦቼ ጣፋጭ ናቸው። ☐ ☐ {ምግብ፣ መጠጥ}
2. የወይን ዝርዝር ሰፋ ያለ እና አስደናቂ ነው። ☐ ☐ {መጠጥ}
3. መጠጡ ለወደፊት ጥሩ ይሁናል። {መጠጥ}

❖ **አገልግሎት(SERVICE)** for comment focusing on the promptness and quality of the customer/kitchen/ service, the staff’s attitude. , the food preparation, the staff’s attitude, and professionalism, the wait time, the options offered (e.g. takeout), etc. Below are some examples:

1. የሆቴሉ አገልግሎት ደስ ይላል። ☐ ☐ {አገልግሎት}

❖ **ድባብ(AMBIENCE)** for opinions focusing on the atmosphere or the environment of the restaurant's interior or exterior space (e.g. terrace, yard, garden), the décor, entertainment options, etc. Here are some examples:

1. ውስጡ በጣም ምቹ እና ሞቃት ነው። ☐ ☐ {ድባብ}
2. የመመገቢያ ክፍሉ ያለ ሙዚቃ የሚያምር ነው ። ☐ ☐ {ድባብ}

❖ **አካባቢ(LOCATION)** for opinions focusing on the location of the reviewed restaurant in terms of its position, the surroundings, the view, etc. Below are some examples:

1. እይታው አስደናቂ ነው ፣ እና ምግቡ በጣም ጥሩ ነው። ☐ ☐ {አካባቢ፣ምግብ}
2. የተመቻቸ ቦታ ነው። ☐ ☐ {አካባቢ}

❖ **ሆቴል(HOTEL)** for opinions evaluating the hotel as a whole.

1. ሽራተን ምርጥ ሆቴል ነው። ☐ {ሆቴል}
2. ምቹ ቦታ ላይ ያለ ሆቴል። ☐ {ሆቴል}

❖ **ክፍያ(PRICES)** for opinions that refer to the prices of the food, the drinks, or the hotel in general e.g.

1. ፒሳው በጣም ውድ ነው። ☐ {ክፍያ}
2. መጠጦቹ በጣም ውድ ናቸው። ☐ {ክፍያ}

Annotation guidelines for Polarity terms

A sentence has to be assigned a polarity, from a set P = {**positive, negative, neutral**}. Adjectives are words that describe or modify other words, making your writing and speaking much more specific, and a whole lot more interesting. There is a positive and negative word list in Appendix E and Appendix F.

1. ይህንን ቦታ ያስወግዱ። ☐ {አሉታዊ}
2. የሚጣፍጥ ምግብ ነው። ☐ {አዎንታዊ}
3. መጠጡ ይምራል። ☐ {አሉታዊ}
4. ምን አይነት ምግብ ነው። {Neutral}

Appendix J: Implementation of Removing Irrelevant Characters

Text Preprocessing Phase

[illegible]

Appendix K: Implementation of Normalization of Amharic Character

```
▶ def clean_text(x):
    h1=['ሀ', 'ሁ', 'ሂ', 'ሃ', 'ሄ', 'ህ', 'ሆ']
    h2=['ሐ', 'ሑ', 'ሒ', 'ሓ', 'ሔ', 'ሕ', 'ሖ']
    h3=['ገ', 'ገሐ', 'ገሂ', 'ገሃ', 'ገሄ', 'ገህ', 'ገሆ']
    a1=['አ', 'አሐ', 'አሂ', 'አሃ', 'አሄ', 'አህ', 'አሆ']
    a2=['ዐ', 'ዐሐ', 'ዐሂ', 'ዐሃ', 'ዐሄ', 'ዐህ', 'ዐሆ']
    s1=['ሰ', 'ሰሐ', 'ሰሂ', 'ሰሃ', 'ሰሄ', 'ሰህ', 'ሰሆ']
    s2=['ሠ', 'ሠሐ', 'ሠሂ', 'ሠሃ', 'ሠሄ', 'ሠህ', 'ሠሆ']
    t1=['ፀ', 'ፀሐ', 'ፀሂ', 'ፀሃ', 'ፀሄ', 'ፀህ', 'ፀሆ']
    t2=['ጸ', 'ጸሐ', 'ጸሂ', 'ጸሃ', 'ጸሄ', 'ጸህ', 'ጸሆ']
    for i in range(len(h1)):
        x=x.replace(h3[i],h1[i])
        x=x.replace(h2[i],h1[i])
        x=x.replace(t2[i],t1[i])
        x=x.replace(s2[i],s1[i])
        x=x.replace(a2[i],a1[i])
    return x
df['comment'] = df['comment'].apply(lambda x: clean_text(x))
df.head()
```

Appendix L: Implementation of Normalizing of Confusion Matrix

```
title_option=["confusion matrix, without normalization of SVM Model Based on Countvectorizer",None),
              ("normalized confusion matrix of SVM Model Based on Countvectorizer",'true')]
for title,normalize in title_option:
    disp=plot_confusion_matrix(sm,X_test,y_test,display_labels=df,cmap=plt.cm.Blues,normalize=normalize)
    disp.ax_.set_title(title)
    print(title)
    print(disp.confusion_matrix)
    plt.show()
```