

Implementation of XGBoost Algorithm for Sentiment Classification of Public Opinions on the Rupiah Redenomination Policy

Andinie Rachmah Basri, Fadhilah Fitri, Dodi Vionanda

Universitas Negeri Padang, Indonesia

andinierb@gmail.com; fadhilahfitri@fmipa.unp.ac.id

Article Info:

Submitted:	Revised:	Accepted:	Published:
Jul 15, 2026	Aug 12, 2026	Aug 24, 2026	Aug 29, 2026

Abstract

Although rupiah redenomination has long featured in Bank Indonesia's monetary policy discourse as a means of simplifying currency denominations without altering the exchange rate or real purchasing power, public concerns about declining purchasing power and price rounding underscore the need for effective policy communication. This study analyzes Indonesian public sentiment toward rupiah redenomination using the Extreme Gradient Boosting (XGBoost) algorithm to classify YouTube comments. Data were collected through the YouTube Data API v3 from 1,763 comments posted on the Tribunnews video titled *Purbaya Targets Rupiah Redenomination Bill to Be Completed in 2027*. Following text cleaning and tokenization, 1,169 comments were retained and transformed using Term Frequency-Inverse Document Frequency (TF-IDF). Manual labeling classified 657 comments as positive and 512 as negative. The XGBoost model, trained using optimized hyperparameters, achieved an accuracy of 73.39%, with F1-scores of 0.772 for positive sentiment and 0.680 for negative sentiment. These results indicate that the model classified positive sentiment more effectively, although ambiguous comments remained

challenging. The findings reveal the distribution of public responses to the proposed policy and emphasize the need for intensive public outreach to mitigate potential resistance and misconceptions. This study contributes a data-driven basis for developing more effective monetary policy communication and anticipating the social dynamics associated with rupiah redenomination.

Keywords: Machine Learning; Monetary Policy Communication; Rupiah Redenomination; Sentiment Analysis; XGBoost

INTRODUCTION

The plan to redenominate the rupiah has long been part of Bank Indonesia's monetary policy agenda. According to Anis Byarwati, a member of Commission XI of the Indonesian House of Representatives from the PKS faction, this discourse is not new because since 2010, Bank Indonesia has designed five stages for the implementation of redenomination (PKS, 2020). According to Permana (2015), redenomination is essentially a simplification of the currency fraction by reducing the number of zero digits without changing the exchange rate or real purchasing power. As an illustration, the price of Premium fuel, which was initially Rp6,500 per liter, will become Rp6.5 per liter after redenomination, with the purchasing power remaining the same.

However, the public often associates redenomination with the currency reform that took place during the Old Order era, leading to concerns about a decrease in purchasing power and the potential rounding up of prices that could disadvantage consumers. Therefore, massive socialization and education are key to the successful implementation of this policy.

Historically, Indonesia implemented a redenomination on December 13, 1965, under President Sukarno by removing three zeros from the old rupiah. The policy initially succeeded in creating a unified monetary system, but the nominal value of the rupiah returned due to economic instability (Puspita & Sulistya, 2024). According to Indriani (2025), the current stages of redenomination designed by Bank Indonesia, as explained by BI Governor Perry Warjiyo, begin with the issuance of the Redenomination Law as a primary requirement before the entire process starts. Without those regulations, the currency simplification policy cannot proceed. The second stage is the formulation of regulations regarding the transparency of the prices of goods traded in Indonesia. This

The introduction should be written descriptively, logically, systematically, and coherently. It should clearly explain the scholarly rationale for conducting the study, position the study within relevant previous research, and highlight the contribution it intends to offer. At a minimum, the introduction should include the following five main components.

According to Hernikawati (2021), sentiment analysis is a process to find meaning from behavior, opinions, views, and emotions from a text, speech, tweets, and database sourced from Natural Language Processing. Sentiment analysis classifies text into positive, negative, or neutral meanings. The application of sentiment analysis can be used in several fields such as consumer information, marketing, politics, and social matters. In government, sentiment analysis can be used to understand public opinion on an issue or policy that has been implemented, allowing the government to create appropriate solutions based on the available data.

With technological advancements, machine learning algorithms like XGBoost excel because they incorporate regulation-based optimization to improve accuracy while reducing the risk of overfitting. This algorithm is known for its high efficiency and excellent accuracy in handling large amounts of data, while also being able to address bias and variance issues more effectively compared to other methods (Wati et al., 2025). The use of the XGBoost algorithm in this research is expected to produce a more accurate and comprehensive sentiment analysis regarding the trends in Indonesian public opinion on the issue of rupiah redenomination based on social media data. The results of this research can subsequently serve as strategic input for the government, Bank Indonesia, and other policymakers in designing more effective public communication and anticipating potential public resistance to the policy.

Several previous studies have applied XGBoost for sentiment analysis on various public policy issues. Research by Widiarta et al. (2023) shows that the application of XGBoost in sentiment analysis regarding the implementation of the PPKM policy on Twitter social media resulted in an accuracy of 85%. Meanwhile, research on public sentiment regarding the redenomination of the rupiah has been conducted using the Naïve Bayes algorithm and achieved an accuracy of 86.3% (Fadli et al., 2026). However, research on the application of XGBoost to classify public sentiment on the issue of rupiah redenomination is still limited. The difference in algorithms used in previous studies serves as the basis for

investigating the capability of XGBoost in identifying public sentiment trends related to the issue.

Therefore, this research aims to analyze the public sentiment in Indonesia regarding the issue of rupiah redenomination using an XGBoost-based machine learning classification approach on social media data. The findings of this research are expected to provide both theoretical and practical contributions in understanding the dynamics of public sentiment toward controversial monetary policies, as well as serving as a data-based reference for formulating more effective and less resistant communication strategies for the rupiah redenomination policy.

METHODS

The data source in this study comes from public comments on the YouTube video titled "Purbaya Targets Rupiah Redenomination Bill to be Completed by 2027" uploaded by the Tribunnews channel. Data was collected using the YouTube Data API v3 through the Python programming language, allowing the comment retrieval process to be automated and resulting in a structured dataset. The research variables consist of input variables in the form of YouTube comment text and output variables in the form of sentiment categories defined as positive, neutral, and negative. All comments go through the text preprocessing stages, including character cleaning, tokenization, normalization, stopword filtering, and stemming to obtain relevant data ready for analysis (Wulla & Rino, 2026). After that, the text data is converted into a numerical representation using Term Frequency–Inverse Document Frequency (TF-IDF) before being used in the classification process.

Extreme Gradient Boosting (XGBoost) is one of the techniques in machine learning for regression and classification analysis based on Gradient Boosting Decision Tree. This algorithm allows for optimization 10 times faster compared to other Gradient Boosting Machines (GBM). Extreme Gradient Boosting (XGBoost) takes the Taylor expansion of the loss function up to the second order and adds a regularization term to find the optimal solution, which is used to balance the reduction of the objective function and the model complexity to avoid overfitting. The Extreme Gradient Boosting (XGBoost) model can be represented by the following equation (Chen & Guestrin, 2016).

$$\hat{y}_i = \sum_{k=1}^K f_k(x_i), f_k \in F \quad (1)$$

Explanation:

K = Number of decision tree

$f_k(x_i)$ = Score generated from k-th tree

$\hat{y}_i$ = predicted value

The data analysis technique in this research is carried out through several main stages as follows:

1. Data collection

Comment data is retrieved using the YouTube Data API v3 and stored in CSV format to facilitate the processing. The collection is carried out automatically through a Python script to ensure that all comments can be consistently acquired.

2. Preprocessing

The preprocessing stage includes cleaning the text from special characters, URLs, numbers, symbols, mentions, and irrelevant hashtags. This process also involves case folding, filtering, tokenizing, and stemming (Ardiani et al., 2020). This step ensures that the generated text is cleaner, more consistent, and ready to be processed in the next stage.

3. Sentiment Labeling

Each comment is categorized into positive, neutral, or negative sentiment classes through a manual labeling process that refers to the sentiment polarity guidelines of the Indonesian language. The annotation process is carried out systematically to reduce subjectivity and ensure consistency in label determination, thereby producing accurate and high-quality training data.

4. Feature Extraction Using TD-IDF

Word representation using the Term Frequency–Inverse Document Frequency (TF-IDF) method to convert text into numerical vectors. Term Frequency reflects the weight of a word or phrase based on how frequently it appears in a document. Inverse Document Frequency (IDF) is a numerical representation that measures the rarity or uniqueness of a term within a collection of documents (Alghazzawi et al., 2025).

$IDF(t, D) = \log \left(\frac{N}{DF + 1} \right)$	(2)
--	-----

$DF(t,D)$ is the Document Frequency, it indicates the number of documents in D that contain the word or term t . D is the collection of documents, and N is the total number of documents in D .

The Term Frequency-Inverse Document Frequency (TF-IDF) statistic measures the importance of a term in a document in comparison to its total occurrence in a collection of documents.

$$tf-idf(t, d, D) = tf(t, d) \times idf(t, D) \quad (3)$$

Where t is the phrase or word, the document is denoted by d , and D is a document collection. The Term Frequency ($TF(t,d)$) represents the number of times the term t appears in the document.

5. Train-Test Split

The dataset is divided into two parts, namely 70% training data and 30% testing data. This division aims to ensure that the model's performance evaluation is conducted objectively using data that was not used during training.

6. Training the XGBoost Model

The model is trained using the XGBoost algorithm, which is a boosting-based ensemble method that builds a series of decision trees incrementally (Zhang et al., 2022). In the initial stage, the model sets the base prediction score before the first tree is formed. Next, several important parameters such as learning rate, maximum depth, gamma, and minimum child weight are adjusted to control the model's complexity and prevent overfitting.

7. Model Evaluation

Model evaluation is conducted using a confusion matrix to compare the predicted results against the actual classes in each sentiment class (Tharwat, 2018). Additionally, metrics such as accuracy, precision, recall, and F1-score are used (Novakovi et al., 2017). This evaluation aims to assess the model's effectiveness in classifying comments.

RESULTS

In this section, the results and discussions are explained, starting from dataset collection, data preprocessing, manual classification or labeling for sentiment, training, and

evaluation of the XGBoost model in classifying YouTube comments regarding Rupiah Redenomination using TF-IDF weighting.

1. Dataset Collection

Data collection in this research was conducted by gathering public comments from the YouTube platform as the basis for sentiment analysis of the public regarding the Rupiah redenomination issue. YouTube was chosen because it has a high level of interaction and allows the public to express their opinions openly on various economic policy issues. Data was obtained from the comments section of the video titled “Purbaya Targetkan RUU Redenominasi Rupiah Rampung Tahun 2027” uploaded by the official Tribun News account, making the data source credible and relevant to the research topic.

The collection of comments was conducted from November 7, 2025, to November 11, 2025, during the period when the issue of rupiah redenomination regained public attention due to increased news coverage and discussions on social media. The collected comments were limited to those directly related to the discussion topic, such as opinions, responses, criticisms, or support for the rupiah redenomination policy. The data collection process was carried out using a Python script, and the data was then stored in Excel format to facilitate further analysis using R software. During the observation period, a total of 1,763 comments were successfully collected, which were subsequently used as the research dataset. Here is an example of the data as shown in Table 1.

Table 1. Data Collection

published_at	author	text
2025-11-11T11:28:11Z	@agasyahbroo7873	Mantap...rupiah jadi gagah... Kayak negara negara maju nilai rupiah...kalo jadi.... (Awesome...the rupiah is becoming strong.. Like the currencies of developed countries...if it happens....)
2025-11-11T06:55:05Z	@azizulzul8937	Makan ja sudah ribu ribuan..ðŸ˜¸,ðŸ˜¸,negara paling banyak 000000000 (Just eating has already cost thousands.. ðŸ˜¸,ðŸ˜¸ the country with the most 000000000)
2025-11-11T02:19:56Z	@moeksssang6188	HARUS dilakukan redenominasi seperti th 1967 DAHULU... agar lebih simple dan efisien.. (It MUST be done a redenomination like in 1967 BEFORE... to be simpler and more efficient..)
2025-11-10T23:16:56Z	@jackysaja4400	Ga ada efek (There is no effect)

1. Preprocessing Data

The preprocessing stage is carried out after all the comment data has been successfully collected. This stage aims to tidy up and simplify the comment text so that it can be optimally processed by the sentiment classification algorithm. Technically, preprocessing steps such as text normalization, removal of irrelevant characters, and standardization of writing have been explained in the methodology section. Through this process, the initially raw and unstructured comment data is transformed into cleaner, more organized, and representative text for analysis purposes.

In the preprocessing process, comments containing emojis, symbols, links (URLs), as well as the use of non-standard or irrelevant words were successfully simplified into a collection of core words with clearer sentiment meanings. Out of a total of 1,763 comments obtained, 1,661 comments were successfully processed and used in the subsequent analysis stage. Meanwhile, some other comments were not included because they contained inappropriate language, deviated from the discussion context, or were identified as duplicate data. To provide a clearer picture of the impact of preprocessing, examples of comment changes before and after this process are presented in Table 2.

Table 2. Comment Transformation

Gpp barang mahal tp ada yg jualnya. Daripada barang murah tp gak ada yg mau jual. Uang jadi gak ada harganya.	tidak apa apa barang mahal tapi ada yang jualnya daripada barang murah tapi tidak ada yang mau jual uang jadi tidak ada harganya. (It's okay if it's expensive as long as someone sells it, rather than being cheap but no one wants to sell it, so it has no value)
ganti pun percuma kalau duitnya banyak di timbun tikus.ðŸ~ðŸ~,	ganti pun percuma kalau duitnya banyak di timbun koruptor (Even if it's replaced, it's useless if the money is hoarded by corrupt people)

2. Labeling

In the preprocessing process, comments containing emojis, symbols, links (URLs), as well as the After going through the preprocessing stage, the dataset is then labeled with sentiment as the basis for the classification process. In the initial stage, labeling is done manually by categorizing comments into three categories: positive, negative, and, neutral based on the content of the message and the emotional nuances contained within. Comments that reflect support or positive judgment are categorized as positive sentiment, while comments that contain criticism, dissatisfaction, or distrust are classified as negative sentiment. Meanwhile, comments that are informative or do not show a tendency toward a particular opinion are classified as neutral sentiment.

Initial test results show that the neutral sentiment class has less clear boundaries and often overlaps with the other two classes, making it difficult for the model to consistently distinguish sentiment patterns and resulting in low accuracy levels. Therefore, this study decided to exclude neutral-labeled data and focus the analysis on the two sentiment classes, namely negative and positive. For the purpose of modeling using the XGBoost algorithm, the sentiment labels were converted into numerical form with a value of 0 for negative sentiment and 1 for positive sentiment. After the selection process, the number of data used decreased from 1,661 comments to 1,169 comments, with examples of the labeling results presented in Table 3.

Table 3. Data Labeling

mantap pak menteri agar kita tidak kebanyakan menghitung uang besar nilai kecil (Great, Minister, so we don't spend too much time counting small amounts of large money)	positive	1
tidak urgensi seperti yang terjadi di venezuela dan lain lain ini akal akalan (not as urgent as what happened in Venezuela and others, this is just a trick)	negative	0

The 1,169 labeled comments, 512 were labeled as negative and 657 as positive. The proportion of each sentiment category is then presented in the form of a pie chart, as shown in Figure 1.

Distribution of Positive and Negative Sentiments (%)

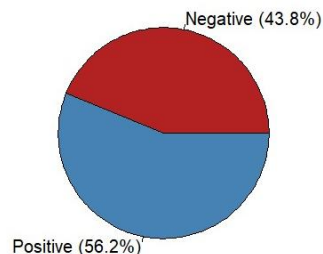


Figure 1. Dataset Proportion

To obtain an initial overview of the dominant word characteristics in each sentiment, a visualization using a word cloud was conducted. This visualization aims to display the words that most frequently appear in comments with negative and positive sentiments based on the results of text preprocessing. The word cloud visualization is shown in Figure 2.

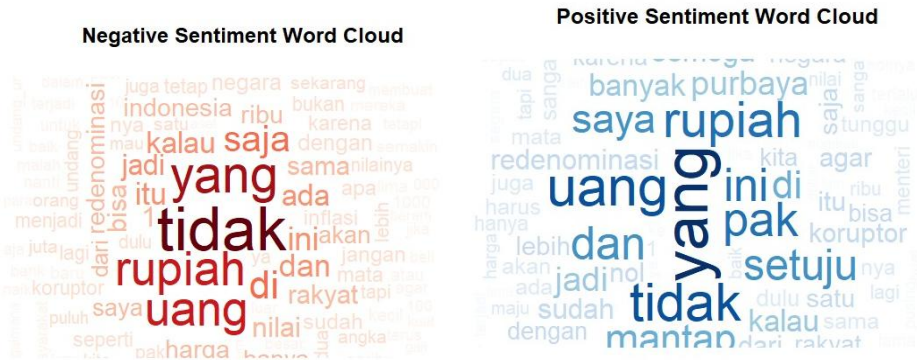


Figure 2. Sentimen Word Cloud

Based on the generated word cloud, a clear difference in word patterns between negative and positive sentiments is evident. In the negative sentiment, words like "tidak" (no), "uang" (money), and "rupiah" appear more dominantly, indicating expressions of dissatisfaction or rejection. Meanwhile, in the positive sentiment, words like "setuju" (agree), "mantap" (good), and "pak" (sir) appear more frequently, reflecting support and appreciation from users. This visualization reinforces the classification results by showing the linguistic context underlying each sentiment.

3. Classification Model

After the comment data goes through the preprocessing and labeling stages, this research continues with the development of a sentiment classification model using the Extreme Gradient Boosting (XGBoost) algorithm. Before the model training process, each comment is transformed into a numerical form through the tokenization process and the creation of features based on Term Frequency–Inverse Document Frequency (TF-IDF) considering unigrams and bigrams.

TF-IDF weighting is performed to represent the importance level of each term in a comment relative to the entire dataset, so that more informative words or phrases have a greater contribution in the classification process. The TF-IDF values are then summarized by calculating the average score for each term as an illustration of the global importance level of the word, with an example of the TF-IDF weighting results shown in Table 4.

Table 4. TF-IDF Score

Term	tfidf_score
mantap	0.0719186
setuju	0.0651122
pak	0.0342651
alhamdulillah	0.0319479
tidak	0.0316092

Table 4 shows how comments originally in text form are transformed into numerical vectors with word weights that reflect their importance in the dataset. The TF-IDF representation is then used as input for the XGBoost classification model to help learn patterns and predict the sentiment of each comment in a more structured manner.

4. Evaluation Model

The evaluation results show that the XGBoost model achieved an accuracy of 73.39%, indicating that most comments were classified correctly. The F1-score for the negative class is 0.680, while for the positive class it reaches 0.772, indicating that the model works more optimally in recognizing positive sentiment compared to negative sentiment. Nevertheless, the results of the confusion matrix still show classification errors in some data, reflecting the similarity in language characteristics between sentiment classes.

The performance was obtained through the hyperparameter tuning process using a cross-validation approach to determine the configuration that best fits the data characteristics. The optimal number of iterations was set to 69 iterations based on the lowest logloss value on the validation data, so that the resulting model has a balance between prediction accuracy and generalization capability. Details of the hyperparameter settings used in this study are presented in Table 5.

Table 5. XGBoost Hyperparameter

Hyperparameter	Value	Definition
learning_rate (eta)	0.15	The rate at which the model adapts to new data
max_depth	7	The maximum depth of each tree
n_estimators (nrounds)	69	The number of boosting rounds or trees to building
subsample	0.8	The fraction of samples to be used for training each tree
colsample_bytree	0.8	The fraction of features to be used for training each tree
min_child_weight	2	The minimum sum of hessian needed in a child node
Gamma	1	Minimum loss reduction is required to make a further partition on a leaf node

Details of the classification evaluation results are presented in the confusion matrix in the figure 3.

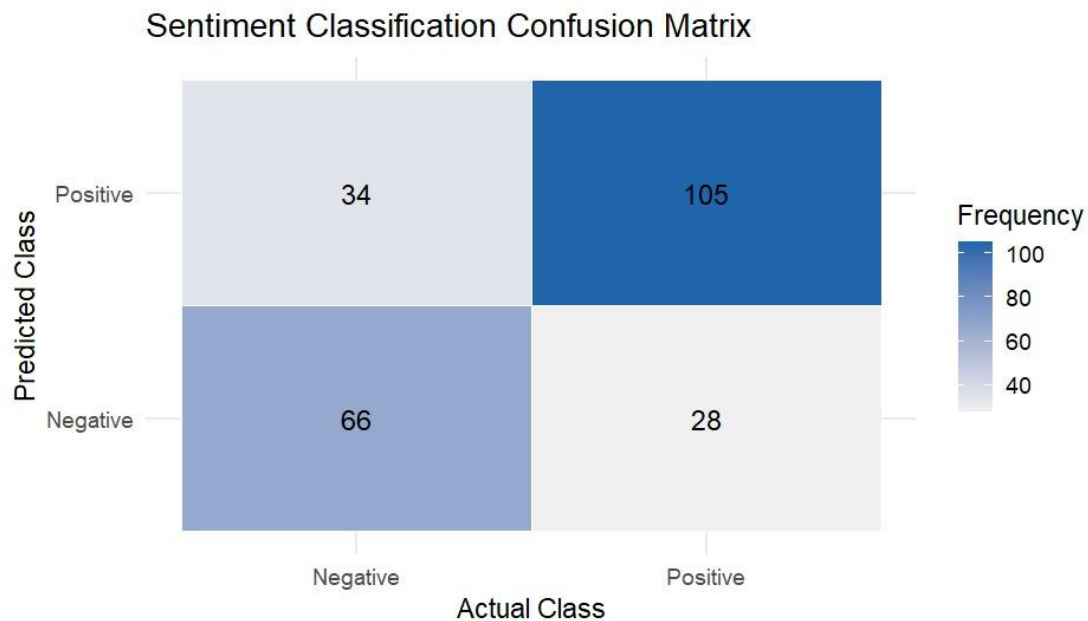


Figure 3. Confusion Matrix

Based on the results of the confusion matrix, the XGBoost model shows a fairly good ability to distinguish between negative and positive sentiments. Most comments were successfully classified correctly, with 66 negative comments and 105 positive comments. However, there were still a number of prediction errors, where 28 negative comments were classified as positive and 34 positive comments were predicted as negative. These findings indicate that the model has generally understood sentiment patterns, but still encounters difficulties in handling comments with ambiguous or unclear language, which are commonly found in social media comment data.

To measure the accuracy of the model in performing classification, several value criteria are used, namely accuracy, sensitivity, precision, specificity, and F1-Score. The test results of these criteria, are presented in Table 6.

Table 6. Classification report

Class	Precision	Recall	F1-Score	Support
Negative	0.660	0.702	0.680	94
Positive	0.789	0.755	0.772	139
Accuracy			0.734	233

The table above shows that the model is able to achieve an accuracy level of 73.39%, indicating a fairly good ability to classify public sentiment into positive and negative categories. In the negative sentiment class, the model obtained a precision value of 0.660, a recall of 0.702, and an F1-score of 0.680, while in the positive sentiment class, it obtained a

precision value of 0.789, a recall of 0.755, and an F1-score of 0.772, which shows the model's relatively more stable performance in recognizing positive sentiment.

DISCUSSION

The data obtained from YouTube comments on the video "Purbaya Targets the Rupiah Redenomination Bill to be Completed by 2027" resulted in 1661 comments for analysis. with 512 comments labeled negative and 657 labeled positive. In positive sentiment, words like 'good' become dominant, and the word 'no' becomes dominant in negative sentiment.

The TF-IDF weighting results show that several words have relatively high values, such as “mantap” and “setuju” This indicates that these words have a significant contribution in representing the characteristics of comments in the dataset. Next, the TF-IDF representation is used as input for XGBoost to learn the relationship patterns between word features and sentiment labels.

The XGBoost model in classifying sentiment achieved an accuracy of 73.39%. Based on the accuracy criteria, this falls into the good category (Pardede et al., 2022). The XGBoost model in classifying sentiment achieved an accuracy of 73.39%. Based on the accuracy criteria, this falls into the good category. The recall value for the negative class is 0.702, indicating that the model can recognize the negative class at 70.2%. Meanwhile, the recall for the positive class is 0.755, indicating that the model can recognize the positive class at 75.5%. The model has a better ability to find positive comments compared to negative comments, as shown by the higher positive recall value.

The evaluation results of the XGBoost model in this study, when compared to previous research, show that the results in this study are higher than those of Purnamasari & Hendrastuty (2026), who achieved an accuracy of 62.11% in sentiment classification of TikTok comments regarding IKN using TF-IDF and XGBoost. However, the results of this study are still lower than those of Maharani et al. (2026), who achieved an accuracy of 78.63% in the sentiment analysis of BPJS Kesehatan, and the research by Diyanulhaq & Saeffurohman (2024), who achieved an accuracy of 82.50% in the sentiment classification of Indonesian-language YouTube comments using TF-IDF.

The difference in results indicates that the use of the XGBoost algorithm and TF-IDF weighting does not always yield the same performance in every study. The model's

performance can be influenced by the characteristics of the dataset, the amount and distribution of data, the source of social media, the topic being analyzed, as well as the linguistic characteristics in the comments.

The process of manual data labeling can also be one of the factors affecting the performance of the classification model. Research by Sofyan et al. (2024), which compared the performance of classification algorithms on data with manual labeling and automatic labeling using the InSet Lexicon, showed that the model with InSet Lexicon labeling achieved an accuracy of 83%, higher than the model with manual labeling which achieved an accuracy of 77%.

CONCLUSION

This research shows that public sentiment analysis on the issue of rupiah redenomination using a machine learning approach based on the XGBoost algorithm can provide a deeper understanding of the Indonesian public's perception of the policy. By collecting data from public comments on YouTube and performing text preprocessing, the XGBoost model achieved an accuracy of 73.39%, indicating a fairly good performance in classifying public sentiment. Although this model shows better performance in identifying positive sentiment compared to negative sentiment, the classification errors in some comments indicate the challenges faced in identifying more subtle nuances of sentiment, such as in ambiguous or unclear comments.

These findings make an important contribution to policy formulation, particularly in designing more effective communication strategies by the government and Bank Indonesia regarding the rupiah redenomination policy. In addition, the results of this research also provide valuable insights into the potential resistance or support from the public toward the monetary policy, which can serve as a basis for anticipating obstacles or challenges in the implementation of the policy in the future. Further research is recommended to explore automatic labeling methods, such as InSet Lexicon, as an alternative in the data labeling process. Additionally, the use and comparison of other classification methods such as naive bayes and SVM can also be conducted to evaluate the possibility of achieving more optimal model performance.

REFERENCES

- Alghazzawi, D., Ullah, H., Tabassum, N., Badri, S. K., & Asghar, M. Z. (2025). Explainable AI-based suicidal and non-suicidal ideations detection from social media text with enhanced ensemble technique. *Scientific Reports*, 15, Article 1111. <https://doi.org/10.1038/s41598-024-84275-6>
- Ardiani, L., Sujaini, H., & Tursina, T. (2020). Implementasi sentiment analysis tanggapan masyarakat terhadap pembangunan di Kota Pontianak. *Jurnal Sistem dan Teknologi Informasi (JUSTIN)*, 8(2), 183. <https://doi.org/10.26418/justin.v8i2.36776>
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794). Association for Computing Machinery. <https://doi.org/10.1145/2939672.2939785>
- Diyanulhaq, M. B., & Saeffurohman. (2026). Klasifikasi sentimen YouTube terkait demonstrasi Agustus 2025 menggunakan metode TF-IDF dan algoritma XGBoost. *INTECOMS: Journal of Information Technology and Computer Science*, 9(3), 613–619. <https://doi.org/10.31539/z74yc182>
- Fadli, M., Sanjaya, I., Surono, M., & Suryono, R. R. (2026). Analisis sentimen isu redominasi rupiah menggunakan lexicon based dan Naïve Bayes. *JUTISI: Jurnal Ilmiah Teknik Informatika dan Sistem Informasi*, 15(2), 757–766. <https://doi.org/10.35889/jutisi.v15i2.3436>
- Hernikawati, D. (2021). Kecenderungan tanggapan masyarakat terhadap vaksin Sinovac berdasarkan lexicon based sentiment analysis. *Jurnal IPTEK-KOM: Jurnal Ilmu Pengetahuan dan Teknologi Komunikasi*, 23(1), 21–31. <https://doi.org/10.17933/iptekkom.23.1.2021.21-31>
- Humas Fraksi PKS. (2020, July 11). RUU denominasi masuk Prolegnas 2020–2024, Anis minta pemerintah prioritaskan program lain. Fraksi Partai Keadilan Sejahtera. <https://fraksi.pks.id/2020/07/11/ruu-denominasi-masuk-prolegnas-2020-2024-anis-minta-pemerintah-prioritaskan-program-lain/>
- Indraini, A. (2025, November 18). Gubernur BI jelaskan tahapan redenominasi, butuh waktu 6 tahun. *detikSumut*. <https://www.detik.com/sumut/bisnis/d-8215997/gubernur-bi-jelaskan-tahapan-redenominasi-butuh-waktu-6-tahun>
- Maharani, M. D., Indradewi, I. G. A. A. D., & Wijaya, I. N. S. W. (2026). Perbandingan performa TF-IDF dan BoW pada analisis sentimen BPJS Kesehatan menggunakan XGBoost. *Jurnal Pendidikan Teknologi dan Kejuruan*, 23(1), 13–24. <https://doi.org/10.23887/jptk-undiksha.v23i1.109604>
- Novaković, J. D., Veljović, A., Ilić, S. S., Papić, Ž., & Tomović, M. (2017). Evaluation of classification models in machine learning. *Theory and Applications of Mathematics & Computer Science*, 7(1), 39–46. <https://www.uav.ro/jour/index.php/tamcs/article/view/2234>
- Pardede, D., Hayadi, B. H., & Iskandar. (2022). Kajian literatur multi layer perceptron: Seberapa baik performa algoritma ini. *Journal of ICT Applications and System*, 1(1), 23–35. <https://doi.org/10.56313/jictas.v1i1.127>
- Permana, S. H. (2015). Prospek pelaksanaan redenominasi di Indonesia. *Jurnal Ekonomi dan Kebijakan Publik*, 6(1), 109–122. <https://doi.org/10.22212/jekp.v6i1.159>
- Purnamasari, N., & Hendrastuty, N. (2026). Perbandingan kinerja XGBoost dan Naive Bayes dalam analisis sentimen komentar TikTok terhadap Ibu Kota Nusantara (IKN) pada data tidak seimbang. *Building of Informatics, Technology and Science (BITS)*, 7(4), 2690–2703. <https://doi.org/10.47065/bits.v7i4.9488>

- Puspita, M. D., & Sulistya, A. R. (2024, December 14). *59 tahun lalu Indonesia lakukan redenominasi rupiah dari Rp1.000 jadi Rp1*. Tempo. <https://www.tempo.co/ekonomi/59-tahun-lalu-indonesia-lakukan-redenominasi-rupiah-dari-rp-1-000-jadi-rp-1--1181265>
- Sofyan, F. M. A., Sulistiyowati, N., & Voutama, A. (2024). Analisis sentimen terhadap respons perubahan nama Twitter menjadi “X” menggunakan metode Naïve Bayes classifier. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(5), 10987–10994. <https://doi.org/10.36040/jati.v8i5.10720>
- Tharwat, A. (2021). Classification assessment methods. *Applied Computing and Informatics*, 17(1), 168–192. <https://doi.org/10.1016/j.aci.2018.08.003>
- Wati, H. L., Wiradinata, M. T. I. R., Anggraeni, N., Kolbiah, S., Hendar, U., & Agustina, N. (2025). Perbandingan algoritma random forest dan XGBoost untuk analisis sentimen publik terhadap Rumah Makan Payakumbuh. *NARATIF: Jurnal Nasional Riset, Aplikasi dan Teknik Informatika*, 7(1), 64–71. <https://doi.org/10.53580/naratif.v7i1.307>
- Widiarta, I. P. A. P., Dwiysaputra, R., & Aranta, A. (2023). Analisis sentimen masyarakat terhadap kebijakan penerapan PPKM di media sosial Twitter dengan menggunakan metode XGBoost. *Jurnal Teknologi Informasi, Komputer, dan Aplikasinya (JTika)*, 5(2), 154–163. <https://doi.org/10.29303/jtika.v5i2.342>
- Wulla, J. D. A., & Rino. (2026). Sentiment analysis of Indonesian tweets on the free lunch and milk program using TF-IDF and Naïve Bayes. *Journal of Computer Science and Intelligent Systems*, 1(1), 29–42. <https://eduscienta.com/index.php/JCSIS/article/view/3>
- Zhang, Y., Liu, J., & Shen, W. (2022). A review of ensemble learning algorithms used in remote sensing applications. *Applied Sciences*, 12(17), Article 8654. <https://doi.org/10.3390/app12178654>