

Korespondensi Deep Learning Approaches for Hate Speech Detection in Resource Constrained Social Media Texts: A Comparative Evaluation

BalasBalas ke se...TeruskanHapusArsipSampahTandaiBerikutnya

[AIA] Submission Acknowledgement

Pengirim Yu Zhang

Pengirim office@bonviewpress.com

Penerima Aditya Hermawan

Tanggal 2025-08-25 16:23

SummaryHeadersTeks mumi

Aditya Hermawan:

Thank you for submitting the manuscript, "Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation" to Artificial Intelligence and Applications. With the online journal management system that we are using, you will be able to track its progress through the editorial process by logging in to the journal web site:

Submission URL: https://ojs.bonviewpress.com/index.php/AIA/authorDashboard/submission/7379

Username: adit

If you have any questions, please contact me. Thank you for considering this journal as a venue for your work.

Artificial Intelligence and Applications

Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation

Published

Workflow

Submission

Review

Review Round 1

Review Round 2

Copyediting

Production

Publication

Title & Abstract

Contributors

Metadata

Galleys

Funding data

WORKFLOW: SUBMISSION

Current Submission Language: **English**

Status

The submission is currently in the Copyediting stage.

Submission Files

Files uploaded at the time of submission

NO	FILE NAME	DATE UPLOADED	TYPE
38250	7379 Submission.docx	2025-10-23	Article Text
38375	7379 Manuscript.docx	2025-08-25	Article Text

Download All Files

Pre-Review Discussions

Add discussion

Name	From	Last Reply	Replies	Closed
Comments for the Editor	adit	2025-08-25 09:23 AM	0	☐

[AIA] EIC Screening Regarding Your Submission #7379

Pengirim

Rosie Huang

Penerima

aditiya.hermawan@ubd.ac.id

Lampiran

doralee

Tanggal

2025-09-01 15:07

Summary

Headers

Teks mumi

Dear Dr. Aditiya Hermawan,

Greetings from *Artificial Intelligence and Applications* (AIA, ISSN: 2811-0854)!

Thank you for your submission #7379, *Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation*. To ensure a smooth transition to the peer review stage, the Editor-in-Chief has provided the following suggestions for revision. The suggestions are as follows:

Dear Authors,

Thank you for your submission. To move your paper forward in the review process, we kindly ask you to address the following points. The architectures in the current version are not presented clearly, so please provide more detailed descriptions. Since the models employed are not entirely new and have been applied to hate speech detection before, it is important to highlight more explicitly how your proposed model differs from state-of-the-art approaches. In addition, we encourage you to include a comparative study with the most recent methods in hate speech detection to strengthen the experimental results.

We encourage you to revise the manuscript with these points in mind before resubmission.

Best regards,
Shivakumara

We kindly request that you address these points to improve your manuscript further. If you have finished the revision, you may send the [revised manuscript\(clean version\)](#) and [revised manuscript\(highlighted version\)](#) directly to me via email.

If you have any questions or need clarifications, please feel free to let me know. I'm willing to provide you with the most efficient author service. Looking forward to hearing from you soon! and wishing you a nice day.

Best wishes,
Rosie Huang

Associate Managing Editor
Artificial Intelligence and Applications
Singapore BON VIEW PUBLISHING PTE. LTD.

Review

[AIA] Editorial Screening Regarding your Submission #7379

Pengirim

Rosie Huang

Penerima

Aditiya Hermawan

Lampiran

doralee, cynthia@aiaceditorialoffice.com

Tanggal

2025-10-10 13:42

Summary

Headers

Teks mumi

7379 Potential AI Generated Rate.pdf (~624 KB)

Dear Dr. Aditiya Hermawan,

Greetings from *Artificial Intelligence and Applications* (AIA, ISSN: 2811-0854)!

Thank you very much for your revisions on submission #7379, *Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation*. We are pleased to inform you that your manuscript has passed the EIC screening and is now under the editorial screening. We found that some content could be improved to meet our standards. Please see the points below:

- Potential AI-Generated Rate:** The AI detection result should be under 40%. Your manuscript shows 43.56%. Please review the highlighted sections in the attached report "[7379 Potential AI Generated Rate](#)" and make the necessary revisions to decrease the rate.
- High-Resolution Images:** To avoid compression issues, please send all high-resolution images separately via email for backup. (Please send me all the high-resolution images in a single compressed folder, preferably in JPG or PNG format.)

7379 Submission Checklist: (Research Article)		
Total number of words: 3500-12000	6723(of.1383)	
Total number of reference: 20+	38	
Similarity rate: < 20% Single Similarity rate: < 8%	20% 2%	
Potential AI generated rate: < 40%	43.56%	Required revision
Self-citations: < 3	0	
Other authors' citations: < 10% of the total references	5%(2/38)	
Non-academic citations: < 3	1(2/9)	
References in the recent 5 years (2021-2025): > 40%	71%(27/38)	
References in the recent 3 years (2023-2025): > 3	19	
Corresponding author uses his/her own organization's email address (instead of ieee, gmail, 126, QQ, etc.)	YES	
The author provides all the high-resolution pictures	NO	

Journal Name: AIA Name: Rosie Huang Time: 2025/10/10

Please make the necessary revisions and send the updated manuscript back to me at your earliest convenience. Additionally, you may refer to the checklist screenshot attached to ensure that all other sections meet the required standards during your revision process.

Thank you again for your time and effort, and I look forward to hearing from you soon.

Best wishes,
Rosie Huang

Associate Managing Editor
Artificial Intelligence and Applications
Singapore BON VIEW PUBLISHING PTE. LTD.

Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation

Library

Published

Workflow

Submission

Review

Review Round 1

Review Round 2

Copyediting

Production

Publication

Title & Abstract

Contributors

Metadata

Galleys

Funding data

WORKFLOW: REVIEW (ROUND 1)

Current Submission Language: English

Status

The submission advanced to the next review round, was accepted, and is currently in the Copyediting stage.

Notifications

[AAJ] Editor Decision2025-10-23 08:08 AM

[AAJ] Editor Decision2026-06-23 01:59 PM

Revisions Uploaded

These files have been submitted by the author after revisions were requested

Upload

NO	FILE NAME	DATE UPLOADED	TYPE
44197	7379 1st Response Letter.docx	2026-01-29	Other
44196	7379 1st Revisions(clear).docx	2026-01-29	Article Text
44195	7379 1st Revisions(clear).docx	2026-01-29	Article Text

Dear Editor,

Thank you for giving me the opportunity to submit a revised draft of my manuscript titled Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation to *Artificial Intelligence and Applications*. We appreciate the time and effort that you and the reviewers have dedicated to providing your valuable feedback on my manuscript. We are grateful to the reviewers for their insightful comments on my paper. We have been able to incorporate changes to reflect most of the suggestions provided by the reviewers. We have highlighted the changes within the manuscript.

Here is a point-by-point response to the reviewers' comments and concerns.

Reviewer F

- **Comment 1:** There is no novelty in the problem statement. Hate and offensive word detection using BERT model (or a composite model that includes a BERT model) has been extensively studied for the last few years. Several Python implementations are available on Github, Kaggle, Medium and so on. Some studies included Indonesian dataset: IndoBERTweet also. Table 2 (Internally Trained models) and Table 3 (previously Trained models) results of comparison of models may not be considered as significant, as better than 87.18% accuracies are available, as presented in the literature review itself.
- **Response:** We thank the reviewer for the constructive feedback. We fully agree that hate and offensive speech detection using BERT and its variants has been widely studied, including for Indonesian datasets such as IndoBERTweet. Accordingly, we do not claim novelty in introducing BERT itself. Instead, the contribution of this study lies in its methodological and empirical focus, namely a controlled and systematic comparison between a baseline multilingual BERT model and multiple BiLSTM variants using different embedding strategies (Word2Vec, FastText, and TF-IDF) under identical experimental settings.

The reported accuracy of 87.18% is the result of our own experimental evaluation, not a value derived from prior literature. This has been clarified in the revised manuscript. Table 2 presents internally trained models evaluated on the same dataset, preprocessing pipeline, data split, and evaluation metrics to ensure fair comparison, while Table 3 is intended solely as an external contextual benchmark rather than a claim of state-of-the-art performance. Although higher accuracies have been reported in previous studies, these often rely on domain-adapted pretrained models, additional features, or ensemble techniques. In contrast, our study isolates the effect of contextual versus non-contextual representations in noisy, low-resource Indonesian social media text. We believe this comparative and diagnostic perspective clarifies model behavior and provides practical insights for researchers and practitioners, thereby strengthening the manuscript's contribution.

Reviewer G

Comments from Reviewer 1

- **Comment 1: Originality:** This paper presents an empirical study examining the efficacy of a BERT-based model in detecting hate speech within low-resource languages, specifically focusing on Indonesian social media content. While the application of BERT for hate speech detection is a relevant topic, it is important to note that similar research utilizing BERT has been extensively documented in existing literature. To strengthen originality, the authors could better articulate the challenges of applying BERT in low-resource language environment and their novel ideas.
- **Response:** We thank the reviewer for the insightful comment regarding originality. We acknowledge that the application of BERT for hate speech detection has been extensively explored in prior studies, including those involving Indonesian language data. Accordingly, this study does not claim novelty in the adoption of BERT itself. Instead, the originality of this work lies in its problem framing and empirical focus on the practical challenges of deploying BERT in low-resource, informal social media environments, where text is highly noisy, abbreviated, and context-dependent.
To address this concern, we have revised the manuscript to more clearly articulate these challenges, including limited domain adaptation, informal linguistic structures, and the absence of specialized pretrained models in many real-world settings. Moreover, we emphasize our novel comparative and diagnostic approach, which systematically evaluates a baseline multilingual BERT model against BiLSTM architectures with different embedding strategies under identical experimental conditions. This design allows us to isolate the impact of contextual modeling in low-resource settings and provides actionable insights for researchers and practitioners who operate under data and computational constraints. We believe these clarifications strengthen the originality and practical relevance of the study.
- **Comment 2: Relationship to Literature:** The paper offers an insightful overview of the existing literature, effectively summarizing research advances in the field. It emphasizes the necessity for a thorough comparative analysis of three distinct modeling approaches: transformer-based models, sequential deep learning models, and mixed models.
- **Response:** We thank the reviewer for highlighting the strength of the literature overview and the motivation for a comprehensive comparison across transformer-based, sequential deep learning, and mixed/hybrid approaches. In response, we have strengthened the "Relationship to Literature" by explicitly mapping prior Indonesian hate speech detection studies into these three paradigms and clarifying how our work addresses fragmentation in earlier evaluations. Specifically, we added a dedicated paragraph in the Literature

Review that (i) explains why prior results are difficult to compare due to differences in pipelines and reporting, and (ii) positions our contribution as a controlled benchmark where dataset, split protocol, and evaluation metrics are held constant to isolate the effect of contextual modeling versus sequential modeling with alternative embeddings. We also revised the discussion around Table 3 to more clearly connect our findings to prior hybrid and classical baselines, emphasizing what is comparable and what remains outside the current experimental scope.

- **Comment 3:** Methodology: The research employs a quantitative methodology that aligns well with established frameworks in the field. The life cycle of comparative analysis regarding hate speech on Indonesian social media content is presented effectively, providing a clear structure for the study. However, the novelty of the work could be articulated more distinctly to better highlight its contributions to existing literature.

Response:

We thank the reviewer for recognizing the clarity and structural coherence of the proposed quantitative methodology. In response to the comment regarding novelty, we have revised the Methodology section to more explicitly articulate the unique contribution of our work. Specifically, we now emphasize that the novelty lies in the design of a unified comparative lifecycle that evaluates transformer-based models, sequential deep learning models, and embedding-driven variants under identical experimental conditions. By fixing the dataset, preprocessing pipeline, data split, training configuration, and evaluation metrics, our methodology isolates the effect of contextual modeling versus sequential and embedding-based representations. This clarification highlights how the proposed methodological framework extends existing literature beyond isolated model evaluations and provides a controlled benchmark for hate speech detection in low-resource Indonesian social media contexts.

- **Comment 4:** The empirical study investigates the effectiveness of a BERT-based model in detecting hate speech in low-resource languages. While BERT-based models are widely recognized in NLP research and have demonstrated impressive performance, it would be beneficial to include a comparative analysis of BERT-based models in both standard and low-resource language environments. Such a comparison would provide valuable insights into the model's adaptability and performance variations across different linguistic contexts.

Response: Thank you for this valuable suggestion. We agree that comparing BERT-based models across standard and low-resource language environments can provide important insights into model adaptability. In response, we have revised the manuscript to explicitly situate our findings within the broader BERT literature by contrasting reported performance trends in high-resource languages (e.g., English benchmark datasets) with those observed in low-resource Indonesian social media data. This comparative discussion has been

added to the Discussion section to highlight performance variations, contextual challenges, and adaptability considerations without introducing additional experiments beyond the scope of the current study.

- **Comment 5: Quality of Communication:** The paper effectively conveys the research. However, there are some repetitions of section number “4.3” and title e.g “4.3. BiLSTM Model Evaluation – TF-IDF”

Response: Thank you for pointing out the issue regarding the repetition of section numbering and titles. We have carefully revised the manuscript to correct the duplicated section number “4.3” and the repeated title “BiLSTM Model Evaluation – TF-IDF.” The section numbering has now been made consistent and sequential throughout the Results and Discussion section to improve clarity and readability. These revisions ensure better structural coherence and enhance the overall quality of communication of the paper.

We look forward to your kind consideration of our submission and would be pleased to respond to any further questions, comments, or requests for clarification that may arise during the review process.

Sincerely,

December, 31 2025
Aditiya Hermawan

Re: [AIA] Editorial Screening Regarding your Submission #7379



Pengirim [Rosie Huang](#)
Penerima [Aditiya Hermawan](#)
Lampiran doralee_cynthia@alaeditorialoffice.com
Tanggal 2025-10-23 14:54
[Summary](#) [Headers](#) [Teks mumi](#)

Dear Dr. Aditiya Hermawan,

Thank you for your emails. We are sincerely sorry for the delayed response. We've completed the check of your revised manuscript, and we're pleased to inform you that your manuscript #7379, *Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation*, has successfully **passed the initial screening by the editorial and EIC screening**.

The submission has now **moved on to the peer review stage**. We will continue to follow up on the review process and provide you with feedback on any suggested revisions.

If you have any questions, please do not hesitate to contact me. I'm here to assist you at any time. Looking forward to hearing from you soon.

Best wishes,
Rosie Huang

Associate Managing Editor
Artificial Intelligence and Applications
Singapore BON VIEW PUBLISHING PTE. LTD.

[AIA] Regarding Your Manuscript #7379 - Revisions Needed Before Acceptance

Pengirim: Rosie Huang
Penerima: Aditya Hermawan
Lampiran: doraalee, arianaxu@aiaeditorial.com
Tanggal: 2026-06-11 14:04
[Summary](#) [Headers](#) [Teks mumi](#)

7379 Potential AI Generated Rate.pdf (~107 KB) 7379 AIA Template.docx (~536 KB)

Dear Dr. Aditya Hermawan,

Greetings from *Artificial Intelligence and Applications* (AIA, ISSN: 2811-0854)!

We are currently conducting a pre-acceptance check and formatting review for submission #7379, *Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation*. During this process, we noticed that a few aspects of the manuscript need to be revised to meet our acceptance standards. Please find the details below:

- Recent References (2021-2026):** At least 80% of the references should be from the last five years. Currently, your manuscript has 71%(27/38) recent 5-year references, so we kindly ask you to make the necessary updates.
- Potential AI-Generated Rate:** The AI detection result should be under 30%. Your manuscript shows 41%. Please review the highlighted sections in the attached report "7379 Potential AI Generated Rate" and make the necessary revisions. Alternatively, you may **provide a report from a reliable detection platform** showing that the potential AI-generated content rate of the full manuscript is below 30%.
- Retracted Reference:** Ref. [16] has been retracted. Please delete or replace it.
- Funding Support:** Please confirm whether the manuscript has any funding support.
- Data Availability Statement:** Please complete this section in accordance with the templates provided in the manuscript.
- Contributors' information:** Please indicate the corresponding affiliation for each author, and provide the author's full name (given name and family name) for Junaedl.

Please send the revised article to me by email. You may directly add the revisions in the file "7379 AIA Template" and **highlight the revised content**. Should you have any further queries, please do not hesitate to contact me at any time. I'm glad to assist you.

Thanks for your patience and attention. We look forward to your response at your earliest convenience.

Best wishes,
Rosie Huang

Associate Managing Editor
Artificial Intelligence and Applications
Singapore BON VIEW PUBLISHING PTE. LTD.

Accepted

[AIA] Notification of Manuscript Acceptance #7379

Pengirim: Rosie Huang
Penerima: Aditya Hermawan
Lampiran: doraalee, arianaxu@aiaeditorial.com
Tanggal: 2026-06-24 14:41
[Summary](#) [Headers](#) [Teks mumi](#)

Bon_View_CC_4.0_OA_licence_agreement.docx (~45 KB)

Dear Dr. Aditya Hermawan,

Greetings from *Artificial Intelligence and Applications* (AIA, ISSN: 2811-0854)!

We are pleased to inform you that your submission #7379, *Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation*, has been successfully **accepted for publication** in our journal. Thank you for your valuable contribution and support.

Next, our team will proceed with the formatting process, and a proof will be sent to you for your review and confirmation. Additionally, an **Open Access Agreement** is attached for your signature. Please kindly **check it and send me the signed file**.

If there are multiple co-authors, each author should sign the agreement, or the corresponding author may sign with the note "on behalf of all authors" to confirm representation.

Should you have any questions or concerns during this time, please don't hesitate to contact me. Once again, thank you for your support of our journal.

Best wishes,
Rosie Huang

Associate Managing Editor
Artificial Intelligence and Applications
Singapore BON VIEW PUBLISHING PTE. LTD.

Artificial Intel

My Submissio

7379

Active s

Revisio

Revisio

Revisio

Schedu

Publicat

Decline

Start A New S

Workflow

Submission

Review

Review R

Review R

Copyediting

Production

Publication

Title & Abs

Contributor

Metadata

Galleys

Funding da

Notifications

[AIA] Editor Decision

2026-06-23 01:59 PM

Yosia Immanuel Bastian, Aditya Hermawan, Ardiane Rossi Kurniawan Maranto, Benny Daniawan, Junaedl:

We have reached a decision regarding your submission to Artificial Intelligence and Applications, "Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation".

Our decision is to: Accept Submission

Copy Editing

[AIA] 2nd Proof for Author Review #7379 (22/07/2026)



Pengirim

Penerima

Lampiran

Tanggal

Rosie Huang

Aditya Hermawan

doralee

2026-07-22 13:31

 Summary

 Headers

 Teks mumi

 AIA62027379_R1.pdf (~828 KB) 

Dear Dr. Aditya Hermawan,

Greetings from *Artificial Intelligence and Applications* (AIA, ISSN: 2811-0854).

The corrections to your manuscript #7379 have been completed, and the proof is now available for your review again. Please take a moment to check the document and confirm if everything is in order. If any revisions are needed, kindly let us know.

Please note that at this stage, you may still make revisions to the manuscript content (except for the author information). However, once the manuscript is published in an official volume, no further changes can be made. Therefore, we kindly ask you to carefully check all content in the manuscript, including spelling, text, data accuracy, and other details.

Thank you for your attention to this matter.

Best wishes,
Rosie Huang

Associate Managing Editor
Artificial Intelligence and Applications
Singapore BON VIEW PUBLISHING PTE. LTD.

[AIA] 1st Proof for Author Review #7379(09/07/2026)



Pengirim

Penerima

Lampiran

Tanggal

Rosie Huang

Aditya Hermawan

doralee

2026-07-09 10:14

 Summary

 Headers

 Teks mumi

 AIA62027379_A1.pdf (~826 KB) 

Dear Dr. Aditya Hermawan,

Greetings from *Artificial Intelligence and Applications* (AIA, ISSN: 2811-0854).

Our production team has completed the typesetting. The first proof of your manuscript #7379, *Deep Learning Approaches for Hate Speech Detection in Resource-Constrained Social Media Texts: A Comparative Evaluation*.

Upon reviewing the document, our production team has noted a few queries on the last page. To address these, we kindly ask that you directly highlight sentences or add comments within the PDF.

Should you require further revisions or have any other feedback, please do not hesitate to reach out. Thank you for your attention to this matter.

Best wishes,
Rosie Huang

Associate Managing Editor
Artificial Intelligence and Applications
Singapore BON VIEW PUBLISHING PTE. LTD.

QUERIES

- AQ1: The sentence starting “Embedding model with Word2Vec...” is not clear. Please check.
- AQ2: Please provide expansion of “BiLSTM” at first mention.
- AQ3: Please expand “TF-IDF” in the abstract and at its first mention in the main text.
- AQ4: Please check if edit made in “To conclude the above matters, this finding...” is okay.
- AQ5: Please check if edit made in “In contrast, BERT has consistently...” is okay.
- AQ6: Please check if edit made in “As a transformer-based model...” is okay.
- AQ7: Please check if edit made in “Even though BERT has been adopted...” is okay.
- AQ8: The phrase “word representation technic developments” is not clear. Please check.
- AQ9: Please check if edit in “Despite this, TF-IDF gives statistical...” is okay.
- AQ10: Please check if edit made in “But comprehensive studies that systematically...” is okay.
- AQ11: Please check if edit made in sentence “This study provides empirical knowledge...” is okay and retains the intended meaning.
- AQ12: Please check if edit in “Hence, these classic models perform...” is okay.
- AQ13: Please provide expansion of “CNN-RNN” at first mention.
- AQ14: Please check if edit made in “Along with the developments, NLP...” is okay.
- AQ15: Please check if edit made in “Given the lack of comprehensive comparative...” is okay.
- AQ16: Please check if the term “sequential recurring architecture” is okay as given.
- AQ17: Please check if edit made in “In this way, the study seeks to...” is okay and retains the intended meaning.
- AQ18: Please check if edit made in “Figure 6(a) indicates that the training accuracy...” is okay.
- AQ19: Please check if edit in “The performance scoring indicates that...” is okay.
- AQ20: The sentence starting “Eventually, this result reflects the...” is not clear. Please check.
- AQ21: The sentence “Along with this dual perspective...” is not clear. Please check.
- AQ22: The sentence “Even though Word2Vec and FastText embedding...” seems to be incomplete. Please check.
- AQ23: Please check if edit in “Second, our BERT training illustrates...” is okay.
- AQ24: Please check if edit in “According to the reposted studies...” is okay.
- AQ25: The phrase “widen feature with synonym based” is not clear. Please check.

