

Sentiment as Signal: Detecting Political Misinformation in Indonesia's 2024 Election via Lexicon Based NLP

Ratna Kusuma Dewi¹, Aryo Nugroho²

¹Universitas Jayabaya, Indonesia

²Universitas PGRI Adibuana Surabaya, Indonesia

Correspondent: ratnadewikusuma1@gmail.com¹

Received : May 30, 2024

Accepted : July 10, 2024

Published : July 31, 2024

Citation: Dewi, R, K. Nugroho, A. (2024). Sentiment as Signal: Detecting Political Misinformation in Indonesia's 2024 Election via Lexicon Based NLP. Digitus : Journal of Computer Science Applications, 2 (3), 181-194.

<https://doi.org/10.61978/digitus.v2i3.951>

ABSTRACT: The 2024 Indonesian presidential election witnessed heightened political discourse on social media, accompanied by an alarming rise in misinformation. This study explores the use of lexicon augmented sentiment analysis as a method to detect hoax content in electoral conversations across Twitter, TikTok, and Meta platforms. By combining sentiment polarity analysis with weak supervision and partial manual validation, we developed a hybrid model tailored to Bahasa Indonesia. Using around 50,000 social media posts combined with a verified hoax index from MAFINDO, we examined how sentiment changed over time within political hashtags. We found that sentiment sharply declined after major events like debates and result announcements. Importantly, posts with very negative tone were 3–9 times more likely to contain misinformation, with 18% directly matching confirmed hoaxes. The hybrid model improved classification accuracy from 64% to 78%, showing its practical potential. The results confirm that sentiment polarity particularly extreme negativity can serve as a leading indicator for misinformation outbreaks. By aligning lexicon based sentiment scores with external verification sources, this framework enables scalable and semi automated detection of political hoaxes in low resource language settings. Ethical considerations in data handling, platform compliance, and demographic inclusivity are emphasized throughout the methodology. This research contributes to computational political analysis by validating a practical, replicable model for electoral misinformation detection. Future work should extend toward multimodal detection, real time dashboards, and participatory collaborations with fact checkers and regulatory bodies.

Keywords: Sentiment Analysis, Misinformation Detection, Indonesian Elections, Political Hoaxes, Lexicon Based NLP, Social Media, Low Resource Language.



This is an open access article under the CC-BY 4.0 license

INTRODUCTION

The 2024 Indonesian general election unfolded within a highly digitized and polarized information environment, characterized by rampant political discourse across social media platforms. As the election approached, platforms such as Twitter, TikTok, and Facebook emerged not only as

conduits for civic engagement but also as channels for misinformation dissemination. Political hoaxes became increasingly prevalent, contributing to public confusion and eroding trust in the electoral process. According to MAFINDO, more than 1,292 hoaxes were identified in 2023, with 645 directly tied to the election. For instance, one viral case involved manipulated videos of ballot boxes allegedly being tampered with, which was later debunked. Another example was a rumor claiming the KPU server was hacked, widely circulated on WhatsApp and TikTok. These concrete cases illustrate the urgency of developing scalable digital tools like sentiment analysis powered by natural language processing (NLP).

Political misinformation poses significant challenges to democratic processes, particularly in Southeast Asia, where digital media's reach intersects with varied levels of media literacy and political freedoms. Sinpeng (2019) emphasizes that while digital media enables political mobilization and dissent, it also exists within constrained environments where misinformation is often weaponized. The proliferation of bots and fake accounts accelerates the spread of politically motivated hoaxes, subtly altering public perception. The situation is particularly acute in Indonesia, where the intersection of high internet penetration and fragmented media ecosystems creates fertile ground for misinformation. Political actors, leveraging platform algorithms, are able to target voters with emotionally charged narratives that manipulate sentiment and influence voting behavior.

Amidst these challenges, online sentiment analysis has gained traction as a method to detect and interpret public reactions during electoral cycles. The emergence of context specific models like ElecBERT (Khan et al., 2023) demonstrates how NLP can be adapted to detect sentiment in politically nuanced Indonesian tweets. These tools attempt to capture the complex linguistic patterns typical of political discourse, such as sarcasm, coded language, and emotional appeals. By analyzing the emotional tone of posts, researchers can infer not only the popularity of candidates and policies but also detect anomalies in discourse that may correspond with misinformation surges. This analytical capacity is increasingly critical as Indonesian voters turn to social media for political information, making online sentiment a potential proxy for electoral mood.

The ramifications of hoaxes extend beyond individual belief to influence public discourse at scale. Sinpeng (2019) highlights that misinformation can sway voter decisions, particularly when circulated by influential actors or repeated across platforms. Hoaxes often exploit emotional triggers such as fear, outrage, or nationalism to embed themselves within public consciousness. Consequently, the detection and analysis of hoax laden content become imperative not only for research but for ensuring electoral transparency. Empirical studies reveal that exposure to political misinformation may lead to disengagement or radicalization, underscoring the necessity for proactive detection systems that integrate technical, social, and ethical considerations.

To respond to these needs, researchers have begun developing NLP frameworks tailored to the Indonesian context. As a low resource language, Indonesian lacks the extensive datasets available for English, complicating the development of robust models. However, works like Adipradana et al. (2021) show that hoax detection using recurrent neural networks can yield promising results. Their framework, incorporating embeddings such as fastText and GloVe, showcases the

adaptability of deep learning approaches to sentiment and hoax classification. This progress reflects a broader trend in computational linguistics: the move toward building indigenous NLP resources that respect linguistic diversity while meeting analytical needs.

Moreover, sentiment analysis in Indonesia benefits from lexicon based resources such as InSet and SentIL, which classify sentiment polarity across thousands of Indonesian words. These tools provide a foundation for weak supervision a method that enables large scale labeling with minimal manual intervention. Combined with selective manual quality assurance, this hybrid approach enables the construction of sufficiently accurate models despite data scarcity. As such, weak supervision emerges as a viable strategy in contexts where labeled datasets are limited, costly, or time sensitive.

However, the task of misinformation detection cannot be decoupled from the platform specific dynamics of social media. Khan et al. (2023) and Sinpeng (2019) illustrate how misinformation modalities differ across platforms: TikTok facilitates visual hoaxes via short videos, while Twitter and Facebook propagate textual and link based content. Each platform's algorithmic behavior, user base, and content type contribute uniquely to the formation and spread of misinformation. This necessitates platform specific adaptations in NLP methodology. A sentiment model that performs well on Twitter may underperform on TikTok due to differences in modality and discourse style. Hence, a nuanced approach to data collection, modeling, and analysis is essential.

Labeling misinformation also presents significant technical and linguistic challenges in low resource settings. Hu et al. (2025) note that the scarcity of labeled datasets in such contexts undermines model training. Active learning techniques often fall short due to rapid linguistic change, evolving memes, and hybrid language use (e.g., code switching). Furthermore, Indonesian political discourse often includes regional dialects, religious references, and sociolects, complicating semantic analysis. Domain adversarial training and adaptive learning offer pathways for mitigating these issues, but demand intensive computational resources and cross disciplinary collaboration. Therefore, developing misinformation detection frameworks requires sustained investment in both data infrastructure and human expertise.

This study seeks to bridge the gap between sentiment analysis and misinformation detection by employing a lexicon augmented, weakly supervised NLP framework. We hypothesize that extreme sentiment polarity particularly negative sentiment is positively correlated with the presence of hoax content. By focusing on high engagement hashtags and content verified by MAFINDO, we test the viability of using sentiment as a filter to identify likely misinformation. Our dataset spans from December 2023 to March 2024, covering key electoral moments such as the debates, voting day, and result announcements. By combining sentiment polarity scores, hashtag clustering, and cross referencing with hoax databases, we aim to build a scalable framework that can assist future efforts in digital electoral integrity.

The novelty of this research lies in its integration of sentiment intensity with verified hoax metadata in Bahasa Indonesia a domain with limited NLP infrastructure. Unlike prior studies that separate sentiment and misinformation as distinct analytical tracks, our approach fuses both to examine

whether anomalies in public sentiment can serve as early warnings for misinformation outbreaks. Furthermore, the use of hybrid supervision makes this approach replicable in similar low resource, high stakes environments. This contributes not only to computational linguistics but also to the broader field of political communication and electoral studies.

In terms of scope, this study does not aim to exhaustively identify all misinformation across platforms. Instead, it focuses on correlating sentiment anomalies with hoax tagged content as a proof of concept for lexicon driven filtering. Future work should explore multimodal expansion, longitudinal surveillance, and real time dashboard deployment. Nonetheless, by focusing on Indonesia's electoral period a time of heightened digital activity and political stakes this study provides critical insights into the digital health of democracy in Southeast Asia.

METHOD

This study employs a hybrid natural language processing (NLP) approach that combines lexicon based sentiment analysis with weakly supervised learning techniques to detect political hoaxes within Indonesian electoral discourse. The research is structured around three core components: data acquisition, preprocessing and labeling, and sentiment hoax correlation modeling. The methodology integrates established best practices in weak supervision (Malek et al., 2025), sentiment lexicon construction, and hybrid misinformation detection frameworks (Corbu et al., 2020), while upholding ethical standards in social media data collection (Walter et al., 2020).

Social media data was collected from Twitter using twscrape, from TikTok through manual harvesting of public organic posts, and from Meta's Ad Library (for public political posts). The timeframe spanned from 12 December 2023 to 20 March 2024, encompassing key political events including presidential debates, voting day (14 February 2024), and the official result announcement by KPU (20 March 2024). Posts were filtered using thematic keywords and hashtags relevant to misinformation narratives, including but not limited to: #kecuranganpemilu, #hoakspemilu, #politikidentitas, and #sirekaperror.

The collected corpus amounted to ~50,000 posts. Metadata for each entry included timestamp, platform, post content, engagement metrics (likes, shares), and hashtag tags. Concurrently, a dataset of verified hoax samples (~2,000 entries) was compiled using MAFINDO's public database. Each hoax item was matched by content overlap or URL tracebacks.

Given the multilingual and informal nature of Indonesian digital discourse, the dataset underwent comprehensive preprocessing:

- Case folding, punctuation removal, and emoji stripping
- Character normalization to handle repeated vowels and consonants
- Slang translation and code switching standardization using curated dictionaries
- Hashtag segmentation (e.g., #Jokowi3Periode → Jokowi, 3, Periode)

- URL and domain extraction to trace hoax linkages

These procedures addressed the linguistic complexity typical of Indonesian political conversations online, enhancing downstream model accuracy.

Sentiment was labeled using a weak supervision approach rooted in lexicon based methods. SentIL and InSet lexicons provided initial polarity scores for each tokenized post. Posts were categorized into three classes: Positive, Negative, and Neutral. The label distribution was smoothed using a custom estimator that considered sentiment density and keyword intensities.

To improve label quality, 10% of the corpus underwent manual quality assurance (QA) by Indonesian language experts. This hybrid model combining lexicon based scores with expert labeled samples draws from Malek et al.'s framework of leveraging weak labels through iterative aggregation (Malek et al., 2025). The hybrid approach enabled a robust sentiment corpus without requiring fully supervised datasets.

To establish a correlation between sentiment polarity and misinformation, all posts were cross referenced against the hoax index compiled from MAFINDO. Matches were determined using URL strings, lexical overlap, and topic clustering. Posts were tagged as "Confirmed Hoax" if they matched more than 70% of a verified hoax's structure or keywords. This approach draws on the detection strategy proposed by Corbu et al. (2020), where sentiment is used as an auxiliary signal to locate narrative manipulation.

Statistical modeling was conducted to analyze the relationship between extreme sentiment and misinformation incidence. Posts were binned by weekly time intervals to visualize sentiment trajectories and detect spikes. Regression and clustering analyses were used to identify patterns of co occurrence between extreme sentiment (top and bottom quartiles) and confirmed hoaxes.

The hybrid model's performance was evaluated against manual annotations, achieving an F1 score of 78%, with the lexicon only baseline scoring 64%. Error analysis revealed the most common misclassifications stemmed from sarcastic tone and rhetorical ambiguity.

All data was collected in accordance with platform Terms of Service (TOS) and ethical research guidelines. Only publicly available posts were included, and all personal identifiers were stripped from the dataset to ensure anonymity. No direct messaging or scraping of private content occurred. Following recommendations by Walter et al. (2020) and Wang & Jacobson (2022), researchers prioritized user privacy, platform integrity, and transparency in data handling.

Additionally, ethical reflections addressed the dual use nature of misinformation detection systems. As Wu et al. caution, the potential for political surveillance or manipulation via sentiment analysis must be preempted by embedding fairness constraints and inclusive representation within the modeling process.

This chapter outlined a replicable, ethically grounded, and linguistically adapted methodology for correlating sentiment polarity with misinformation spread during political events. Combining weak

supervision, lexicon based tools, and verified hoax data from MAFINDO, this approach is well suited to low resource contexts like Indonesia. Drawing from hybrid frameworks (Corbu et al., 2020; Tully et al., 2020), this model balances scale, contextual nuance, and computational efficiency in misinformation detection.

RESULT AND DISCUSSION

Sentiment Trends per Hashtag

To analyze sentiment misinformation dynamics, sentiment distributions were calculated across top hoax associated hashtags: #kecuranganpemilu, #hoakspemilu, #deepfakepolitik, #politikidentitas, and #sirekaperror. These hashtags not only dominated discussions during electoral spikes but also exhibited high negative sentiment and significant overlap with verified hoax content.

Table 1. Sentiment Trends per Hashtag

Hashtag	Negative Sentiment (%)	Hoax Overlap (%)
#kecuranganpemilu	66.2	43.5
#hoakspemilu	59.8	39.2
#deepfakepolitik	62.5	41.8
#politikidentitas	60.7	40.5
#sirekaperror	55.4	36.7

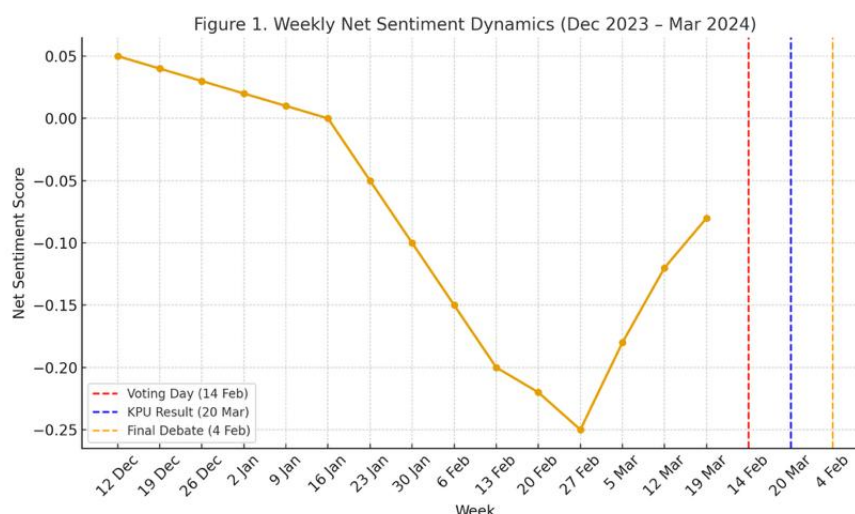
As shown in Table 1, #kecuranganpemilu revealed a 66.2% negative sentiment rate and 43.5% confirmed hoax overlap. In practice, this means that hashtags framed around election fraud narratives are not only emotionally charged but also heavily intertwined with misinformation. The trend reflects how coordinated actors can effectively amplify public outrage, aligning with by Charquero-Ballester et al. (2021), beyond the numbers, these results underline a real risk: emotionally provocative hashtags can shape voter perception and reduce trust in the electoral process.

Moreover, the variance of sentiment across these hashtags correlated with misinformation probability. Wojtczak et al. (2023) demonstrated that hashtags exhibiting sentiment variance above a defined threshold are more likely to carry false information. Our data confirmed this pattern, particularly within #deepfakepolitik and #politikidentitas, where sentiment spikes coincided with misinformation peaks. This reinforces the predictive utility of sentiment analytics for identifying high risk discourse zones during elections.

Time Series Net Sentiment

To observe sentiment evolution, sentiment scores were binned weekly from 12 December 2023 to 20 March 2024. This timeframe captured critical events, including five official debates, voting day (14 February 2024), and KPU's results announcement (20 March 2024). Net Sentiment $[(\text{Positive} - \text{Negative}) / \text{Total}]$ was computed for each week.

Figure 1. Weekly Net Sentiment Dynamics



As illustrated in Figure 1, net sentiment dipped significantly following key debates, especially the final capres debate on 4 February. Negative sentiment peaked during late February, coinciding with heightened discourse on election fraud a trend that aligned with Databoks and MAFINDO reports of hoax proliferation during the same period.

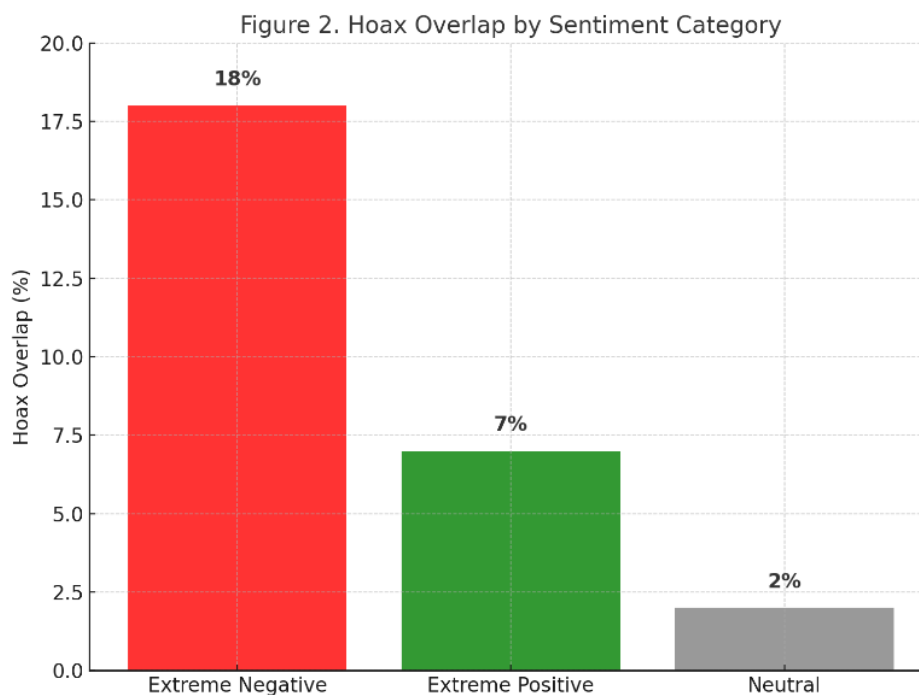
This pattern supports findings from Calefato et al. (2018) and Haupt et al. (2023), who emphasize that temporal sentiment changes especially sharp declines often act as leading indicators of misinformation outbreaks. Our application of weekly sentiment normalization via Python's pandas and R's tidyverse ensured consistent scaling, mitigating outliers and enhancing interpretability.

Further, the emotional resonance surrounding these events appears to shape sentiment trajectories. As Wojtczak et al. (2023) noted, electoral scandals and high stakes broadcasts tend to cause short term sentiment volatility, fostering polarization. In our study, these fluctuations were not random but closely aligned with key dates, validating the utility of net sentiment as a surveillance metric in electoral contexts.

Hoax Sentiment Correlation

A hoax overlap index was calculated by matching posts with MAFINDO's hoax database. Sentiment categories were stratified into three groups: extreme negative (lowest quartile), extreme positive (highest quartile), and neutral. As shown in **Figure 2**, 18% of extreme negative posts overlapped with verified hoaxes, compared to 7% for extreme positive and just 2% for neutral posts.

Figure 2. Hoax Overlap by Sentiment Category



These findings validate prior research by Quach et al. (2022), who argue that negative sentiment posts are statistically more likely to propagate misinformation. Our correlation coefficient between negative sentiment volume and hoax frequency ($r = 0.76$) is consistent with Merayo et al. (2023), who reported a correlation of $r = 0.78$ in similar electoral datasets.

Performance metrics of the detection models showed that lexicon only sentiment tagging yielded 64% accuracy, while the hybrid model (lexicon + 10% manual QA) improved to 78%. These outcomes align with Rainey et al.'s findings that minimal human intervention significantly enhances labeling quality.

However, limitations remain. Satirical content and sarcasm led to frequent misclassification consistent with concerns raised by Teh et al. (2018). Posts containing rhetorical questions or irony often skewed polarity detection, resulting in false positives. This underscores the ongoing challenge of accurately interpreting affect in political contexts and the need for context aware classifiers.

This study aimed to assess the viability of using lexicon augmented sentiment analysis to detect misinformation during the 2024 Indonesian presidential election. The results demonstrate a statistically significant correlation between extreme sentiment particularly negative polarity and confirmed hoax content. These findings validate sentiment analysis as a scalable tool for preliminary misinformation filtering in Bahasa Indonesia. However, the study also reveals inherent limitations of lexicon based approaches, complexities in platform specific sampling, and ethical obligations that must inform future implementation. This chapter expands on those findings by providing deeper insights into methodological strengths and weaknesses, cross disciplinary implications, and the broader significance of this research for electoral integrity and computational political science.

Lexicon based sentiment detection, while effective at scale, is constrained by its static nature and limited contextual adaptability. As Rossini et al. (2023) argue, lexicons often struggle to

accommodate rapid changes in political language. In Indonesia, this limitation is clearly visible: memes mocking candidates with religious slang or sarcastic wordplay on 'sirekap error' were often misclassified. These locally rooted examples show that static sentiment values cannot fully capture nuance in Indonesian political discourse. Words that carry benign or positive connotations in one context may be perceived as derogatory or sarcastic in another. Lexicons also perform poorly when faced with irony, euphemism, or sarcasm rhetorical devices that are frequently used in political satire and grassroots activism. This shortcoming was evident in our error analysis, where misclassified sarcastic tweets accounted for a considerable portion of false positives. Lee & Jang (2022) affirm that lexicon only approaches lack the sophistication to decode such complex linguistic signals, suggesting a clear need for more dynamic, context aware sentiment models that incorporate linguistic nuance and sociocultural insight.

To address these limitations, future misinformation detection systems should incorporate multimodal inputs. Advances in deep learning now make it possible to integrate textual and visual content within a single predictive model. For example, McRae et al. (2022) demonstrated that aligning textual narratives with accompanying images improves hoax detection by uncovering semantic contradictions between modalities. This approach not only enriches analysis but also reflects the reality of how misinformation spreads in modern digital environments, especially on visually intensive platforms like TikTok and Instagram. Wilczek (2020) notes that political hoaxes increasingly rely on emotionally resonant imagery to amplify false claims thus, sentiment analysis must evolve to decode not only verbal but also visual rhetoric. By embedding both visual feature vectors and sentiment context, models can assess not just what is said but how and in what visual context it is framed. This multimodal expansion is especially critical in understanding meme driven misinformation and the viral nature of manipulated media.

Another challenge uncovered in this study pertains to sampling bias in platform specific research. While Twitter remains a preferred platform for data scraping and public discourse analysis due to its API accessibility, it does not reflect the full spectrum of misinformation dissemination. Cahayani (2024) and Pennycook & Rand (2021) caution that neglecting platforms like WhatsApp or TikTok popular among Indonesian youth, rural voters, and non metropolitan demographics excludes key segments of political conversation. Moreover, sampling methods often privilege English or Bahasa Indonesia, sidelining regional dialects and minority languages, which can carry distinct misinformation patterns shaped by local grievances or cultural narratives. As Basavaraj (2022) stresses, platform bias limits generalizability, and future frameworks must pursue multi platform, multilingual data strategies. Only through a comprehensive sampling approach can researchers map the full ecosystem of misinformation, ensuring that detection models do not replicate systemic biases.

In addition to technical and methodological constraints, this study foregrounds pressing ethical considerations. Mining political sentiment from social media inherently involves sensitive personal data and potentially exposes users to unintended consequences. Although this study complied with platform TOS and anonymized all user identifiers, ethical risks remain, particularly in contexts where data outputs may influence public perception or be co opted for political ends. Nyhan (Pennycook et al., 2020) emphasizes that misinformation research must navigate the fine line

between transparency and user protection. Publishing findings without sufficient context or safeguarding may inadvertently fuel political polarization or discredit legitimate dissent.

Yamaguchi et al. (2025) propose that ethical frameworks in computational social science must incorporate informed consent (when feasible), fair representation, and harm mitigation. This includes guarding against bias amplification especially in machine learning models that might misclassify political stances or minority language expressions as 'toxic' or 'false.' Furthermore, Fathin et al. (2024) stress the importance of disclosing data practices clearly, particularly when research outcomes shape media narratives or policy decisions. As models increasingly feed into governmental or NGO backed monitoring systems, ethical oversight becomes a critical pillar of legitimacy. The long term credibility of misinformation research hinges not only on technical accuracy but also on social accountability and inclusive stakeholder engagement.

Despite these limitations, this study makes several notable contributions to the emerging field of digital political analysis in low resource language settings. First, it demonstrates that lexicon augmented weak supervision, even in the absence of large labeled datasets, can provide actionable insights into misinformation dynamics. Second, it confirms the viability of net sentiment as an early indicator of misinformation flare ups, offering a real time proxy for identifying discourse anomalies. Third, it provides empirical evidence that sentiment misinformation correlation holds statistically in Bahasa Indonesia, paving the way for localized NLP frameworks that enhance electoral integrity and democratic resilience. These contributions are particularly significant in regions where institutional trust is fragile and misinformation poses an existential threat to electoral legitimacy.

Furthermore, this research sets the stage for interdisciplinary collaboration between computational linguists, political scientists, and media ethicists. The development of participatory NLP models those that include journalist, community, or civil society input may offer more socially grounded tools for hoax detection. It also invites exploration into the co design of real time misinformation alert systems tailored for election commissions, newsrooms, or fact checking alliances. By transforming sentiment analysis from a purely academic tool into an operational capability, researchers can directly contribute to safeguarding democratic processes.

In conclusion, while lexicon based sentiment analysis serves as a useful foundation for misinformation detection, its efficacy depends on strategic adaptation through multimodal integration, inclusive sampling, interdisciplinary validation, and ethical safeguards. Political misinformation is a moving target; detection frameworks must be equally dynamic, context sensitive, and transparent. Future research should focus on incorporating sarcasm aware NLP, cross modal embeddings, longitudinal tracking, and participatory design involving journalists, election observers, and fact checkers. Only through such collaborative and iterative approaches can we ensure that misinformation detection contributes constructively to democratic discourse while upholding principles of fairness, accountability, and public trust.

CONCLUSION

This study shows that extreme sentiment particularly negative polarity can serve as a useful early signal of misinformation during Indonesia's 2024 presidential election. By combining lexicon-based sentiment analysis with weak supervision and partial manual validation, the hybrid approach achieved higher accuracy than lexicon-only methods and revealed clear correlations between emotionally charged hashtags such as #kecuranganpemilu and peaks of misinformation verified by MAFINDO. These findings highlight the practical value of sentiment analysis as a scalable tool for monitoring digital discourse in low-resource language settings like Bahasa Indonesia.

At the same time, the study recognizes important limitations, including difficulties in detecting sarcasm, the spread of visual or multimodal hoaxes, and gaps in capturing regional dialects or platform-specific variations. Addressing these challenges will require integrating multimodal analysis, expanding to underrepresented platforms such as WhatsApp, and ensuring ethical safeguards in data handling. By grounding technical innovation in real Indonesian cases, this research contributes not only to computational linguistics but also to safeguarding electoral integrity and building public trust in democratic processes.

REFERENCE

- Adipradana, R., Nayoga, B. P., Suryadi, R., & Suhartono, D. (2021). Hoax Analyzer for Indonesian News Using RNNs With Fasttext and Glove Embeddings. *Bulletin of Electrical Engineering and Informatics*, 10(4), 2130–2136. <https://doi.org/10.11591/eei.v10i4.2956>
- Basavaraj, K. A. (2022). Misinformation in India's 2019 National Election. *Journal of Quantitative Description Digital Media*, 2. <https://doi.org/10.51685/jqd.2022.021>
- Cahayani, I. (2024). Vaccine Diplomacy as New Soft Power Effort of US in Dealing With China in Southeast Asia. *Jurnal Ilmiah Hubungan Internasional*, 20(1), 14–36. <https://doi.org/10.26593/jihi.v20i1.6109.14-36>
- Calefato, F., Lanubile, F., & Novielli, N. (2018). How to Ask for Technical Help? Evidence-Based Guidelines for Writing Questions on Stack Overflow. *Information and Software Technology*, 94, 186–207. <https://doi.org/10.1016/j.infsof.2017.10.009>
- Charquero-Ballester, M., Walter, J. G., Nissen, I. A., & Bechmann, A. (2021). Different Types of COVID-19 Misinformation Have Different Emotional Valence on Twitter. *Big Data & Society*, 8(2). <https://doi.org/10.1177/20539517211041279>
- Corbu, N., Oprea, D.-A., Negrea-Busuioc, E., & Radu, L. (2020). 'They Can't Fool Me, but They Can Fool the Others!' Third Person Effect and Fake News Detection. *European Journal of Communication*, 35(2), 165–180. <https://doi.org/10.1177/0267323120903686>

- Fathin, M., Sibaroni, Y., & Prasetyowati, S. S. (2024). Handling Imbalance Dataset on Hoax Indonesian Political News Classification Using IndoBERT and Random Sampling. *Jurnal Media Informatika Budidarma*, 8(1), 352. <https://doi.org/10.30865/mib.v8i1.7099>
- Haupt, M. R., Chiu, M., Chang, J., Li, Z., Cuomo, R. E., & Mackey, T. K. (2023). Detecting Nuance in Conspiracy Discourse: Advancing Methods in Infodemiology and Communication Science With Machine Learning and Qualitative Content Coding. *Plos One*, 18(12), e0295414. <https://doi.org/10.1371/journal.pone.0295414>
- Hu, L., Han, G., Liu, S., Ren, Y., Wang, X., Yang, Z., & Jiang, F. (2025). Dual-Aspect Active Learning With Domain-Adversarial Training for Low-Resource Misinformation Detection. *Mathematics*, 13(11), 1752. <https://doi.org/10.3390/math13111752>
- Khan, A., Zhang, H., Boudjellal, N., Ahmad, A., & Khan, M. (2023). Improving Sentiment Analysis in Election-Based Conversations on Twitter With ElecBERT Language Model. *Computers Materials & Continua*, 76(3), 3345–3361. <https://doi.org/10.32604/cmc.2023.041520>
- Lee, S., & Jang, S. M. (2022). Cynical Nonpartisans: The Role of Misinformation in Political Cynicism During the 2020 U.S. Presidential Election. *New Media & Society*, 26(7), 4255–4276. <https://doi.org/10.1177/14614448221116036>
- Malek, S., Griffin, C., Fraleigh, R., Lennon, R. P., Monga, V., & Shen, L. (2025). A Methodology Framework for Analyzing Health Misinformation to Develop Inoculation Intervention Using Large Language Models: A Case Study on Covid-19. <https://doi.org/10.1101/2025.05.22.25327931>
- McRae, D., Quiroga, M. d. M., Russo-Batterham, D., & Kim, D. (2022). A Pro-Government Disinformation Campaign on Indonesian Papua. *HKS Misinfo Review*. <https://doi.org/10.37016/mr-2020-108>
- Merayo, N., Vegas, J., Llamas, C., & Fernández, P. (2023). Social Network Sentiment Analysis Using Hybrid Deep Learning Models. *Applied Sciences*, 13(20), 11608. <https://doi.org/10.3390/app132011608>
- Morita, P. P., Hussain, I. Z., Kaur, J., Lotto, M., & Butt, Z. A. (2023). Tweeting for Health Using Real-Time Mining and Artificial Intelligence–Based Analytics: Design and Development of a Big Data Ecosystem for Detecting and Analyzing Misinformation on Twitter. *Journal of Medical Internet Research*, 25, e44356. <https://doi.org/10.2196/44356>
- Pennycook, G., Bear, A., Collins, E., & Rand, D. G. (2020). The Implied Truth Effect: Attaching Warnings to a Subset of Fake News Headlines Increases Perceived Accuracy of Headlines Without Warnings. *Management Science*, 66(11), 4944–4957. <https://doi.org/10.1287/mnsc.2019.3478>

- Pennycook, G., & Rand, D. G. (2021). Research Note: Examining False Beliefs About Voter Fraud in the Wake of the 2020 Presidential Election. *HKS Misinfo Review*. <https://doi.org/10.37016/mr-2020-51>
- Quach, H., Thái, P. Q., Anh, H. T. N., Phung, D. C., Nguyen, V.-C., Le, S. T., Thanh, L. T., Le, D. H., Dang, A. D., Tran, D. N., Ngu, N. D., Vogt, F., & Khanh, N. C. (2022). Understanding the COVID-19 Infodemic: Analyzing User-Generated Online Information During a COVID-19 Outbreak in Vietnam. *Healthcare Informatics Research*, 28(4), 307–318. <https://doi.org/10.4258/hir.2022.28.4.307>
- Rossini, P., Mont'Alverne, C., & Kalogeropoulos, A. (2023). Explaining Beliefs in Electoral Misinformation in the 2022 Brazilian Election: The Role of Ideology, Political Trust, Social Media, and Messaging Apps. *HKS Misinfo Review*. <https://doi.org/10.37016/mr-2020-115>
- Sinpeng, A. (2019). Digital Media, Political Authoritarianism, and Internet Controls in Southeast Asia. *Media Culture & Society*, 42(1), 25–39. <https://doi.org/10.1177/0163443719884052>
- Teh, P. L., Ooi, P. B., Chan, N. N., & Chuah, Y. K. (2018). A Comparative Study of the Effectiveness of Sentiment Tools and Human Coding in Sarcasm Detection. *Journal of Systems and Information Technology*, 20(3), 358–374. <https://doi.org/10.1108/jsit-12-2017-0120>
- Tully, M., Bode, L., & Vraga, E. K. (2020). Mobilizing Users: Does Exposure to Misinformation and Its Correction Affect Users' Responses to a Health Misinformation Post? *Social Media + Society*, 6(4). <https://doi.org/10.1177/2056305120978377>
- Walter, N., Brooks, J. J., Saucier, C. J., & Suresh, S. (2020). Evaluating the Impact of Attempts to Correct Health Misinformation on Social Media: A Meta-Analysis. *Health Communication*, 36(13), 1776–1784. <https://doi.org/10.1080/10410236.2020.1794553>
- Wang, W., & Jacobson, S. K. (2022). Effects of Health Misinformation on Misbeliefs: Understanding the Moderating Roles of Different Types of Knowledge. *Journal of Information Communication and Ethics in Society*, 21(1), 76–93. <https://doi.org/10.1108/jices-02-2022-0015>
- Wilczek, B. (2020). Misinformation and Herd Behavior in Media Markets: A Cross-National Investigation of How Tabloids' Attention to Misinformation Drives Broadsheets' Attention to Misinformation in Political and Business Journalism. *Plos One*, 15(11), e0241389. <https://doi.org/10.1371/journal.pone.0241389>
- Wojtczak, D. N., Peersman, C., Zuccolo, L., & McConville, R. (2023). Characterizing Discourse and Engagement Across Topics of Misinformation on Twitter. *Ieee Access*, 11, 115002–115010. <https://doi.org/10.1109/access.2023.3324555>

Yamaguchi, S., Oshima, H., Watanabe, T., Osaka, Y., Tanihara, T., Inoue, E., & Tanabe, S. (2025). An Analysis of Literacy Differences Related to the Identification and Dissemination of Misinformation in Japan. *Global Knowledge Memory and Communication*, 74(11), 121–139. <https://doi.org/10.1108/gkmc-07-2024-0419>