

Sentiment Analysis of Nuclear Waste on Twitter and Instagram Data

Chao Bing ^a, Fatimah Sidi ^{a,1}, Iskandar Ishak ^{a,2}, Lili Nurliyana Abdullah ^b, Kurmashev Ildar Gusmanovich ^c

^a Department of Computer Science, Faculty of Computer Science and Information Technology, University Putra Malaysia, Serdang, Malaysia

^b Department of Multimedia, Faculty of Computer Science and Information Technology, University Putra Malaysia, Serdang, Malaysia

^c Department of Information and Communication Technologies, Faculty of Engineering and Digital Technology, M. Kozybayev North Kazakhstan University, Kazakhstan

Corresponding author: ¹fatimah@upm.edu.my; ²iskandar_i@upm.edu.my

Abstract—This study investigates public opinion and sentiment toward Japan's Fukushima nuclear wastewater discharge as expressed on social media, specifically Twitter and Instagram. The goal is to perform sentiment analysis to provide stakeholders with a reliable foundation for informed decision-making. The main issue of this study is the absence of a comprehensive analysis of public opinion and sentiment on social media regarding nuclear wastewater discharge from the Fukushima Daiichi Nuclear Power Plant. Thus, capturing and understanding these sentiments is essential, as they provide crucial insights into public opinion. Hence, a dataset of 12,195 posts from January to March 2024 was collected using Python's Tweepy and Instagrapi libraries. Data preprocessing involved cleaning non-English text, hashtags, usernames, URLs, and emojis, followed by stop word removal. Key topics such as "nuclear," "waste," and "Fukushima" were identified through LDA topic modeling, focusing on posts highly relevant to the issue. Sentiment analysis was conducted using the BERT (uncased) model, categorizing sentiments into "very positive," "positive," "neutral," "negative," and "very negative." VADER was also applied to validate and compare the results, enhancing the robustness of the sentiment classification. Findings show that negative sentiment consistently outweighs positive sentiment, indicating an overall unfavorable public opinion toward the discharge. This analysis provides policymakers and stakeholders with valuable insights into and guidance for addressing global concerns about Japan's nuclear wastewater strategy. Future studies could include additional social media platforms and incorporate geolocation data to analyze regional sentiment trends more precisely.

Keywords—Sentiment analysis; nuclear waste; social media; NLP.

Manuscript received 29 Dec. 2024; revised 30 Nov. 2025; accepted 12 Jan. 2026. Date of publication 30 Jun. 2026.
IJASEIT is licensed under a Creative Commons Attribution-Share Alike 4.0 International License.



I. INTRODUCTION

On the 11th of March 2011, an earthquake with a magnitude of 9.0 struck off the east coast of Japan. A massive tsunami damaged the Fukushima Daiichi Nuclear Power Plant's cooling system, causing three nuclear reactors to overheat and parts of their cores to melt [1]. Since melted fuel fragments were found in the three destroyed reactors, water had to be continuously injected to cool them. Japanese authorities continued to inject seawater into the reactors, along with groundwater and rainwater. This water became contaminated when it came into contact with melted fuel, fuel fragments, and even radioactive fallout. Japanese authorities stored this radioactive water in storage tanks and then used a multi-nuclide processing system (ALPS) to eliminate substances of radioactive except carbon 14 and tritium. The Fukushima Daiichi Nuclear Power Plant's water tanks stored 1.3 million tons of nuclear wastewater, enough to fill 500 Olympic-sized

swimming pools [2]. As of November 2022, Japan has spent about 12.1 trillion yen (82 billion U.S. dollars) to address the incident (including compensation and reactor decommissioning costs), and this figure continues to grow [3]. To speed up the reactor's decommissioning, given Japan's economic and technological maturity and current events, Japan decided to dilute the wastewater and discharge it into the Pacific Ocean.

On the 22nd of August 2023, Japanese Prime Minister Fumio Kishida announced that the Fukushima nuclear wastewater would be discharged into the sea starting on August 24, 2023. On August 24, 2023, at 1 pm Tokyo time, Tokyo Electric Power Company officially began to discharge treated water. The International Atomic Energy Agency (IAEA) reviewed the treated water and found that it met international safety standards. According to Tokyo Electric Power Company's plan, the radioactive impact of ALPS-treated water discharged into the sea on humans and the environment could be ignored [4]. The first stage of discharge

lasted 17 days, during which 7,800 tons of nuclear-treated water were discharged at a rate of 460 tons per day [5]. This matter has sparked discussion among people from all over the world on social media, especially on mainstream platforms such as Twitter and Instagram.

The power of public opinion on social media is powerful. It not only resonates with people but, when it becomes a social phenomenon, can even influence changes in national policy. For example, Swedish environmental activist Greta Thunberg uses Twitter and Instagram to raise awareness of climate change worldwide and to organize global student climate strikes. Her activities have promoted the younger generation's attention and participation in climate change issues [6]. They have sparked extensive discussions on environmental protection policies and sustainable development, and have put pressure on policy-making in many countries.

With the implementation of Japan's Fukushima plan to dump nuclear wastewater into the Pacific Ocean, there has been a significant increase in posts on Twitter and Instagram. This surge in social media activity reflects public reactions and sentiments regarding the event. However, there is a notable gap in systematically analyzing and understanding these sentiments.

The primary issue is the lack of a comprehensive analysis of public sentiment on social media regarding the discharge of Fukushima nuclear wastewater. Understanding these sentiments is crucial, as they provide insights into public opinion that can guide policymakers and relevant stakeholders in addressing public concerns and improving communication strategies [7].

The significance of this research lies in its potential to inform decision-makers about the societal impact of the Fukushima plan. By systematically analyzing social media posts, this study can reveal the public's predominant emotions and opinions. This understanding can aid in anticipating public reactions to future similar events and in formulating effective response strategies [8].

To address this problem, this study aims to conduct a sentiment analysis of social media posts related to the Fukushima nuclear wastewater discharge. Specifically, it will employ text mining methods to analyze data collected from Twitter and Instagram using specific hashtags as keywords. The data cleaning and sentiment analysis processes will be carried out using Python, with advanced NLP techniques such as LDA for topic modeling and the pre-trained BERT model for sentiment analysis.

By integrating advanced NLP techniques, the research aims to provide a comprehensive understanding of public opinion on the discharge of Fukushima nuclear wastewater. The use of LDA will help reduce noise and enhance the relevance of the text by identifying key topics [9], while BERT will offer a nuanced sentiment analysis of the posts [10]. Ultimately, the goal is to obtain a detailed picture of public sentiment and the main concerns surrounding the event, thus contributing valuable insights for future policy and communication efforts.

Section I of this paper covers the background, the problem statement, and the study's significance. Section II delves into the tools and methods applied for sentiment analysis of tweets and Instagram posts, including the detailed methodology.

Section III presents the results and discussion, while Section IV draws the conclusion and outlines potential future work.

II. MATERIALS AND METHOD

Over the past two decades, social media has transformed how people communicate, acquire, and share information, with platforms like Facebook, Twitter, and Instagram fostering more connected and dynamic interactions [11], [12]. Whether sharing daily moments or real-time updates of major events, social media enables users to exchange views and stay connected with friends, family, and colleagues globally.

In political and social activities, social media enhances citizen participation by providing a platform for real-time expression and interaction. For example, during elections, voters can learn about candidates, join discussions, and interact directly with politicians. In social movements, it empowers organizers to mobilize and spread information, amplifying public engagement and influence [13], [14]. With its vast data resources, social media has become an ideal platform for analyzing behavior and psychology. Sentiment analysis, a key data analysis method, helps reveal public attitudes, emotional tendencies, and satisfaction levels, providing crucial insights for decision-making [15], [16], [17].

During crises, sentiment analysis [18] and machine learning [19] play a critical role in helping businesses and governments quickly gauge public reactions, address negative emotions, and manage crises effectively. For instance, understanding public sentiment during product recalls, natural disasters, or public health emergencies enables timely and efficient responses [20].

Sentiment analysis methodologies have evolved significantly over time, beginning with lexicon-based methods in the early 2010s that relied on predefined sentiment dictionaries. These methods increased vocabulary coverage and did not require labeled data, but they were prone to noise from irrelevant words. In the mid-2010s to early 2020s, machine learning-based methods such as Naïve Bayes and SVMs emerged, offering better context understanding. However, their effectiveness depended heavily on the clarity of the training data labels, and they struggled with mixed sentiments.

From the late 2010s onwards, deep learning-based approaches such as MF-CNN-BLSTM and BERT-based models marked a significant advancement by automatically learning complex feature representations, thereby enhancing sentiment analysis accuracy. These methods, however, required large amounts of labeled data and involved complex parameter tuning. Recently, hybrid models like PLASA have combined various techniques to improve sentiment analysis, particularly for short texts, by incorporating punctuation and extensive dictionary information. Despite these improvements, challenges remain in effectively analyzing longer texts and filtering out non-human content.

Cho et al. [21] introduced a data-driven method for adapting sentiment dictionaries to various domains. It involves merging multiple dictionaries, filtering non-contributory words, and adjusting word polarities based on empirical ratios of positive to negative occurrences in training data. The integrated sentiment dictionary, refined through these operations, demonstrates robust performance in sentiment analysis across domains such as smartphones, movies, and books. Comparative evaluations highlight its

superiority over individual dictionaries. The study discusses the effectiveness of this automated approach for domain-specific sentiment dictionary adaptation, leveraging abundant labeled web opinion data and existing lexical resources. This method offers a straightforward, data-driven solution to improve sentiment classification accuracy without manual intervention [21].

Sano et al. [22] introduced a visualized comparative review analysis model using the Naïve Bayes classifier for the tourism domain. The study aims to enhance understanding of tourist experiences by visualizing their opinions across different destinations. The methodology includes data collection from primary interviews and secondary online reviews, data preprocessing (cleaning, case folding, tokenization, stop-word removal, stemming, and TF-IDF weighting), and k-fold cross-validation for model evaluation.

The model was tested on reviews from two Balinese beaches, Jimbaran and Kuta. The accuracy was 80% for Jimbaran and 64% for Kuta. The results are visualized to facilitate easy comparison and informed decision-making for various stakeholders in the tourism industry.

The visual comparison aids decision-making. However, the accuracy for Kuta beach was lower, potentially due to ambiguous sentiments in the training data. The study suggests that the model can be extended to other domains requiring comparative analysis for decision-making [22].

Wang and Zhao [23] used a Support Vector Machine (SVM) to predict investor sentiment based on news texts. The study uses a psychological line (PSY) as a sentiment indicator. Data is collected via web crawlers from financial news media, specifically from the medical industry. News texts are processed and vectorized using doc2vector, and PSY is calculated from stock data.

The SVM model achieves high accuracy in predicting investor sentiment. However, including industry-specific news can lead to overfitting and reduce accuracy. The study finds that varying amounts of news data affect prediction accuracy, with less data leading to lower accuracy. Despite these challenges, the model effectively predicts stock information and investor sentiment [23].

Aslan [24] collected 192,915 tweets related to the Ukraine–Russia conflict using specific hashtags. The tweets were preprocessed to remove special characters, URLs, and non-English words, and were then tokenized and lemmatized. The Valence Aware Dictionary and Sentiment Reasoner (VADER) was used to classify the tweets into positive, negative, and neutral sentiments. Latent Dirichlet Allocation (LDA) was employed to identify dominant topics in the tweets.

A new deep learning-based sentiment classification model, MF-CNN-BiLSTM, was proposed. This model combines Convolutional Neural Networks (CNN) and Bidirectional Long Short-Term Memory (BiLSTM) to capture both local and long-term dependencies in the data. The FastText word embedding technique was used to convert tweets into vectors. The analysis revealed that 50.38% of the tweets were negative, 30.95% were positive, and 18.67% were neutral. The LDA identified seven dominant topics in the tweets, which were visualized using word clouds.

The proposed MF-CNN-BiLSTM model achieved an accuracy of 95.32%, outperforming CNN, LSTM, and BiLSTM models. The FastText word embedding technique

proved more effective than GloVe. The study found that the majority of tweets expressed negative sentiment towards the Ukraine–Russia conflict, a trend that was consistent across countries, though intensity varied. The dominant topics identified by LDA revealed the main concerns and discussions among Twitter users regarding the conflict, including political reactions, humanitarian issues, and military actions.

The MF-CNN-BiLSTM model's high accuracy demonstrates the effectiveness of combining CNN and BiLSTM for sentiment analysis. The use of FastText embeddings further enhanced the model's performance by capturing sub-word information and handling out-of-vocabulary words [24].

Ashayeri and Abbasabadi [25] used social media data from the X platform (formerly Twitter) and urban building data from NYC's PLUTO and the Department of Energy's LEAD databases. Keywords like "stay-at-home," "energy bill," and "HVAC" were used to filter relevant tweets. The data underwent preprocessing, tokenization, and summarization. Sentiment analysis was performed using the RoBERTa transformer deep learning model. Co-occurrence matrices and advanced visualization techniques were employed to understand word relationships and model attention patterns.

The RoBERTa model, pre-trained on sentiment analysis tasks, was used to classify tweets into positive, negative, or neutral sentiments. The analysis included generating word clouds, embedding visualizations, and attention maps. The study found varying sentiment distributions across NYC boroughs for the terms "stay-at-home," "energy bill," and "HVAC." Manhattan had more positive sentiments, while Brooklyn and the Bronx showed higher negative sentiments.

Detailed attention patterns from the RoBERTa model highlighted the relationships between words in tweets, showing how the model prioritized different parts of the text. The study visualized neuron activations in the RoBERTa model, revealing intricate word associations during the COVID-19 era. The research highlighted significant disparities in energy burdens and urban building characteristics across NYC boroughs, emphasizing the need for equitable urban development and energy policies.

The findings can inform data-driven policy development, targeting energy subsidies and promoting energy-efficient systems. The study also suggests the importance of ongoing monitoring of public sentiment to refine energy policies [25].

Li and Zou [26] proposed the PLASA model, which utilizes the pre-trained BERT model to produce contextually sensitive embeddings and semantically rich comments, transforming words and sentences into vectors. This model incorporates lexicon-query and punctuation-location attentions. Lexicon-query attention used a specialized lexicon to adjust the relationship between background words and sentiment-related keywords, while punctuation-location attention assigns weights to words based on their proximity to punctuation marks.

The refined representations are fed into various classifiers, including machine-learning methods (RF, SVM, XGBoost) and deep-learning methods (RNN, BiLSTM), to classify sentiment polarity. The PLASA model demonstrated significant performance improvements over baseline and commonly used models, achieving micro-F1 scores of

93.94%, 90.34%, and 88.79% on the SentiTikTok-2023, SentiBilibili-2023, and SentiWeibo-2023 datasets, respectively.

The study showed that both lexicon-query and punctuation-location attentions contributed to the model's performance. The punctuation-location attention was more effective for shorter texts, while the lexicon-query attention was crucial for longer texts. The optimal parameters for the collaborative attention module were identified as a window dimension of 5 and an intensity coefficient of 1, providing stability and strong performance across datasets.

The study emphasized that both punctuation and lexicon are crucial factors in sentiment analysis. Punctuation consistently made up approximately 13% of comments, indicating its vital role in language expression. Lexicon features also differed significantly between positive and negative segments. The PLASA model effectively combined state-of-the-art NLP techniques with traditional sentiment analysis methods, enhancing performance and interpretability. The collaborative attention module was particularly effective in refining word representations for sentiment classification [26].

Jia et al. [27] proposed the BERT-CBLBGA model, which integrates BERT with BiGRU, BiLSTM, TextCNN, and a self-attention mechanism to enhance text sentiment classification. BERT generates word embeddings, which are then processed by TextCNN for local features and BiLSTM/BiGRU for global features. The self-attention mechanism fuses these features for better sentiment classification.

TextCNN captures local features via convolution and pooling, while BiGRU and BiLSTM extract global features by modeling sequential information. The self-attention mechanism dynamically assigns different features varying importance. The model uses pre-trained BERT models for text vectorization. For Chinese datasets, preprocessing includes segmentation and text cleaning, whereas for English datasets, it involves stemming, stop-word removal, and text cleaning.

The BERT-CBLBGA model was evaluated on datasets: (i) weibo_senti_100k, (ii) IMDB, (iii) waimai_10k, and (iv) jidian_10k. It outperformed baseline models (BERT, BERT-BiLSTM, BERT-LSTM, BERT-BiGRU, and BERT-CNN) in terms of accuracy, recall, and F1-score. The model achieved significant improvements in accuracy across all datasets, demonstrating its ability to capture both local and global sentiment features.

The BERT-CBLBGA model outperformed models using Word2vec and fastText embeddings, highlighting the advantages of BERT's bidirectional encoding mechanism. The integration of BERT, TextCNN, BiLSTM, BiGRU, and self-attention mechanisms allows the model to capture a wide range of sentiment-related features. This comprehensive feature extraction leads to better sentiment classification performance.

The BERT-CBLBGA model effectively combines local and global feature extraction with dynamic feature weighting, making it robust for various sentiment analysis tasks. The self-attention mechanism enhances the model's ability to focus on important sentiment features. The paper suggests further exploration of multilingual support, model

interpretability, and domain-specific adaptations to improve the model's generalization capabilities and applicability in different contexts [27].

Uthirapathy and Sandanam [28] pre-processed raw Twitter data to remove noise and irrelevant information, including stopword removal, URL removal, pattern removal, tokenization, and vectorization. LDA was used to identify hidden topics within the tweets, generating a document-topic matrix and a topic-word matrix to map relevant words to topics. The Bidirectional Encoder Representation from Transformers (BERT) model was used for sentiment classification, processing the tokenized tweets and classifying them into anti, neutral, news, and support.

The LDA model identified the most discussed topics in climate change tweets, categorizing them into anti-, neutral-, news-, and support-related categories. The BERT model achieved high performance in sentiment classification, with a recall of 89.65%, a precision of 91.35%, and an accuracy of 93.50%. The analysis revealed that a significant number of tweets supported the idea of man-made climate change, while a smaller portion opposed it. The BERT model outperformed other methods in terms of precision, recall, and accuracy, demonstrating its effectiveness in sentiment classification [28].

The exploration of big data through sentiment analysis on platforms like Twitter and Instagram offers profound insights into public perceptions of nuclear waste. As technology evolves, so does the potential to harness this information to inform and shape environmental policy effectively. Continued advancements in NLP and machine learning will have a significant role in refining the accuracy and applicability of sentiment analysis in environmental science.

When conducting research, a clear process design is key to ensuring a smooth process and high-quality results. Fig. 1 shows a complete process flow from the data collection to the final result.

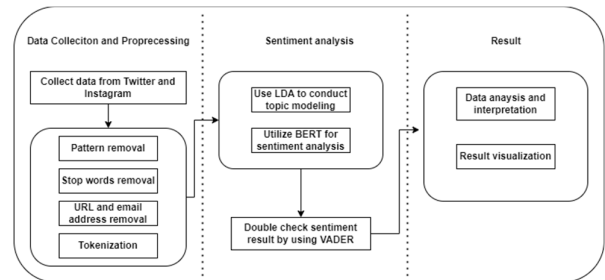


Fig. 1 Research Methodology

A. Data Collection

In this study, users' Twitter tweets and Instagram posts are considered as data sources. These tweets are collected using the Tweepy and Instagrapi libraries in Python. Both libraries can obtain corresponding posts based on keywords. The time span of the collected Twitter dataset is from January 20, 2024 to February 15, 2024, with a total of 9153 posts. The collected IG dataset covers March 20, 2024, and includes 3042 posts. During this period, the news media did not make special reports on the incident, ensuring consistent public performance. The dataset was collected using specific keywords (hashtags) to ensure relevance to the topic. For each hashtag, the corresponding number of posts collected is detailed in Table I (for Twitter) and Table II (for Instagram).

TABLE I

Keyword	Amount	Keyword	Amount
#fukushima water	99	#nuclear waste	3092
#fukushima	99	#nuclear water	2089
#FukushimaWaterRelease	17	#nuclearsafety	19
#japan radiation	227	#nuclearwaste	58
#nuclear disease	97	#Pacific radiation	214
#nuclear leakage	64	#radiation	482
#nuclear leaking	185	#radioactive waste	1500
#nuclear pollute	59	#radioactive water	568
#nuclear seafood	5	#sellafield nuclear	279

TABLE II

Keyword	Amount	Keyword	Amount
#fukushimanclearplant	104	#nuclearwater	115
#fukushimawater	182	#radioactivewaste	999
#nuclearwaste	999	#radioactivewater	416

B. Data Preprocessing

Tweets and IG posts are platforms for expressing users' perspectives through brief messages containing hashtags, emojis, text, and special symbols. Consequently, proper preprocessing of tweet data is essential. This work employs several preprocessing techniques, including pattern removal, stop word removal, URL removal, tokenization, and vectorization:

- **Pattern removal:** Symbols and special characters such as @, #, %, and others should be eliminated from the dataset which are not necessary for sentiment analysis and topic modeling.
- **Stop words removal:** Stop words should be eliminated from the dataset such as various forms of the verb 'be', personal pronouns, 'and', 'or', and similar terms, as they carry no meaningful information for sentiment analysis.
- **URL and email address remove:** If the URLs and email addresses in the dataset do not contribute meaningful information and can negatively impact classification accuracy, they should be eliminated.
- **Tokenization:** Splitting sentences into individual words or tokens. For example, the sentence Stop Japan's reckless nuclear contamination in our oceans' will be tokenized as follows. 'Stop', 'Japans', 'reckless', 'nuclear', 'contamination', 'in', 'our', 'oceans'.

C. Experimental Settings

The purpose of this article is to clarify public sentiment regarding Japan's discharge of nuclear wastewater and to provide stakeholders with a basis for decision-making. After data preprocessing, this article applies the LDA topic modeling algorithm to analyze the dataset's topics. Topic modeling, a probabilistic generative model for analyzing text data, is designed to uncover the latent topic structure in large-scale text and classify documents into these topics [29]. LDA identifies a set of topics in text data, where each topic consists of a collection of related words. Subsequently, the BERT algorithm is employed to perform sentiment analysis on each identified topic. BERT leverages a sophisticated input representation, including token embeddings, segment

embeddings, and position embeddings, to effectively handle both single-sentence and sentence-pair tasks [30], [31].

The sentiment analysis results will be visualized in a bar graph and further validated using the VADER data dictionary. The VADER, introduced by C.J. Hutto and Eric Gilbert in 2014, is a rule- and dictionary-based natural language processing tool particularly suited for sentiment analysis of informal texts, such as social media content and news articles [32]. VADER employs a sentiment dictionary where each word or phrase is assigned a sentiment score, reflecting its positive or negative intensity. It calculates an overall sentiment score by analyzing text components such as words, phrases, modifiers, and punctuation [33].

III. RESULTS AND DISCUSSIONS

This dataset contains 9,794 tweets and 2,815 Instagram posts. After data preprocessing, 8375 posts remain. A total of 197,631 words were identified from these posts. The following figure is the word cloud of the posts and the ten most frequently appearing words, as shown in Fig. 2. The larger the font in the word cloud, the higher the frequency of its occurrence. As shown in Fig. 3, the most frequently appearing words are "nuclear, waste, water, radioactive, power, Fukushima, plant, Japan, energy, radiation". The term "nuclear" appears as many as 3,700 times, "waste" around 2,300 times, and "water" more than 2,000 times.



Fig. 2 Word Cloud of the Dataset

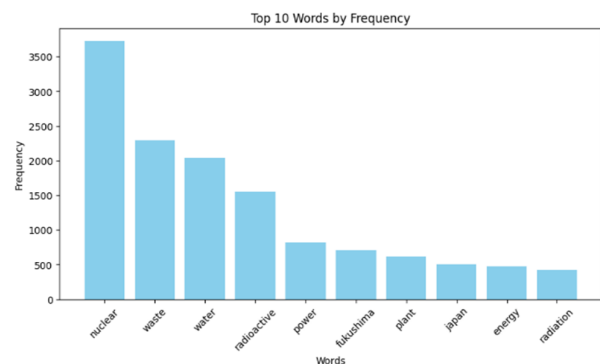


Fig. 3 Word Cloud of the Dataset

In social media data, there is a lot of noise and irrelevant information. LDA can help identify the main topics, thereby filtering out irrelevant content and reducing noise interference in sentiment analysis. Using LDA techniques, three topics were identified, as shown in Fig. 4. To mitigate noise

interference, the top 10 most relevant terms for each topic were selected to filter the dataset.

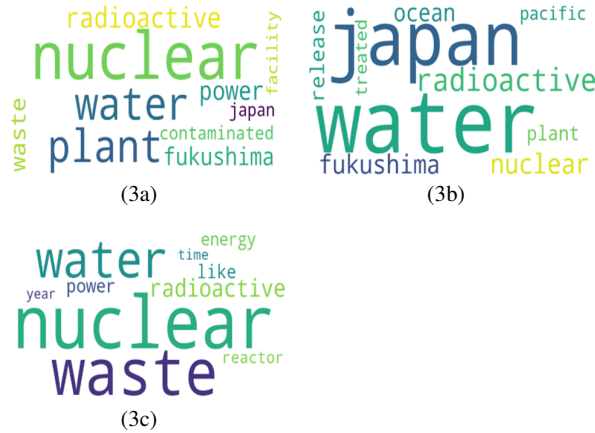


Fig. 4 Main Topics. (3a), (3b), (3c) are respectively Topic 1, 2, 3

Based on these terms as keywords, sentiment analysis was conducted using the BERT-base-multilingual-uncased-sentiment model on comments within the dataset. These topic-related terms are shown in Table III.

TABLE III
TOPIC DETAILS

Topic number	Theme	Topic words
1	Radioactive Waste Management in Fukushima	radioactive, nuclear, water, power, Japan, facility, waste, contaminated, Fukushima
2	Ocean Discharge of Treated Radioactive Water	ocean, Pacific, Japan, release, treated, radioactive, water, plant, Fukushima, nuclear
3	Long-term Impact of Nuclear Energy and Waste	water, energy, time, like, power, year, radioactive, nuclear, waste, reactor

The analysis results include 'very negative', 'negative', 'neutral', 'positive', and 'very positive'. The count of sentiment classification results for each topic is presented in Fig. 5 to 7. (The labels on the X-axis, from left to right, are "very positive," "positive," "neutral," "negative," and "very negative." The Y-axis represents the number of posts.)

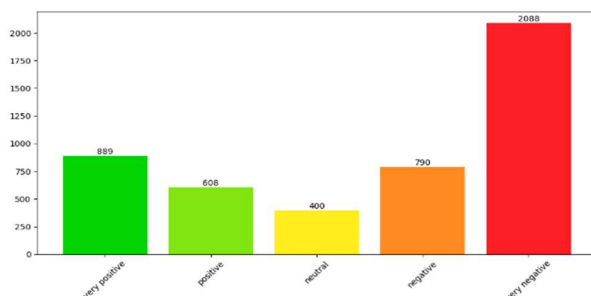


Fig. 5 Sentiment Analysis of Topic 1

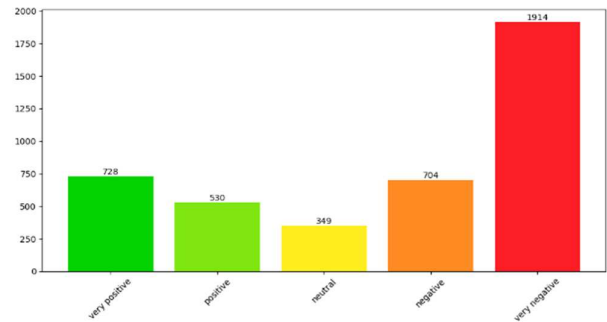


Fig. 6 Sentiment Analysis of Topic 2

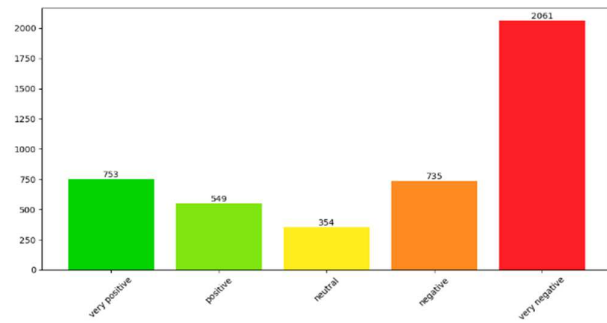


Fig. 7 Sentiment Analysis of Topic 3

The VADER results are used to further verify the BERT results shown in Table IV.

TABLE IV
SENTIMENT ANALYSIS BY VADER

Topic number	Positive number	Neutral number	Negative number
1	1404	775	2596
2	1172	713	2340
3	1205	708	2539

It can be observed that, for each topic, negative comments outnumber positive ones, indicating that netizens' overall attitude towards Japan's nuclear wastewater discharge is negative. In addition to analyzing topic distributions and overall sentiment patterns, this study performed a sentence-level analysis to better capture the nuances of public sentiment expressed in posts. Table V provides examples of sentence-level sentiment analysis results, illustrating how individual sentences are categorized across different sentiment classes. This approach ensures that the context and tone of each post are preserved, offering a more reliable sentiment classification.

TABLE V
SENTENCE-LEVEL SENTIMENT ANALYSIS RESULTS

Sentiment	Posts
negative	Over 1 million tons of treated wastewater will be released by Japan over the next thirty years.
positive	The National Report describes the legislative and regulatory measures taken by the UAE to meet its obligations under the Joint Convention. FANR issued and drafted six regulations containing requirement on managing radioactive waste.
negative	Japan began releasing treated nuclear water this Thursday. it will reach Vancouver, eventually
neutral	Tensions over the Fukushima water release.

IV. CONCLUSION

This paper investigates public opinion surrounding Japan's controversial plan to dispose of the decision on treated nuclear wastewater from the Fukushima Daiichi Nuclear Power Plant, which is being discharged into the Pacific Ocean. The analysis focuses on social media data from Twitter and Instagram, utilizing these platforms as primary channels for understanding global public views on this contentious issue. Using BERT and VADER to analyze the sentiment of 12,195 posts reveals that public sentiment toward this topic is predominantly negative.

This paper identifies several limitations that offer opportunities for future research. Firstly, the dataset used in this study was not sufficiently large to fully capture the complexity of the subject matter, thereby limiting the robustness and generalizability of our findings. Expanding the dataset in future studies to include a larger, more diverse sample could address this limitation and provide a more in-depth understanding of the topic. Furthermore, the absence of geographical data represents another limitation, as opinions may vary significantly across regions. Incorporating such data in future studies could facilitate the examination of regional differences and their impact on the results.

Additionally, this study relied on data from only two social media platforms, which limits the representativeness of the findings. Future studies should broaden the scope to analyze data from a wider range of platforms to better capture the diversity of online social interactions and opinions.

Finally, sarcasm detection in short social media posts remains a challenging aspect of this research. Sarcasm often depends on context, tone, and non-verbal cues, which are frequently absent or ambiguous in short text snippets, thereby affecting the reliability of sentiment analysis. Future research could address this limitation by employing advanced methods such as incorporating contextual data, multimodal analysis (e.g., combining text with images or emojis), or more sophisticated NLP techniques. Exploring tone, sentiment progression, or user interaction history can also mitigate the challenges posed by the lack of non-verbal cues and broader context. By addressing these limitations, future studies can refine sentiment analysis methodologies and yield more reliable, nuanced interpretations of sarcasm and sentiment in online discourse.

ACKNOWLEDGMENT

This work was sponsored by the Faculty of Computer Science and Information Technology, Universiti Putra Malaysia. The corresponding authors can be reached via email at Fatimah Sidi and Iskandar Ishak {fatimah; iskandar_i}@upm.edu.my.

REFERENCES

- [1] IAEA, "Fukushima Nuclear Accident Update Log," Apr. 28, 2011. [Online]. Available: <https://www.iaea.org/newscenter/news/fukushima-nuclear-accident-update-log-43>.
- [2] BBC News, "Fukushima nuclear disaster: Japan to release treated water in 48 hours." Accessed: Dec. 16, 2024. [Online]. Available: <https://www.bbc.com/news/world-asia-66578158>.
- [3] The Asahi Shimbun, "12.1 trillion yen spent so far on Fukushima nuclear disaster." Accessed: Dec. 16, 2024. [Online]. Available: <https://www.asahi.com/ajw/articles/14762193>.
- [4] IAEA, "IAEA Comprehensive Report On The Safety Review Of The Alps-Treated Water At The Fukushima Daiichi Nuclear Power Station." Accessed: Dec. 16, 2024. [Online]. Available: https://www.iaea.org/sites/default/files/iaea_comprehensive_alps_report.
- [5] TEPCO, "The Methods and Results of Treated Water Discharge and Monitoring Based on Scientific Evidence." [Online]. Available: <https://www.tepco.co.jp/en/decommission/progress/treated-water-lan/index-e.html>.
- [6] BBC News, "Greta Thunberg: Who is the climate activist and what has she achieved?" Accessed: Dec. 16, 2024. [Online]. Available: <https://www.bbc.com/news/world-europe-49918719>.
- [7] Z. Qu, J. Wang, and M. Zhang, "Mining and analysis of public sentiment during disaster events: The extreme rainstorm disaster in megacities of China in 2021," *Heliyon*, vol. 9, no. 7, p. e18272, Jul. 2023. doi: 10.1016/j.heliyon.2023.e18272.
- [8] P. SV, S. Sourav, and D. Kasilingam, "Critique: Sentiment-topic dynamic collaborative analysis-based public opinion mapping in aviation disaster management: A case study of the MU5735 air crash," *Int. J. Disaster Risk Reduct.*, vol. 108, p. 104546, Jun. 2024. doi: 10.1016/j.ijdr.2024.104546.
- [9] M. N. P. Ma'ady, A. F. A. Rahim, T. S. N. Syahda, A. F. Rizqi, and M. C. A. Ratna, "Malaysia Citizen Sentiment on Government Response Towards Covid-19 Disaster Management: Using LDA-based Topic Visualization on Twitter," *Procedia Comput. Sci.*, vol. 234, pp. 561–569, 2024. doi: 10.1016/j.procs.2024.03.040.
- [10] E. Rosenberg et al., "Sentiment analysis on Twitter data towards climate action," *Results Eng.*, vol. 19, p. 101287, Sep. 2023. doi:10.1016/j.rineng.2023.101287.
- [11] R. Bodhi, Y. Joshi, and A. Singh, "How does social media use enhance employee's well-being and advocacy behaviour? Findings from PLS-SEM and fsQCA," *Acta Psychol.*, vol. 251, p. 104586, Nov. 2024. doi:10.1016/j.actpsy.2024.104586.
- [12] K. Kuralová, K. Zychová, L. Kvasničková Stanislavská, L. Pilařová, and L. Pilař, "Work-life balance Twitter insights: A social media analysis before and after COVID-19 pandemic," *Heliyon*, vol. 10, no. 13, p. e33388, Jul. 2024. doi: 10.1016/j.heliyon.2024.e33388.
- [13] A. Khan, N. Boudjellal, H. Zhang, A. Ahmad, and M. Khan, "From Social Media to Ballot Box: Leveraging Location-Aware Sentiment Analysis for Election Predictions," *Comput. Mater. Contin.*, vol. 77, no. 3, pp. 3037–3055, 2023. doi: 10.32604/cmc.2023.044403.
- [14] P. Chauhan, N. Sharma, and G. Sikka, "The emergence of social media data and sentiment analysis in election prediction," *J. Ambient Intell. Humaniz. Comput.*, vol. 12, no. 2, pp. 2601–2627, Aug. 2020. doi:10.1007/s12652-020-02423-y.
- [15] T. Bhowmik, N. Eluru, S. Hasan, A. Culotta, and K. C. Roy, "Predicting hurricane evacuation behavior synthesizing data from travel surveys and social media," *Transp. Res. Part C Emerg. Technol.*, vol. 165, p. 104753, Aug. 2024. doi: 10.1016/j.trc.2024.104753.
- [16] L. Charmaraman, H. Hojman, J. Q. Auqui, and Z. Bilyalova, "Understanding Adolescent Self-esteem and Self-image Through Social Media Behaviors," *Pediatr. Clin. North Am.*, vol. 72, no. 2, pp. 189–201, Apr. 2025. doi: 10.1016/j.pcl.2024.07.034.
- [17] Y. Badulescu, F. Cañas, and N. Cheikhrouhou, "Judgmental adjustment of demand forecasting models using social media data and sentiment analysis within industry 5.0 ecosystems," *Int. J. Inf. Manag. Data Insights*, vol. 4, no. 2, p. 100272, Nov. 2024. doi:10.1016/j.ijime.2024.100272.
- [18] M. S. Md Suhaimin et al., "Social media sentiment analysis and opinion mining in public security: Taxonomy, trend analysis, issues and future directions," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 35, no. 9, p. 101776, Oct. 2023. doi: 10.1016/j.jksuci.2023.101776.
- [19] M. T. Almuqati, F. Sidi, S. N. M. Rum, M. Zolkepli, and I. Ishak, "Challenges in Supervised and Unsupervised Learning: A Comprehensive Overview," *Int. J. Adv. Sci. Eng. Inf. Technol.*, vol. 14, no. 4, pp. 1449–1455, Aug. 2024. doi: 10.18517/ijaseit.14.4.20191.
- [20] W. Liu, X. Zhao, M. M. Zhan, and S. Hernandez, "Streaming disasters on TikTok: Examining social mediated crisis communication, public engagement, and emotional responses during the 2023 Maui wildfire," *Public Relat. Rev.*, vol. 50, no. 5, p. 102512, Dec. 2024. doi:10.1016/j.pubrev.2024.102512.
- [21] H. Cho, S. Kim, J. Lee, and J.-S. Lee, "Data-driven integration of multiple sentiment dictionaries for lexicon-based sentiment classification of product reviews," *Knowl.-Based Syst.*, vol. 71, pp. 61–71, Nov. 2014. doi: 10.1016/j.knsys.2014.06.001.
- [22] A. V. D. Sano et al., "Proposing a visualized comparative review analysis model on tourism domain using Naïve Bayes

- classifier," *Procedia Comput. Sci.*, vol. 227, pp. 482–489, 2023. doi:10.1016/j.procs.2023.10.549.
- [23] D. Wang and Y. Zhao, "Using News to Predict Investor Sentiment: Based on SVM Model," *Procedia Comput. Sci.*, vol. 174, pp. 191–199, 2020. doi: 10.1016/j.procs.2020.06.074.
- [24] S. Aslan, "A deep learning-based sentiment analysis approach (MF-CNN-BILSTM) and topic modeling of tweets related to the Ukraine–Russia conflict," *Appl. Soft Comput.*, vol. 143, p. 110404, Aug. 2023. doi: 10.1016/j.asoc.2023.110404.
- [25] M. Ashayeri and N. Abbasabadi, "Unraveling energy justice in NYC urban buildings through social media sentiment analysis and transformer deep learning," *Energy Build.*, vol. 306, p. 113914, Mar. 2024. doi: 10.1016/j.enbuild.2024.113914.
- [26] Z. Li and Z. Zou, "Punctuation and lexicon aid representation: A hybrid model for short text sentiment analysis on social media platform," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 36, no. 3, p. 102010, Mar. 2024. doi: 10.1016/j.jksuci.2024.102010.
- [27] K. Jia, F. Meng, J. Liang, and P. Gong, "Text sentiment analysis based on BERT-CBLBGA," *Comput. Electr. Eng.*, vol. 112, p. 109019, Dec. 2023. doi: 10.1016/j.compeleceng.2023.109019.
- [28] S. E. Uthirapathy and D. Sandanam, "Topic Modelling and Opinion Analysis On Climate Change Twitter Data Using LDA And BERT Model," *Procedia Comput. Sci.*, vol. 218, pp. 908–917, 2023. doi:10.1016/j.procs.2023.01.071.
- [29] M. Venugopalan and D. Gupta, "An enhanced guided LDA model augmented with BERT based semantic strength for aspect term extraction in sentiment analysis," *Knowl.-Based Syst.*, vol. 246, p. 108668, Jun. 2022. doi: 10.1016/j.knosys.2022.108668.
- [30] S. Shashank and R. K. Behera, "Factors influencing recommendations for women's clothing satisfaction: A latent dirichlet allocation approach using online reviews," *J. Retail. Consum. Serv.*, vol. 81, p. 104011, Nov. 2024. doi: 10.1016/j.jretconser.2024.104011.
- [31] J. Liu, R. Long, H. Chen, M. Wu, W. Ma, and Q. Li, "Topic-sentiment analysis of citizen environmental complaints in China: Using a Stacking-BERT model," *J. Environ. Manage.*, vol. 371, p. 123112, Dec. 2024. doi: 10.1016/j.jenvman.2024.123112.
- [32] C. Hutto and E. Gilbert, "VADER: A Parsimonious Rule-Based Model for Sentiment Analysis of Social Media Text," *Proc. Int. AAAI Conf. Web Soc. Media*, vol. 8, no. 1, pp. 216–225, May 2014. doi:10.1609/icwsm.v8i1.14550.
- [33] V. Vought, R. Vought, A. S. Lee, I. Zhou, M. Gameni, and S. A. Greenstein, "Application of sentiment and word frequency analysis of physician review sites to evaluate refractive surgery care," *Adv. Ophthalmol. Pract. Res.*, vol. 4, no. 2, pp. 78–83, May 2024. doi:10.1016/j.aopr.2024.03.002.