

Stock Sentiment Prediction of LQ-45 Based on News Articles Using LSTM

Kristina^{1*}, I Made Agus Dwi Suarjaya^{2*}, Anak Agung Ketut Agung Cahyawan Wiranatha^{3*}

*Teknologi informasi, Universitas Udayana

kristina@student.unud.ac.id¹, agussuarjaya@it.unud.ac.id², agung.cahyawan@gmail.com³

Article Info

Article history:

Received 2025-06-02

Revised 2025-06-24

Accepted 2025-07-03

Keyword:

LQ-45,
Realtime Sentiment Prediction,
LSTM,
Financial News RSS,
Stock Market.

ABSTRACT

The growth in the number of investors in the financial market indicates that the investment world is currently experiencing rapid development. One of the long-term investment instruments that has experienced significant growth in the financial market is the stock market. Growth data as of September 2024 sourced from the Indonesia Stock Exchange report reveals that the number of stock market investors has reached more than 6 million single investor identification (SID). The share price of a company can be influenced by two main factors, namely internal factors and external factors. Internal factors come from within the company itself, while external factors come from conditions outside the company. Model development uses the Long Short-Term Memory (LSTM) method to predict daily stock sentiment in realtime. Labeling is done based on the history of stock price changes taken from Yahoo Finance. Stock market news data is obtained automatically every day through Really Simple Syndication (RSS) with the help of cronjob. The results of the LSTM model showed good performance, with a macro F1-Score of 0.73, a macro precision of 0.72, and a macro recall of 0.75. When compared to baseline models such as Logistic Regression, Naive Bayes, and Random Forest which only achieve a macro F1-Score of 0.58, 0.54, and 0.65, respectively, it can be concluded that the developed LSTM model has superior performance. This research can provide new considerations to investors, so as to reduce the risk of loss due to errors in choosing companies to invest in.



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

I. INTRODUCTION

Capital markets play a key role in liquidity and economic development. They also reflect a country's economic conditions and serve as an investment tool that drives broad economic growth[1]. One of the long-term investment instruments that has experienced significant growth in the financial market is the stock market. This significant growth can be seen from the report of the Indonesia Stock Exchange, which revealed that the number of stock market investors had reached more than 6 million single investor identification (SID), or more precisely, 6,001,573 SID based on data as of September 2024. This marks a growth of more than 744 thousand new stock investors[2].

The fluctuation of a company's stock price can be influenced by two main factors: internal factors and external

factors[3]. Internal factors come from within the company itself, such as rights issue policies, total assets, total liabilities, and public sentiment towards information or news related to the company. Meanwhile, external factors come from conditions outside the company, such as fluctuations in global oil prices, global gold prices, the exchange rate of the rupiah, and global market sentiment. Both of these factors play an important role in determining the rise and fall of stock prices in the capital market. Company sentiment plays a role in influencing stock prices.

A study in 2020 by Yinghao Ren, Fangqing Liao, and Yongjing Gong proposed a stock price trend prediction model using Deep Bidirectional LSTM (DBLSTM) combined with an automatic labeling method based on the BIAS indicator [4]. The BIAS indicator calculates the percentage difference between the closing price of a stock on the day the news is

published and the average closing price for several days afterward (in this study for 6 days). The labeling process is done automatically (auto-labeling) by looking at whether the BIAS value is positive or negative, so that the news is classified as “Up” or “Down” without the need for human intervention. The experimental results show that the built model is able to achieve an accuracy of 0.64, which indicates that the use of historical stock price data in labeling can improve the accuracy of predicting market trends based on financial news.

A 2024 study by Khonde explored the use of sentiment analysis on financial news headlines to predict stock performance by applying techniques such as Term Frequency-Inverse Document Frequency (TF-IDF), Long Short-Term Memory (LSTM), and SentiWordNet for sentiment classification. The approach achieved a prediction accuracy of 86% by integrating sentiment analysis with a regression-based stock prediction model[5].

Another study by Lakshya conducted in 2022 focused on stock price prediction using news sentiment analysis through deep learning approaches, including Recurrent Neural Network (RNN) and Deep Neural Network (DNN)[6]. News data were processed using tools such as VADER and the Natural Language Toolkit (NLTK) to generate sentiment polarity. The findings show that combining RNN with VADER sentiment polarity yields more accurate prediction results than other methods, with an average Mean Absolute Percentage Error (MAPE) of 2.1 for major companies such as Apple and Tesla. These results highlight the significant influence of news sentiment on stock prices and confirm its potential as a powerful tool in stock market prediction.

A 2022 study by Ibnu, Syafaah and Lestandy aimed to classify emotions in text using the TF-IDF method for word weighting and LSTM for classification[7]. The dataset used contained 18,000 data points classified into six emotion categories. The results showed that the LSTM method combined with TF-IDF achieved an accuracy of 97.50%, while the LinearSVC method with TF-IDF achieved only 89%. These findings underscore the effectiveness of LSTM in capturing temporal patterns in emotional text, and highlight the significant performance gap between LSTM and conventional classification methods such as LinearSVC.

A study conducted in 2025 by Saputra highlights the role of fundamental analysis in predicting stock price movements by incorporating macroeconomic variables such as the exchange rate, interest rate, and the composite stock price index (CSPI)[8]. These indicators reflect broader economic conditions that significantly influence corporate performance and market trends. The findings show that integrating fundamental data with technical indicators and sentiment analysis in an LSTM-based model can improve prediction accuracy, as indicated by the low RMSE values achieved by certain variable combinations.

A study conducted in 2024 by Bouchra shows that the use of Word2Vec can enhance the accuracy of the Bi-LSTM model in sentiment analysis[9]. The study utilized data from

the social media platform Twitter, combining the Word2Vec word embedding technique with the Bi-LSTM architecture. The resulting model achieved a high accuracy of 96.4%, demonstrating the effectiveness of this combination in capturing word context and sentence structure in text data. The findings also revealed that neutral sentiment dominated public perception of eight national television channels, with positive sentiment outweighing negative sentiment on most channels. Overall, the study offers insights into public satisfaction trends regarding national television broadcasts and reinforces the role of Word2Vec in improving deep learning-based text classification performance.

A study conducted in 2020 by Galih Pradana shows that sentiment on social media can influence stock price movements. By applying the K-Nearest Neighbor (KNN) method to classify public sentiment toward telecommunications companies in Indonesia, it was found that positive sentiment had a correlation of 0.659 with stock price increases. These findings indicate that public sentiment can serve as an important indicator for predicting stock movement trends[10].

Based on previous research, stock sentiment has been proven to have a significant influence on stock price movements. However, most previous studies still use manual sentiment labeling methods or sentiment dictionaries, which may contain subjectivity and are unable to capture actual market dynamics. Therefore, this study innovates by applying an automatic sentiment labeling process based on historical stock price data from Yahoo Finance. This approach is considered more objective because the labels Naik, Turun, and Netral are determined based on changes in the LQ-45 stock price after the article is published, making the sentiment representation more relevant to the actual market reaction. This innovation not only reduces dependence on human annotators but also opens up opportunities to build adaptive and sustainable systems in the face of ever-evolving news data. Additionally, this research develops an article-based sentiment prediction system using the Long Short-Term Memory (LSTM) algorithm combined with two feature representation methods: TF-IDF and GloVe word embedding. TF-IDF is used to measure word weights based on their frequency across the corpus, while GloVe is used to capture semantic relationships between words in a sentence. The integration of these various methods is implemented in an automated service system that updates predictions daily in real-time, providing more accurate and contextual market sentiment information. The system offers innovation by utilizing RSS feeds with real-time data collection, ensuring sentiment predictions are always aligned with current and outstanding news. With this approach, users not only gain insights from financial news but can also understand sentiment trends and patterns that impact more precise investment decision-making.

II. METHODS

The methodology used in this research aims to predict daily sentiment based on articles. This research consists of several stages, namely data collection, data labeling, data preprocessing, feature extraction, prediction model Building LSTM, Model Evaluation and Vaidation, and service development.

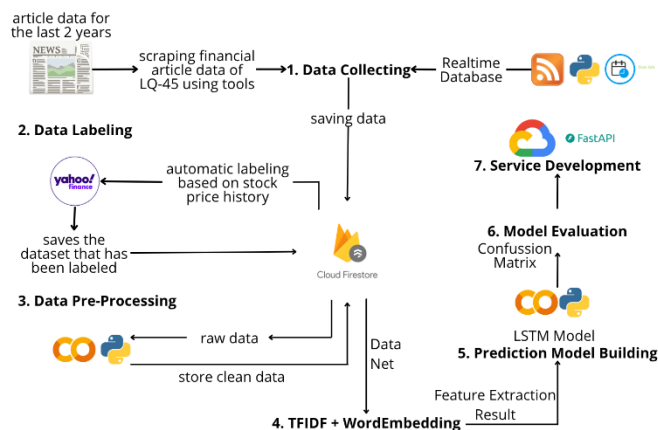


Figure 1. System Overview

A. Data Collection

This research begins with the process of collecting article data which consists of two types, namely primary data and secondary data. Primary data is data that is collected directly by the researcher from the first source for the specific purpose of the research being conducted. Secondary data, on the other hand, is data that has been previously available, usually collected by others for different purposes, and then reused by the researcher in a new research context[11].

Primary data is obtained from RSS feeds of various Indonesian online media platforms, such as CNBC Indonesia, Kompas, Kumparan, Detik.com, JPNN, Tempo, Viva, Kontan, Bisnis.com, and Stockwatch.

Primary data collection was conducted through two methods. First, archived article data from the last two years was collected using tools such as Webscraper and BeautifulSoup to build the initial dataset. The amount of data collected from this period reached 35,000 articles. This dataset was updated daily through an automated scraping process scheduled twice a day using cronjob.

Secondly, data for daily prediction input was collected in real-time through the same RSS feed. The focus of the article collection is on news related to the member stocks of the LQ-45 index. This research specifically limits the scope of the analysis to only the stocks included in the LQ-45 index, and does not cover stocks outside the index.

All articles used are sourced from Indonesian-language news outlets obtained via RSS feeds, so foreign-language news is not included in the scope of this research. Additionally, this study does not process complex figurative or emotional language elements, such as sarcasm. The

information extracted from the articles includes the title, article content, links, publication date, and publication time.

The secondary data in this study consists of historical stock price data obtained from the Yahoo Finance platform. The data used is the closing price of the stock, which serves to analyze the relationship between the sentiment contained in the news articles and the movement of the related stock prices.

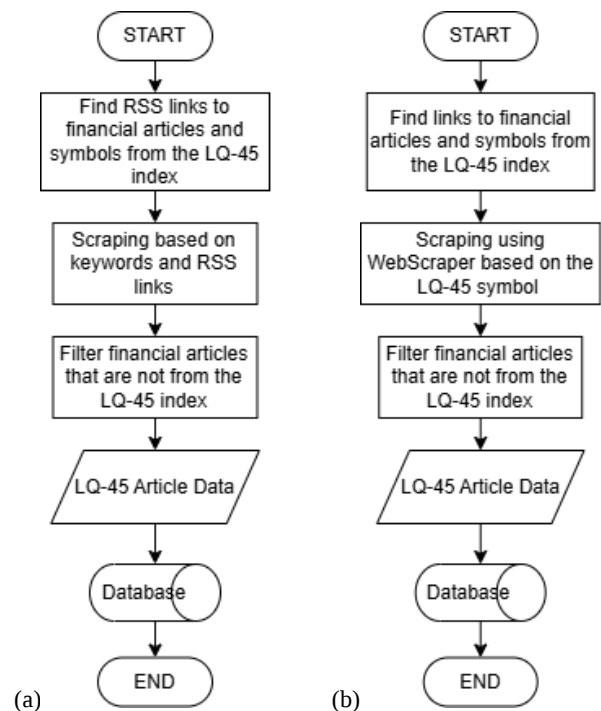


Figure 2 (a). Initial Data Scraping Flow (b). Automated Data Scraping Flow

B. Data Labelling

The labeling in this study is a development of previous studies that have shown a relationship between financial news sentiment and stock price movements[4]. The previous study used a labeling method based on the BIAS indicator, which is the percentage difference between the closing price of the stock on the day the news is published and the average closing price for several days after (usually for 6 days). Based on the BIAS value, if $BIAS > 0$ then the stock is considered to be trending up, and if $BIAS < 0$ then it is considered to be trending down. To simplify analysis, a label of “1” is assigned to uptrends, and a label of “0” to downtrends. This labeling process is done automatically (auto-labeling) based on historical stock price data, without involving manual annotation from humans. Thus, this research develops a labeling method by comparing the closing price of the stock on the day before the news is published with the closing price on the day the news is published. If the closing price increases, the news is labeled as “Naik”; if it decreases, it is labeled as “Turun”; and if there is no change, it is labeled as “Netral”. According to Del Corro and Hoffart (2020), autolabeling techniques that utilize structured financial data offer a scalable and objective way to build sentiment datasets,

reduce subjectivity, and align more closely with actual market behavior[12].

TABEL I
LABELING DATA USING CLOSE PRICE

Label	Description
Naik	Close price before issue date < Close price on issue date
Turun	Close price before issue date > Close price on issue date
Netral	Close price before issue date = Close price on issue date

Table I is the sentiment labeling scheme used in this study. The labeling is determined based on the date and time of publication of the article, as well as the closing price of the corresponding stock on that date. Specifically, each news article published between 00:00 and 23:59 on a given day is labeled by comparing the closing price of the stock on the day before publication with the closing price on the publication date.

Labeling is done by comparing the closing price (Close) of the stock on the day before the article is published with the closing price on the day the article is published. If the closing price of the stock on the publication date is higher than the previous day, then the article is labeled “Naik.” If the closing price is lower, the article is labeled “Turun.” If the closing price is the same as the previous day, the article is labeled “Netral.” This labeling scheme is used to link articles with historical stock price data retrieved from Yahoo Finance, allowing further analysis of the relationship between news sentiment and stock price movements. A Netral label is given if there is no change in the closing price (0% movement). An Naik label is given when the closing price increases by more than 0%, while a Turun label is given when the closing price decreases by less than 0%.

C. Data Preprocessing

Data pre-processing aims to improve the quality and consistency of the text data so that the model can understand the patterns better.

TABEL II
DATA PREPROCESSING

Before Preprocessing	
Date	Selasa, 19 Maret 2024 09.40
Article Title	Saham BMRI, AKRA, BRPT, CPIN Dorong Indeks Bisnis-27
Label	Naik
After Preprocessing	
Change Date Format	2024-03-2024
Encode Label	0
Lowecase	saham bmri, akra, brpt, cpin dorong indeks bisnis-27
Remove alphanumeric characters	saham bmri akra brpt cpin dorong indeks bisnis

This process starts by cleaning the text in the Article Title column by removing non-alphanumeric characters such as punctuation marks and symbols, then converting all text into lowercase letters to avoid duplication of the same word in different uppercase forms. Sentiment category labels that were originally in text form such as “Naik”, “Turun”, and “Netral” were then converted into numeric form using encoding so that they could be processed by the model. The imbalance in the amount of data between classes is overcome by oversampling the class with less data so that the class distribution becomes balanced.

D. Feature Extraction

Feature extraction is essential in the analysis of unstructured behavioral sequences. It includes the retrieval of potential information using data-driven methods, as well as the transformation of the sequence into structured data that can be understood by a computer. This process has a significant impact on the accuracy and reliability of subsequent model analysis[13]. Feature extraction can be done by various methods such as TFIDF, WordEmbedding, One hot Encoding, etc. Term Frequency-Inverse Document Frequency (TF-IDF) aims to assign a weight to each word to improve the effectiveness of sentiment analysis in text mining[14]. Word embedding using GloVe (Global Vectors for Word Representation) converts each word into a low-dimensional vector that represents the semantic and syntactic relationships between words based on co-occurrence statistics in a large corpus. GloVe uses a count-based model that utilizes a global co-occurrence matrix, thus capturing word meaning relations in a broader and more stable manner[15].

TABEL III
FEATURE EXTRACTION

Method	Feature	Value
Tokenizer + Padding	Article Title	300 (Sequence Length)
TF-IDF	Article Title	15.000 (Word Features)
One-hot Encoding	Stock Symbol	45 (Number of Symbol)
GloVe Pre-trained Embedding	Embedding Layer	100 (Vector Dimensions)
LSTM	Output from Embedding	128(Units)

E. Preprocessing NLP

Preprocessing NLP begins with cleaning up the article title by removing punctuation and changing all letters to lowercase to make the text more uniform. Next, the text is converted into a sequence of tokens using a Tokenizer which is limited to the 15,000 most frequently occurring words. The tokenization results are then padded to a maximum length of 300 tokens. In addition, feature extraction is also performed using TF-IDF with a maximum of 15,000 features to represent the importance of words in each article. An additional feature used is one-hot encoding of stock symbols. For word

representation, a pretrained GloVe with a vector dimension of 100 is loaded into the Embedding Layer and the weights are non-trainable to retain the semantic meaning of the pretrained model. The embedding result is then processed using LSTM architecture with two layers and continued with concatenation with TF-IDF and one-hot symbol features before entering the fully connected layer and producing sentiment classification output.

F. Model Building

Model building using LSTM is done after the feature extraction process is performed by combining three types of feature representations to improve the accuracy of stock sentiment prediction.

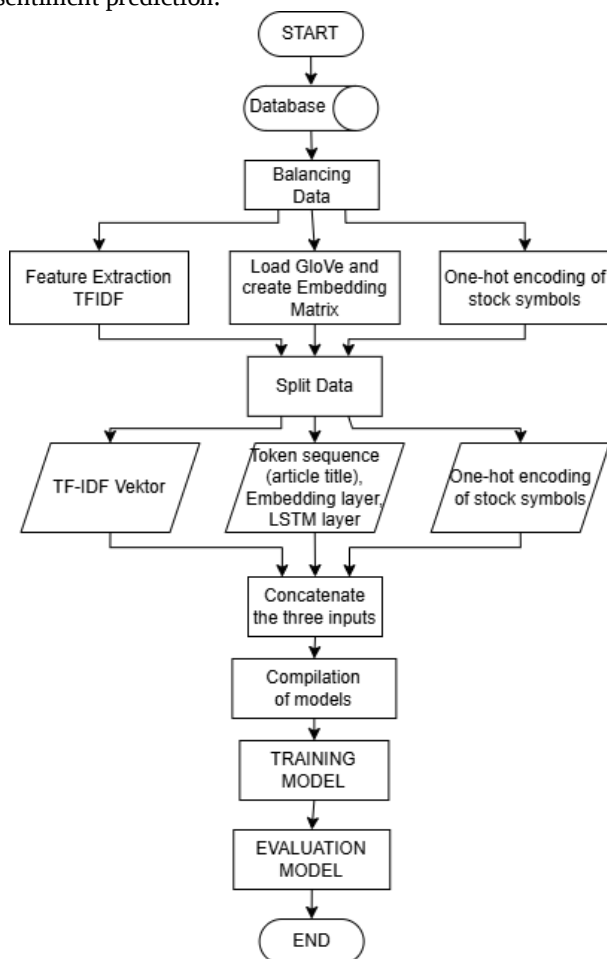


Figure 2. Model implementation flow

In the first stage, article titles are processed using a Tokenizer that converts the text into a sequence of numbers (word index) and condensed into a fixed length of 300 words; this representation is then passed to an Embedding Layer that utilizes word vector weights from GloVe (100 dimensions per word) to capture semantic context. In the second stage, textual features are also extracted using the TF-IDF Vectorizer with the top 15,000 words, resulting in a numerical vector reflecting the importance of each word in the document. In

the third stage, the stock symbols of each article are converted into binary vectors using one-hot encoding to introduce company context information. These three features (LSTM output, TF-IDF, and stock symbols) are then combined into a single input vector for further processing by the dense neural network layer to generate sentiment prediction. This multi-feature approach can enable the model to learn from the text structure, word frequency, pattern, and context of each stock entity simultaneously.

The model training process begins with preparing input data that has undergone feature extraction. The data is then divided into training and testing data according to a specified ratio. The LSTM model is trained using training data with specific parameter configurations for several epochs, where each epoch consists of several batches of data. After each epoch is complete, the model is validated to monitor performance and prevent overfitting. This process continues until the specified number of epochs is reached and the model achieves optimal performance based on the evaluation metrics used.

TABEL IV
MODEL TRAINING PARAMETERS

Parameter	Value
Epochs	20
Batch Size	64
Optimizer	Adam
Learning Rate	0.0001 (1e-4)
Split Data	80:20

The combined features are then processed through multiple fully connected (dense) layers equipped with dropout and batch normalization techniques to reduce overfitting and improve stability during training. To optimize the training process, Adam's optimizer is used with a learning rate scheduler and early stopping to automatically stop training when validation no longer shows improvement.

TABEL V
DAILY SENTIMENT PREDICTION SCENARIOS

Session	Prediction Time	Article Data Time Range
1	00:01–12.00 WIB	Past 3 months up to 00:00 of the previous day
2	12:01–00.00 WIB	Past 3 months + articles published from 00:01 to 12:00 today

Sentiment prediction in this study is done twice a day, at 00.01 WIB for the morning session (Session 1) and at 12.01 WIB for the afternoon session (Session 2). In session 1, the system uses news articles from the last three months until 00:00 WIB the previous day as input to predict stock price movements in the morning session. While in session 2, the system utilizes articles from the last three months plus articles published between 00:01 to 12:00 WIB on the same day to predict stock price movements in the afternoon session. This article data acts as the main input in the sentiment prediction

process, allowing the system to identify stock market sentiment. This approach allows the system to respond to the latest information faster and produce more accurate predictions according to market dynamics in the two daily trading sessions.

G. Retrain Model

The implementation of retraining can improve the accuracy of the model in predicting sentiment. By periodically retraining using the latest data from the real-time database, the model becomes more adaptive to the latest news patterns and trends.

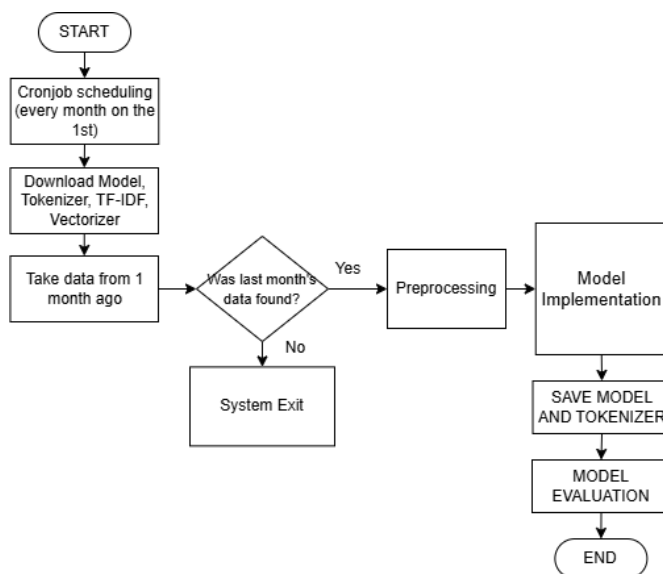


Figure 3. Model retrain flow

Each retrained model is stored in a separate version every month. The goal is to record the history of the model's development in a structured manner, allowing for performance analysis between versions and facilitating the rollback process if there is a decrease in accuracy in the latest version[16]. This versioning also provides an overview of the improvement of the model's capabilities as the data is updated and retrained regularly. The retrained model is not directly integrated into the CID system. Each retrained version goes through further evaluation to ensure its feasibility and performance before being published as an active version. This evaluation process is a crucial step to maintain the quality of the system's predictions and prevent a decline in accuracy[17].

H. Model Evaluation

The evaluation metric used in this study is the confusion matrix, which provides a comprehensive overview of the model's performance in classifying data into the correct classes. From the confusion matrix, several important evaluation metrics can be obtained, such as accuracy, precision, recall, and F1-score[18]. Precision measures the proportion of positive predictions that are actually correct,

indicating how accurate the model is when predicting a sentiment. Recall evaluates the model's ability to identify all relevant examples of a particular class, indicating how well the model captures sentiments that truly belong to each category. These metrics help assess the effectiveness of the model in predicting sentiment and identifying specific classes that may still require improvement.

I. Service Development

Service development in a microservice architecture allows each component to be developed, deployed, and scaled independently, driving modularity, agility, and resilience in modern software systems. This independent deployment model supports cross-functional teams working independently, leading to faster iterations and better system scalability[19].

In this research, a sentiment prediction service was developed and deployed using Google Cloud Platform (GCP) so that it could be accessed by an integrated stock price prediction application. This application combines two main approaches, namely sentimental analysis and technical analysis, with the developed system focusing on the sentiment aspect through news sentiment analysis. The service is designed to automate the sentiment prediction process based on real-time data. FastAPI was used as a lightweight, high-performance framework to build an efficient, easy-to-maintain, and scalable RESTful API. This API functions to receive article data as input, process it into sentiment predictions, and automatically update the prediction results on a daily basis. As a result, the system is capable of responding quickly and accurately to market information dynamics. The API exposes 45 path-parameterized endpoint each corresponding to a specific stock symbol allowing users to retrieve sentiment predictions simply by specifying the desired symbol in the URL path.

III. RESULT AND DISCUSSION

This model was trained using a dataset containing 35,000 articles from the past two years. Each article has been labeled with a stock symbol based on predefined keywords. These keywords represent company names or stock codes used for news searches. The dataset was then classified into three sentiment classes, Turun (16,535 articles), Naik (15,438 articles), and Netral (3,027 articles). Model evaluation in this study was carried out using the single holdout method, where the dataset was divided into 80% training data and 20% testing data.

TABLE VI
MODEL EVALUATION BASELINE

Method	Metric	Result
Logistic Regression	Precision(macro)	0.59
	Recall(macro)	0.59
	F1-Score(macro)	0.58
	Precision(weighted)	0.59
	Recall(weighted)	0.59
	F1-Score(weighted)	0.58
Naïve Bayes	Precision(macro)	0.55
	Recall(macro)	0.55
	F1-Score(macro)	0.54
	Precision(weighted)	0.55
	Recall(weighted)	0.55
	F1-Score(weighted)	0.54
Random Forest	Precision(macro)	0.65
	Recall(macro)	0.65
	F1-Score(macro)	0.65
	Precision(weighted)	0.65
	Recall(weighted)	0.65
	F1-Score(weighted)	0.65

The three baseline models used, namely Logistic Regression, Naïve Bayes, and Random Forest, produce macro F1-Score of 0.58, 0.54, and 0.65, respectively. Based on these results, a model was developed using a deep learning approach (LSTM) to obtain better prediction performance. To improve the model's ability to understand textual data, feature extraction techniques are applied, including Term Frequency-Inverse Document Frequency (TF-IDF) and pre-trained word embedding using GloVe. These techniques allow the model to capture both statistical and semantic representations of the input text, thus improving its ability to detect patterns and predict stock movement sentiment more accurately[20].

TABLE VII
MODEL EVALUATION RESULTS BASED ON FEATURE EXTRACTION

Method	Metric	Result
LSTM dengan fitur ekstraksi (TF-IDF+ Word Embedding(Glove))	Precision(macro)	0.72
	Recall(macro)	0.75
	F1-Score(macro)	0.73
	Precision(weighted)	0.73
	Recall(weighted)	0.73
	F1-Score(weighted)	0.73
LSTM dengan fitur ekstraksi (Word Embedding(Glove))	Precision(macro)	0.67
	Recall(macro)	0.70
	F1-Score(macro)	0.70
	Precision(weighted)	0.71
	Recall(weighted)	0.70
	F1-Score(weighted)	0.70
LSTM dengan fitur ekstraksi (TF-IDF+ Fine Tuning Word Embedding(Glove))	Precision(macro)	0.71
	Recall(macro)	0.76
	F1-Score(macro)	0.73
	Precision(weighted)	0.74
	Recall(weighted)	0.73
	F1-Score(weighted)	0.73

In this study, a single holdout method is used to evaluate the performance of the model, in which the dataset is divided

into 80% training data and 20% test data with a fixed proportion. Based on the evaluation results of three scenarios of feature extraction methods in the LSTM model, it can be concluded that the use of a combination of TF-IDF and Word Embedding GloVe provides the most consistent and superior performance. The model integrating both representations recorded an F1-Score macro of 0.73, with Precision macro of 0.72 and Recall macro of 0.75. These results indicate that the combination of statistical and semantic representations is able to capture text context more effectively and equally among all sentiment classes (Naik, Turun, dan Netral).

The model using only Word Embedding (GloVe) without TF-IDF shows a decline in performance, with the macro's F1-Score only reaching 0.70. This suggests that embedding-based representations alone are not sufficient to capture relevant word variations and frequencies in stock market news texts, which are often dense with technical and economic terms.

GloVe embedding was fine-tuned during training and combined with TF-IDF, the model performance remained competitive with a macro F1-Score of 0.73 and a slight increase in macro recall to 0.76. This shows that fine-tuning embedding can improve the model's sensitivity to minority classes, although it does not significantly improve precision.

The results of these three scenarios can be summarized as follows: the combination of statistical and semantic feature representations has a positive impact on text sentiment classification performance. By enriching the model's input using a multi-representation approach, the LSTM model is able to generate more accurate predictions and better generalization on complex and contextual data such as stock news. The LSTM model that combines TF-IDF and pretrained GloVe word embeddings demonstrated the best overall performance, with the highest macro F1-score and stable metrics indicating good generalization on multi-category data.

When compared to traditional baseline models such as Logistic Regression, Naïve Bayes, and Random Forest-which only produce macro F1-Score of 0.58, 0.54, and 0.65 respectively-the LSTM model shows a significant performance improvement. This comparison further strengthens the effectiveness of deep learning approaches in capturing complex patterns in financial news sentiment.

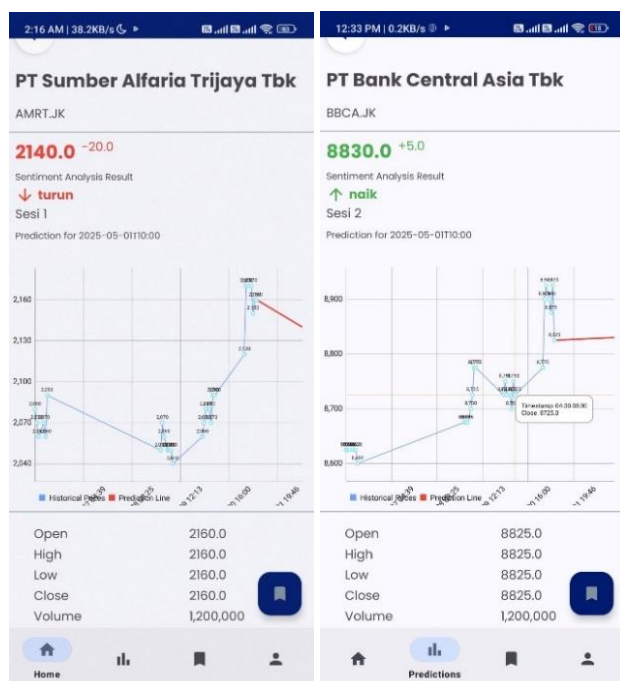


Figure 4. Integration System

Figure 4 shows the results of the API implementation that has been integrated with the Android application for technical stock price prediction. The API provides responses in the form of real-time daily sentiment predictions, which include company name, sentiment prediction, date, and prediction session. The system will display sentiment predictions that are relevant to the time session when the application is accessed. If the user accesses the application at 00.01-12.00, the system will display sentiment predictions for session 1. Conversely, if the application is accessed at 12.01-00.00, it will display sentiment predictions for session 2. With this integration, users can obtain sentiment information that is accurate and in accordance with stock price movements in the ongoing session. This shows that the system is able to run in real-time, adjust sentiment predictions to the access time, and support the investment decision-making process by users. This integration provides added value for users in determining which stocks to invest in, because it provides two points of view of analysis, namely from the sentiment and technical sides.

IV. CONCLUSION

This study aims to develop a daily stock sentiment prediction system based on LSTM using real-time updated economic news articles that are processed daily. The system is built using a multi-input approach, which combines various text feature representations and stock symbol information to improve prediction accuracy. The results show that integrating text representations based on embedding (GloVe) and statistical representations via TF-IDF, combined with stock symbol information using one-hot encoding,

significantly improves model performance. Among the three feature extraction methods tested, the combination of TF-IDF and pretrained GloVe Word Embedding yielded the best performance. The model achieved a macro F1-Score of 0.73, with a macro precision of 0.72 and a macro recall of 0.75. Based on this research, it can be concluded that the combination of methods is more effective in capturing the context and variation of words in stock news texts rich in technical and economic terms. Additionally, fine-tuning the embedding during training further enhances the model's sensitivity to minority classes, as evidenced by the increase in macro recall to 0.76 without compromising Precision stability. This research still opens opportunities for further development, particularly in the use of methods like BERT and the addition of stock recommendation systems.

REFERENCES

- [1] C. Chikwira and J. I. Mohammed, 'The Impact of the Stock Market on Liquidity and Economic Growth: Evidence of Volatile Market', *Economies*, vol. 11, no. 6, May 2023, doi: 10.3390/economies11060155.
- [2] BEI, 'Detail Siaran Pers Jumlah Investor Saham di Indonesia Lampau 6 Juta SID', Bursa Efek Indonesia.
- [3] F. Dwi Riyanto and A. Farhan Habibie, 'Effect of Internal and External Factors on Stock Price in Companies That Conduct Rights Issue', *Jurnal Ekbis : Analisis, Prediksi dan Informasi*, vol. 24, no. 2, Sep. 2023.
- [4] Y. Ren, F. Liao, and Y. Gong, 'Impact of News on the Trend of Stock Price Change: an Analysis based on the Deep Bidirectional LSTM Model', in *Procedia Computer Science*, Elsevier B.V., 2020, pp. 128–140. doi: 10.1016/j.procs.2020.06.068.
- [5] S. R. Khonde, S. S. Virnoddar, S. B. Nemade, M. A. Dudhedia, B. Kanawade, and S. H. Gawande, 'Sentiment Analysis and Stock Data Prediction Using Financial News Headlines Approach', *Revue d'Intelligence Artificielle*, vol. 38, no. 3, pp. 999–1008, Jun. 2024, doi: 10.18280/ria.380325.
- [6] Lakshya, Prateek, and D. Sethia, 'Stock Price Prediction Using News Sentiment Analysis', in *2022 2nd International Conference on Intelligent Technologies, CONIT 2022*, Institute of Electrical and Electronics Engineers Inc., Jun. 2022. doi: 10.1109/CONIT55038.2022.9847747.
- [7] M. Ibnu Alfarizi, L. Syafaah, and M. Lestandy, 'Emotional Text Classification Using TF-IDF (Term Frequency-Inverse Document Frequency) And LSTM (Long Short-Term Memory)', *JUITA: Jurnal Informatika*, vol. 10, no. 2, pp. 225–232, Nov. 2022, [Online]. Available: https://atapdata.ai/dataset/192/HIMPUNAN_DATA_E
- [8] M. I. Saputra, E. Nurawati, and R. Abyasa, 'Predicting Stock Price Movements with Technical, Fundamental, and Sentiment Analysis Using the LSTM Model', *Jurnal Informatika*, vol. 12, no. 1, pp. 1–9, Mar. 2025, doi: 10.31294/inf.v12i1.22299.
- [9] F. Bouchra, I. M. Agus Dwi Suarjaya, and N. K. DwiRusjanyanthi, 'Analisis Sentimen Masyarakat terhadap Tayangan Televisi Nasional menggunakan Metode Deep Learning', *Jurnal Buana Informatika*, vol. 15, pp. 89–99, Oct. 2024.
- [10] M. Galih Pradana, A. Christian Nurcahyo, P. Hari Saputro, U. Alma Ata Yogyakarta, S. Shanti Bhuana Bengkulu, and K. Barat, 'Pengaruh Sentimen di Sosial Media dengan Harga Saham Perusahaan', May 2020.
- [11] 'Comparison between Primary Data and Secondary Data Basis for Comparison Primary Data Secondary Data'. [Online]. Available: [www.ej-edu.orgDOI:http://dx.doi.org/19810.21091/](http://dx.doi.org/19810.21091/)
- [12] L. Del Corro and J. Hoffart, 'From Stock Prediction to Financial Relevance: Repurposing Attention Weights to Assess News

- Relevance Without Manual Annotations', Feb. 2021, doi: 10.48550/arXiv.2001.09466.
- [13] J. Zhou, Z. Ye, S. Zhang, Z. Geng, N. Han, and T. Yang, 'Investigating response behavior through TF-IDF and Word2vec text analysis: A case study of PISA 2012 problem-solving process data', *Heliyon*, vol. 10, no. 16, Aug. 2024, doi: 10.1016/j.heliyon.2024.e35945.
- [14] O. I. Gifari, M. Adha, I. Rifky Hendrawan, F. Freddy, and S. Durrand, 'Analisis Sentimen Review Film Menggunakan TF-IDF dan Support Vector Machine', *Jifotech (Journal Of Information Technology)*, vol. 2, no. 1, 2022.
- [15] K. Das, Kamlish, and F. Abid, 'Advancements in Word Embeddings: A Comprehensive Survey and Analysis', Sep. 23, 2024, *Pakistan Academy of Sciences*. doi: 10.53560/PPASA(61-3)842.
- [16] P. Liang, B. Song, X. Zhan, Z. Chen, and J. Yuan, 'Automating the Training and Deployment of Models in MLOps by Integrating Systems with Machine Learning'.
- [17] T. Antoniou and M. Mamdani, 'Evaluation of machine learning solutions in medicine', *CMAJ*, vol. 193, no. 36, pp. E1425–E1429, Sep. 2021, doi: 10.1503/cmaj.210036.
- [18] S. Panggabean and A. Junika, 'JOISTECH: Journal of Information System and Technology Sentiment Analysis on Public Opinions Regarding the 2024 Regional Elections Using Long Short-Term Memory (LSTM), Random Forest, and Naive Bayes', *JOISTECH: Journal of Information System and Technology*, vol. 1, no. 2, pp. 67–75, Dec. 2024.
- [19] H. Rodrigues, R. Silva, A. Avritzer, and A. R. Silva, 'Highlights Assessment of Performance and its Scalability in Microservice Architectures: Systematic Literature Review Assessment of Performance and its Scalability in Microservice Architectures: Systematic Literature Review', *Journal of Systems and Software*, pp. 1–19, 2025, doi: <https://doi.org/10.1016/j.jss.2025.112500>.
- [20] M. George and R. Murugesan, 'Improving sentiment analysis of financial news headlines using hybrid Word2Vec-TFIDF feature extraction technique', in *Procedia Computer Science*, Elsevier B.V., 2024, pp. 1–8. doi: 10.1016/j.procs.2024.10.172.