

BACHELOR OF SCIENCE IN COMPUTER SCIENCE AND ENGINEERING

ASPECT-BASED SENTIMENT ANALYSIS OF HAUSA LANGUAGE

Saleem Ahmed Assan

200041257

Mounbagna Abdella Abasse

200041261

Ibrahim Nuhu

200041263

Oumar Guidado

180041250

Department of Computer Science and Engineering

Islamic University of Technology

APRIL, 2025

ASPECT-BASED SENTIMENT ANALYSIS OF HAUSA LANGUAGE

Saleem Ahmed Assan

200041257

Mounbagna Abdella Abasse

200041261

Ibrahim Nuhu

200041263

Oumar Guidado

180041250

Department of Computer Science and Engineering

Islamic University of Technology

APRIL, 2025

Declaration of Candidate

This is to certify that the work presented in this thesis is the outcome of the analysis and experiments carried out by **Saleem Ahmed Assan, Mounbagna Abdella Abasse, Ibrahim Nuhu, and Oumar Guidado** under the supervision of **Dr. Md Azam Hossain**, Associate Professor, Department of Computer Science and Engineering, Islamic University of Technology, Dhaka, Bangladesh. It is also declared that neither this thesis nor any part of it has been submitted anywhere else for any degree or diploma. Information derived from the published and unpublished work of others have been acknowledged in the text and a list of references is given.

Dr. Md Azam Hossain

Associate Professor

Department of Computer Science and Engineering

Islamic University of Technology (IUT)

Date: APRIL 29, 2025

Saleem Ahmed Assan

Student ID: 200041257

Date: APRIL 29, 2025

**Mounbagna Abdella
Abasse**

Student ID: 200041261

Date: APRIL 29, 2025

Ibrahim Nuhu

Student ID: 200041263

Date: APRIL 29, 2025

Oumar Guidado

Student ID: 180041250

Date: APRIL 29, 2025

Contents

List of Abbreviations	vi
1 Introduction	1
1.1 Motivation and Scope	3
1.2 Problem Statement	4
1.2.1 Objectives	5
1.2.2 Research Challenges	5
1.2.3 Research Gap	6
1.3 contribution	6
2 Related Works	8
2.1 Aspect-Based Sentiment Analysis (ABSA) in High-Resource Language (English)	8
2.2 Aspect-Based Sentiment Analysis for Low-Resource Languages	11
2.3 Comparative Overview: High-Resource vs. Low-Resource ABSA	15
3 Dataset Construction	17
3.1 Dataset Acquisition	17
3.2 Manual Annotation	18
3.3 Dataset Details	21
3.4 Dataset Statistics	21
4 Methodology	24
4.1 Models	24
4.2 Overview of the Model	25
4.2.1 BiLSTM	25
4.2.2 HauBERT	25
4.3 Dataset Preprocessing and Preparation	28
5 Results and Discussion	32

5.1	Experimental Setup	32
5.2	Results Presentation	34
5.2.1	Bi-LSTM Model Results	34
5.2.2	HauBERT Model Results	34
5.2.3	Average Evaluation Metrics	35
5.3	Examining the Findings	36
5.4	Difficulties and Restrictions	36
5.5	Implications and Significance of Findings	37
5.6	Chapter Conclusion	37
6	Conclusion	39
6.1	Restating Research Objectives and Questions	39
6.2	Discussion of Implications	40
6.3	Summary of Key Findings	40
6.4	Contributions of the Research	41
6.5	Acknowledging Limitations	42
6.6	Future Work	42
	References	44

List of Figures

1.1	Difference in Classification Results of SA and ABSA	3
1.2	Map showing the distribution of Hausa-speaking countries in Africa .	4
3.1	Our new annotations for students' feedback compared to annotations from <i>Rakhmanov et al.[27]</i>	20
3.2	Distribution of sentiment polarities and Aspects in our dataset	23
4.1	Model Architecture	31
5.1	Training Loss and Accuracy Curves for HauBERT. The left plot shows aspect extraction, and the right shows sentiment classification.	35

List of Abbreviations

ABSA	Aspect-Based Sentiment Analysis
ACTE	Aspect Category and Term Extraction
AE	Aspect Extraction
AfriBERTa	African Bidirectional Encoder Representations from Transformers
AfroXLM-R	African Cross-Lingual Language Model RoBERTa-based
AI	Artificial Intelligence
APD	Aspect Polarity Detection
BERT	Bidirectional Encoder Representations from Transformers
BiLSTM	Bidirectional Long Short-Term Memory
BIO	Beginning, Inside, Outside (annotation tagging scheme)
CoConABSA	ConcisePrompt and Contrastive Alignment Aspect-Based Sentiment Analysis
CORSA	Conditional Relation and Sentiment Analysis (multi-task framework)
CRD	Conditional Relation Detection
CRF	Conditional Random Field
CrosGrpsABS	Cross Graph Attention-Based Sentiment Analysis Model
DualGCN	Dual Graph Convolutional Network (SynGCN + SemGCN)
EN	English (language code)
FN	False Negative
FP	False Positive
FR	French (language code)
HauBERT	Hausa Bidirectional Encoder Representations from Transformers
HESAC	Hausa-English Sentiment Analysis Corpus
HESAC-ABSA	Hausa-English Sentiment Analysis Corpus for Aspect-Based Sentiment Analysis
HP-LSTM	Hierarchical Parallel Long Short-Term Memory
IAN	Interactive Attention Network
ICL	In-Context Learning
LLaMA	Large Language Model Meta AI
LLM	Large Language Model
LSTM	Long Short-Term Memory
MAMS	Multi-Aspect Multi-Sentiment Dataset
ML	Machine Learning
mBERT	Multilingual BERT

NLP	Natural Language Processing
PLM	Pretrained Language Model
P-SUM	Parallel Summarization (aggregation method with BERT+CRF)
REC	Relational Extraction Condition
RNN	Recurrent Neural Network
SA	Sentiment Analysis
SAfriSenti	South African Sentiment Dataset (Twitter corpus)
SemEval	Semantic Evaluation (International NLP Benchmark Competition)
SemGCN	Semantic Graph Convolutional Network
SERENGETI	South and Eastern African Language Representation Model
SynGCN	Syntactic Graph Convolutional Network
TASD	Target-Aspect Sentiment Detection
TN	True Negative
TP	True Positive
UzABSA	Uzbek Aspect-Based Sentiment Analysis Dataset
VOL	Volatility (task within CORSA framework)
XGBoost	Extreme Gradient Boosting
YOLOv8	You Only Look Once, version 8 (object detection model)

Acknowledgement

We are profoundly grateful to our parents for their unconditional love and support, for giving us the opportunity to continue our studies. We are truly grateful.

To our supervisor, Dr. Azam Hossain , for his endless patience, insightful critiques, and for always guiding us. His expertise and encouragement were instrumental in completing this thesis.

A heartfelt appreciation to all.

Abstract

Aspect-Based Sentiment Analysis (ABSA) enables fine-grained understanding of textual data by identifying specific aspects of a target entity and determining the sentiment expressed toward each aspect. Despite significant advances in high-resource languages such as English, low-resource languages like Hausa remain underrepresented, particularly in educational contexts. This study addresses this gap by developing an annotated Hausa-English ABSA dataset derived from an existing Hausa-English sentiment dataset and enhanced for aspect-based analysis. The dataset was manually annotated with aspect categories and sentiment polarities to improve analytical depth and applicability. Five key aspect categories were defined to reflect the educational domain: lecturer behavior, course content, course difficulty, teaching quality, and teaching performance. Each instance was annotated with an aspect category and the corresponding sentiment polarity (positive, neutral, or negative) providing a comprehensive resource for detailed sentiment interpretation.

The objectives of this study include: (i) developing an annotated dataset for ABSA in the Hausa language using student comments; (ii) implementing machine learning models capable of extracting aspect terms and classifying their sentiment polarities; (iii) evaluating the performance of the models in low-resource settings for ABSA tasks; and (iv) contributing to the development of NLP tools and resources for underrepresented African languages, with a focus on educational feedback analysis. Model performance was assessed using standard metrics, including accuracy, precision, recall, and F1-score. The results highlight challenges in handling explicit and implicit aspect mentions in low-resource languages and underscore the importance of robust annotation frameworks. This work establishes the first structured Hausa ABSA dataset in the educational domain, providing a foundation for future multilingual and cross-lingual ABSA research in African languages.

Chapter 1

Introduction

With the proliferation of textual data in the digital age, Natural Language Processing (NLP) has emerged as a crucial area of study in computer science and artificial intelligence [11]. Machines can now comprehend and produce responses that are similar to those of humans thanks to their capacity to process, analyze, and interpret human language. Large volumes of unstructured text have been produced as a result of the growth of digital platforms like social media, online reviews, and user feedback systems; therefore, natural language processing (NLP) is crucial for gleaning valuable information from this deluge of data. Sentiment analysis (SA), one of the most influential subfields of natural language processing (NLP), is concerned with examining text to ascertain the emotional tone that a piece of information (digital, text, etc) conveys. This falls into one of three general categories: neutral, negative, or positive [2], [23], [28].

A key component of sentiment analysis is sentiment polarity, which establishes whether the sentiment conveyed in a text is neutral, negative, or positive. In order to better capture the subtleties of emotion in textual data, sentiment analysis can be more granular and involve multiple classes (e.g., strong positive, positive, neutral, negative, and strong negative) [31]. In its most basic form, sentiment analysis can be binary, identifying only positive or negative sentiment. Businesses, politicians, and organizations can better understand social media trends, consumer feedback, and public opinion by using this polarity classification. Sentiment analysis is used, for instance, in customer service to evaluate product reviews, monitor public opinion during political campaigns, and even identify possible public relations crises on social media [37].

However, the subtleties that emerge when several sentiments are expressed toward various aspects of a good, service, or event within a single text are frequently missed

by conventional sentiment analysis methods, which concentrate on general sentiment polarity. A simple sentiment analysis system might, for example, only return a neutral sentiment for the entire sentence in the sentence "The food at the restaurant is delicious, but the service is bad," ignoring the different sentiments toward "food" (which is positive) and "service" (which is negative). This illustrates a significant drawback of conventional sentiment analysis, which is that it can produce unduly simplistic interpretations when several sentiments coexist in a single sentence. Figure 1.1 provides a clear illustration of this problem by visually separating aspect-based sentiment analysis (ABSA) from traditional sentiment analysis.

A more advanced method within the larger sentiment analysis field, aspect-based sentiment analysis (ABSA), was introduced in response to these limitations. Beyond evaluating the sentiment of a text as a whole, ABSA determines the sentiment associated with particular features or aspects that are mentioned in the text. For instance, in the aforementioned sentence, ABSA would acknowledge that the sentiment regarding product quality is positive, but the sentiment regarding delivery is negative. Absa offers more in-depth, fine-grained insights by breaking the text down into different parts and examining the sentiment connected to each, which enables a deeper comprehension of the attitudes and feelings that users have expressed [4], [6], [24], [35]. This method is especially useful in applications like opinion mining, customer feedback analysis, and product reviews, where knowing how people feel about particular features can provide useful information for targeted marketing and product enhancements.

The development of sophisticated NLP techniques like Aspect-Based Sentiment Analysis (ABSA) has been highly skewed toward high-resource languages such as English and Mandarin, creating a significant digital linguistic divide. In contrast, many low-resource languages, particularly in South Asia and sub-Saharan Africa, remain critically underserved. This disparity is exacerbated by a scarcity of annotated corpora and foundational linguistic tools, which are prerequisites for data-intensive model training.

Although recent research has initiated efforts to bridge this gap through cross-lingual transfer learning and the creation of benchmark datasets, substantial challenges persist. Significant methodological and application gaps continue to hinder the development of robust, equitable sentiment analysis technologies for these diverse linguistic communities.

Using aspect-based sentiment analysis (ABSA) techniques on the Hausa language, a low-resource language spoken mostly in West Africa, this study aims to close this

gap. Because of its intricate syntax and scarcity of high-quality annotated datasets, the Hausa language poses particular difficulties for sentiment analysis. In order to improve the accessibility of sentiment analysis tools in these communities, this research intends to advance NLP for underrepresented languages by adapting ABSA techniques to Hausa. This will enhance the linguistic resources available for further study in addition to facilitating more equitable participation in the digital world.

In addition to laying the groundwork for future research in low-resource languages and showcasing the value of ABSA in more complex sentiment analysis, this study seeks to advance the larger objective of linguistic inclusivity in NLP by concentrating on a language with few available resources.

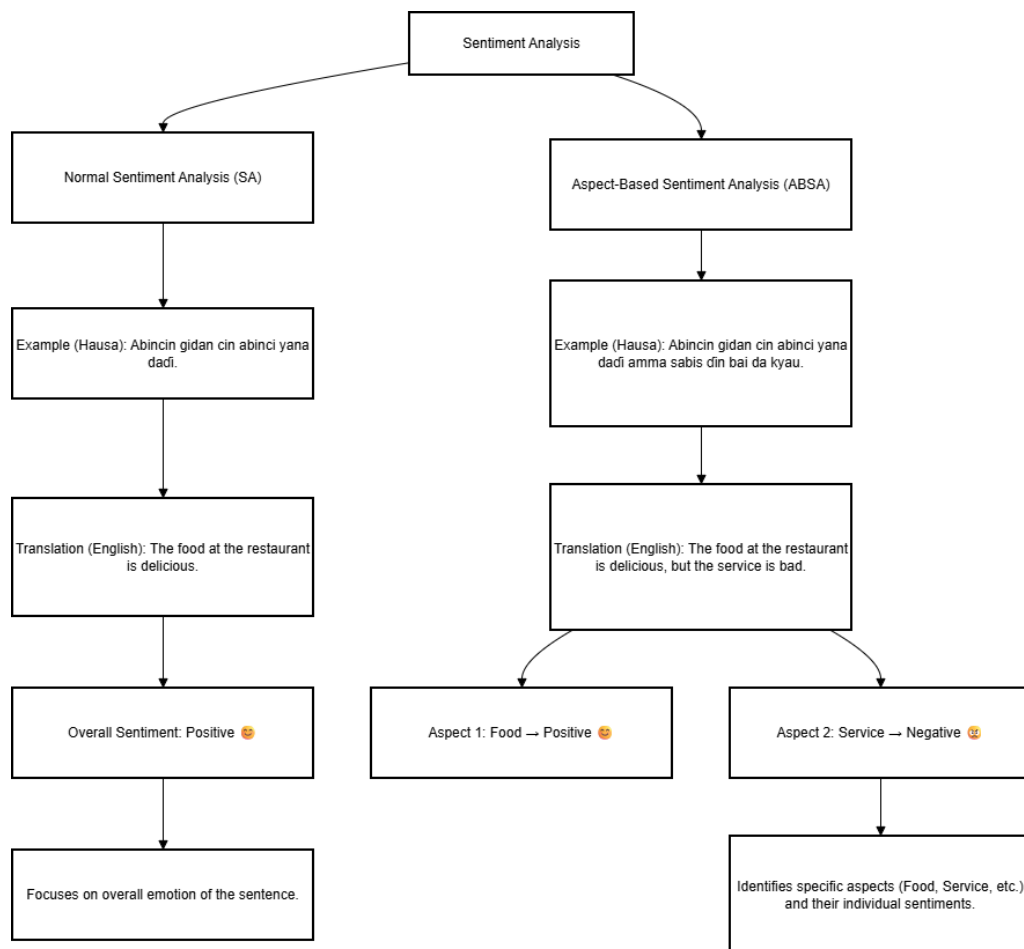


Figure 1.1: Difference in Classification Results of SA and ABSA

1.1 Motivation and Scope

Hausa is one of the predominantly spoken languages in West Africa Figure 1.2. It is used in many aspects of day-to-day life in the area: business, education, governance,

etc. Despite this, it's still quite underrepresented in the domain of natural language processing research. Given that it has wide usage across many fields, there is a need for tools that can understand fine-grained sentiments across them. This study aims to improve feedback analysis in education and expand Hausa's NLP capabilities, supporting applications in healthcare, governance, business, and beyond.

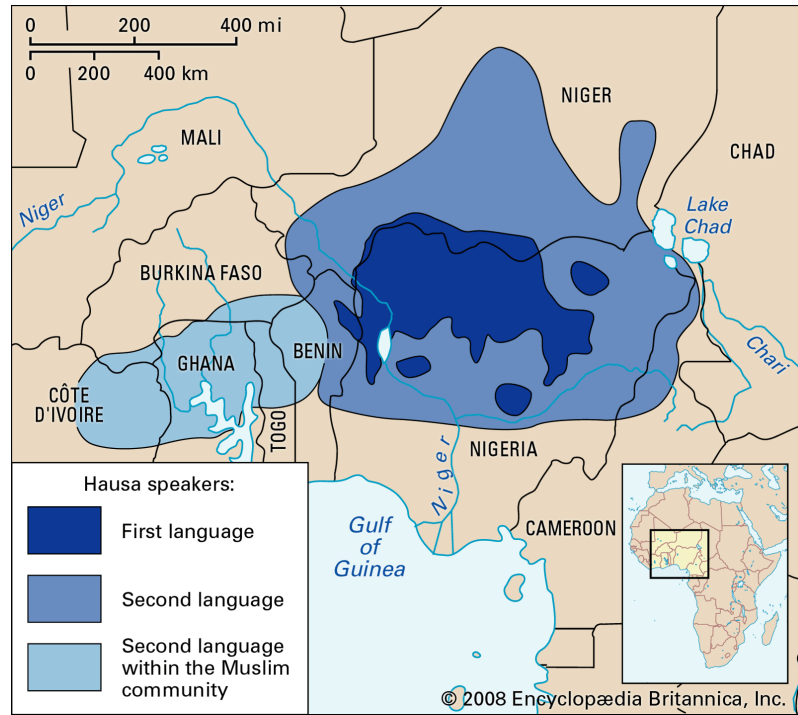


Figure 1.2: Map showing the distribution of Hausa-speaking countries in Africa

^a<https://www.britannica.com/topic/Hausa-language>

1.2 Problem Statement

Despite the increasing need for opinion mining tools in educational and other feedback-based systems, the Hausa language, spoken by millions of West Africans, continues to be underrepresented in the research of natural language processing (NLP). Current sentiment analysis initiatives in Hausa focus mainly on general polarity classification and cannot identify sentiments directed toward specific aspects within a text. This shortcoming hinders educational institutions and other sectors from obtaining detailed insights essential for targeted enhancements. Additionally, the lack of annotated datasets and models for aspect-based sentiment analysis (ABSA) in Hausa presents a considerable obstacle to creating precise, intelligent systems that comprehend the nuanced opinions conveyed in the language.

1.2.1 Objectives

1. To develop an annotated dataset for Aspect-Based Sentiment Analysis in the Hausa language using student comments.
2. To design and implement machine learning models capable of extracting aspect terms and classifying their corresponding sentiment polarities in Hausa text.
3. To evaluate the performance of different language models (e.g., AfroXLM-R, AfriBERTa) in low-resource settings for ABSA tasks.
4. To contribute to the development of NLP tools and resources for underrepresented African languages, with a focus on educational feedback analysis.

1.2.2 Research Challenges

1. **Lack of Annotated Datasets:** Hausa lacks large, high-quality datasets annotated with aspect terms and sentiment polarities, which are essential for training and evaluating ABSA models[6].
2. **Low-Resource Language Status:** Hausa is considered a low-resource language in NLP, with limited language models, lexicons, and linguistic tools available compared to high-resource languages like English or Chinese[20].
3. **Complex Morphology:** Hausa exhibits rich morphological structures, such as affixation and reduplication, making aspect term extraction more challenging for standard models[24].
4. **Code-Switching and Dialectal Variations:** Hausa speakers often mix English or other local languages in their texts, especially in student comments, complicating aspect and sentiment detection[23].
5. **Scarcity of Pretrained Models:** Few pretrained language models exist that are optimized for fine-grained tasks like ABSA in Hausa[24].
6. **Limited Sentiment Lexicons:** The absence of comprehensive sentiment lexicons in Hausa hinders lexicon-based sentiment analysis approaches[4].
7. **Ambiguity in Aspect Identification:** Aspects in student feedback are frequently implied or context-dependent, making automatic aspect extraction more difficult.
8. **Annotation Difficulties:** Developing ABSA datasets requires native speakers and domain experts, making the annotation process costly and time-consuming[35].

1.2.3 Research Gap

Although there have been numerous investigations into sentiment analysis within the Hausa language, most concentrate on **general polarity classification (positive, negative, or neutral)** at the levels of sentences or documents. These methods fail to capture the **aspect-level sentiments** present in a text, where various segments of a sentence may convey differing opinions on distinct topics.

Furthermore, although notable advances have been made in developing general sentiment analysis models for Hausa using data from Twitter and news sources, there is a lack of significant research focused on creating aspect-based sentiment analysis (ABSA) models for Hausa, particularly in the context of education. The absence of readily available annotated datasets, tools for aspect extraction, and techniques for fine-grained analysis in the Hausa language further emphasizes the limited resources in this domain.

Consequently, a distinct gap exists in current research: the need for detailed aspect-based sentiment analysis methodologies for Hausa texts, supported by specifically developed datasets and models that can identify and classify sentiments related to specific aspects, particularly in feedback and other formal settings.

The objective of this research is to address that gap by establishing an ABSA framework for Hausa, creating annotated resources, and assessing contemporary NLP techniques tailored to low-resource environments.

1.3 contribution

- **Aspect-Based Sentiment Analysis (ABSA) Study for Hausa:** This research is among the first to apply Aspect-Based Sentiment Analysis specifically to the Hausa language, moving beyond general sentiment polarity detection to fine-grained aspect-level sentiment classification.
- **Creation of an Annotated Hausa ABSA Dataset:** A new dataset of student comments in Hausa was manually annotated with aspect terms and their associated sentiment polarities, providing a foundational resource for future Hausa ABSA research.
- **Adaptation and Fine-Tuning of Pretrained Multilingual Models for Hausa ABSA:** Multilingual pretrained models such as mBERT and LSTM were fine-tuned to perform aspect extraction and sentiment classification on Hausa text, addressing the low-resource nature of the language.

- **Handling Linguistic Challenges Unique to Hausa:** Techniques were developed to manage Hausa-specific challenges such as complex morphology, code-switching with English, and dialectal variations within the student comment dataset.
- **Framework for Future Low-Resource ABSA Research:** The research presents a replicable methodology that can guide the development of ABSA systems for other low-resource African languages, supporting broader efforts to make NLP more inclusive.

Chapter 2

Related Works

To place this research within a larger academic framework, it is essential to understand the landscape of Aspect-Based Sentiment Analysis (ABSA) in languages with limited resources. This chapter offers a critical analysis of the most relevant research on sentiment analysis, ABSA methodologies, and the particular difficulties faced by languages like Hausa. It also compares the two environments after analyzing Aspect-Based Sentiment Analysis (ABSA) in high-resource and low-resource languages. This analysis reveals important gaps that underscore the need for language-specific approaches to ABSA.

2.1 Aspect-Based Sentiment Analysis (ABSA) in High-Resource Language (English)

In this section, we evaluate aspect-based sentiment analysis for high-resource languages, with the English language being the focus. With an emphasis on the datasets, annotation techniques, models, and other elements of each study, this section assesses a variety of ABSA for English studies. All this can be found in table 2.1.

Table 2.1: Summary of ABSA Research Including New Contributions

Paper	Dataset	Annotation	Model	Domain	Polarity	Aspects
Karimi et al. (2021) [15]	SemEval 2014 Laptop, Restaurant	Manual (BIOES-style) format for AE	BERT + CRF with Parallel (P-SUM) and Hierarchical (H-SUM) Aggregation	Product Reviews	Positive, Negative, Neutral	Laptop, Food, Service, Restaurant

Paper	Dataset	Annotation	Model	Domain	Polarity	Aspects
<i>Bao et al.</i> (2025) [5]	ACOS (Restaurant, Laptop), en-Phone, Rest15/16	Manual (Annotated quadruples)	LLaMA-2-7B + Hybrid Sampling (Element, Quadruple, Neighbor, Structural Rollback)	Product Reviews	Positive, Negative, Neutral	Food, Price, Service, Ambience, Menu
<i>Ma et al.</i> (2017) [18]	SemEval 2014 (Restaurant, Laptop)	Manual	Interactive Attention Networks (IAN) with LSTM	Restaurant, Laptop	Positive, Neutral, Negative	Food, Service, Ambience, Price, Environment
<i>Liu et al.</i> (2025) [17]	Twitter-15, Twitter-17	Auto (Conditional relation via REC; object bounding boxes via YOLOv8)	CORSA: Multi-task (CRD + VOL + multimodal BART)	Social Media (Multimodal)	Positive, Negative, Neutral	Product Features, Brand, Service, Customer Experience
<i>Li et al.</i> (2021) [16]	SemEval 2014 (Restaurant, Laptop), Twitter	Manual (Standard ABSA labels)	DualGCN: SynGCN + SemGCN with BiAffine, Orthogonal	Differential Regularizers	Positive, Negative, Neutral	Laptop, Restaurant, Food, Service, Ambience
<i>Ruder et al.</i> (2016) [32]	SemEval 2016 (11 multilingual datasets)	Manual (Aspect-level polarity per sentence)	HP-LSTM: H-LSTM with pre-trained multilingual embeddings	Restaurants, Hotels, Laptops, Phones, Cameras	Positive, Negative, Neutral	Hotel, Phone, Laptop
<i>Chen et Al</i> (2018) [7]	BeerAdvocate (unsupervised), SemEval 2014 (Restaurant, Laptop)	Unsupervised (latent aspect-sentiment trees)	DOT: Discrete Opinion Tree (VAE with Gumbel-Softmax)	Product Reviews	Positive, Negative, Neutral	Food, Service, Price
<i>Akhtar et al.</i> (2019) [3]	SemEval 2016 (EN, ES, FR, DU), GermEval 2017 (DE), Akhtar et al. Hindi Dataset	Manual (BIO for AE, polarity labels)	Language-Agnostic Bi-LSTM (A1A4 architectures) + Handcrafted Features	Multilingual Reviews	Positive, Negative, Neutral	Product, Service
<i>Neveditsin et al.</i> (2025) [25]	SemEval-14, MAMS, Twitter, Synthetic Sports Feedback	Hybrid (LLM-generated + human-verified)	Mistral 7B, LLaMA-3 8B (ICL, fine-tuned, blended)	Product Reviews, Sports Feedback	Positive, Negative, Neutral	Service, Customer Experience
<i>Wang et al.</i> (2024) [10]	SemEval 2014 (Restaurant, Laptop), Twitter, Co-ConABSA Benchmark	Manual (for supervised), Unsupervised (pretraining), Synthetic Augmentation	CoConABSA: Transformer with ConcisePrompt and Contrastive Alignment	Product	Social Media Reviews	Food, Service, Restaurant, Price, Product

Product reviews, social media posts, and domain-specific texts are the usual sources of datasets used in ABSA research. Aspect-Based Sentiment Analysis (ABSA) research landscape has been profoundly shaped by the SemEval shared task series. Benchmarks established in tasks such as SemEval-2014 and SemEval-2016 [26] provided

foundational annotated datasets for domains like restaurant and laptop reviews, which have since become standard resources for model development and evaluation. These datasets are manually annotated with sentiment polarity and aspect terms; aspect extraction (AE) frequently use BIO-format annotations.

Karimi et al [15], for instance, used the SemEval 2014 datasets for restaurants and laptops, where reviews were manually annotated in BIO-format for the purpose of extracting aspects. They used BERT-enhanced CRF models to study sentiment and aspect-based classification. Similar to this, Li et al [16] applied standard ABSA labels for elements like food, service, and ambience using the SemEval 2014 dataset and Twitter reviews.

Other datasets, such as ACOS and BeerAdvocate, are also important. Bao et al [5] worked with the ACOS dataset, which contains annotations for product reviews in quadruple format (entity, aspect, sentiment, and opinion). Chen et al [7] used the BeerAdvocate dataset to investigate unsupervised aspect extraction using latent aspect-sentiment trees.

In addition, since deep learning and transformer-based models were introduced, ABSA has undergone significant changes. Traditional machine learning algorithms like Conditional Random Fields (CRF) and Support Vector Machines (SVM) were frequently used in early methods. To improve performance, more recent research has used transformer models like BERT, LLaMA, and multi-task learning frameworks, particularly when working with large datasets.

Bidirectional Encoder Representations from Transformers, or BERT, is now the foundation of many ABSA systems. For aspect extraction and sentiment classification, Karimi et al.[15] combined BERT and CRF with hierarchical aggregation techniques. With high F1 scores for aspect extraction and precision for sentiment classification, this method demonstrated excellent performance on the SemEval 2014 datasets. Similarly, Ruder et al. [32] addressed ABSA tasks across multiple languages and domains by utilizing pre-trained BERT embeddings in their hierarchical LSTM model.

Bao et al [5] and Neveditsin et al [25] have employed LLaMA (Large Language Model Meta AI), another cutting-edge model in the ABSA landscape, for aspect and sentiment classification. These models employ hybrid sampling methods to handle different levels of granularity in aspect extraction. LLaMA-2-7B was specifically used by Bao et al [5], who concentrated on quadruple-based annotations and achieved good results in a variety of domains, such as food, price, and service.

Other models, such as CORSA Liu et al [17], integrate textual sentiment analysis with

visual features in multimodal tasks. The model’s capacity to examine elements like product attributes and customer experience from social media data is improved by this multimodal approach.

Another significant contribution is DualGCN (Graph Convolutional Networks) by Li et al [16], which employs a hybrid graph-based method for aspect extraction. Their model beats conventional techniques on the SemEval 2014 and Twitter datasets by utilizing syntactic and semantic graph convolutions.

Rich datasets, sophisticated machine learning methods, and the use of transformer models like BERT and LLaMA have all contributed significantly to the development of ABSA in English. These studies show how flexible ABSA models are in a variety of contexts, including social media and product reviews, and how well-suited contemporary models are for extracting fine-grained details like food, service, and product attributes. Future studies will probably concentrate on improving aspect extraction and sentiment classification in various languages and contexts, as well as honing these models to handle noisy and multimodal data, as the field develops.

These studies all offer insightful perspectives on ABSA, and their application of transformer-based models has expanded the realm of what is feasible in sentiment analysis for high-resource languages such as English.

2.2 Aspect-Based Sentiment Analysis for Low-Resource Languages

In this section, we take a look at aspect-based sentiment analysis for low-resource. We focus on how different models have been used to achieve aspect extraction and sentiment classification [21], [24]. Here we are going to take a look at languages from Asia, Europe, and Africa. We’ll look at their datasets, annotations, language, aspects, models. This is summarized in Table 2.2

Table 2.2: Summary of ABSA Research on Low-Resource Languages

Paper	Dataset	Annotation	Aspects	Polarity	Language
Bhattacharya et al. (2021) [6]	Translated SemEval-2014 (Hindi)	Manual (native speakers; inter-annotator agreement reported)	Laptop/Restaurant aspects (e.g., food, service, price)	(focus on extraction only)	Hindi
Regatte et al. (2020) [29]	Telugu movie reviews	Manual	Aspect terms, categories (generic review settings)	Positive, Negative, Neutral	Telugu

Paper	Dataset	Annotation	Aspects	Polarity	Language
<i>Dewangan et al.</i> (2025) [8]	Odia ABSA Benchmark (Ethnic food, Handloom, Handicrafts, Hotels, Electronics, Books, Movies)	Manual (native Odia speakers with expert validation; augmented with back-translation and T5 paraphrasing)	Food, Service, Product Quality, Price, Ambience, Cultural elements	Positive, Neutral	Negative, Odia
<i>Neveditzin et al.</i> (2025) [25]	Synthetic Sports Feedback + Composite (SemEval-14, MAMS, Twitter)	Hybrid (LLM-generated + human refined)	Player Performance, Coaching, Teamwork, Game Strategy, Refereeing	Positive, Neutral	Negative, English (sports domain)
<i>Matlatipov et al.</i> (2024) [21]	Uzbek ABSA Dataset	Manual (native Uzbek annotators)	Food, Service, Price, Ambience	Positive, Neutral	Negative, Uzbek
<i>Koto et al.</i> (2024) [4]	Multiple (Urdu, Hindi, Arabic, Bengali, Tamil, Vietnamese, etc.)	Mixed (manual, translation, weak supervision)	Product, Service, Experience	Positive, Neutral	Negative, Multiple (20+ low-resource languages)
<i>Akhtar et al.</i> (2016) [2]	Hindi ABSA Dataset (5417 sentences across 12 domains: laptops, mobiles, tablets, cameras, home appliances, apps, movies, etc.)	Manual (native Hindi speakers, SemEval-style, Cohens Kappa 95%)	Audio Quality, Weight, Screen, Price, Camera, Travel aspects	Positive, Neutral, Conflict	Negative, Hindi
<i>Hossain et al.</i> (2025) [13]	Bengali ABSA (Car, Mobile, Movie, Restaurant) + SemEval 2014 EN (Laptop, Restaurant)	Manual (Bengali); Standard Benchmark (English)	Product, Service, Experience	Positive, Neutral	Negative, Bengali, English
<i>Sotibe et al.</i> (2025) [30]	SAfriSenti (Sepedi, Sesotho, Setswana, isiXhosa, isiZulu Twitter Corpus)	Semi-automatic (emoji/lexicon-based distant supervision + manual verification)	Product, Service, Experience	Positive, Neutral	Negative, Sepedi, Sesotho, Setswana, isiXhosa, isiZulu
<i>Salahudeen et al.</i> (2023) [33]	AfriSenti-SemEval 2023 (Hausa Tweets)	Manual (native speakers, SemEval guidelines)	Service, Politics, Entertainment, General social topics	Positive, Neutral	Negative, Hausa
<i>Musa et al.</i> (2025) [24]	Hausa Movie Reviews	Manual (native Hausa speakers)	Movie, Acting, Direction, Storyline	Positive, Neutral	Negative, Hausa

Because low-resource ABSA datasets are usually small, domain-skewed, and expensive to label, authors rely on native-speaker annotators, careful schema design, and

augmentation or hybrid (LLM+human) pipelines whenever feasible. One recent example is the Odia benchmark by Dewangan et al [8], which uses native Odia speakers with expert validation to label both aspects and sentence-level sentiment in seven culturally salient domains: ethnic food, handloom, handicrafts, hotels, electronics, books, and movies. In order to combat sparsity, they experiment with data augmentation techniques like back-translation and paraphrasing, while maintaining a clean, manually verified gold set for evaluation. The task framing is SemEval-style (three polarities). In order to preserve inter-annotator consistency, the paper documents annotation guidelines and explicitly motivates domain choice by covering commonplace Odia consumer discourse.

Where human labeling is hard to scale, several teams explored hybrid supervision. Neveditsin et al [25] build an ABSA-like resource for sports feedbacka domain under-represented in classic review corporaby drafting large language models to propose candidate spans/labels and then using human raters to refine them. They also blend this with composites of existing English datasets (SemEval-14, MAMS, Twitter) to study generalization. The result is a pipeline that yields aspect terms such as player performance, coaching, refereeing, and tactics, while preserving standard three-way sentiment. Their analysis stresses the gains (coverage, speed) and the risks (LLM bias), and reports quality control via human verification thresholds.

In other papers a number of teams investigated hybrid supervision in situations where human labeling is difficult to scale. By creating extensive language models to suggest potential spans or labels and then employing human raters to refine them, Neveditsin et al. [25] create an ABSA-like resource for sports feedback, a field that is under-represented in traditional review corpora. In order to investigate generalization, they also combine this with composites of preexisting English datasets (SemEval-14, MAMS, Twitter). The end product is a pipeline that maintains the typical three-way sentiment while producing aspect terms like player performance, coaching, refereeing, and tactics. In addition to reporting quality control through human verification thresholds, their analysis highlights the risks (LLM bias) and the benefits (speed, coverage).

In Uzbek, Matlatipov et al [21] present UzABSA, the first publicly available ABSA dataset with native Uzbek annotators that adhere to SemEval-style standards. The corpus focuses on traditional review categories with three polarities, such as food, service, price, and ambience in restaurants or product scenarios. The resource is positioned by the authors as a platform for cross-lingual transfer from more robust English models, and they document label definitions and inter-annotator procedures. Their

publication stresses reproducibility and provides baselines, which is especially crucial for under-resourced languages with little previous research.

A broader perspective is offered by survey-style initiatives. Koto et al [4] summarize how manual labeling, translation-projection, and weak supervision are combined to obtain aspect categories like product, service, and user experience across social media and review domains. They do this by synthesizing methods and data across more than 20 low-resource languages (South/Southeast Asian and beyond). In addition to listing common modeling options (mBERT/XLM-R transfer, adapters), the survey identifies common problems with multilingual ABSA, such as domain shift, code-mixing, and divergent aspect taxonomies.

One of the first Hindi ABSA datasets in South Asia is provided by Akhtar et al [2] and consists of 5,417 sentences across 12 domains (e.g., laptops, mobiles, tablets, cameras, appliances, apps, movies, and travel). Annotation was conducted by native speakers under SemEval-style guidelines, with reported high agreement (Cohens 0.95). Crucially, the aspect schema encompasses tangible categories like *audio quality*, *weight*, *screen*, *price*, and *camera*, and the Hindi labels include the *conflict* class in addition to the typical *positive/negative/neutral*. This resource is now used as a benchmark for later low-resource and cross-lingual ABSA projects.

Hossain et al. [13] use CrosGrpsABS, a model that employs cross-attention on top of a Transformer backbone and overlays syntactic and semantic graphs, to target both sentiment classification and aspect extraction for Bengali. Their assessment includes the English SemEval-2014 laptop/restaurant benchmarks, maintaining three polarities, and four Bengali domains: car, mobile, movie, and restaurant (manually annotated). The study describes their inventories of aspect categories and demonstrates how graph structure aids in the deciphering of aspect terms that are incorporated into a flexible Bengali word order.

There are two noteworthy complementary threads regarding African languages. First, Sotibe et al [4] compare fine-tuned multilingual PLMs (AfriBERTa, AfroXLM-R), occasionally ensembling with XGBoost, and analyze SAfriSenti, a multi-language Twitter collection (Sepedi, Sesotho, Setswana, isiXhosa, and isiZulu). The authors outline useful heuristics and verification to maintain labels in reliable social data, even though the shared task datasets are mostly sentiment (three-class) and not fully labeled for aspect terms. Their analysis highlights the influence of domain-specific pre-training as well as language-family effects (Nguni vs. Sotho-Tswana). Secondly, the primary Hausa Twitter sentiment sets utilized by numerous teams, such as Salahudeen et al [33] (HausaNLP), are provided by the AfriSenti SemEval-2023 shared task. Teams

frequently add Hausa-specific tokenization and adapter tuning, and Task A covers more than 12 languages with positive, negative, and neutral labels. Datasets are curated from Twitter with privacy-preserving preprocessing. Many system papers discuss aspect surrogates (hashtags or entity cues) to steer models, even though the official task is sentiment-only. The data creation for the shared task is mentioned in the HausaNLP paper, which also details their pipeline.

Last but not least, Musa et al [24] focus on Hausa film reviews and specifically frame ABSA using aspect categories like film, acting, direction, and plot, which are annotated by native Hausa speakers with three polarities. They provide two contributions: (i) a Hausa ABSA dataset that is domain-focused and enhances overall sentiment on Twitter, and (ii) HauBERT, a Hausa-adapted Transformer optimized for sentiment classification and aspect extraction, which is helpful in situations where standard multilingual PLMs perform poorly on genre-specific vocabulary.

2.3 Comparative Overview: High-Resource vs. Low-Resource ABSA

In terms of data resources, annotation techniques, and modeling strategies, aspect-based sentiment analysis (ABSA) in high-resource and low-resource languages differs significantly. The foundation for extensive experimentation with sophisticated models, such as attention-based LSTMs [18], graph convolutional networks [16], and BERT-based systems with advanced aggregation mechanisms [15], is provided by standardized and large-scale benchmarks like the SemEval datasets [26] in high-resource languages like English. Research on prompt-driven transformers [10] and multimodal frameworks [17] has also been made possible by the abundance of annotated data, allowing for reliable and repeatable evaluation across a variety of domains, including social media, laptops, and restaurants. Low-resource ABSA, on the other hand, is subject to far stricter regulations. Hausa movie reviews are an example of a dataset that is smaller and frequently domain-specific.[24], the Hindi ABSA corpus covering twelve product categories [2] and the Odia benchmark covering seven cultural domains [8]. Annotation quality and transparency are crucial in these situations: While Neveditsin et al. [25]. (2025) investigate hybrid annotation pipelines where large language models propose drafts that are refined by human annotators, Dewangan et al. [8]. (2025) report inter-annotator agreement and use back-translation and paraphrasing for augmentation [25]. These limitations are also reflected in model development, as researchers mainly rely on multilingual pretrained models like XLM-R and

AfriBERTa [21], [33]. These models are frequently improved with parameter-efficient tuning or domain-specific ongoing pretraining, as shown by HauBERT for Hausa. Whereas high-resource ABSA is driven by architectural innovations and benchmark-driven competition, low-resource ABSA emphasizes careful dataset design, culturally grounded aspect taxonomies, augmentation methods, and cross-lingual transfer. This contrast highlights that advancing ABSA for languages like Hausa requires adopting the best practices of high-resource settings in terms of annotation and evaluation, while tailoring modeling approaches to the unique constraints and linguistic realities of low-resource contexts.

Chapter 3

Dataset Construction

A crucial first step in developing Aspect-Based Sentiment Analysis (ABSA) is creating high-quality datasets, particularly for low-resource languages. Low-resource languages encounter difficulties such as a lack of standardized aspect categories, inconsistent sentiment labeling, and a lack of well-annotated datasets, in contrast to high-resource languages, which have numerous benchmark datasets. These restrictions limit the possibility of conducting fruitful cross-lingual research and impede the creation of reliable ABSA models.

The approach used to create and annotate a Hausa ABSA dataset in the educational field is presented in this chapter. After describing how comparable efforts in other low-resource languages have influenced our design decisions, we go on to describe how raw data is obtained, preprocessed, and then annotated. The chapter also describes our annotation framework, which includes the definition procedure. It also describes our chosen annotation framework, which includes polarity assignment, aspect category definitions, and annotated data examples. Our goal is to provide a dataset that supports Hausa research and makes cross-lingual comparisons with other ABSA resources, like the SemEval datasets, easier by adhering to this structured process.

3.1 Dataset Acquisition

We aim to build a Hausa dataset for ABSA within the education domain, utilizing the format consistent with Semeval-2015 [26] dataset, whose annotation format is very useful in addressing more complex ABSA tasks like Aspect Polarity Polarity Detection (APD), Aspect Category and Term Extraction (ACTE), Target-Aspect Sentiment

Detection(TASD). Additionally, following this format allows future cross-lingual experiments between Hausa, English and other languages [35].

Our newly created dataset utilized *Rakhamov et al [27]* existing publicly available dataset, the Hausa-English Sentiment Analysis Corpus for Educational Environments (HESAC), which contains 40,000 English student comments collected from the 2018\2019 course evaluation database at Nile University in Nigeria. The dataset was originally labeled with three sentiment classes (negative, neutral, positive). As with other annotated data collections, the labels were cross-checked for accuracy [27]. To produce comments in Hausa, each comment was first machine-translated and then corrected by three PhD students from Nile University with excellent proficiency in both Hausa and English [27].

The HESAC dataset was not initially optimized for ABSA, as it lacked aspect-specific annotations. To address this, we converted the dataset into an ABSA-optimized dataset by manually annotating it for specific aspects while retaining the original sentiment polarity(positive-3, neutral-2, negative-1).

The new Hesac-Absa dataset comprises 10,000 manually annotated student comments.

3.2 Manual Annotation

Four final-year Hausa-speaking native students manually annotated the dataset. The original HausaEnglish parallel dataset from *Rakhmanov et al. (2022) [27]* was used by the annotators, who modified it to meet the needs of Aspect-Based Sentiment Analysis (ABSA). Finding pertinent aspect categories and allocating sentiment polarities were the main annotation tasks because the dataset was already in Hausa with aligned English translations.

The annotators developed five core aspect categories-**lecturer behavior, course content, course difficulty, teaching quality, and teaching performance**. Through an iterative review process that captures the most important aspects of student feedback in an educational setting. Each of these categories is explained below, along with an example in Hausa and an English translation.

- **Lecturer Behavior:** This category describes the lecturer’s professional and interpersonal behavior, including their attitudes toward students, approachability, and deference. *Translation (English): "The lecturer is very respectful." Example (Hausa):*“Malam yana da mutunci sosai.”
- **Course Content:** This sums up the course materials’ organization, relevance,

and clarity. It shows how well-structured, current, and helpful the students think the course material is. *Example (Hausa): "Abubuwan da ake koyarwa suna da amfani."* *Translation (English): "The course material is beneficial."*

- **Course Difficulty:** This category deals with the perceived degree of difficulty presented by the course, including workload, topic complexity, and general manageability. *Example (Hausa): "Darasin nan yana da wuya sosai."* *Translation (English): "This course is really challenging."*
- **Teaching Quality:** This component assesses how well the lecturer explains concepts and promotes understanding, as well as how clear and effective the teaching process is. *Example (Hausa): "Yadda malam yake bayyana darasi yana da sauki a fahimta."* *Translation (English): "The lecturer's explanation of the lesson is simple to comprehend".*
- **Teaching Performance:** This describes how well the lecturer delivers the course overall, including preparation, timeliness, and students satisfaction. *Example (Hausa): "Malam yana koyarwa cikin kwarewa "* *Translation (English): "The lecturer is a professional teacher."*

This categorization provides a structured foundation for ABSA in the Hausa educational context, enabling fine-grained sentiment analysis across dimensions that are meaningful to both students and instructors.

After successfully identifying the aspect categories, the annotators had the following tasks to accomplish:

1. **Identify Objective Reviews:** Objective reviews and reviews without any expressed sentiment had to be deleted from the data set. Example: *"Babu Sharhi."* ("No comment").
2. **Identify Aspect Terms:** Word expressions that name a specific aspect of the target entity. e.g *"Hanya ce mai kyau"* ("it is a nice course"). Some feedback came with implicitly referred aspect terms e.g *"Kuna murmushi da yawa"* ("You smile a lot").
3. **Assign Aspect Category:** The annotators had to assign aspect categories for each feedback based on the identified aspect terms. It must be chosen from the predefined aspect categories mentioned above (lecturer behavior, course content, course difficulty, teaching quality, teaching performance. Example: *"Ina zaune hanya kuma yana taimakawa wajen sanin sabbin yaruka."* ("I love the course and it helps in knowing new languages").

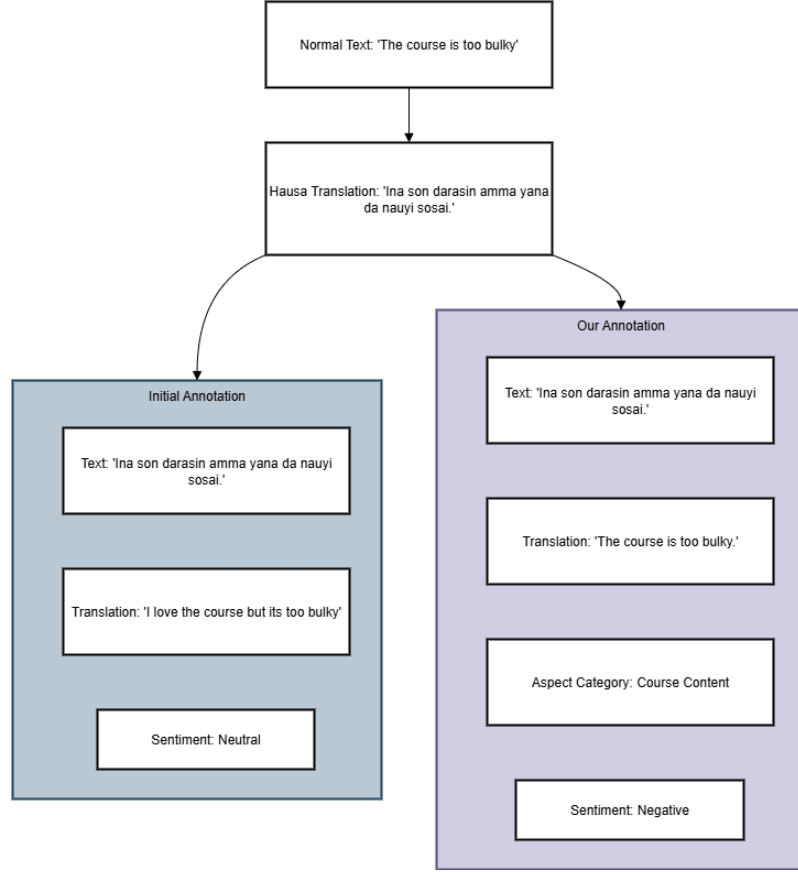


Figure 3.1: Our new annotations for students’ feedback compared to annotations from *Rakhmanov et al.*[27]

4. **Assign Polarity:** Finally, for each each aspect category tuple, the annotators had to assign the sentiment polarity from the following predefined values: *neutral*, *negative* and *positive*. The *neutral polarity* applies to neither positive nor negative sentiment.

An example of the annotation duplets compared to the annotation format used in *Rakhmanov et al* [27] is shown in Figure 3.1. The figure provides a comparative example, contrasting the original annotation from Rakhmanov et al. with the new, fine-grained annotations. For instance, the Hausa sentence ‘Ina son darasin amma yana da nauyi sosai’ was initially translated as “I love the course but its too bulky” and assigned a general Neutral sentiment. However, our ABSA framework re-analyzes this same sentence, more accurately translating it as “The course is too bulky,” and identifies the Course Content aspect with a Negative sentiment polarity. This comparison underscores the enhanced analytical depth and precision achieved by our annotation scheme, which successfully disentangles mixed feedback and assigns specific sentiment to explicit aspects.

3.3 Dataset Details

In Pontiki et al. [26] semEval-2015, for a given sentence X , one or multiple triplets(a, c, p), where a represents the aspect term, c denotes the aspect category, and p indicates the sentiment expressed toward the aspect. The annotators did not assign aspect terms to the sentences in the dataset; they assigned an aspect category and the sentiment directed towards it.

Five aspect categories were identified in the original dataset by Rakhmanov et Al [27]. Using this, a pair of aspect-category and sentiment polarity was assigned. Table -3.1 shows a detailed example.

Table 3.1: Annotated examples for Aspect categories with Hausa translations

Aspect Category	Duplets	Sentence	Translation (Hausa)
Teaching Quality	("TEACHING-QUALITY", "POSITIVE")	Interactive teaching method.	Hanyar koyarwa ta hula da alibai.
Course Content	("COURSE-CONTENT", "POSITIVE")	The course is really interesting.	Karatun yana da ban sha'awa sosai.
Lecturer Behaviour	("LECTURER-BEHAVIOUR", "POSITIVE")	Dr. Sadiq Thomas is a very friendly and nice lecturer anybody will like to meet.	Dr. Sadiq Thomas malami ne mai kirki saukin kai wanda kowa zai so ya hau shi.
Teaching Performance	("TEACHING-PERFORMANCE", "POSITIVE")	Commendable teaching skills.	Kwarewar koyarwa mai yabawa.
Course Difficulty	("COURSE-DIFFICULTY", "NEGATIVE")	The course is difficult.	Karatun yana da wahala.

3.4 Dataset Statistics

The data set was significantly reduced from its original size of 40,000 to 10,000 student comments. This is because to optimize the original Hesac data set for the ABSA tasks, the annotators had to :

1. Delete Sentences without aspect category, for example sentences such as "*Babu matsala*" "*alright*"
2. Removed repetitive comments like "*lafiya*" "*good*"

The made were to add terms to adjust other statements to get an aspect category: for example, repetitive comments such as "*lafiya*"

Table 3.2: Examples of Modifications done to repetitive comments with no aspect category

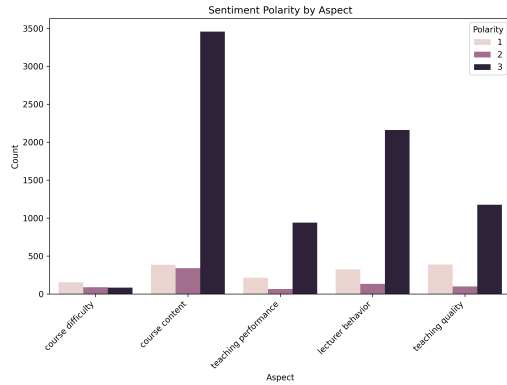
Comment	Modification	Aspect Category	Translation (English)
mai Kyau	("Malami Mai kyau")	Lecturer behavior	A good teacher.
mai kyau	("Hanyoyin koyarwa masu kyau")	Teaching quality	Good teaching Methods
mai kyau	("Aiki mai kyau ranka ya dae")	Teaching Performance	Good Work Sir
mai kyau	("Kwas mai kyau, yana da sauin fahimta")	Course difficulty	Good course easy to understand

Table 3.3: Dataset Distribution statistics, Sentence count, positive, negative, neutral

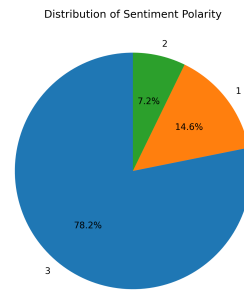
Aspect Category	Number of Sentences	Classification		
		Positive	Negative	Neutral
Teaching Quality	1660	1175	386	99
Teaching Performance	1219	940	214	65
Lecturer Behavior	2616	2159	325	132
Course Content	4181	3457	385	339
Course Difficulty	324	84	153	87

By clarifying the relationships between the aspect categories and their corresponding sentiments, the updated dataset improves the original HESAC dataset [27] and makes more complex ABSA tasks possible. As described in *salahudeen et al* [24], we conducted an exploratory data analysis to better understand the structure, distribution, and features of our dataset. The results of the analysis showed an uneven distribution of neutral, negative, and positive sentiments, which could affect how well the model performs. This disparity is depicted in Fig -3.2. Figure 3.2(a), which displays the sentiment polarity analysis across various aspects, identifies clear trends in student feedback. The 'Course Difficulty' component has a remarkably high number of annotations, indicating that it is a major issue in student assessments. Additionally, an initial analysis suggests that this category contains a concentration of negative sentiment. A more balanced or even positive sentiment distribution is seen in aspects pertaining to the instructor, such as "Teaching Performance" and "Lecturer Behavior." This discrepancy emphasizes how useful ABSA is at separating feedback; general course discontent, which is frequently caused by perceived difficulty, can be distinguished from admiration for the caliber of instruction.

A significant skew in the dataset is confirmed by the overall sentiment distribution, which is shown in Figure 3.2(b). With 78.2% of the annotations, positive sentiments



(a) Sentiment polarity distribution in the dataset



(b) Overall Sentiment distribution in the dataset

Figure 3.2: Distribution of sentiment polarities and Aspects in our dataset

make up the majority class. At 14.6% and 7.2%, respectively, negative and neutral sentiments are much less prevalent. This imbalance is a crucial feature of the dataset because machine learning models trained on it may become biased toward predicting the majority (positive) class, which would make it harder to identify the subtle but significant instances of neutral and negative feedback. Creating a strong ABSA model will require mitigation techniques for this imbalance.

Chapter 4

Methodology

In this chapter, we describe the research methodology employed to achieve the objectives of this study. The methodology includes everything from the research design, dataset preparation, dataset annotation and data analysis, and techniques.

4.1 Models

Amongst the highly sought-after models employed in aspect-based sentiment analysis due to their high accuracy and performance are LSTM and BERT. This section provides a comprehensive overview of how these models are applied for ABSA and references recent research that has been able to utilize these techniques successfully.

LSTM, a family of recurrent neural networks (RNNs), is capable of learning long-term dependencies in sequence data. In ABSA tasks, LSTM is generally employed to capture the interaction between words in a sentence and the given aspect. LSTMs' extensions, such as attention-based LSTMs, allow the model to focus on words carrying sentiment within the context of a particular aspect. For example, Ma et al. [20] presented an attentive LSTM model that well integrates commonsense knowledge to improve aspect-level sentiment classification.

Conversely, transformer models like BERT are the current state-of-the-art solution for ABSA since they can represent deep bidirectional context. Musa et al. [24] developed HauBERT, a transformer tailored to aspect-based sentiment analysis of Hausa-language film reviews, to demonstrate the effectiveness of BERT architectures in low-resource language. Mabokela et al. [30] also addressed the use of pretrained language models for sentiment analysis in African languages, further cementing the effectiveness of BERT-based models for low-resource multilingual environments.

Other works such as Abdullah et al. [1] and Aliyu et al. [4] also provide extensive reviews and hybrid approaches that employ an ensemble of both shallow and deep learning models in sentiment analysis of poverty-stricken languages. These works collectively highlight the superior performance of BERT models and the ongoing relevance of LSTM-based architectures for aspect-based sentiment analysis tasks.

4.2 Overview of the Model

The proposed methodology of the ABSA of Hausa is shown in the 4.1. The proposed architecture has been created for **Aspect-Based Sentiment Analysis (ABSA)**, on **Hausa student comments**, allowing for rich sentiment classification along various dimensions of education.

4.2.1 BiLSTM

BiLSTM architecture is a form of recurrent neural network that reads text from both directions using Bidirectional Long Short-Term Memory layers [12], [34]. The input words are first converted into dense vector representations using an embedding layer [22]. Then, the embeddings are passed to the BiLSTM layer for extracting contextual information from both directions [14]. For aspect term extraction, the output from BiLSTM is passed through a dense layer that performs token-level classification following the BIO tagging scheme [36]. For sentiment classification, aggregated hidden states are used for sequence-level sentiment classification [19]. The model employs sequential-based learning and is trained with a cross-entropy loss function [9].

4.2.2 HauBERT

HauBERT is based on the pre-trained bert-base-multilingual-cased and is fine-tuned for Aspect-Based Sentiment Analysis (ABSA) on Hausa [24]. It addresses two tasks: aspect term extraction with token-level classification with the BIO tagging scheme, and sentiment classification with sequence-level prediction. It employs subword tokenization via the WordPiece tokenizer and generates contextual embeddings via several transformer layers with self-attention. The output embeddings are fed to task-specific classification heads: one for aspect labels and one for sentiment labels. Both tasks are trained using a cross-entropy loss function. HauBERT’s property of multilingualism aids in transcending the low-resource constraint of Hausa language by knowledge transfer from a higher resource language.

The fine-tuning procedure for ABSA of HauBERT is as shown in Algorithm 1. Capitalizing on a pre-trained multilingual BERT (mBERT) model, a task-agnostic classification head is added to enable token-level aspect extraction and sequence-level sentiment classification. The model is optimized using a cross-entropy loss function, propagating parameters through backpropagation over a number of epochs. Training data is split into batches, shuffled each epoch, and passed into the model to compute predictions and losses. The computed gradients are then updated by the optimizer in order to update model weights. Periodically, its performance is evaluated on a validation set for accuracy and F1-score. This ensures that the mBERT weights are appropriately adapted to the Hausa ABSA task. As such, this approach leverages the multilingual capabilities of mBERT to compensate for the lack of resources in Hausa [24].

Algorithm 1 Fine-Tuning mBERT for ABSA in Hausa Language

Require: Preprocessed dataset D with inputs in Hausa; pre-trained mBERT model

M ; optimizer O ; number of epochs E ; batch size B ; and loss function L .

- 1: Initialize mBERT model M with pre-trained weights.
 - 2: Add a task-specific classification head to M for ABSA.
 - 3: Split dataset D into training and validation sets.
 - 4: Divide the training set into batches of size B .
 - 5: **for** epoch $e = 1$ to E **do**
 - 6: Shuffle the training dataset.
 - 7: **for** each batch (X, y) in the training set **do**
 - 8: Unpack the batch: input IDs X_{ids} , attention mask X_{mask} , and labels y .
 - 9: Transfer the batch to the computation device (e.g., GPU or CPU).
 - 10: Zero the gradients of O .
 - 11: Compute predictions $\hat{y}$ using mBERT:
 - 12: $\hat{y} = M(X_{ids}, X_{mask})$
 - 13: Compute loss ℓ :
 - 14: $\ell = L(\hat{y}, y)$
 - 15: Perform backpropagation to compute gradients:
 - 16: $\nabla_{\theta} \ell$
 - 17: Update model parameters using the optimizer:
 - 18: $\theta \leftarrow \theta - O(\nabla_{\theta} \ell)$
 - 19: **end for**
 - 20: Evaluate the model on the validation set to compute metrics such as accuracy and F1-score.
 - 21: **end for**
 - 22: **return** the fine-tuned mBERT model M .
-

Algorithm 1 describes how to fine-tune a pre-trained mBERT model to conduct Hausa Aspect-Based Sentiment Analysis (ABSA). In order to provide a solid multilingual foundation that can be utilized for low-resource languages like Hausa, the algorithm starts by initializing the mBERT model with pre-trained weights. The model is then supplemented with task-specific classification heads: a sequence-level head for sentiment classification and a token-level head for aspect term extraction using the BIO tagging scheme. To enable effective computation, the dataset is separated into training and validation sets, and the training data is further subdivided into batches. To lessen overfitting, the dataset is shuffled at the beginning of each training epoch. The model generates predictions for each batch based on input IDs and attention masks,

which are then compared to the true labels using a cross-entropy loss function. Back-propagation is used to calculate gradients, and the designated optimizer is used to update the model parameters. Metrics like accuracy and F1-score are used to periodically assess the model’s performance on the validation set and make sure it is successfully learning aspect extraction and sentiment classification. The algorithm does not cover some implementation details, such as the precise architecture of the classification heads, hyperparameter values (e.g., learning rate, batch size, number of epochs), dataset characteristics, or performance outcomes, even though it offers a clear, step-by-step framework for fine-tuning. However, the process successfully illustrates how the low-resource constraints of the Hausa language for ABSA tasks can be addressed by utilizing mBERT’s multilingual capabilities.

4.3 Dataset Preprocessing and Preparation

The suggested architecture has been crafted for **Aspect-Based Sentiment Analysis (ABSA)**, focusing on **Hausa student comments**, allowing detailed sentiment classification across various educational dimensions.

1. Data Preprocessing

Procedure:

- **Text Cleaning:** Eliminate unnecessary characters (emojis, punctuation, URLs). Fix spelling errors (common in casual student feedback).
- **Tokenization:** Break down sentences into words/subwords (e.g., “Darasi yana da wuya” → [“Darasi”, “yana”, “da”, “wuya”]).
- **Stop Word Removal:** Remove non-significant Hausa words (e.g., “da” (and), “a” (in)).
- **Lemmatization:** Convert words to their root forms (e.g., “karatu” (reading) → “kara” (read)).
- **Annotation:** Manually tag aspects (e.g., “Teaching Quality”) and polarities (Positive/Negative/Neutral) for supervised learning.

2. Aspect Extraction

- **Rule-Based + ML Hybrid Approach:** Utilize predefined aspect categories (e.g., *Lecturer Behavior*, *Course Content*) combined with keyword matching. Train an **BiLSTM** to identify implicit aspects (e.g., “Malam

ya kasance mai hauri” → “Lecturer Behavior”).

- **Normalize Polarity:** Standardize sentiment labels (e.g., “yana da kyau” → Positive, “ba ya da kyau” → Negative).
- **Polarity Encoding:** Transform sentiments into numerical codes (Positive: 3, Neutral:2 , Negative: 1).

3. Sentiment Classification Model

Models:

- **HauBERT (Fine-Tuned mBERT):** Utilizes pretrained multilingual BERT for better comprehension of the Hausa context. Fine-tuned on annotated student comments for aspect-specific sentiment analysis.
- **BiLSTM + Attention:** Captures long-term dependencies present in student feedback. The attention mechanism emphasizes sentiment-driven words (e.g., “kyau” (good) → Positive).

Output: Sentiment predictions for each aspect (e.g., “Course Difficulty” → Negative).

4. Model Evaluation

Metrics:

- **Aspect Extraction:** Precision, Recall, F1-score.
- **Sentiment Classification:** Accuracy, Confusion Matrix.
- **Cross-Validation:** Ensure reliability on previously unseen student comments.

Tools: Scikit-learn, Hugging Face Transformers.

5. Key Benefits of the Architecture

- **Language Adaptability:** HauBERT and BiLSTM effectively address the morphological nuances of Hausa.
- **Detailed Analysis:** Dissects sentiments relating to specific educational aspects (e.g., “Teaching Quality” versus “Course Content”).
- **Scalability:** The rule-based aspect extraction can be expanded by adding new keywords (e.g., “ilimi” (knowledge) → “Course Content”).

This architecture facilitates **precise, understandable and useful insights** from the feedback provided by Hausa students.

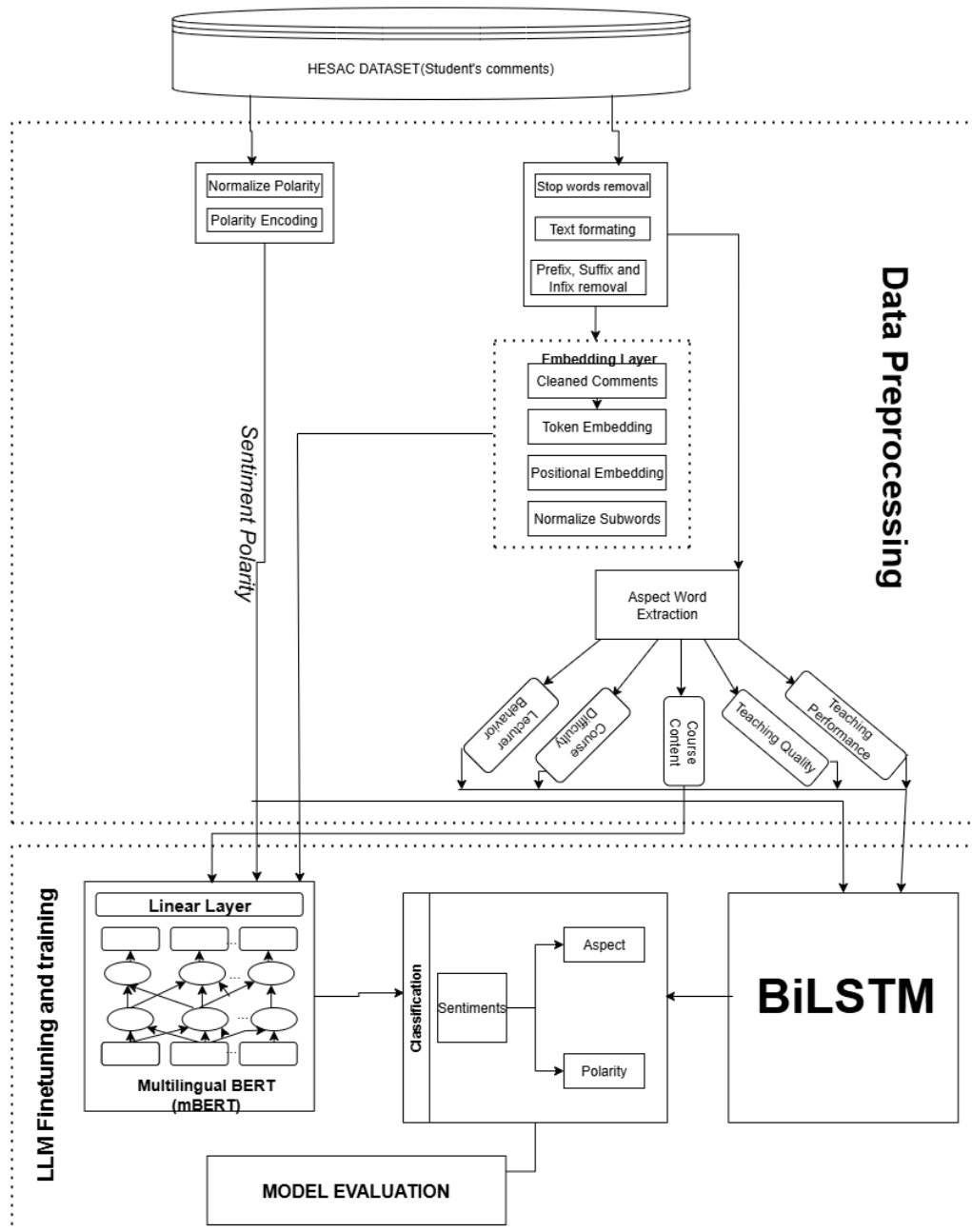


Figure 4.1: Model Architecture

Chapter 5

Results and Discussion

In this chapter, we dive into the experimental setups and the experiments we did, present and analyze the results, challenges, limitations, implications, and significance of findings.

5.1 Experimental Setup

The experimental setup used in this study is a transformer model called HauBERT and we also used Bi-LSTM to train and evaluate both models on the Hesac-ABSA dataset. To test the quality of the dataset, the models were evaluated based on the following tasks:

1. Aspect Category Detection: Detection of Aspect categories.
2. Aspect-term-category extraction: Extraction of (Aspect term, category) tuples.
3. Aspect Polarity Detection: Detection of (Aspect Category, sentiment Polarity) duplets.

Common classification metrics are used to assess the performance of our Aspect-Based Sentiment Analysis (ABSA) models. The following elements are used to derive these metrics:

The evaluation metrics rely on four fundamental classification outcomes, defined as follows with examples from our HESAC-ABSA dataset:

- **True Positive (TP):** The model correctly assigns the correct sentiment polarity or recognizes an aspect. *For Instance:* If the model correctly predicts that "course content" is an aspect in the sentence "*The course is very interesting*"

(*Karatun yana da ban shaawa sosai*) and assigns it a positive sentiment, this counts as a TP.

- **False Positive (FP):** The model incorrectly identifies an aspect or assigns a sentiment when it should not. *For Example:* If the model wrongly predicts "lecturer behaviour as an aspect in the sentence *The course is very interesting* (*Karatun yana da ban shaawa sosai*), even though no lecturer-related aspect is annotated, this counts as an FP.
- **True Negative (TN):** The model correctly ignores a non-aspect term or correctly classifies a sentiment as neutral/absent. *Example:* In the sentence *No comment* (*Babu Sharhi*), if the model does not assign any aspect or sentiment polarity, this is a TN.
- **False Negative (FN):** The model fails to identify an actual aspect or assigns an incorrect sentiment polarity. *For Example:* In the sentence *The lecturer was not helpful* (*Malamin bai taimaka ba*), if the model fails to detect the lecturer behaviour aspect or misclassifies the sentiment as neutral instead of negative, this counts as an FN.

The evaluation metrics listed below are used in accordance with these definitions:

1. **Aspect Detection Accuracy:** calculates the proportion of accurately identified aspect terms. Since aspect detection is the pipeline's initial step and mistakes there spread to subsequent tasks, this is critical for ABSA.

$$\text{Aspect_Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (5.1)$$

2. **Sentiment Polarity Accuracy:** evaluates the proportion of accurately predicted sentiment for the aspects that have been identified. This measure assesses the model's accuracy in identifying positive, negative, or neutral polarities.

$$\text{Sentiment_Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (5.2)$$

3. **Precision:** Indicates the proportion of predictions that are correct out of all instances the model labeled as positive. High precision means the model avoids false alarms, which is important when the cost of incorrect predictions is high.

$$\text{Precision} = \frac{TP}{TP + FP} \quad (5.3)$$

4. **Recall:** Represents the percentage of correctly identified positive instances out of all actual positives. High recall ensures that fewer pertinent aspects or sentiments are missed.

$$\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}} \quad (5.4)$$

5. **F1-Score:** The harmonic mean of Precision and Recall, balancing both metrics. It is especially useful when dealing with imbalanced datasets, where accuracy alone may be misleading.

$$\text{F1-Score} = \frac{2 \cdot \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (5.5)$$

5.2 Results Presentation

Below are separate presentations of the experimental results for the HauBERT and Bi-LSTM models on the HESAC dataset.

5.2.1 Bi-LSTM Model Results

For both aspect and polarity detection, the Bi-LSTM model was trained and assessed. Table 5.1 provides a summary of the findings.

Table 5.1: Efficacy of Bi-LSTM Model on HESAC Dataset

Task	Accuracy	Precision	Recall	F1-Score
Aspect Detection	0.667	0.657	0.667	0.655
Polarity Detection	0.860	0.850	0.860	0.853

Observation:

The model performed moderately well in aspect detection (66.7% accuracy), meaning that while it could correctly identify some aspect categories, it had trouble with others. The model demonstrated its ability to accurately classify sentiment for detected aspects with an accuracy of 86% for polarity detection.

5.2.2 HauBERT Model Results

For both aspect and polarity detection, the HauBERT transformer model was trained and assessed. Table 5.2 provides a summary of the findings.

Observation:

Table 5.2: HauBERT Model Performance on HESAC Dataset

Task	Accuracy	Precision	Recall	F1-Score
Aspect Detection	0.972	0.944	0.972	0.958
Polarity Detection	0.783	0.613	0.783	0.688

- HauBERT’s ability to identify aspect categories in Hausa text was demonstrated by its outstanding aspect detection performance (97.2% accuracy).
- Haubert obtained 78.3% accuracy and 61.3% precision for polarity detection, suggesting some misclassifications that were probably caused by the dataset’s class imbalance.

Figure 5.1 shows the training loss and accuracy curves over epochs for both tasks. These plots illustrate how the model converged during training and highlight differences in difficulty between aspect extraction and sentiment classification.

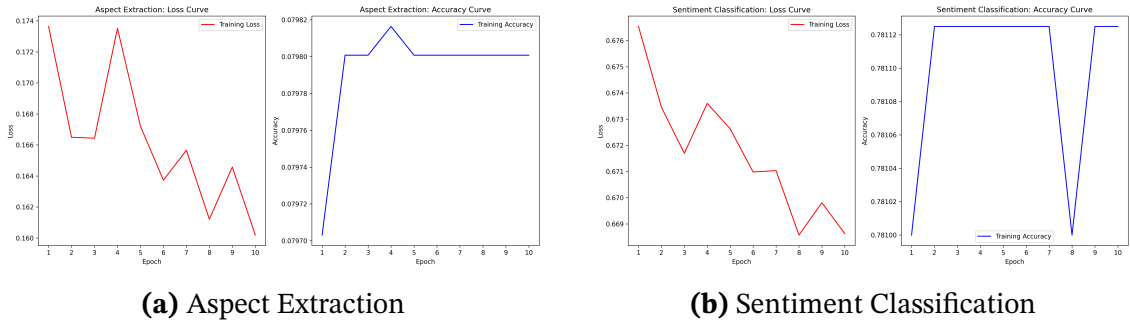


Figure 5.1: Training Loss and Accuracy Curves for HauBERT. The left plot shows aspect extraction, and the right shows sentiment classification.

5.2.3 Average Evaluation Metrics

The average metrics for aspect and polarity tasks are shown in Table 5.3 to give a general comparison of model performance.

Table 5.3: Average Performance Metrics Across Both Tasks

Model	Average Accuracy	Average Precision	Average Recall	Average F1-Score
Bi-LSTM	0.764	0.754	0.764	0.754
HauBERT	0.878	0.779	0.878	0.823

5.3 Examining the Findings

The outcomes of the HauBERT and Bi-LSTM models offer important new information about how well Aspect-Based Sentiment Analysis works with the Hausa language dataset.

- **Aspect Recognition:** The accuracy of HauBERT was substantially higher (97.2%) than that of Bi-LSTM (66.7%). This implies that even in a low-resource environment, transformer-based models are very successful at recognizing aspect categories and capturing contextual information in Hausa text.
- **Detection of Polarity:** In terms of sentiment polarity detection, Bi-LSTM outperformed HauBERT (78.3%) by a small margin (86%). This suggests that, especially when working with smaller or unbalanced datasets, sequential models such as Bi-LSTM are still competitive for sentiment detection.
- **Precision and Recall:** While HauBERT's lower precision in polarity detection indicates a propensity to misclassify certain sentiments, its high recall in aspect detection indicates that it can accurately identify the majority of aspect categories. When it came to polarity detection, Bi-LSTM maintained a more balanced precision-recall performance.
- **Overall Average Metrics:** With a higher average accuracy and F1-score, HauBERT performed better overall than Bi-LSTM, demonstrating the usefulness of transformer architectures in Hausa ABSA tasks.

5.4 Difficulties and Restrictions

A number of difficulties and restrictions were noted throughout the experimentation process:

- **Data Sparsity:** The HESAC dataset contained relatively few instances of certain aspect categories, which hindered the models' capacity to efficiently learn these uncommon categories.
- **Unbalanced Classes:** Model precision suffered as a result of the imbalanced sentiment polarity and aspect category classes, particularly for minority classes.
- **Hausa Language Complexity:** Tokenization and embedding representation were complicated by morphological richness, spelling variations, and dialect differences.

- **Computational Resources:** Significant memory and processing time were needed for the Haubert transformer’s training, which might have limited its scalability for bigger datasets.
- **Generalization:** Without further fine-tuning, models trained on the HESAC dataset might not generalize well to other Hausa text domains.

5.5 Implications and Significance of Findings

The study’s conclusions have numerous significant implications:

- **Transformer Efficiency:** HauBERT’s superior aspect detection performance shows how well transformer-based models work with low-resource languages like Hausa.
- **Sequential Model Relevance:** Traditional sequential models can still be useful for sentiment classification, particularly in situations where resources are scarce, as demonstrated by Bi-LSTM’s competitive polarity detection performance.
- **Useful Applications:** The findings can direct the creation of useful ABSA systems for Hausa-language applications like feedback analysis, social media monitoring, and customer review analysis.
- **Future Research:** The shortcomings noted point to areas that require further development, such as handling class imbalance, constructing multilingual or domain-adapted models, and augmenting data.

5.6 Chapter Conclusion

The experimental setup, evaluation metrics, and comprehensive results for the HauBERT and Bi-LSTM models on the HESAC-ABSA dataset were presented in this chapter. The efficacy of transformer-based architectures for Aspect-Based Sentiment Analysis in the Hausa language was demonstrated by the HauBERT model’s superior aspect detection performance. Sequential models are still useful for some sentiment classification tasks, as demonstrated by the Bi-LSTM model’s competitive results in sentiment polarity detection. Along with the implications of these findings for future research and real-world applications, the chapter also covered important challenges such as class imbalance, data sparsity, and the complexity of Hausa language morphology. All things considered, the results offer a thorough grasp of the model’s performance, its

limitations, and the importance of the discoveries for creating reliable ABSA systems in Hausa.

Chapter 6

Conclusion

By going over the goals and research questions again, summarizing the main conclusions, talking about their wider ramifications, highlighting the contributions made, and admitting the study’s limitations, this chapter brings the research to a close. This work fills a significant gap in low-resource language research in Natural Language Processing (NLP) by offering one of the first attempts to create Aspect-Based Sentiment Analysis (ABSA) resources and models for Hausa in the educational domain.

6.1 Restating Research Objectives and Questions

Establishing a framework for ABSA in Hausa, a language spoken by millions of people but underrepresented in computational linguistics, was the main goal of this study. In order to do this, the study concentrated on creating datasets, developing and improving models and assessing contemporary language technologies in low-resource environments. The following goals and related queries served as the research’s compass:

1. **Dataset Creation:** How can a reliable Hausa ABSA dataset be developed from student feedback that reflects meaningful aspect categories and sentiment polarities?
2. **Model Development:** Which machine learning and pretrained language models are effective in extracting aspect terms and classifying sentiment polarities in Hausa?
3. **Model Comparison:** How do multilingual pretrained models such as HauBERT, AfroXLM-R, perform in comparison with traditional models like Bi-LSTM?

4. **Linguistic Adaptation:** What strategies can be employed to address Hausa-specific challenges such as complex morphology, code-switching, and dialectal variations?
5. **Broader Applicability:** In what ways can this research contribute to NLP resources and methodologies for other low-resource African languages?

6.2 Discussion of Implications

The implications of this research extend beyond Hausa and offer lessons for low-resource NLP in general:

- **Educational Applications:** By enabling fine-grained sentiment analysis of student feedback, the Hausa ABSA framework can help educational institutions identify specific areas of strength and weakness, informing targeted improvements in teaching, course design, and institutional performance.
- **Advancement of African NLP:** The creation of a Hausa ABSA dataset contributes to the growing body of resources for African languages. It demonstrates that high-quality, manually annotated datasets can be developed even in under-represented languages, paving the way for more inclusive NLP research.
- **Model Development for Low-Resource Settings:** The performance of HauBERT underscores the promise of pretrained multilingual transformers for ABSA in low-resource contexts, while the results of Bi-LSTM highlight the continued value of lightweight models that can perform well with limited computational resources.
- **Methodological Significance:** The framework adopted in this research dataset creation, annotation, model benchmarking, and adaptation to linguistic challenges serves as a replicable methodology for future ABSA studies in other low-resource languages.

6.3 Summary of Key Findings

The findings of this study can be summarized as follows:

- **Annotated Hausa ABSA Dataset:** Using student feedback that was manually annotated by native Hausa speakers, a novel dataset (HESAC-ABSA) was produced. Lecturer behavior, course content, course difficulty, teaching qual-

ity, and teaching performance were the five main aspect categories that were found. The first benchmark dataset for Hausa-ABSA in an educational setting was created by labeling each aspect with sentiment polarity (positive, negative, and neutral).

- **Model Performance:** According to the experimental results, transformer-based models in particular, HauBERT performed better than Bi-LSTM (66.7%) in aspect detection, with an accuracy of 97.2%. In polarity detection, however, Bi-LSTM performed better than HauBERT (86.0% vs. 78.3%), demonstrating the ongoing competitiveness of sequential models for specific tasks.
- **Precision-Recall Trade-off:** HauBERT demonstrated a high recall in aspect detection, effectively identifying the majority of aspect categories, but a lower precision in polarity detection, indicating difficulties with class imbalance. For sentiment tasks, Bi-LSTM preserved a more balanced precision and recall.
- **Overall Results:** On average, HauBERT outperformed Bi-LSTM across tasks, with higher accuracy and F1-scores. This indicates the effectiveness of transformer architectures in handling Hausa ABSA, while also confirming the usefulness of simpler models for polarity classification.
- **Insights for Low-Resource NLP:** The results suggest that a hybrid approach—leveraging transformers for aspect detection and Bi-LSTM for polarity detection—may provide the most effective solution in low-resource contexts like Hausa.

6.4 Contributions of the Research

The research makes several key contributions:

1. Development of the first annotated Hausa ABSA dataset in the educational domain.
2. Systematic evaluation and benchmarking of traditional (Bi-LSTM) and pretrained transformer models (HauBERT) for Hausa ABSA tasks.
3. Introduction of annotation categories specifically relevant to Hausa educational feedback, thereby contextualizing ABSA to local needs.
4. Adaptation of models and preprocessing strategies to address linguistic features of Hausa, including morphology, code-switching, and dialectal variations.
5. Provision of a methodological framework that can guide the creation of ABSA

resources and models for other African low-resource languages.

6.5 Acknowledging Limitations

While this study has made significant progress, several limitations should be acknowledged:

- **Dataset Size:** The dataset, though novel, is relatively small compared to established ABSA benchmarks in high-resource languages. This limits generalization and scalability.
- **Domain Restriction:** The dataset focuses solely on student feedback, which may reduce its applicability to other domains such as product reviews or social media discourse.
- **Annotation Constraints:** Despite careful annotation by native speakers, some aspect categories are inherently subjective, and cultural differences may influence labeling.
- **Model Limitations:** Only a limited number of pretrained models were tested. Broader experimentation with Afro-centric models (e.g., AfroXLM-R, AfriBERTa) could yield further insights.

Future work should expand the dataset across multiple domains, incorporate semi-supervised and data augmentation techniques, and evaluate additional pretrained African language models. These directions will strengthen the robustness and applicability of ABSA for Hausa and contribute to the broader goal of advancing NLP for low-resource languages.

6.6 Future Work

While this study provides an initial step towards developing an ABSA dataset and models for Hausa in the educational domain, several directions can further enhance the quality and impact of this work:

- **Expand the Dataset Size:** In addition to improving model robustness, increasing the number of annotated student feedback instances will enable more representative coverage of viewpoints from a variety of lecturers and courses.
- **Triplet-Based Annotation:** Future research should incorporate triplets of *aspect term*, *aspect category*, and *sentiment polarity* in addition to duplet annota-

tion. In order to capture more nuanced feedback, polarity could also be extended from three classes (positive, neutral, and negative) to four or more, for instance, by adding a "mixed" or "conflict" or "more negative" polarities.

- **Handling Explicit and Implicit Aspects:** Current annotations focus primarily on explicit aspect terms. Future models should be trained to effectively capture both implicit and explicit opinions (e.g., indirect references to teaching style or course workload), thereby improving coverage of subtle sentiment expressions.
- **Model Development and Evaluation:** Pre-trained multilingual language models, domain-adapted transformers, and hybrid models that blend lexicon-based and neural approaches are just a few examples of the broader variety of machine learning and deep learning architectures that should be investigated. Better performance on aspect detection and sentiment classification tasks as well as a more thorough evaluation will be made possible by this.

References

- [1] R. Abdullah, R. Sarno, et al., “Aspect based sentiment analysis for explicit and implicit aspects in restaurant review using grammatical rules, hybrid approach, and senticircle.,” *International Journal of Intelligent Engineering & Systems*, vol. 14, no. 5, 2021.
- [2] M. S. Akhtar, A. Ekbal, and P. Bhattacharyya, “Aspect based sentiment analysis in Hindi: Resource creation and evaluation,” in *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC’16)*, N. Calzolari et al., Eds., Portoro, Slovenia: European Language Resources Association (ELRA), May 2016. [Online]. Available: <https://aclanthology.org/L16-1429/>.
- [3] M. S. Akhtar, A. Kumar, A. Ekbal, C. Biemann, and P. Bhattacharyya, “Language-agnostic model for aspect-based sentiment analysis,” in *Proceedings of the 13th International Conference on Computational Semantics - Long Papers*, S. Dobnik, S. Chatzikyriakidis, and V. Demberg, Eds., Gothenburg, Sweden: Association for Computational Linguistics, May 2019. [Online]. Available: <https://aclanthology.org/W19-0413/>.
- [4] Y. Aliyu, A. Sarlan, K. Danyaro, A. Rahman, and M. Abdullahi, “Sentiment analysis in low-resource settings: A comprehensive review of approaches, languages, and data sources,” *IEEE Access*, vol. PP, pp. 1–1, Jan. 2024. DOI: 10.109/ACCESS.2024.3398635.
- [5] X. Bao, M. Qiang, J. Gu, Z. Wang, and C.-R. Huang, “Exploring hybrid sampling inference for aspect-based sentiment analysis,” in *Findings of the Association for Computational Linguistics: NAACL 2025*, L. Chiruzzo, A. Ritter, and L. Wang, Eds., Albuquerque, New Mexico: Association for Computational Linguistics, Apr. 2025. [Online]. Available: <https://aclanthology.org/2025.findings-naacl.236/>.
- [6] A. Bhattacharya, A. Debnath, and M. Shrivastava, “Enhancing aspect extraction for Hindi,” in *Proceedings of the 4th Workshop on e-Commerce and NLP*, S. Malmasi, S. Kallumadi, N. Ueffing, O. Rokhlenko, E. Agichtein, and I. Guy, Eds., Association for Computational Linguistics, Aug. 2021. [Online]. Available: <https://aclanthology.org/2021.ecnlp-1.17/>.

- [7] C. Chen, Z. Teng, Z. Wang, and Y. Zhang, “Discrete opinion tree induction for aspect-based sentiment analysis,” in *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Dublin, Ireland: Association for Computational Linguistics, 2022, pp. 2051–2064. [Online]. Available: <https://aclanthology.org/2022.acl-long.145/>.
- [8] L. Dewangan, Z. A. Sayeed, and C. Maurya, “Benchmark creation for aspect-based sentiment analysis in low-resource Odia language and evaluation through fine-tuning of multilingual models,” in *Proceedings of the 31st International Conference on Computational Linguistics*, O. Rambow, L. Wanner, M. Apidianaki, H. Al-Khalifa, B. D. Eugenio, and S. Schockaert, Eds., Abu Dhabi, UAE: Association for Computational Linguistics, Jan. 2025. [Online]. Available: <https://aclanthology.org/2025.coling-main.391/>.
- [9] I. Goodfellow, Y. Bengio, and A. Courville, *Deep learning*. MIT Press, 2016.
- [10] E. Häglund and J. Björklund, “Opinion units: Concise and contextualized representations for aspect-based sentiment analysis,” in *Proceedings of the Joint 25th Nordic Conference on Computational Linguistics and 11th Baltic Conference on Human Language Technologies (NoDaLiDa/Baltic-HLT 2025)*, R. Johansson and S. Stymne, Eds., Tallinn, Estonia: University of Tartu Library, Mar. 2025. [Online]. Available: <https://aclanthology.org/2025.nodalida-1.24/>.
- [11] M. A. Hedderich, L. Lange, H. Adel, J. Strötgen, and D. Klakow, “A survey on recent approaches for natural language processing in low-resource scenarios,” K. Toutanova et al., Eds.
- [12] S. Hochreiter and J. Schmidhuber, “Long short-term memory,” *Neural computation*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [13] M. M. Hossain, M. S. Hossain, S. Chaki, M. R. Hossain, M. S. Rahman, and A. B. M. S. Ali, *Crosgrpsabs: Cross-attention over syntactic and semantic graphs for aspect-based sentiment analysis in a low-resource language*, 2025. arXiv: 2505.19018 [cs.CL]. [Online]. Available: <https://arxiv.org/abs/2505.19018>.
- [14] Z. Huang, W. Xu, and K. Yu, “Bidirectional lstm-crf models for sequence tagging,” in *arXiv preprint arXiv:1508.01991*, 2015.
- [15] A. Karimi, L. Rossi, and A. Prati, “Improving BERT performance for aspect-based sentiment analysis,” in *Proceedings of the 4th International Conference on Natural Language and Speech Processing (ICNLSP 2021)*, M. Abbas and A. A. Freihat, Eds., Trento, Italy: Association for Computational Linguistics, Dec. 2021. [Online]. Available: <https://aclanthology.org/2021.icnls-1.5/>.
- [16] R. Li, H. Chen, F. Feng, Z. Ma, X. Wang, and E. Hovy, “Dual graph convolutional networks for aspect-based sentiment analysis,” in *Proceedings of the*

- 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, C. Zong, F. Xia, W. Li, and R. Navigli, Eds., Online: Association for Computational Linguistics, Aug. 2021. [Online]. Available: <https://aclanthology.org/2021.acl-long.494/>.
- [17] X. Liu, R. Li, S. Ye, G. Zhang, and X. Wang, “Multimodal aspect-based sentiment analysis under conditional relation,” in *Proceedings of the 31st International Conference on Computational Linguistics*, O. Rambow, L. Wanner, M. Apidianaki, H. Al-Khalifa, B. D. Eugenio, and S. Schockaert, Eds., Abu Dhabi, UAE: Association for Computational Linguistics, Jan. 2025. [Online]. Available: <https://aclanthology.org/2025.coling-main.22/>.
 - [18] D. Ma, S. Li, X. Zhang, and H. Wang, *Interactive attention networks for aspect-level sentiment classification*, 2017. arXiv: 1709.00893 [cs.AI]. [Online]. Available: <https://arxiv.org/abs/1709.00893>.
 - [19] X. Ma and E. Hovy, “End-to-end sequence labeling via bi-directional lstm-cnns-crf,” in *arXiv preprint arXiv:1603.01354*, 2016.
 - [20] Y. Ma, H. Peng, and E. Cambria, “Targeted aspect-based sentiment analysis via embedding commonsense knowledge into an attentive lstm,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 32, no. 1, Apr. 2018. [Online]. Available: <https://ojs.aaai.org/index.php/AAAI/article/view/12048>.
 - [21] S. G. Matlatipov, J. Rajabov, E. Kuriyozov, and M. Aripov, “UzABSA: Aspect-based sentiment analysis for the Uzbek language,” in *Proceedings of the 3rd Annual Meeting of the Special Interest Group on Under-resourced Languages @ LREC-COLING 2024*, Torino, Italia: ELRA and ICCL, May 2024. [Online]. Available: <https://aclanthology.org/2024.sigul-1.47/>.
 - [22] T. Mikolov, K. Chen, G. Corrado, and J. Dean, “Efficient estimation of word representations in vector space,” in *Proceedings of Workshop at ICLR*, 2013.
 - [23] S. H. Muhammad et al., *Naijasenti: A nigerian twitter sentiment corpus for multilingual sentiment analysis*, 2022. arXiv: 2201.08277 [cs.CL]. [Online]. Available: <https://arxiv.org/abs/2201.08277>.
 - [24] A. Musa, F. M. Adam, U. Ibrahim, and A. Y. Zandam, “Haubert: A transformer model for aspect-based sentiment analysis of hausa-language movie reviews,” *Engineering Proceedings*, vol. 87, no. 1, 2025, ISSN: 2673-4591. [Online]. Available: <https://www.mdpi.com/2673-4591/87/1/43>.
 - [25] N. Neveditsin, P. Lingras, and V. K. Mago, “From annotation to adaptation: Metrics, synthetic data, and aspect extraction for aspect-based sentiment analysis with large language models,” in *Proceedings of the 2025 Conference of the*

- Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 4: Student Research Workshop)*, A. Ebrahimi, S. Haider, E. Liu, S. Haider, M. Leonor Pacheco, and S. Wein, Eds., Albuquerque, USA: Association for Computational Linguistics, Apr. 2025. [Online]. Available: <https://aclanthology.org/2025.naacl-srw.14/>.
- [26] M. Pontiki, D. Galanis, H. Papageorgiou, S. Manandhar, and I. Androutsopoulos, "SemEval-2015 task 12: Aspect based sentiment analysis," in *Proceedings of the 9th International Workshop on Semantic Evaluation (SemEval 2015)*, P. Nakov, T. Zesch, D. Cer, and D. Jurgens, Eds., Denver, Colorado: Association for Computational Linguistics, Jun. 2015. [Online]. Available: <https://aclanthology.org/S15-2082/>.
 - [27] O. Rakhmanov and T. Schlippe, "Sentiment analysis for Hausa: Classifying students' comments," in *Proceedings of the 1st Annual Meeting of the ELRA/ISCA Special Interest Group on Under-Resourced Languages*, M. Melero, S. Sakti, and C. Soria, Eds., Marseille, France: European Language Resources Association, Jun. 2022. [Online]. Available: <https://aclanthology.org/2022.sigul-1.13/>.
 - [28] N. Raychawdhary, A. Das, G. Dozier, and C. D. Seals, "Seals_Lab at SemEval-2023 task 12: Sentiment analysis for low-resource African languages, Hausa and Igbo," in *Proceedings of the 17th International Workshop on Semantic Evaluation (SemEval-2023)*, A. K. Ojha, A. S. Doruöz, G. Da San Martino, H. Tayyar Madabushi, R. Kumar, and E. Sartori, Eds., Toronto, Canada: Association for Computational Linguistics, Jul. 2023. [Online]. Available: <https://aclanthology.org/2023.semeval-1.208/>.
 - [29] Y. R. Regatte, R. R. R. Gangula, and R. Mamidi, "Dataset creation and evaluation of aspect based sentiment analysis in Telugu, a low resource language," in *Proceedings of the Twelfth Language Resources and Evaluation Conference*, N. Calzolari et al., Eds., Marseille, France: European Language Resources Association, May 2020. [Online]. Available: <https://aclanthology.org/2020.lrec-1.617/>.
 - [30] K. Ronny Mabokela, M. Primus, and T. Celik, "Advancing sentiment analysis for low-resourced african languages using pre-trained language models," *PLOS ONE*, vol. 20, Jun. 2025. [Online]. Available: <https://doi.org/10.1371/journal.pone.0325102>.
 - [31] S. Rosenthal, N. Farra, and P. Nakov, "SemEval-2017 task 4: Sentiment analysis in Twitter," in *Proceedings of the 11th International Workshop on Semantic Evaluation (SemEval-2017)*, S. Bethard, M. Carpuat, M. Apidianaki, S. M. Mohammad, D. Cer, and D. Jurgens, Eds., Vancouver, Canada: Association for Computational Linguistics, Aug. 2017. [Online]. Available: <https://aclanthology.org/S17-2088/>.

- [32] S. Ruder, P. Ghaffari, and J. G. Breslin, “A hierarchical model of reviews for aspect-based sentiment analysis,” in *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, Austin, Texas: Association for Computational Linguistics, 2016, pp. 999–1005. [Online]. Available: <https://aclanthology.org/D16-1103/>.
- [33] S. A. Salahudeen et al., *Hausanlp at semeval-2023 task 12: Leveraging african low resource tweetdata for sentiment analysis*, 2023. arXiv: 2304.13634 [cs.CL]. [Online]. Available: <https://arxiv.org/abs/2304.13634>.
- [34] M. Schuster and K. K. Paliwal, “Bidirectional recurrent neural networks,” in *IEEE Transactions on Signal Processing*, IEEE, vol. 45, 1997, pp. 2673–2681.
- [35] J. míd, P. Pibá, O. Prazak, and P. Kral, “Czech dataset for complex aspect-based sentiment analysis tasks,” in *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, N. Calzolari, M.-Y. Kan, V. Hoste, A. Lenci, S. Sakti, and N. Xue, Eds., Torino, Italia: ELRA and ICCL, May 2024. [Online]. Available: <https://aclanthology.org/2024.lrec-main.384/>.
- [36] E. F. Tjong Kim Sang and F. De Meulder, “Introduction to the conll-2003 shared task: Language-independent named entity recognition,” in *Proceedings of CoNLL*, 2003, pp. 142–147.
- [37] W. Zhang, Y. Deng, B. Liu, S. Pan, and L. Bing, “Sentiment analysis in the era of large language models: A reality check,” in *Findings of the Association for Computational Linguistics: NAACL 2024*, K. Duh, H. Gomez, and S. Bethard, Eds., Mexico City, Mexico: Association for Computational Linguistics, Jun. 2024. [Online]. Available: <https://aclanthology.org/2024.findings-naacl.246/>.