

Analysis of Public Concerns Towards COVID-19 Vaccines on X Post-Peak Pandemic Phase

Luthfi Atikah¹, Novrindah Alvi Hasanah², Binti Kholifah³

¹ Management Informatic: Politeknik Astra, Kabupaten Bekasi, Indonesia

² Departement of Electrical Engineering: Universitas Islam Maulana Malik Ibrahim, Malang, Indonesia

³ Management Informatic Universitas Negeri Surabaya, Surabaya, Indonesia

luthfiatikah@polytechnic.astra.ac.id¹, bintikholifah@unesa.ac.id², novrindah@uin-malang.ac.id³

Accepted 03 June 2026

Approved 28 June 2026

Abstract— The rapid development of social media, particularly platform X, has created both opportunities and challenges in public health discourse. This study aims to analyze public concerns and perceptions toward COVID-19 vaccines in the post-peak pandemic phase. A dataset of 4.519 Indonesian-language tweets collected between September 2024, and September 2025 was processed through four stages of preprocessing and subsequently labeled. The labeled data were then evaluated using two experimental scenarios, namely with and without data augmentation, and analyzed using a seven-class sentiment classification framework. The results indicate that negative sentiment overwhelmingly dominates the discourse at 96.29%, with strongly negative sentiment accounting for 91.88%. Meanwhile, neutral sentiment is only 0.36%, and positive sentiment remains relatively low at 3.35%. These findings suggest that public distrust toward COVID-19 vaccines persists even in the post-peak phase of the pandemic. Therefore, more effective, targeted, and evidence-based health communication strategies are necessary to address public concerns and improve vaccine confidence.

Index Terms— BERT; IndoBert; COVID-19 Vaccine; Sentiment Analysis; X; Vaccine Hesitancy.

I. INTRODUCTION

In March 2020, the World Health Organization (WHO) declared COVID-19 a global pandemic with widespread impacts on public health, the environment, human psychology, the global socioeconomy, and education [1]; [2]. The surge in mortality due to the virus placed immense pressure on healthcare systems worldwide [3]. Various measures were implemented to curb the uncontrolled spread of the virus, including mobility restrictions. However, these policies led to increased unemployment and reduced household income due to significant losses faced by many companies. Additionally, major shifts in social interaction patterns widened social inequalities and made mental health a critical concern for policymakers. In fact, widespread psychological distress resulted in increased cases of depression globally [4].

Effective pandemic management became a crucial step in maintaining global health system stability. [5] stated that delayed responses and weak healthcare systems contributed to approximately 14.8 million deaths during the pandemic. In response, experts developed various solutions, including vaccination programs as an effective preventive strategy. Several studies highlight vaccination as a key solution to reducing COVID-19 transmission. For instance, [6] demonstrated that vaccines significantly reduce symptomatic infections. Similarly, B [7] reported 94.1% efficacy in preventing COVID-19 in large-scale trials of the Moderna vaccine.

The success of vaccination programs heavily depends on public trust. Vaccine hesitancy, often driven by misinformation and distrust in health authorities, can significantly reduce vaccination coverage. Social media has become a key platform for expressing public concerns regarding vaccination programs [8]. Studies also show that hesitancy is not limited to the public but is also present among healthcare workers, with vaccine side effects being a primary concern [9]. Public rejection of government vaccination programs often stems from low trust in health institutions. This is reinforced by [10] who emphasized the importance of trust and confidence in determining vaccine acceptance. Therefore, successful vaccination programs require not only vaccine availability but also comprehensive communication, wide coverage, equitable distribution, and healthcare system preparedness [11].

If these aspects are neglected, public concern will increase, becoming a major issue for health authorities. Social media serves as an effective tool to monitor these concerns. During the pandemic, platforms like X became central to public opinion formation due to their accessibility and compatibility with social distancing measures. However, the rapid spread of information on such platforms also accelerates the dissemination of misinformation. [12], in “The COVID-19 Infodemic: Twitter versus Facebook,” showed that social media can amplify misinformation, significantly influencing

public perception and decision-making. Thus, while social media has great potential as a health communication tool, it also requires strategies to mitigate misinformation risks.

Social media data analysis enables real-time measurement of public sentiment, as these platforms serve as primary outlets for expressing opinions and emotions. Twitter is widely used to detect public perception trends and monitor reactions to government policies [13]. However, the exponential growth of digital data presents significant challenges for manual analysis. [14] highlighted that current data volumes exceed human processing capacity, necessitating automated and algorithmic approaches. During the pandemic, millions of daily posts made manual sentiment analysis impractical [15]. Therefore, integrating machine learning approaches becomes essential for handling large-scale data efficiently [16].

Text mining techniques allow researchers to extract valuable insights from massive textual datasets [17]. Sentiment analysis plays a key role in understanding public emotions, with deep learning models achieving high accuracy. Machine learning approaches have also proven effective for real-time monitoring of public sentiment during crises [18]. Consequently, integrating natural language processing (NLP) and machine learning is crucial for efficient and scalable analysis of textual data.

Artificial intelligence (AI), especially NLP, has proven effective in automatically understanding and classifying text. Studies such as [19] and [20] demonstrated high accuracy in text classification tasks using deep learning. Transformer-based models like BERT have significantly advanced NLP by improving semantic and syntactic understanding [21]. These models enable tasks such as topic classification, entity recognition, and sentiment analysis. Additionally, large-scale datasets like COVID Senti further demonstrate the

effectiveness of NLP in analysing pandemic-related sentiment.

BERT (Bidirectional Encoder Representations from Transformers) is a state-of-the-art NLP model known for its bidirectional context understanding. It outperforms traditional models in capturing nuanced emotions and hidden concerns in text. Studies have shown its effectiveness in analysing COVID-19-related tweets and detecting public concerns [22]. Despite its strong performance, previous studies in Indonesia have mostly focused on simple three-class sentiment classification or non-Indonesian datasets.

Based on this gap, this study focuses on analysing public sentiment toward COVID-19 vaccines in the post-pandemic phase using BERT, by classifying sentiment into positive, negative, and neutral categories. This research aims to provide deeper insights into public perception and emotional responses toward vaccination programs in Indonesia, particularly as expressed on Twitter (X).

II. METHOD

A. Research Design

At this stage, several data preparation steps are carried out based on Figure 1. Including data collection, Preprocessing, Deep Learning based classification, and Evaluation. The figure illustrates the overall workflow of this study.

This study begins with data collection from X (formerly Twitter) over a one-year period, spanning from September 2024 to September 2025. Some of the collected raw data were preprocessed before proceeding to the next stage, including the removal of irrelevant components. After the preprocessing stage, the dataset was manually labeled into seven main

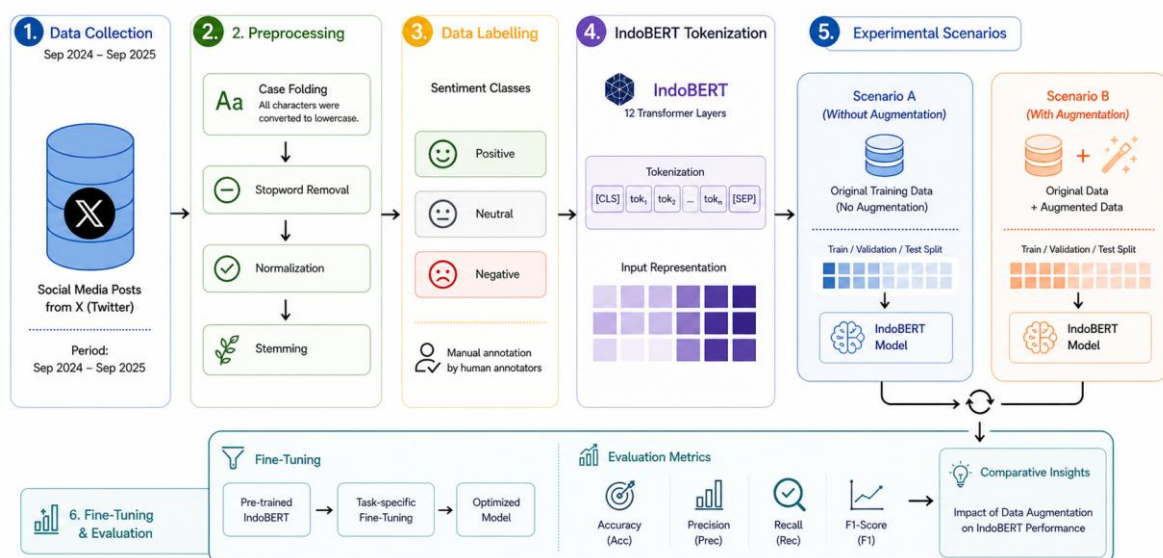


Figure 1. Scenario Research Pipeline

classes: strongly negative, negative, weakly negative, neutral, weakly positive, positive, and strongly positive.

The labeled data are then utilized as input for the IndoBERT model, a pre-trained transformer-based language model designed to capture deep contextual representations of Indonesian text. The output of the IndoBERT model is obtained as a pooled representation with a dimensionality of 768. This representation is then passed through a dropout layer to reduce overfitting, followed by a softmax layer to generate probability distributions over the target classes.

The model was evaluated using several evaluation metrics, including accuracy, precision, recall, and F1-score. The sentiment results were classified into seven classes: strongly negative, negative, weakly negative, neutral, weakly positive, positive, and strongly positive.

The classification into seven sentiment classes in this study was intended to capture the intensity of sentiments expressed regarding the COVID-19 vaccine, thereby providing a more optimal representation of sentiment analysis. For example, tweets expressing mild hesitation have a different representation from tweets expressing strong rejection, even though both would generally be categorized as negative in a three-class classification scheme.

These seven categories were developed from the original three-class manual annotations through a probability-based mapping approach. The mapping process from the three initial classes was further divided into strongly negative (softmax probability > 0.85), negative (0.65–0.85), and weakly negative (0.50–0.65). For neutral tweets, the confidence score threshold was set at > 0.70 , while tweets with lower confidence scores were categorized as either weakly positive or weakly negative. A similar approach was also applied to the positive class. The seven-class categorization adopted in this study was based on previous research on fine-grained sentiment analysis [1]. A detailed explanation of each step is provided in the following section.

B. Data Collection

Data were collected from X (formerly Twitter) over one year, spanning from September 2024 to September 2025. The data collection process was conducted using the Twitter Academic Research API v2, which enables systematic and large-scale retrieval of publicly available tweets for research purposes. This approach ensures that the data acquisition process is both structured and reproducible.

During the data collection process, a total of 4,519 tweets in the Indonesian language were collected. The tweets were then filtered based on predefined keywords. In addition to the tweet text, the dataset also consisted of several important attributes, including publication date, number of likes, number of retweets, and other related metadata. Furthermore, bot and spam filtering was conducted by manually removing irrelevant tweets. This stage was carried out to reduce noise and improve the representativeness of the dataset.

All collected data were stored in Comma-Separated Value (CSV) format to facilitate efficient data handling, processing, and reproducibility of the study. The use of a standardized data format also ensures compatibility with various data analysis tools and machine learning frameworks.

Table 1. Keyword Categories

Theme	Keywords (Indonesian)
Side Effects (Efek Samping)	efek samping, reaksi alergi, sakit kepala, demam tinggi
Long-term Effects (Efek Jangka Panjang)	infertilitas, autoimun, risiko kanker
Safety (Keamanan)	risiko kesehatan, anak-anak dan ibu hamil
Efficacy (Efektivitas)	efektif atau tidak, varian baru, perlindungan menurun
Distribution (Distribusi)	ketidakadilan, akses terbatas, negara berkembang
New Technology (Teknologi Baru)	mRNA, teknologi belum teruji, manipulasi genetik

C. Preprocessing

The raw data obtained through the crawling technique are then preprocessed to improve data quality before being used in the modeling process. This preprocessing stage aims to remove noise, simplify text representation, and ensure data consistency, thereby improving the model's ability to understand language context.

In this study, preprocessing consists of four main stages: case folding, stopword removal, normalization, and stemming. In the case folding stage, all characters in the text are converted to lowercase to avoid differences in word representation due to capitalization. Next, in the stopword removal stage, common words that do not contribute meaningful information to the sentiment analysis process, such as conjunctions and prepositions, are removed from the text.

The next stage is normalization, which aims to convert non-standard words, such as slang, abbreviations, or informal spelling variations, into their standard forms in accordance with Indonesian language rules. This is followed by stemming, which involves reducing affixed words to their base forms using the Sastrawi stemmer algorithm, allowing morphological variations of a word to be represented uniformly.

In addition to these four main steps, further cleaning is performed by removing punctuation, URLs, and mention tags that do not contribute to the sentiment analysis. After completing all preprocessing steps, the dataset is reduced from 4,519 tweets to 3,815 tweets, making it ready for use in the subsequent modeling stage.

D. Data labeling

The model was evaluated using several evaluation metrics, including accuracy, precision, recall, and F1-score. The sentiment results were classified into seven categories: strongly negative, negative, weakly negative, neutral, weakly positive, positive, and strongly positive. The classification into seven sentiment categories in this study was intended to capture the intensity of sentiments expressed regarding the COVID-19 vaccine, thereby providing a more optimal representation of sentiment analysis. For example, tweets expressing mild doubt have a different representation from tweets expressing strong rejection, even though both would generally be categorized as negative in a three-class sentiment scheme.

These seven categories were developed from three manually annotated classes through probability-based mapping. The mapping process grouped the three classes into strongly negative (softmax probability > 0.85), negative (0.65–0.85), and weakly negative (0.50–0.65). For neutral tweets, the confidence score threshold was set at > 0.70, while tweets that did not meet this threshold were categorized as either weakly positive or weakly negative. A similar approach was applied to the positive class. The seven sentiment categories were adopted based on previous research in fine-grained sentiment analysis [1].

Table 2 presents representative examples from each sentiment class, illustrating the characteristics and linguistic patterns associated with positive, neutral, and negative labels in the dataset.

Table 2. Sample Data Labeled

Sentiment Label	Example Tweet (Original Indonesian)
Negative	imun gw emg agak kurang tapi semenjak gw vaksin covid gw kena neurodermatitis tiba2 bete bgt.

Neutral	Vaksin covid dan alat2 testnya cba cek skalian
Positive	Covid di RI lu dapet vaksin China, aman.

E. Model Architecture and Classification

Sentiment classification in this study was performed using the IndoBERT model (indobenchmark/indobert-base-p1) [19], a transformer-based language model pre-trained to understand the context of the Indonesian language. In the initial stage, tokenization was performed to split tweet text into subword tokens, with a classification token [CLS] added at the beginning of each sequence to represent the overall text.

Next, a padding process was applied to standardize the input length to a maximum of 128 tokens. If a sequence exceeded this limit, truncation was performed to ensure it conformed to the specified maximum length.

In this study, two experimental scenarios were conducted: data augmentation and non-data augmentation. These scenarios were designed to address the issue of class imbalance in the dataset. The dataset was highly imbalanced, with 94% of the tweets labeled as negative. Data augmentation was applied to the minority classes, namely positive and neutral, to improve the model’s generalization capability. The augmentation techniques employed included synonym replacement, random insertion, random swap, and random deletion based on the approach proposed in [21]. For each tweet in the minority classes, three augmented versions were generated, resulting in an augmentation ratio of 1:4. The model output generated probabilities for seven sentiment classes, namely strongly negative, negative, weakly negative, neutral, weakly positive, positive, and strongly positive. The final sentiment label was determined based on the highest probability score. The overall stages of the classification process are illustrated in Figure 2.

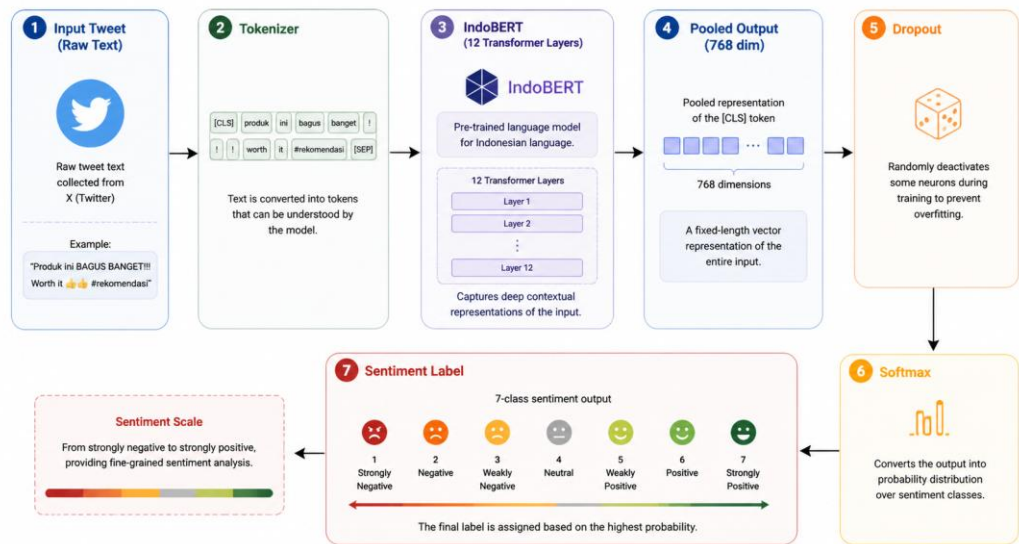


Figure 2. IndoBERT-Based Sentiment Classification Architecture

F. Evaluation Metric

During the evaluation stage, the model was assessed using several evaluation metrics, including accuracy, precision, recall, and F1-score. Accuracy was used to measure the proportion of correct predictions over the entire dataset, while precision and recall were used to evaluate the model's ability to classify data accurately and comprehensively. The F1-score served as a combined metric that balances precision and recall. In addition, a confusion matrix was employed to provide a more in-depth analysis of the classification results in order to identify the strengths and weaknesses of the model for each class.

III. RESULT AND DISCUSSION

In the results and discussion stage, the model was analyzed through the following aspects:

A. Dataset Statistic

After the preprocessing stage, a total of 3,366 tweets were obtained. The preprocessed tweets were then labeled using a semi-supervised approach and categorized into seven sentiment classes: strongly negative, negative, weakly negative, neutral, weakly positive, positive, and strongly positive.

The dataset distribution showed a significant class imbalance, with approximately 90% of the tweets classified as negative sentiment. This imbalance indicates that public opinion was predominantly characterized by concerns, skepticism, and doubts regarding vaccination. The imbalance in label distribution may affect the model's performance; therefore, more in-depth model training and evaluation are required, particularly for the minority classes.

Table 3. Primary Sentimen Distribution

Sentiment Categories	Total Tweets	Percentage (%)
Positive	178	5,29
Neutral	19	0,56
Negative	3.169	94,14
Total	3.366	100,00

B. Sentiment Classification

In the sentiment analysis stage, classification was performed using the IndoBERT model. The results of the classification are presented in Table 4. The IndoBERT model calculated the probability for each predefined sentiment class. The final distribution results were then used to determine the sentiment tendency within each class.

Table 4. Seven Class Sentiment Distribution

Sentiment Categories	Percentage (%)
Strongly Negative	91,88
Negative	3,09

Slightly Negative	1,32
Neutral	0,36
Slightly Positive	0,79
Positive	0,51
Stongly Positive	2,05

Based on the analysis Table 4, it was found that the majority of tweets were classified as negative sentiment. These tweets were dominated by six major issues, namely vaccine side effects, long-term health risks, vaccine safety, the effectiveness of vaccines against new variants, vaccine distribution to remote areas, and concerns regarding mRNA technology. The analysis results indicate that public concern toward vaccination was relatively high.

C. Discussion

Based on the confusion matrix analysis in Scenario B (using data augmentation), the model demonstrated a good classification capability. This was indicated by the successful prediction of Tabel 5 testing data instances correctly. However, the model still experienced difficulties in distinguishing between the negative and weakly negative classes, as well as between the positive and strongly positive classes.

The implementation of data augmentation in Scenario B was considered effective in improving the classification performance for minority classes compared to Scenario A. This can be observed from the increase in the F1-score from 0.9487 to 0.9512, as well as the improvement in overall accuracy from 96.13% to 96.48%, as illustrated in Figure 4.

Based on the experimental results, the sentiment distribution was dominated by negative opinions, accounting for 91.88% of all tweets. The high proportion of negative sentiment in this study may have been influenced by several factors. First, the collected dataset was obtained during the peak phase of the pandemic (September 2024 – September 2025), when vaccination mandates had already ended, leading to an increase in criticism and negative sentiment. Second, the labeling scheme using seven sentiment classes allowed tweets containing negative expressions to be detected more easily compared to neutral categories. Third, numerous negative expressions emerged following reports of vaccine side effects, which further strengthened the spread of negative sentiment within society.

Based on the comparison among the three sentiment categories, a significant difference was observed between the strongly negative sentiment class, which accounted for 91.88%, and the negative sentiment class, which accounted for 1.32%. This indicates that the public not only expressed negative opinions but also showed strong rejection toward vaccination. In contrast, the proportion of positive sentiment in this study was only 3.35%, which is considered significantly lower compared to the positive sentiment proportion reported in previous studies, ranging from approximately 20% to 35%. This indicates the

existence of systematic differences in public trust toward vaccination.

Table 5. Model Performance Evaluation Result

Metric	Scenario A (Without Augmentation)		Scenario B (With Augmentation)	
Class	Precision	Recall	Precision	Recall
Strongly Negative	0.97	0.99	0.97	0.99
Negative	0.41	0.20	0.55	0.35
Slightly Negative	0.38	0.17	0.52	0.30
Neutral	0.00	0.00	0.20	0.14
Slightly Positive	0.33	0.19	0.48	0.32
Positive	0.44	0.23	0.61	0.40
Strongly Positive	0.50	0.31	0.64	0.44
Accuracy	96.13 %		96.48%	

Based on the analysis, this study still has several limitations, particularly the high level of dataset imbalance, which may introduce bias during both the labeling process and model prediction. Therefore, future studies are expected to utilize larger datasets and include more diverse data variations in order to reduce label distribution imbalance.

IV. CONCLUSION

This study shows that COVID-19 vaccination remains a topic of discussion, particularly after the pandemic situation began to subside. In this research, sentiment analysis was conducted on Indonesian-language tweets using the IndoBERT method. The results indicated that public conversations were dominated by negative sentiment, accounting for 96.29% of the overall sentiment classes. Tweets within the negative sentiment categories reflected various emotional expressions, including fear, skepticism, criticism toward vaccination, concerns about vaccine side effects, uncertainty regarding the long-term impact of vaccination, and distrust toward mRNA technology. However, this study also encountered the issue of data imbalance; therefore, data augmentation was applied to the dataset used for modeling. Future studies are expected to utilize larger and more diverse datasets in order to improve model performance and reduce dataset imbalance.

REFERENCES

- A. Kumar et al., "Waste Classification Using Deep Learning," V. P. Kalanjati et al., "Sentiment analysis of Indonesian tweets on COVID-19 and COVID-19 vaccinations," *F1000Res.*, vol. 12, p. 1007, Aug. 2023, doi: 10.12688/f1000research.130610.1.
- [2] Y. Miyah, M. Benjelloun, S. Lairini, and A. Lahrichi, "COVID-19 Impact on Public Health, Environment, Human Psychology, Global Socioeconomy, and Education," 2022, *Hindawi Limited*. doi: 10.1155/2022/5578284.
- [3] H. Wang et al., "Estimating excess mortality due to the COVID-19 pandemic: a systematic analysis of COVID-19-related mortality, 2020–21," *The Lancet*, vol. 399, no. 10334, pp. 1513–1536, Apr. 2022, doi: 10.1016/S0140-6736(21)02796-3.
- [4] H. Offi cials and S. Galea, "Public Health COVID-19 Impact Assessment: Lessons Learned and Compelling Needs About the NAM series on Emerging Stronger After," 2021. [Online]. Available: <https://astho.org/Research/Data-and-Analysis/State-and-Local-Governance->
- [5] O. J. Watson, G. Barnsley, J. Toor, A. B. Hogan, P. Winskill, and A. C. Ghani, "Global impact of the first year of COVID-19 vaccination: a mathematical modelling study," *Lancet Infect. Dis.*, vol. 22, no. 9, pp. 1293–1302, Sep. 2022, doi: 10.1016/S1473-3099(22)00320-6.
- [6] E. J. Haas et al., "Impact and effectiveness of mRNA BNT162b2 vaccine against SARS-CoV-2 infections and COVID-19 cases, hospitalisations, and deaths following a nationwide vaccination campaign in Israel: an observational study using national surveillance data," *The Lancet*, vol. 397, no. 10287, pp. 1819–1829, May 2021, doi: 10.1016/S0140-6736(21)00947-8.
- [7] L. R. Baden et al., "Efficacy and Safety of the mRNA-1273 SARS-CoV-2 Vaccine," *New England Journal of Medicine*, vol. 384, no. 5, pp. 403–416, Feb. 2021, doi: 10.1056/nejmoa2035389.
- [8] S. L. Wilson and C. Wiysonge, "Social media and vaccine hesitancy," *BMJ Glob. Health*, vol. 5, no. 10, Oct. 2020, doi: 10.1136/bmjgh-2020-004206.
- [9] A. A. Dror et al., "Vaccine hesitancy: the next challenge in the fight against COVID-19," *Eur. J. Epidemiol.*, vol. 35, no. 8, pp. 775–779, Aug. 2020, doi: 10.1007/s10654-020-00671-y.
- [10] L. S. Rozek, P. Jones, A. Menon, A. Hicken, S. Apsley, and E. J. King, "Understanding Vaccine Hesitancy in the Context of COVID-19: The Role of Trust and Confidence in a Seventeen-Country Survey," *Int. J. Public Health*, vol. 66, p. 636255, 2021, doi: 10.3389/ijph.2021.636255.
- [11] S. L. Wilson and C. Wiysonge, "Social media and vaccine hesitancy," *BMJ Glob. Health*, vol. 5, no. 10, pp. 1–7, 2020, doi: 10.1136/bmjgh-2020-004206.
- [12] M. Cinelli et al., "The COVID-19 social media infodemic," *Sci. Rep.*, vol. 10, no. 1, Dec. 2020, doi: 10.1038/s41598-020-73510-5.
- [13] N. El-Bassel, K. R. Hochstatter, M. N. Slavina, C. Yang, Y. Zhang, and S. Muresan, "Harnessing the Power of Social Media to Understand the Impact of COVID-19 on People Who Use Drugs During Lockdown and Social Distancing," *J. Addict. Med.*, vol. 16, no. 2, pp. E123–E132, Mar. 2022, doi: 10.1097/ADM.0000000000000883.
- [14] Viktor. Mayer-Schönberger and Kenneth. Cukier, *Big data : a revolution that will transform how we live, work, and think*. Houghton Mifflin Harcourt, 2013.
- [15] S. Sharma, A. Kundu, S. Basu, N. P. Shetti, and T. M. Aminabhavi, "Sustainable environmental management and related biofuel technologies," Nov. 01, 2020, *Academic Press*. doi: 10.1016/j.jenvman.2020.111096.
- [16] M. M. Najafabadi, F. Villanustre, T. M. Khoshgoftaar, N. Seliya, R. Wald, and E. Muharemagic, "Deep learning applications and challenges in big data analytics," *J. Big Data*, vol. 2, no. 1, Dec. 2015, doi: 10.1186/s40537-014-0007-7.
- [17] X. Dou et al., "Research progress on chitosan and its derivatives in the fields of corrosion inhibition and

- antimicrobial activity,” May 01, 2024, *Springer*. doi: 10.1007/s11356-024-33351-5.
- [18] “icns_20_1-3_53”.
- [19] N. A. Hasanah, N. Suciati, and D. Purwitasari, “Identifying degree-of-concern on covid-19 topics with text classification of twitters,” *Register: Jurnal Ilmiah Teknologi Sistem Informasi*, vol. 7, no. 1, pp. 50–62, 2021, doi: 10.26594/register.v7i1.2234.
- [20] L. Atikah, D. Purwitasari, and N. Suciati, “DETEKSI KEJADIAN LALU LINTAS PADA TEKS TWITTER DENGAN PENDEKATAN KLASIFIKASI MULTI-LABEL BERBASIS DEEP LEARNING”, doi: 10.25126/jtiik.202295206.
- [21] W. J. Li, “Letter for Qiu et al. (2021) regarding ‘The distribution and behavioral characteristics of plateau pikas (*Ochotona curzoniae*)’,” 2022, *Pensoft Publishers*. doi: 10.3897/zookeys.1102.79148.
- [22] F. Abbas and A. Taeihagh, “Unmasking deepfakes: A systematic review of deepfake detection and generation techniques using artificial intelligence,” Oct. 15, 2024, *Elsevier Ltd*. doi: 10.1016/j.eswa.2024.124260.

