

Aspect-Based Sentiment Analysis of Indonesian Electric Vehicles on Media Social

Muh. Hanafi Halik ^{1*}, Hazriani ^{2*}, Nasrullah ^{3*}, Andani Achmad ^{4**}, Ingrid Nurtanio ^{5**},
Wardi ^{6**}

* Department of Computer Systems, Handayani University of Makassar

** Department of Electrical Engineering, Hasanuddin University, Makassar

muhhanafihalik@gmail.com ¹, hazriani@handayani.ac.id ², nasrullahstimik@handayani.ac.id ³, andani@unhas.ac.id ⁴,
ingrid@unhas.ac.id ⁵, wardi@unhas.ac.id ⁶

Article Info

Article history:

Received 2026-05-21

Revised 2026-07-02

Accepted 2026-08-07

Keywords:

*Aspect-Based Sentiment
Analysis,
Electric Vehicles,
Hybrid Model,
Media Social.*

ABSTRACT

The increasing adoption of electric vehicles (EVs) in Indonesia has generated diverse and unstructured public opinions across social media platforms. This study applies Aspect-Based Sentiment Analysis (ABSA) to classify Indonesian EV opinions across six predefined aspects: battery, design, price, performance, infrastructure, and general perception. The ABSA process was conducted using manual multi-label aspect-sentiment annotation rather than fully automatic aspect extraction, allowing one review to contain multiple EV aspects with corresponding sentiment labels. A total of 3,000 reviews were collected from YouTube, Instagram, and TikTok and used to evaluate three deep learning architectures: Bi-LSTM, IndoBERT, and Hybrid IndoBERT-BiLSTM. In the Hybrid architecture, IndoBERT was used as a Transformer-based contextual feature extractor, and the resulting contextual embeddings were passed into a Bi-LSTM layer to capture bidirectional sequential dependencies before final sentiment classification. The dataset was divided into training, validation, and testing sets using an 80:10:10 split. Model performance was measured using accuracy, precision, recall, and F1-score. The results show that IndoBERT achieved the highest average accuracy of 80.50%, followed by Hybrid IndoBERT-BiLSTM with 78.83% and Bi-LSTM with 71.67%. Although IndoBERT performed best overall, the Hybrid model showed competitive performance and better validation-loss stability in selected aspect-level contexts. These findings indicate the effectiveness of Transformer-based models for Indonesian EV sentiment analysis, while hybrid sequential modeling can provide a stable alternative for handling informal social media text.



This is an open access article under the [CC-BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.

I. INTRODUCTION

Climate change is currently the biggest worldwide obstacle to sustainable development. One of the most challenging mitigation targets is achieving carbon neutrality, where natural systems may safely absorb the resultant emissions[1]. The high utilization of fossil fuel vehicles is directly proportional to the increasing level of air pollution caused by exhaust emissions [2]. In Indonesia, the transportation sector is the largest contributor to greenhouse gas emissions. Consequently, transitioning to electric

vehicles (EVs) emerges as an effective solution, as EVs offer numerous advantages, including being environmentally friendly, energy-efficient, and having lower operational costs[3]. Indonesian automotive sector is still dominated by foreign original equipment manufacturers (OEMs) since indigenous businesses, such as EV startup producers, are vulnerable [4].

On the other hand, according to market appeal and expected use, EVs still have a tiny market share [5]. EVs have a lot of advantages, but their market penetration is quite limited, especially in developing nations [6]. The proportion

of EVs is positively impacted by all government policies, particularly those that interact with charging infrastructure, such as financial incentives, traffic laws that encourage EVs, and charging infrastructure. The policies supporting electric mobility reflect consumer preferences and impact high buying power and the greater uptake of EVs [7].

Globally, EVs perform well in urban areas where charging stations are more accessible and emissions are more concentrated [7]. However, in rural areas or regions with colder climates, the performance of EVs can be less reliable [8]. During this current EV adoption transition, several manufacturers have released their products to the market. Upon introducing a product to the public, manufacturers naturally seek to gauge public response toward their offerings [9].

Electric vehicles have become a major trend in the modern automotive world. With growing environmental concerns and the need to reduce greenhouse gas emissions, EVs are viewed as a potential solution to the negative environmental impacts of transportation. However, the implementation of electric vehicles depends not only on technical and economic factors but also on user perceptions toward these vehicles [10]. The current state of social media indicates that the discourse regarding electric vehicles in Indonesia is no longer confined to government policy spaces but has become a highly dynamic primary topic on social media. This development creates a fragmented public opinion ecosystem across various platforms, where each social media presents unique interaction characteristics. On the YouTube platform, public opinion is dominated by in-depth reviews from automotive experts and direct test drives. This triggers discussions in the comments section that tend to be technical and rational, where the audience critically questions battery durability efficiency, long-term maintenance costs, and objective performance comparisons between electric and conventional vehicles.

On the other hand, the Instagram platform presents a different perspective by emphasizing visual aspects, lifestyle, and brand prestige. User opinions on this platform are heavily influenced by the aesthetic design of both the vehicle's interior and exterior, as well as the image of modernity attached to EV ownership. Meanwhile, TikTok has emerged as a space for more spontaneous and viral opinions through short video content. On this platform, real-time issues such as queues at charging stations, real experiences facing technical obstacles on the highway, and emotional responses to subsidy price fluctuations become the center of attention. The diversity of data sources from these three platforms is crucial as they represent a vast spectrum of public sentiment, ranging from technical concerns to lifestyle enthusiasm.

Summarizing the dynamics of these three social media platforms, a common thread can be drawn: the perception of the Indonesian people is currently in a crucial transition period. There is a polarization of opinion between groups optimistic about environmental sustainability and groups still

skeptical of supporting infrastructure readiness and the long-term economic value of electric vehicles. Therefore, integrating and analyzing data from YouTube, Instagram, and TikTok as a unified research background is essential. This aims to capture a complete and accurate representation of public opinion. Therefore, integrating and analyzing data from YouTube, Instagram, and TikTok is important to capture a broader representation of social-media-based public opinion and to map the obstacles and opportunities for EV adoption in Indonesia.

Another challenge in analyzing EV-related social media comments is the informal and context-dependent nature of user-generated language. Comments from TikTok and Instagram often contain slang, abbreviations, Indonesian-English code-mixing, emojis, repeated characters, and sarcastic expressions. These linguistic features may change or obscure the intended sentiment, making aspect-level classification more difficult. Therefore, Indonesian EV sentiment analysis requires not only contextual language modeling but also careful preprocessing and manual annotation to preserve the intended meaning of informal expressions.

Sentiment analysis on Big Data requires specialized techniques for managing and extracting relevant information. Previous research indicates that data preprocessing is a vital step in sentiment analysis, as it can significantly affect the accuracy of classification results [11]. To overcome these limitations, a more precise approach is required, namely Aspect-Based Sentiment Analysis (ABSA).

Aspect-Based Sentiment Analysis (ABSA) is a form of sentiment analysis that considers sentiment based on specific aspects. Various methods can be used to analyze topics; deep learning was chosen as the method for this research [12]. ABSA aims to identify specific aspects discussed in a text (e.g., battery, design, price, etc.) and then determine the sentiment associated with each aspect. Thus, ABSA is capable of providing a far more detailed and comprehensive picture of public perception.

Sentiment analysis is a field of Natural Language Processing (NLP) capable of studying community comments or reviews in text form. The learning conducted in sentiment analysis provides text polarity trained into positive, neutral, and negative classes. An innovation in NLP, particularly in sentiment analysis, is the BERT pre-trained model. Furthermore, BERT is highly compatible with dataset labeling [13]. The presence of transformer models such as IndoBERT and DeBERTa serves as a new alternative to overcome existing limitations. IndoBERT, specifically designed and trained on an Indonesian language corpus, has shown promising performance in various NLP tasks [14].

The LSTM method is a variation of the Recurrent Neural Network (RNN) with several advantages, including the ability to store long-term information, read, and update previous information, as well as handle the vanishing

gradient problem in training that typically occurs in other RNN variations. LSTM generally consists of a cell, input gate, output gate, and forget gate [15]. Bidirectional Long Short-Term Memory (Bi-LSTM) is a deep learning combination method consisting of two LSTM layers. Thus, Bi-LSTM is a development that allows additional training by traversing input data twice: from left to right and from right to left [16].

Referring to the challenges in sentiment analysis described, this research is conducted to fill the gap in comparative studies between Recurrent Neural Network-based models (LSTM) and Transformer-based models (BERT) within the context of Aspect-Based Sentiment Analysis (ABSA). By utilizing data sourced from various social media platforms with diverse linguistic characteristics, this research is expected to provide empirical evidence regarding the most optimal and accurate model.

In this study, the aspect categories were determined by combining findings from previous EV-related studies, EV adoption considerations, and exploratory observation of social media comments. The final aspects include battery, design, price, performance, infrastructure, and general perception, representing both technical product attributes and broader consumer acceptance issues.

TABLE I
COMPARISON OF RELATED STUDIES

Study	Method	Data / Domain	Analysis Level	Research Gap
Widagdo et al. [2]	LSTM	YouTube EV comments	Sentence-level	No ABSA and no Transformer comparison
Ernawati et al. [3]	SVM, Naive Bayes	EV sentiment	General sentiment	Limited contextual understanding
Anwar et al. [9]	VADER	Indonesian EV opinion	General sentiment	Lexicon-based and not aspect-level
Indrayana et al. [10]	Naive Bayes	Twitter/X EV comments	General sentiment	No multi-aspect sentiment analysis
Nugroho et al. [14]	IndoBERT, DeBERTa	EV incentive policy	Multiclass sentiment	Focused on policy, not EV product aspects
This study	Bi-LSTM, IndoBERT, Hybrid IndoBERT-BiLSTM	YouTube, Instagram, TikTok EV opinions	ABSA	Six EV aspects with validated labels

Based on Table I, previous studies have mainly focused on general or sentence-level sentiment classification using classical machine learning, lexicon-based approaches, sequential deep learning models, or Transformer-based models. Although these studies have contributed to electric vehicle sentiment analysis, most of them have not explicitly modeled public sentiment at the aspect level across multiple social media platforms. In addition, the role of a Hybrid

IndoBERT-BiLSTM architecture for combining Transformer-based contextual representation and sequential dependency learning remains insufficiently explored in Indonesian electric vehicle opinion analysis.

Although previous EV sentiment studies have used classical machine learning approaches such as SVM and Naive Bayes, this study focuses on evaluating deep learning architectures within a multi-aspect ABSA framework. Therefore, classical machine learning models were not re-implemented as experimental baselines in this study. The comparison with classical approaches is discussed through related studies, while direct experimental comparison with TF-IDF-based SVM or Logistic Regression is acknowledged as a limitation and future research direction.

Therefore, the novelty of this study lies not merely in comparing IndoBERT and Bi-LSTM for general sentiment classification, but in evaluating their behavior within a multi-aspect ABSA framework for Indonesian electric vehicle opinions across heterogeneous social media platforms. This study constructs and evaluates an aspect-sentiment dataset covering six EV-related aspects: battery, design, price, performance, infrastructure, and general perception. Furthermore, this study proposes and evaluates a Hybrid IndoBERT-BiLSTM architecture, where IndoBERT functions as a contextual feature extractor and Bi-LSTM is added to capture sequential dependencies from the extracted contextual embeddings. This design enables the study to examine whether sequential modeling after Transformer-based representation can improve aspect-level sentiment classification stability on informal and noisy social media comments.

The main contributions of this study are summarized as follows:

1. This study provides a multi-platform ABSA analysis of Indonesian electric vehicle opinions collected from YouTube, Instagram, and TikTok.
2. This study applies multi-label aspect-sentiment annotation to capture multiple EV-related aspects within a single social media comment.
3. This study compares Bi-LSTM, IndoBERT, and Hybrid IndoBERT-BiLSTM models using aspect-level evaluation metrics, including accuracy, precision, recall, and F1-score.
4. This study provides empirical evidence that IndoBERT achieves the best overall accuracy, while the Hybrid model offers competitive performance and better validation-loss stability in selected aspect-level contexts.

II. METHODS

This research was conducted through several systematic stages, starting from collecting raw data from social media to evaluating the performance of the proposed deep learning models. Broadly speaking, the workflow of these research stages can be seen in Fig. 1.

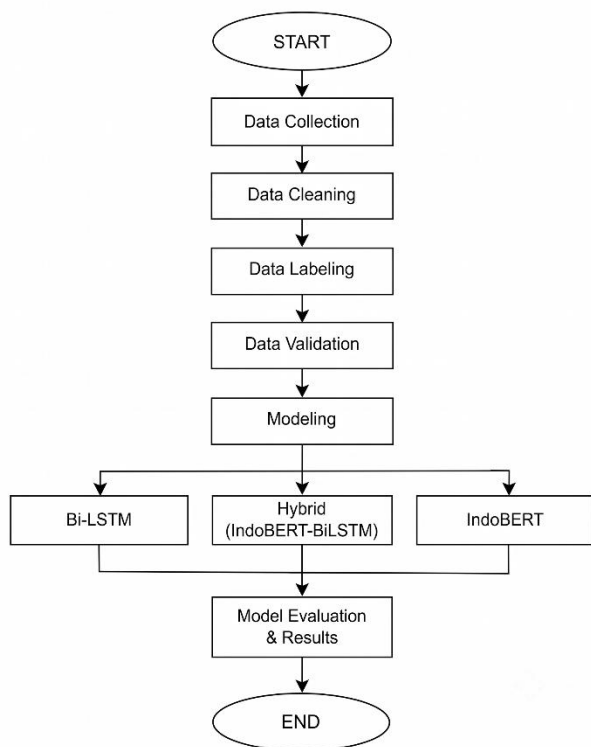


Fig. 1. Flowchart of the electric vehicle sentiment analysis research stages.

A. Data Collection

The initial stage of this research is the collection of public opinion data regarding electric vehicles in Indonesia. Raw data was extracted from three social media platforms with the highest user interaction rates: YouTube, Instagram, and TikTok. The extracted text focuses on the comment sections and reviews of content discussing electric vehicles. This data collection process yielded a corpus containing 3,000 reviews. This dataset size is considered representative for training natural language processing (NLP) models to recognize everyday conversational language patterns.

B. Data Preprocessing

The primary characteristics of social media text data are high noise levels, irregular sentence structures, and the frequent use of slang, abbreviations, emojis, hashtags, mentions, and code-mixed expressions. Therefore, the raw data underwent several preprocessing steps before annotation and model training.

- 1) *Cleansing*: Text elements that did not directly contribute to sentiment interpretation were removed, including URLs, hashtags, account mentions, numbers, punctuation marks, repeated symbols, emojis, and emoticons.
- 2) *Case Folding*: All text was converted to lowercase to ensure consistency and avoid duplicate representations caused by capitalization differences.

- 3) *Normalization*: Non-standard words, slang, and common abbreviations in Indonesian social media comments were normalized into standard Indonesian expressions when their meanings could be clearly identified.
- 4) *Stopword Removal*: Stopword removal was applied to common function words. Because negation words may affect sentiment polarity, their treatment was carefully checked during preprocessing and annotation.
- 5) *Tokenizing*: Cleaned sentences were split into token units. These tokens were then converted into numerical representations during the modeling stage.

In handling informal social media language, Indonesian-English code-mixed terms were retained when they carried important sentiment or aspect information, such as “fast charging,” “worth it,” “overprice,” or “maintenance.” Emojis and emoticons were removed during cleansing because the models in this study focused on textual input. However, this decision may remove some emotional cues from the comments. Sarcastic expressions were not automatically detected during preprocessing; instead, they were considered during manual annotation and discussed as one of the sources of classification difficulty.

C. Aspect Selection Rationale

The six EV-related aspects used in this study were selected based on a combination of previous studies, EV adoption considerations, and exploratory observation of the collected social media comments. These aspects were not selected randomly, but were determined to represent the main dimensions commonly discussed by users when evaluating electric vehicles.

Battery, price, and design were selected because these aspects are frequently discussed in previous EV-related sentiment and consumer perception studies. Battery represents concerns regarding driving range, charging duration, durability, safety, and replacement cost. Price represents affordability, subsidy, resale value, and perceived economic value. Design represents visual appeal, interior and exterior aesthetics, modernity, and product identity.

In addition, performance, infrastructure, and general perception were added based on exploratory observation during data collection and preliminary review of comments from YouTube, Instagram, and TikTok. Performance was included because users frequently discussed acceleration, driving comfort, handling, power, and technological capability. Infrastructure was included because many comments referred to charging station availability, charging access, after-sales service, and long-distance travel readiness. General perception was added to accommodate overall opinions that did not explicitly mention a specific technical aspect but still reflected acceptance, rejection, interest, or skepticism toward electric vehicles.

TABLE II
BASIS OF ASPECT SELECTION

Aspect	Selection Basis	Main Coverage
Battery	Previous studies and EV adoption consideration	Driving range, charging duration, durability, safety, replacement cost
Price	Previous studies and EV adoption consideration	Affordability, subsidy, resale value, ownership cost, economic value
Design	Previous studies and product perception	Interior, exterior, aesthetics, modernity, visual appeal
Performance	Exploratory observation of social media comments	Acceleration, comfort, handling, power, driving experience
Infrastructure	Exploratory observation and EV adoption barrier	Charging stations, charging access, after-sales service, travel readiness
General	Exploratory observation of holistic comments	Overall perception, brand opinion, purchase interest, acceptance or rejection

Based on this rationale, the six aspects were used as predefined categories in the ABSA annotation process. Each review was then examined to determine whether it contained sentiment toward one or more of these aspects.

D. Aspect and Sentiment Labeling

In this study, aspect extraction was conducted using six predefined EV-related aspects rather than fully automatic aspect extraction. The predefined aspects were battery, design, price, performance, infrastructure, and general perception. Each review was manually annotated using a multi-label aspect-sentiment scheme. For each aspect, annotators assigned one of four labels: positive, negative, neutral, or none. The “none” label indicates that the corresponding aspect was not discussed in the review. This annotation scheme allows a single review to contain multiple aspect-sentiment pairs, since social media users often discuss several EV attributes in one comment.

Before model training, the distribution of aspect-sentiment labels was examined to identify class imbalance. Because the dataset uses a multi-label ABSA structure, each review was evaluated independently for each of the six aspects. The “none” label was included to indicate that a specific aspect was not discussed in a review. This structure allows the model to learn both explicit aspect sentiment and the absence of aspect information.

E. Annotation Reliability and Data Validation

The annotation process was conducted to construct a reliable ground-truth dataset for the ABSA task. In this study, each review was annotated using a manual multi-label aspect-sentiment scheme. The annotation was performed based on six predefined EV-related aspects: battery, design, price, performance, infrastructure, and general perception. For each aspect, one of four labels was assigned: positive,

negative, neutral, or none. The “none” label indicates that the corresponding aspect was not discussed in the review.

To support annotation consistency, an annotation guideline was prepared before the labeling process. The guideline defined the scope of each aspect, sentiment polarity criteria, and examples of typical expressions found in Indonesian social media comments. The annotation process also considered informal language, abbreviations, slang, code-mixing, and implicit expressions commonly found in YouTube, Instagram, and TikTok comments.

To improve label reliability, an expert validation process was conducted. A total of 300 samples, representing 10% of the dataset, were reviewed by an expert validator. The validation results showed that 282 samples were consistent with the assigned labels, while 18 samples required revision due to ambiguous context and informal language. Based on these results, the observed agreement reached 94%, indicating a high level of labeling reliability for the constructed ABSA dataset.

The revised labels from the validation process were used as the final ground-truth labels for model training and evaluation. This validation step was important because label quality directly affects the validity of ABSA model performance.

TABLE III
ANNOTATION AND VALIDATION SUMMARY

Item	Description
Annotation scheme	Manual multi-label aspect-sentiment annotation
Aspects	Battery, design, price, performance, infrastructure, general
Sentiment labels	Positive, negative, neutral, none
Annotation unit	Review-level aspect-sentiment pair
Validation sample	300 reviews or 10% of dataset
Consistent labels	282 samples
Revised labels	18 samples
Agreement measure	Observed agreement
Observed agreement	94%
Final label use	Revised labels were used as ground truth

F. Data Splitting

The validated dataset was divided into three subsets: training data, validation data, and testing data using an 80:10:10 ratio. From the total of 3,000 reviews, 2,400 reviews were used as training data, 300 reviews were used as validation data, and 300 reviews were used as testing data.

The training set was used to learn model parameters, while the validation set was used to monitor model performance during training and support hyperparameter adjustment. The testing set was kept separate and used only for final model evaluation. This separation was applied to reduce data leakage and ensure that the final performance was measured on unseen data.

The same train-validation-test split was used for all modeling scenarios, including Bi-LSTM, IndoBERT, and Hybrid IndoBERT-BiLSTM, to ensure a fair comparison across models. The data were divided using a stratified random sampling strategy to maintain the proportional distribution of aspect-sentiment labels across the training, validation, and testing subsets.

TABLE IV
DATA SPLITTING AND VALIDATION STRATEGY

Subset	Ratio	Number of Reviews	Purpose
Training set	80%	2,400	Used to train model parameters
Validation set	10%	300	Used to monitor model performance and support hyperparameter adjustment
Testing set	10%	300	Used only for final model evaluation on unseen data
Total	100%	3,000	Complete validated dataset

This splitting strategy was applied consistently across all models to improve reproducibility and ensure that Bi-LSTM, IndoBERT, and Hybrid IndoBERT-BiLSTM were evaluated under the same data partition.

G. Modeling Scenarios

Three deep learning model architectures were developed and trained to compare their classification performance in the ABSA task :

- 1) *Bi-LSTM*: The Bidirectional Long Short-Term Memory model processes word token sequences simultaneously from both directions (forward and backward). This model is designed to retain the contextual information of long sentences and capture dependency relationships between words that are widely separated within the sentence structure.
- 2) *IndoBERT*: A Transformer-based language model specifically pre-trained with an Indonesian language corpus. In this study, the base IndoBERT model underwent a fine-tuning process. The self-attention mechanism in this architecture allows it to assign different importance weights to each word, resulting in a highly rich understanding of semantic context.
- 3) *Hybrid (IndoBERT + Bi-LSTM)*: The Hybrid IndoBERT-BiLSTM model was designed to combine Transformer-based contextual representation with sequential dependency learning. In this architecture, IndoBERT is placed in the first stage to generate contextual embeddings from Indonesian social media text. Instead of directly using the pooled output for classification, the last hidden states generated by IndoBERT are passed into a Bi-LSTM layer. The Bi-LSTM layer processes the contextual embeddings in both forward and backward directions to capture sequential dependencies that may still be relevant in informal comments, especially when sentiment expressions appear far from the target aspect. The final representation is then passed

through a dropout layer and a dense classification layer to predict the sentiment label for each aspect. The purpose of this Hybrid architecture is not to assume that sequential modeling will always outperform standalone IndoBERT, but to empirically examine whether adding Bi-LSTM after contextual embedding extraction can improve stability and aspect-level classification performance in selected conditions.

H. Evaluation Metrics

The performance of each model was evaluated using confusion matrix-based metrics, namely accuracy, precision, recall, and F1-score. These metrics were calculated for each aspect to provide a detailed evaluation of the ABSA task. In this study, the evaluated aspects include battery, design, price, performance, infrastructure, and general perception.

Accuracy measures the overall proportion of correctly predicted labels compared with the total number of testing samples. However, in Aspect-Based Sentiment Analysis (ABSA), accuracy alone may be insufficient because the distribution of aspect-sentiment labels is often imbalanced. In particular, the “none” label may dominate several aspects because not every review discusses all EV-related attributes.

Precision measures the proportion of correctly predicted labels among all labels predicted by the model for a particular class. High precision indicates that the model produces fewer false positive predictions. Recall measures the proportion of actual labels that are successfully identified by the model. High recall indicates that the model is able to retrieve more relevant instances of a particular sentiment class.

F1-score is the harmonic mean of precision and recall. This metric provides a more balanced evaluation because it considers both false positives and false negatives. Therefore, F1-score is important for evaluating ABSA performance, especially when some sentiment classes are less represented than others. In this study, accuracy, precision, recall, and F1-score were reported for each aspect to avoid relying only on aggregate accuracy.

III. RESULTS

A. Dataset Distribution

The dataset used in this study consists of 3,000 Indonesian-language reviews collected from YouTube, Instagram, and TikTok. Although the dataset size is relatively limited for a multi-aspect deep learning task, each review was annotated using a multi-label ABSA scheme across six predefined EV-related aspects: battery, design, price, performance, infrastructure, and general perception. Therefore, each review may contain more than one aspect-sentiment pair.

To provide a more transparent description of the dataset, the distribution of reviews across platforms, cumulative sentiment classes, and aspect-sentiment labels is reported in

this section. This distribution is important because ABSA datasets often contain class imbalance, especially when some aspects are discussed less frequently than others.

TABLE V
DATASET DISTRIBUTION BY SOCIAL MEDIA PLATFORM

Platform	Number of Reviews	Percentage
YouTube	1.000	33.33%
Instagram	1.000	33.33%
TikTok	1.000	33.33%
Total	3.000	100%

As shown in Table V, the dataset was balanced across the three platforms, with each platform contributing 1,000 reviews. This balanced platform composition was used to reduce source dominance and provide broader coverage of public opinion from different social media environments.

TABLE VI
ASPECT-SENTIMENT LABEL DISTRIBUTION

Aspect	Positive	Negative	Neutral	None	Total
Battery	74 (2.47%)	377 (12.57%)	92 (3.07%)	2,457 (81.90%)	3,000
Design	227 (7.57%)	64 (2.13%)	37 (1.23%)	2,672 (89.07%)	3,000
Price	431 (14.37%)	502 (16.73%)	274 (9.13%)	1,793 (59.77%)	3,000
Performance	157 (5.23%)	179 (5.97%)	76 (2.53%)	2,588 (86.27%)	3,000
Infrastructure	98 (3.27%)	368 (12.27%)	146 (4.87%)	2,388 (79.60%)	3,000
General	920 (30.67%)	939 (31.30%)	286 (9.53%)	855 (28.50%)	3,000

Table VI shows that the aspect-sentiment labels are not evenly distributed across aspects. The “none” label dominates several aspects, especially design, performance, battery, and infrastructure, indicating that not all reviews discuss every EV attribute. This condition is common in social media-based ABSA because users usually focus only on certain aspects in a single comment.

The general aspect has the largest number of explicit sentiment labels, with 920 positive, 939 negative, and 286 neutral labels. This indicates that many users expressed overall opinions toward electric vehicles without referring to a specific technical attribute. Price is also frequently discussed, with 431 positive, 502 negative, and 274 neutral labels, showing that affordability, resale value, and economic value are major issues in public EV discourse.

Battery and infrastructure show more negative than positive labels. Battery contains 377 negative labels compared to 74 positive labels, while infrastructure contains 368 negative labels compared to 98 positive labels. This indicates that public concerns are concentrated around battery durability, replacement cost, charging duration, charging station availability, and long-distance travel readiness. In contrast, design has more positive than negative

labels, suggesting that visual appearance and modern product image are relatively well received by users.

This imbalanced distribution justifies the use of precision, recall, and F1-score in addition to accuracy. Accuracy alone may overestimate model performance when the majority label, particularly “none,” dominates the dataset. Therefore, aspect-level evaluation metrics are necessary to provide a fairer interpretation of model performance in the ABSA task.

B. Training Dynamics and Convergence Evaluation

The training dynamics were analyzed to observe the learning behavior of each model during training. Table VII summarizes the training configuration and final loss values. These results are used only to interpret validation-loss behavior and should not be interpreted as a comprehensive computational efficiency evaluation.

TABLE VII
TRAINING CONFIGURATION AND FINAL LOSS METRICS

Model Architecture	Epochs	Batch Size	Final Training Loss	Final Validation Loss
IndoBERT	10	16	1.8	5.2
Bi-LSTM	50	32	2.2	4.9
Hybrid (IndoBERT - Bi-LSTM)	5	16	4.0	4.8

In-depth analysis of the loss curves provides visual confirmation of these learning behaviors.

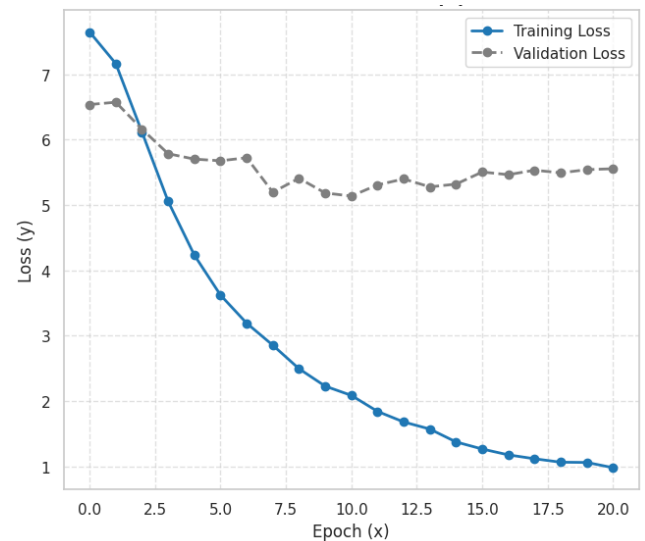


Figure 3. Bi-LSTM Loss Chart.

As shown in Fig. 3, the Bi-LSTM model experiences a steady decline in training loss, but its validation loss plateaus and begins to fluctuate after the 10th epoch. This indicates that while the standalone sequential model effectively learns the training data, it struggles to generalize the complex and noisy social media vocabulary to unseen data, capping its overall performance.

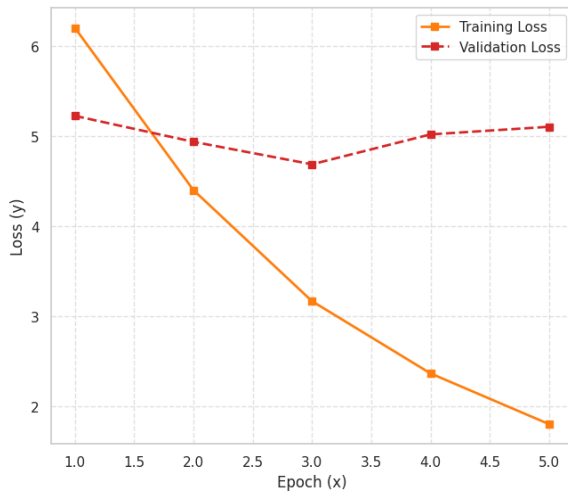


Figure 4. IndoBERT Loss Chart.

In contrast to the sequential model, Fig. 4 shows that IndoBERT is highly aggressive in learning training patterns, achieving the lowest training loss (1.8). However, its validation loss curve begins to rise sharply after the 5th epoch. This indicates that the Transformer model started "memorizing" noise within social media reviews rather than learning generalized sentiment patterns, leading to mild overfitting.

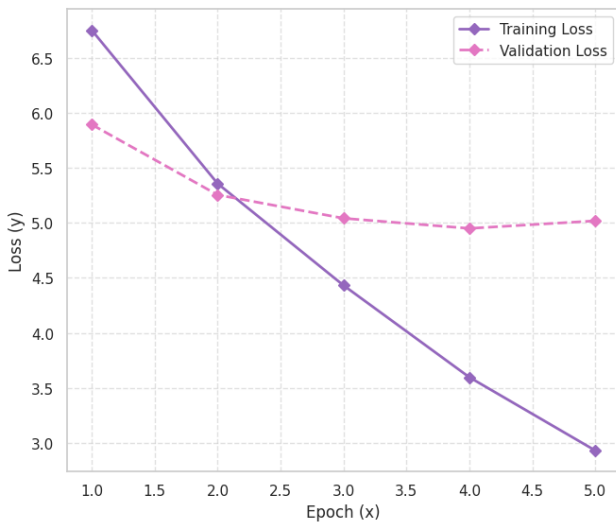


Figure 5. Hybrid (IndoBERT - Bi-LSTM) Loss Chart.

Figure 5 presents the validation-loss behavior of the Hybrid IndoBERT-BiLSTM model. Figure 5 shows that the Hybrid IndoBERT-BiLSTM model obtained the lowest final validation loss of 4.8 compared with IndoBERT and Bi-LSTM. This result indicates better validation-loss stability during training. However, this finding does not indicate that the Hybrid model is computationally more efficient or more stable in terms of training time, inference speed, memory consumption, model size, or GPU utilization. Therefore, the Hybrid model is interpreted only as a validation-loss-stable alternative in selected aspect-level contexts.

C. Confusion Matrix and Error Analysis

To further dissect the prediction accuracy of each architecture across all sentiment classes, including instances where no specific aspect was identified (labeled as "none"), a confusion matrix was generated for the IndoBERT, Bi-LSTM, and Hybrid models. The main diagonal of each matrix represents the True Positive (TP) values, indicating correct predictions, while values outside the diagonal represent misclassifications.

Actual (y)	none	1042	107	69	78
	negative	51	138	22	28
	neutral	20	21	39	16
	positive	31	45	21	72
		none	negative	neutral	positive
		Predicted (x)			

Figure 6. Confusion Matrix of the Bi-LSTM model.

Fig. 6 illustrates the confusion matrix for the baseline Bi-LSTM model. While it successfully identified 1042 instances of the "none" class, its performance significantly drops when classifying explicit sentiments. For example, it only correctly predicted 138 negative instances and 72 positive instances. A notable weakness is seen in its tendency to misclassify true positive (31 instances) and true negative (51 instances) reviews as "none," indicating that the sequential model struggles to capture sentiment features without strong contextual embeddings.

Actual (y)	Predicted (x)			
	none	negative	neutral	positive
none	1123	57	74	42
negative	24	149	34	32
neutral	7	9	66	14
positive	13	27	18	111

Figure 7. Confusion Matrix of the IndoBERT model.

The performance improvement is clearly visible in the IndoBERT confusion matrix (Fig. 7). The True Positive

values on the main diagonal are consistently higher across all classes. It correctly identified 1123 "none" instances, 149 negative instances, and significantly improved positive sentiment detection with 111 correct predictions. The bidirectional self-attention mechanism of BERT allows it to better distinguish between sentiment classes. However, it still occasionally misclassifies "none" instances into negative (57 instances) or neutral (74 instances), showing a slight oversensitivity to certain keywords.

Actual (y)	none	1084	66	80	66
	negative	22	157	32	28
	neutral	6	12	63	15
	positive	14	22	17	116
		none	negative	neutral	positive
		Predicted (x)			

Figure 8. Confusion Matrix of the Hybrid (IndoBERT - Bi-LSTM) model.

Figure 8 presents the confusion matrix for the Hybrid model. The model achieved the highest correct predictions for the negative class and improved positive detection compared with standalone IndoBERT. It also reduced several misclassifications of explicit sentiment labels into the "none" category. This result suggests that adding a Bi-LSTM layer after IndoBERT feature extraction may help retain sequential information in some complex or ambiguous sentences. However, the improvement is not consistent across all aspects; therefore, the Hybrid model should be interpreted as a competitive alternative rather than a consistently superior architecture.

Based on the three confusion matrices, the most frequent classification errors occurred between explicit sentiment labels and the "none" label. This indicates that the models sometimes failed to distinguish aspect-specific sentiment from general opinion. Bi-LSTM showed a stronger tendency to misclassify explicit sentiment as "none," while IndoBERT reduced this error but still misclassified some "none" instances into negative or neutral classes. The Hybrid model reduced several misclassifications of explicit sentiment labels into "none," but errors remained in neutral, implicit, and context-dependent expressions. These patterns indicate that Battery, Price, Infrastructure, and General are more difficult to classify because they often contain technical concerns, ownership cost, charging accessibility, and broad or implicit opinions.

D. Detailed Aspect-Based Classification Results

The core performance of the Aspect-Based Sentiment Analysis (ABSA) is detailed in TABLE VIII. To avoid relying only on aggregate model performance, this study reports the performance of each model separately for each EV aspect. This aspect-level evaluation is important because different aspects, such as battery and infrastructure, may contain different sentiment patterns and classification difficulties.

TABLE VIII
DETAILED ASPECT-LEVEL PERFORMANCE USING PRECISION, RECALL, F1-SCORE, AND ACCURACY

Aspect	Model	Precision	Recall	F1-Score	Accuracy
Battery	Bi-LSTM	0.87	0.84	0.85	0.84
	IndoBERT	0.90	0.88	0.88	0.88
	Hybrid	0.90	0.88	0.89	0.88
Design	Bi-LSTM	0.88	0.87	0.87	0.87
	IndoBERT	0.95	0.95	0.94	0.95
	Hybrid	0.93	0.93	0.93	0.92
Price	Bi-LSTM	0.67	0.64	0.65	0.64
	IndoBERT	0.81	0.75	0.77	0.75
	Hybrid	0.80	0.71	0.74	0.71
Performance	Bi-LSTM	0.83	0.80	0.81	0.80
	IndoBERT	0.91	0.91	0.91	0.91
	Hybrid	0.91	0.89	0.90	0.89
Infrastructure	Bi-LSTM	0.76	0.71	0.73	0.71
	IndoBERT	0.88	0.86	0.87	0.86
	Hybrid	0.86	0.84	0.85	0.84
General	Bi-LSTM	0.50	0.44	0.41	0.44
	IndoBERT	0.67	0.48	0.42	0.48
	Hybrid	0.66	0.49	0.41	0.49

Table VIII presents the aspect-level performance of each model using precision, recall, F1-score, and accuracy. This evaluation is important because ABSA data are often imbalanced, especially due to the dominance of the "none" label in several aspects. Therefore, accuracy alone is insufficient to represent model performance, and F1-score is needed to provide a more balanced interpretation across sentiment classes.

Overall, IndoBERT achieved the strongest observed performance across most aspects. It obtained particularly high scores in Design and Performance, with accuracy values of 0.95 and 0.91, respectively. These results indicate that Transformer-based contextual representation is effective in capturing explicit evaluative expressions related to visual appearance, comfort, acceleration, and technological capability. Although IndoBERT also performed better than the other models in Price and Infrastructure, the scores in these aspects were lower than those in Design and Performance, suggesting that economic and infrastructure-related opinions are more complex and context-dependent.

Battery and Infrastructure require special attention because both aspects show different sentiment characteristics compared with Design and Performance. Battery-related comments often involve technical and long-term ownership concerns, such as charging duration,

durability, safety, warranty, and replacement cost. Infrastructure-related comments frequently refer to charging station availability, charging access, after-sales service, and long-distance travel readiness. These characteristics make both aspects more context-dependent and more difficult to classify using aggregate evaluation alone.

The General aspect produced the lowest F1-score across all models, indicating that broad and implicit opinions are more difficult to classify than aspect-specific comments. Similarly, Battery, Price, and Infrastructure showed more variable performance because they often involve technical concerns, ownership cost, charging accessibility, and implicit criticism. These findings indicate that aspect-level evaluation is necessary in ABSA, since each EV aspect has different linguistic and sentiment characteristics.

The Hybrid IndoBERT-BiLSTM model achieved the highest F1-score in the Battery aspect, with a score of 0.89. It also slightly outperformed IndoBERT in recall and accuracy for the General aspect. However, the performance gap was very small. Thus, Hybrid IndoBERT-BiLSTM should be interpreted as a competitive and validation-loss-stable alternative in selected contexts, not as a model that consistently outperforms IndoBERT.

In addition, the overall accuracy difference between IndoBERT and Hybrid IndoBERT-BiLSTM was only 1.67%. Since this study did not conduct statistical significance testing, this difference should be interpreted descriptively rather than as statistically significant evidence of superiority. Therefore, IndoBERT is reported as the best-performing model based on observed accuracy, while the Hybrid model remains a competitive alternative in selected aspect-level contexts.

IV. DISCUSSION

This study applied Aspect-Based Sentiment Analysis (ABSA) to Indonesian electric vehicle opinions collected from YouTube, Instagram, and TikTok. The analysis was conducted using six predefined aspects: battery, design, price, performance, infrastructure, and general perception. The results show that IndoBERT achieved the highest overall average accuracy of 80.50%, followed by Hybrid IndoBERT-BiLSTM with 78.83% and Bi-LSTM with 71.67%. These findings indicate the effectiveness of Transformer-based models for Indonesian EV sentiment analysis.

However, the performance difference between IndoBERT and Hybrid IndoBERT-BiLSTM was relatively small and was not statistically tested in this study. Therefore, the comparison should be interpreted descriptively. IndoBERT can be considered the best-performing model based on observed average accuracy, while Hybrid IndoBERT-BiLSTM should be interpreted as a competitive and validation-loss-stable alternative in selected aspect-level contexts, not as a model that consistently outperforms IndoBERT. It should be noted that the stability of the Hybrid

model in this study refers only to validation-loss behavior, not computational efficiency, because training time, inference speed, memory consumption, model size, and GPU utilization were not measured.

The aspect-level results also show that public sentiment toward electric vehicles is multidimensional. Positive sentiment tends to appear in the design and performance aspects, while price, battery, and infrastructure remain major public concerns. This indicates that ABSA is more suitable than general sentiment classification for analyzing Indonesian EV opinions because users may express different sentiments toward different vehicle attributes within the same comment.

From a practical perspective, the ABSA results can help electric vehicle stakeholders translate public sentiment into more targeted actions. Positive sentiment toward Design and Performance can be used by EV manufacturers and marketers to emphasize visual appeal, modern product image, driving comfort, acceleration, and technological features in promotional campaigns. Meanwhile, negative sentiment toward Price, Battery, and Infrastructure indicates that public concerns are concentrated around affordability, ownership cost, battery durability, replacement cost, charging duration, and charging station availability. These findings can support product improvement, after-sales communication, consumer education, and EV infrastructure policy.

TABLE IX
PRACTICAL IMPLICATIONS OF ABSA FINDINGS

Aspect	Main Finding	Practical Use
Design	More positive sentiment than negative sentiment	Emphasize exterior, interior, modernity, and visual identity in marketing campaigns
Performance	Positive response toward driving experience and technology	Highlight acceleration, comfort, handling, and smart features
Price	High negative sentiment and frequent discussion	Communicate total cost of ownership, subsidies, financing options, and resale value
Battery	More negative than positive sentiment	Strengthen battery warranty transparency, durability information, and replacement-cost explanation
Infrastructure	More negative than positive sentiment	Expand charging access, improve SPKLU visibility, and strengthen after-sales support
General	Mixed overall perception	Build public trust through education, user testimonials, and evidence-based EV ownership information

Table IX shows that ABSA findings are not only useful for model evaluation, but also for identifying which EV attributes should be promoted, improved, or supported through policy intervention.

Informal social media language remains an important challenge in this study. Comments from TikTok and Instagram are often short, spontaneous, and highly contextual. Users frequently use slang, abbreviations,

Indonesian-English code-mixing, emojis, and sarcastic expressions to convey their opinions. These expressions can make sentiment classification more difficult because the polarity is not always explicitly stated through formal words.

For example, code-mixed expressions such as “worth it,” “overprice,” “fast charging,” or “maintenance” may carry important sentiment and aspect information in EV discussions. Similarly, sarcastic expressions may appear neutral at the lexical level but actually imply negative sentiment toward battery reliability, price, or charging infrastructure. This condition is particularly challenging for ABSA because a comment may not explicitly mention an aspect but still refer to it implicitly. Therefore, classification errors may occur when the model fails to distinguish between explicit aspect sentiment and general informal opinion.

Although IndoBERT provides stronger contextual representation than Bi-LSTM, informal and sarcastic expressions remain difficult to classify because their meaning often depends on platform context, shared user assumptions, or pragmatic interpretation. The Hybrid model may help in selected cases by processing contextual embeddings sequentially, but it does not fully solve the challenges of slang, code-mixing, emoji-based sentiment, and sarcasm. These findings indicate that future Indonesian ABSA research should consider slang normalization dictionaries, sarcasm detection, emoji-aware embeddings, and platform-specific modeling.

Future studies should expand the dataset size, involve multiple independent annotators, include classical baselines such as TF-IDF-based Logistic Regression and Linear SVM, apply statistical significance testing, and evaluate computational efficiency using training time, inference speed, memory consumption, and GPU utilization.

The use of YouTube, Instagram, and TikTok may also introduce potential platform bias. Although the dataset was balanced across platforms, with each platform contributing 1,000 reviews, each platform has different user behavior, content format, and linguistic style. YouTube comments tend to be longer and more technical because users often respond to automotive reviews or test-drive videos. Instagram comments are more visually oriented and may emphasize design, lifestyle, and brand image. Meanwhile, TikTok comments are usually shorter, more spontaneous, informal, and emotionally expressive. These differences may influence aspect frequency, sentiment polarity, and classification difficulty. Therefore, the findings should be interpreted as social-media-based EV opinion rather than a fully generalizable representation of all Indonesian consumers.

V. LIMITATIONS

This study has several limitations that should be acknowledged for a more accurate interpretation of the results. First, the dataset consists of 3,000 reviews collected only from YouTube, Instagram, and TikTok within a specific timeframe. Although these platforms represent active social

media discussions about electric vehicles, the dataset may not fully capture the diversity of Indonesian public opinion or long-term changes in sentiment as the electric vehicle market continues to evolve. In addition, different platform characteristics may introduce platform bias because YouTube, Instagram, and TikTok have different user behaviors, content formats, and linguistic patterns.

Second, the dataset size remains limited for a multi-aspect deep learning task. In addition, the aspect-sentiment distribution shows a strong imbalance, particularly because the “none” label dominates several aspects such as battery, design, performance, and infrastructure. Therefore, the findings should be interpreted as an empirical evaluation on a manually annotated Indonesian EV opinion dataset rather than as a universal benchmark for all EV-related social media discourse.

Third, although the dataset was manually annotated and validated by an expert reviewer, this study used observed agreement rather than advanced inter-annotator reliability metrics such as Cohen’s Kappa. Future studies should involve multiple independent annotators and report inter-annotator agreement to further strengthen the reliability of ABSA labeling.

Fourth, this study did not conduct statistical significance testing to confirm whether the 1.67% accuracy difference between IndoBERT and Hybrid IndoBERT-BiLSTM was statistically meaningful. Therefore, the comparison between these two models is interpreted descriptively based on the observed evaluation metrics. Future research should apply paired significance tests, such as McNemar’s test or bootstrap resampling, to validate whether small differences in model performance are statistically significant.

Fifth, the high prevalence of informal language, slang, abbreviations, Indonesian-English code-mixing, and implicit expressions in social media comments remains a challenge for ABSA. Emojis and non-textual symbols were removed during preprocessing, which may eliminate some emotional cues. In addition, this study did not include a dedicated sarcasm detection module or emoji-aware sentiment representation. Future studies should consider slang normalization, sarcasm detection, and emoji-aware embeddings to improve sentiment interpretation on short-form platforms such as TikTok and Instagram.

Sixth, this study did not include classical machine learning baselines such as Logistic Regression or Support Vector Machine. Therefore, the comparison focuses on deep learning architectures only, namely Bi-LSTM, IndoBERT, and Hybrid IndoBERT-BiLSTM. Future studies should include TF-IDF-based Logistic Regression and Linear SVM as baseline models to evaluate whether deep learning models provide substantial performance improvement over classical text classification approaches.

Finally, this study focused on classification performance and validation-loss behavior, but did not conduct a comprehensive computational efficiency analysis. Training

time, inference speed, memory consumption, model size, and GPU utilization were not measured. Therefore, the claim regarding the Hybrid IndoBERT-BiLSTM model is limited to validation-loss stability, not computational efficiency. Future research should include computational benchmarking to evaluate whether the Hybrid architecture is practical for real-time or large-scale deployment.

VI. CONCLUSION

This study applied Aspect-Based Sentiment Analysis (ABSA) to Indonesian electric vehicle opinions collected from YouTube, Instagram, and TikTok. Six predefined aspects were analyzed, namely battery, design, price, performance, infrastructure, and general perception. The results show that IndoBERT achieved the highest overall average accuracy of 80.50%, followed by Hybrid IndoBERT-BiLSTM with 78.83% and Bi-LSTM with 71.67%. These findings indicate that Transformer-based contextual representation is effective for Indonesian EV sentiment classification, especially in handling informal and context-dependent social media text.

However, the difference between IndoBERT and Hybrid IndoBERT-BiLSTM was relatively small and was not statistically tested. Therefore, the comparison should be interpreted descriptively. IndoBERT can be considered the best-performing model based on observed accuracy, while Hybrid IndoBERT-BiLSTM should be interpreted as a competitive and validation-loss-stable alternative in selected aspect-level contexts.

The theoretical contribution of this study lies in showing that Indonesian EV sentiment is not adequately represented by general sentiment polarity alone. Public opinion toward electric vehicles is aspect-dependent, where positive sentiment toward design and performance may coexist with negative sentiment toward price, battery, and infrastructure. This finding strengthens the relevance of ABSA for Indonesian social media analysis, especially in domains where users express multidimensional opinions within a single comment.

Methodologically, this study contributes to Indonesian ABSA research by applying a manual multi-label aspect-sentiment annotation scheme across six EV-related aspects and evaluating deep learning models at the aspect level rather than only through aggregate performance. The findings also show that adding sequential modeling after Transformer-based representation does not automatically outperform standalone IndoBERT, but may provide validation-loss stability in selected contexts. This provides a more cautious empirical basis for future hybrid ABSA model development in Indonesian-language social media data.

Practically, the findings can support EV marketing, product improvement, ownership-cost communication, battery warranty transparency, and charging infrastructure policy. Future studies should include classical baselines such as Logistic Regression and SVM, apply statistical

significance testing, expand the dataset, involve multiple independent annotators, and evaluate computational efficiency.

ACKNOWLEDGMENT

The authors would like to express their sincere gratitude to the Master's Program in Computer Systems, Universitas Handayani Makassar, for providing the academic environment and support necessary to complete this research. Special thanks are also extended to the research supervisors for their valuable guidance, insights, and technical advice during the development of the hybrid models.

REFERENCES

- [1] J. Li, S. He, S. Member, and Q. Yang, "A Comprehensive Review of Second Life Batteries Toward Sustainable Mechanisms: Potential, Challenges, and Future Prospects," *IEEE Trans. Transp. Electr.*, vol. 9, no. 4, pp. 4824–4845, 2023, doi: 10.1109/TTE.2022.3220411.
- [2] Sri Widagdo Adika Nuresa Qodri, Krisna Edi Nugroho, Fachruddin S, Akbar Rizky, and Nisrina P, "Analisis Sentimen Mobil Listrik di Indonesia Menggunakan Long-Short Term Memory (LSTM)," *J. Fasilkom*, vol. 13, no. 3, pp. 416–423, 2023, doi: <https://doi.org/10.37859/jf.v13i3.6303>.
- [3] N. W. Ernawati, I. N. Satya Kumara, and W. Setiawan, "Perbandingan Metode Klasifikasi Support Vector Machine Dan Naïve Bayes Pada Analisis Sentimen Kendaraan Listrik," *J. SPEKTRUM*, vol. 10, no. 3, p. 106, 2023, doi: <https://doi.org/10.24843/SPEKTRUM.2023.v10.i04.p12>.
- [4] T. W. Sasongko, U. Ciptomulyono, B. Wirjodirdjo, and A. Prastawa, "Identification of electric vehicle adoption and production factors based on an ecosystem perspective in Indonesia," *Cogent Bus. Manag.*, vol. 11, no. 1, p., 2024, doi: 10.1080/23311975.2024.2332497.
- [5] Tu, J.-C., & Yang, C. (2019). Key Factors Influencing Consumers' Purchase of Electric Vehicles. *Sustainability*, 11(14), 3863. <https://doi.org/10.3390/su11143863>
- [6] R. Hidayat and J. Cowie, "ScienceDirect ScienceDirect A framework to explore policy to support the adoption of electric vehicles in developing nations: A case study of Indonesia," *Transp. Res. Procedia*, vol. 70, pp. 364–371, 2023, doi: 10.1016/j.trpro.2023.11.041.
- [7] N. Rietmann and T. Lieven, "How policy measures succeeded to promote electric mobility – Worldwide review and outlook," *J. Clean. Prod.*, vol. 206, pp. 66–75, 2019, doi: <https://doi.org/10.1016/j.jclepro.2018.09.121>.
- [8] J. Lee, F. Baig, M. A. H. Talpur, and S. Shaikh, "Public intentions to purchase electric vehicles in Pakistan," *Sustainability*, vol. 13, no. 10, p. 5523, 2021.
- [9] M. Taufiq Anwar, D. Riandhita Arief Permana, P. STMI Jakarta, P. Sistem Informasi Industri Otomotif, J. Letjen Suprpto No, and J. Pusat, "Analisis Sentimen Masyarakat Indonesia Terhadap Produk Kendaraan Listrik Menggunakan VADER," *Tek. Inform. dan Sist. Inf.*, vol. 10, no. 1, pp. 783–792, 2023, doi: <https://doi.org/10.35957/jatisi.v10i1.3406>.
- [10] R. Indrayana, M. Solikhin, and M. Si, "Analisis Sentimen Kendaraan Listrik Pada Media Sosial Twitter Menggunakan Algoritma Naive Bayes Classifier," *Univ. Negeri Malang*, vol. 8, p. 2023, 2023, doi: <https://doi.org/10.1080/09638288.2019.1595750%0A>.
- [11] R. talita Trista and E. tri Asmoro, "Analisis Sentimen dalam Data Big Data," *J. Multidiscip. Inq. Sci. Technol. Educ. Res.*, vol. 2, no. 1, pp. 700–706, 2025, doi: <https://doi.org/10.32672/mister.v2i1.2518>.

-
- [12] R. Naquitasia, D. H. Fudholi, and L. Iswari, "Analisis Sentimen Berbasis Aspek pada Wisata Halal dengan Metode Deep Learning," *J. Teknoinfo*, vol. 16, no. 2, p. 156, 2022, doi: 10.33365/jti.v16i2.1516.
- [13] N. A. R. Putri and Ardiansyah, "Analisis Sentimen Terhadap Kemajuan Kecerdasan Buatan di Indonesia Menggunakan BERT dan RoBERTa," *J. Sains dan Inform.*, vol. 9, no. 2, pp. 136–145, 2023, doi: <https://doi.org/10.34128/jsi.v9i2.649>.
- [14] M. B. Nugroho, A. K. Zyen, and N. A. Widiastuti, "Multiclass Sentiment Analysis of Electric Vehicle Incentive Policies Using IndoBERT and DeBERTa Algorithms," vol. 9, no. 3, pp. 910–919, 2025, doi: <http://jurnal.polibatam.ac.id/index.php/JAIC>.
- [15] A. R. Isnain, H. Sulistiani, B. M. Hurohman, and A. Nurkholis, "Analisis Perbandingan Algoritma LSTM dan Naive Bayes untuk Analisis Sentimen," vol. 8, no. 2, pp. 299–303, 2022, doi: <https://doi.org/10.33365/jti.v16i2.1751>.
- [16] D. I. Afidah, S. F. Handayani, R. W. Pratiwi, and S. N. Sari, "Sentimen Ulasan Destinasi Wisata Pulau Bali Menggunakan Bidirectional Long Short Term Memory Sentiment of Balis Touristic Destination Reviews Using Bidirectional Long Short Term Memory," vol. 21, no. 3, 2022, doi: 10.30812/matrik.v21i3.1402.