



Artikel Ilmiah Cek Plagiasi Ft Umsida

18%
Suspicious
texts



14% Similarities
0 % similarities between quotation marks
2 % among the sources mentioned
5% Unrecognized languages
33% Texts potentially generated by AI (ignored)

Document name: Artikel Ilmiah Cek Plagiasi Ft Umsida.pdf
Document ID: 32683ad4d752daf69c376fb0361b8fe29939769b
Original document size: 817.87 KB

Submitter: fst umside
Submission date: 1/30/2026
Upload type: interface
analysis end date: 1/30/2026

Number of words: 9,140
Number of characters: 70,505

Location of similarities in the document:



Sources of similarities

Main sources detected

No.	Description	Similarities	Locations	Additional information
1	 ejournal.unikama.ac.id Analisis Tingkat Berfikir Kreatif Peserta Didik Dalam Me... http://ejournal.unikama.ac.id/index.php/jtst/article/download/3225/2111 34 similar sources	11%		 Identical words: 11% (1,152 words)
2	 archive.umsida.ac.id https://archive.umsida.ac.id/index.php/archive/preprint/download/4336/30985/34944 33 similar sources	10%		 Identical words: 10% (1,074 words)
3	 KARYA ILMIAH Agustina Kurniawati.pdf KARYA ILMIAH Agustina Kurni... #769aca Comes from my group 33 similar sources	10%		 Identical words: 10% (1,073 words)
4	 archive.umsida.ac.id https://archive.umsida.ac.id/index.php/archive/preprint/download/4621/33078/37498 2 similar sources	7%		 Identical words: 7% (770 words)
5	 doi.org Detecting cyberbullying using deep learning techniques: a pre-trained gl... https://doi.org/10.7717/peerj-cs.1961 11 similar sources	2%		 Identical words: 2% (191 words)

Sources with incidental similarities

No.	Description	Similarities	Locations	Additional information
1	 doi.org Enhancing Cyberbullying Detection on Platform 'X' Using IndoBERT and ... https://doi.org/10.52436/1.jutif.2025.6.2.4321	< 1%		 Identical words: < 1% (31 words)
2	 doi.org Klasifikasi Jenis Kejahatan berdasarkan Teks Amar Putusan Pengadilan ... https://doi.org/10.29408/edumatic.v9i2.30326	< 1%		 Identical words: < 1% (29 words)
3	 doi.org Deteksi Komentar Cyberbullying Pada YouTube Dengan Metode Convol... https://doi.org/10.34148/teknika.v12i3.677	< 1%		 Identical words: < 1% (23 words)
4	 doi.org Penyeimbangan Kelas SMOTE dan Seleksi Fitur Ensemble Filter pada Su... https://doi.org/10.25126/jtiik.1067234	< 1%		 Identical words: < 1% (15 words)
5	 dx.doi.org Penerapan Kalman Filter Pada Pembacaan Sensor Loadcell Berbasis ... http://dx.doi.org/10.33795/elkolind.v11i3.6158	< 1%		 Identical words: < 1% (14 words)

Referenced sources (without similarities detected) These sources were cited in the paper without finding any similarities.

1	 https://t.co/NLTPgeBuka
2	 https://t.co/ewhKdZIGGa
3	 https://doi.org/10.3390/make6010009
4	 https://doi.org/10.18280/isi.280410
5	 https://doi.org/10.1109/ICAECA56562.2023.10201149

Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.

Hybrid CNN-LSTM for Indonesian Cyberbullying Text Classification on Social Media X

Muhammad Hafizh Fattah1),



Mochamad Alfian Rosid2) *, Sukma Aji3) & Suprianto4)

1,2,3,4) Program

Studi Informatika, Fakultas Sains dan Teknologi, Universitas Muhammadiyah
Sidoarjo, Indonesia

*Corresponding Email: [alfanrosid@umsida.ac](mailto:alfanrosid@umsida.ac.id)

.id

Abstrak

Perundungan siber di platform media sosial X telah menjadi ancaman digital yang kritis dan memerlukan mekanisme deteksi otomatis untuk memitigasi dampak psikologis pada korban. Penelitian ini mengusulkan arsitektur deep learning hibrida yang menggabungkan Convolutional Neural Network (CNN) untuk ekstraksi fitur lokal dan Long Short-Term Memory (LSTM) untuk pemahaman konteks sekuensial dalam mengklasifikasikan komentar perundungan siber berbahasa Indonesia. Penelitian ini mengevaluasi kinerja model menggunakan dataset sebanyak 13.677 komentar dari media sosial X melalui serangkaian skenario pengujian sistematis, termasuk dampak regularisasi, pemanfaatan embedding FastText, dan studi komparatif terhadap model mutakhir. Hasil eksperimen menunjukkan bahwa mekanisme Early Stopping merupakan faktor kritis dalam arsitektur ini, di mana tanpa mekanisme tersebut model mengalami degradasi akurasi hingga 32%. Model CNN-LSTM yang diusulkan mencapai akurasi 88,38% dan F1-Score 88%, meningkat menjadi AUC 0,9559 dengan integrasi FastText. Model ini mencapai lebih dari 97% kinerja IndoBERTweet dengan kompleksitas komputasi 22 kali lipat lebih rendah (4,97 juta berbanding 110,88 juta parameter) dan mengungguli metode machine learning seperti SVM dengan margin akurasi lebih dari 10 poin persentase. Penelitian ini menyimpulkan bahwa arsitektur CNN-LSTM menawarkan solusi robust dan efisien untuk deteksi perundungan siber, khususnya untuk lingkungan dengan sumber daya terbatas.



Kata Kunci: Perundungan Siber, CNN-LSTM, Deep Learning, Media Sosial X,

Bahasa Indonesia.

Abstract

Cyberbullying on social media platform X has become a critical digital threat and requires automatic detection mechanisms to mitigate psychological impacts on victims. This study proposes a hybrid deep learning architecture that combines Convolutional Neural Network (CNN) for local feature extraction and Long Short-Term Memory (LSTM) for sequential context understanding in classifying Indonesian language cyberbullying comments. This study evaluates model performance using a dataset of 13,677 comments from social media X through a series of systematic testing scenarios, including the impact of regularization, utilization of FastText embeddings, and comparative studies against state-of-the-art models. Experimental results demonstrate that the Early Stopping mechanism is a critical factor in this architecture, where without this mechanism the model experiences accuracy degradation of up to 32%. The proposed CNN-LSTM model achieves 88.38% accuracy and 88.00% F1-Score, improving to 0.9559 AUC with FastText integration. This model achieves over 97% of IndoBERTweet's performance with 22 times lower computational complexity (4.97 million versus 110.88 million parameters) and outperforms machine learning methods such as SVM with an accuracy margin of more than 10 percentage points. This study concludes that the CNN-LSTM architecture offers a robust and efficient solution for cyberbullying detection, particularly for resource-constrained environments.

Keywords: Cyberbullying, CNN-LSTM, Deep Learning, Social Media X, Indonesian Language.



I. PENDAHULUAN

Era digital kontemporer ditandai dengan penetrasi media sosial yang masif ke dalam berbagai aspek kehidupan masyarakat. Platform seperti X (sebelumnya dikenal sebagai Twitter) telah bertransformasi dari sekadar jejaring sosial menjadi arena utama untuk wacana publik, diseminasi informasi, dan interaksi sosial berskala global. Statistik terbaru dari tahun 2024 (Shubham Singh, 2025) menunjukkan bahwa X memiliki sekitar 611 juta pengguna aktif bulanan di seluruh dunia dan menempati peringkat di antara platform media sosial teratas secara global. Skala besar ini menjadikan X sebagai salah satu ekosistem komunikasi paling berpengaruh. Platform ini memainkan peran penting sebagai ruang untuk ekspresi pendapat yang cepat dan dinamis secara khusus di Indonesia, dengan pengguna secara aktif mendiskusikan isu-isu sosial, politik, dan budaya populer. Basis pengguna X Indonesia menunjukkan tingkat keterlibatan yang tinggi, sehingga menjadikannya sangat signifikan untuk mempelajari pola perilaku daring dan fenomena sosial.

Kemudahan akses dan kebebasan berekspresi yang ditawarkan oleh media sosial X juga menghasilkan dampak negatif berupa fenomena perundungan siber. Perundungan siber merujuk pada tindakan agresif dan disengaja yang dilakukan oleh individu atau kelompok menggunakan media elektronik, yang dapat terjadi kapan saja dan di mana saja, melampaui batasan fisik lingkungan tradisional. Studi terkini mengindikasikan bahwa perundungan siber memengaruhi sekitar 36% pengguna media sosial Indonesia,

3

ejournal.unikama.ac.id | Analisis Tingkat Berfikir Kreatif Peserta Didik Dalam Menyelesaikan Soal Matematika Ditinjau Dari Adversity Quotient (AQ)
<http://ejournal.unikama.ac.id/index.php/jst/article/download/3225/2111>

Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.

4

KARYA ILMIAH Agustina Kurniawati.pdf | KARYA ILMIAH Agustina Kurniawati
♥ Comes from my group

dengan korban mengalami konsekuensi psikologis yang parah termasuk gangguan kecemasan, depresi, dan isolasi sosial (El Koshiry et al., 2024). Karakteristik media sosial X seperti potensi anonimitas, kemampuan diseminasi konten viral, dan kesingkatan pesan (keterbatasan karakter) menciptakan lingkungan yang sangat subur untuk perkembangan perilaku verbal yang toksik dan intimidatif. Perundungan siber di X dapat menjangkau audiens yang masif secara instan dan meninggalkan jejak digital permanen, berbeda dengan perundungan tatap muka, yang memperkuat efek berbahayanya. Kondisi ini menjadikan perundungan siber di X sebagai masalah sosial mendesak yang memerlukan solusi teknologi yang efektif.

5

doi.org | Klasifikasi Jenis Kejahatan berdasarkan Teks Amar Putusan Pengadilan Hukum Pidana KUHP menggunakan IndoBERT
<https://doi.org/10.29408/edumatic.v9i2.30326>

Kecerdasan Buatan (Artificial Intelligence/AI), khususnya Pemrosesan Bahasa Alami (Natural Language Processing/NLP),

menawarkan kerangka solusi yang menjanjikan untuk mengatasi masalah perundungan siber. NLP berfokus pada pengembangan sistem komputasi yang mampu secara otomatis memahami, menginterpretasi, dan memproses bahasa manusia. Teknik klasifikasi teks dapat dimanfaatkan dalam konteks perundungan siber untuk membangun model yang mampu secara otomatis mengidentifikasi dan memfilter konten yang mengandung elemen perundungan. Tinjauan sistematis komprehensif pada tahun 2023 mengonfirmasi bahwa penelitian terkini secara konsisten menunjukkan keunggulan model berbasis deep learning dibandingkan algoritma machine learning konvensional dalam menangani data yang besar dan kompleks, serta kemampuannya mengekstrak fitur secara otomatis tanpa rekayasa fitur manual yang ekstensif (Hasan et al., 2023).



Berbagai studi empiris telah membuktikan keunggulan model deep learning dalam konteks deteksi perundungan siber berbahasa Indonesia. Penelitian yang menerapkan arsitektur deep learning sekuensial menunjukkan bahwa model Bidirectional Long Short-Term Memory (BiLSTM) berhasil mencapai akurasi tinggi sebesar 81,82% pada data uji untuk klasifikasi komentar perundungan siber berbahasa Indonesia di platform Instagram (Iryani et al., 2025). Temuan ini menggarisbawahi kekuatan model berbasis LSTM dalam

memahami urutan kalimat dan konteks, yang merupakan elemen krusial dalam mengidentifikasi pola perilaku perundungan yang tertanam dalam teks percakapan. Penelitian lain menunjukkan kinerja tinggi yang konsisten dari model Bi-LSTM dengan mencapai akurasi hingga 97% dalam mendeteksi berbagai jenis pelanggaran UU ITE, termasuk perundungan siber, di platform media sosial berbahasa Indonesia (Maulana & Aditya, 2025). Hasil-hasil ini memvalidasi efektivitas jaringan saraf rekuren dalam menangkap sifat sekuensial bahasa alami.

Model tunggal menghadapi keterbatasan inheren dalam menangkap fitur lokal dan pola sekuensial global secara bersamaan, meskipun arsitektur deep learning individual telah menunjukkan hasil yang menjanjikan. Efektivitas arsitektur CNN dalam klasifikasi teks berbahasa Indonesia telah divalidasi dalam berbagai konteks analisis bahasa natural. Penelitian menunjukkan bahwa arsitektur CNN dengan



embedding layer, convolutional layer, pooling, dan dense layer dapat mencapai akurasi 99,10% dalam

analisis sentimen komentar YouTube berbahasa Indonesia, dengan penekanan pada pentingnya preprocessing komprehensif yang mencakup text cleaning,



normalization, tokenization, stopword removal, dan stemming (Cahyo et al., 2025).

Dalam domain identifikasi konten problematik lainnya, CNN juga menunjukkan performa tinggi untuk identifikasi berita hoaks berbahasa Indonesia dari Facebook, mencapai akurasi 92,53% pada data pelatihan dan 81,09% pada data pengujian dengan pendekatan preprocessing serupa serta pemanfaatan Word2Vec untuk representasi vektor kata (Ramadhan et al., 2024). Penggunaan pre-trained word embeddings seperti Word2Vec terbukti meningkatkan pemahaman konteks semantik, aspek yang relevan untuk membedakan antara komentar bullying dan non-bullying dalam media sosial yang sering mengandung bahasa informal, slang, dan pencampuran kode.



Kesadaran akan keterbatasan model tunggal telah mendorong tren penelitian menuju arsitektur hibrida yang lebih canggih, yang kini dianggap sebagai salah satu pendekatan paling mutakhir dalam klasifikasi teks. Arsitektur hibrida Convolutional Neural Network - Long Short-Term Memory (CNN-LSTM) secara khusus menunjukkan potensi luar biasa karena kemampuannya menggabungkan kekuatan kedua model secara sinergis. CNN unggul dalam mengekstrak fitur lokal yang signifikan seperti frasa kunci, pola n-gram, atau kombinasi kata-kata toksik, sementara LSTM unggul dalam memahami konteks sekuensial, dependensi jangka panjang antar kata, dan alur sentimen keseluruhan dalam sebuah kalimat. Beberapa studi terkini telah memvalidasi efektivitas pendekatan hibrida ini pada data berbahasa Indonesia. Penelitian pada tahun 2023 berhasil menerapkan model hibrida CNN-LSTM pada data berbahasa Indonesia dari media sosial X dan mencapai akurasi 79,48% (Asqolani & Setiawan, 2023). Pada tahun yang sama, penelitian lain memperkuat temuan ini dengan menggunakan arsitektur hibrida CNN-BiLSTM yang ditingkatkan dengan ekspansi fitur FastText dan berhasil mencapai akurasi 80,55% pada data tweet Indonesia (Nasution & Setiawan, 2023). Konsistensi performa tinggi CNN dalam ekstraksi fitur lokal, dikombinasikan dengan kemampuan LSTM dalam memahami konteks sekuensial, menjadi fondasi kuat untuk pengembangan arsitektur hibrida CNN-LSTM dalam penelitian ini.



ejournal.unikama.ac.id | Analisis Tingkat Berfikir Kreatif Peserta Didik Dalam Menyelesaikan Soal Matematika Ditinjau Dari Adversity Quotient (AQ)
<http://ejournal.unikama.ac.id/index.php/jst/article/download/3225/2111>

Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use,

distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original

publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply

with these terms.

Beberapa gap penelitian masih tersisa dan perlu diatasi, meskipun model hibrida CNN-LSTM telah menunjukkan hasil yang menjanjikan. Pertama, penelitian terdahulu yang menerapkan arsitektur hibrida CNN-LSTM sering kali terbatas pada platform atau konteks data spesifik dan kurang memiliki evaluasi komprehensif khusus pada media sosial X, di mana pola linguistik dan manifestasi perundungan siber yang unik terjadi (Hafiza & Setiawan, 2025). Kedua, implementasi yang ada sering kali belum menginvestigasi secara sistematis dampak teknik regularisasi (seperti Early Stopping) terhadap generalisasi model, meskipun overfitting merupakan tantangan yang dikenal luas dalam deep learning. Ketiga, peran embedding

kata yang telah dilatih sebelumnya (seperti FastText) dalam meningkatkan kinerja model hibrida untuk bahasa Indonesia belum diteliti secara menyeluruh. Keempat, studi komparatif komprehensif yang memposisikan hibrida CNN-LSTM terhadap metode machine learning maupun model berbasis Transformer mutakhir masih langka. Penelitian menunjukkan bahwa kinerja model hibrida CNN-LSTM sangat dipengaruhi oleh proses optimasi hiperparameter dan pilihan desain arsitektural, yang tanpanya model mungkin tidak mencapai kinerja puncaknya (Faritha Banu et al., 2022). Kondisi ini mengindikasikan kebutuhan untuk mengembangkan model dasar yang kuat dengan kinerja yang diukur secara sistematis dan metodologi implementasi yang robust khusus untuk deteksi perundungan siber berbahasa Indonesia di media sosial X.



Penelitian ini bertujuan untuk mengisi gap tersebut melalui implementasi sistematis dan evaluasi yang ketat terhadap arsitektur hibrida CNN-LSTM. Penelitian ini secara spesifik berupaya menjawab pertanyaan-pertanyaan penelitian berikut: (1) Seberapa efektif Early Stopping dalam mencegah overfitting pada model CNN-LSTM untuk deteksi perundungan siber berbahasa Indonesia? (2) Sejauh mana embedding FastText yang telah dilatih sebelumnya meningkatkan kinerja klasifikasi model? (3) Bagaimana perbandingan model hibrida CNN-LSTM terhadap metode machine learning, arsitektur deep learning tunggal, dan model Transformer mutakhir dalam hal akurasi dan efisiensi komputasi? (4) Apa keseimbangan optimal antara kinerja model dan kompleksitas komputasi untuk deployment praktis?

Penelitian ini memiliki empat kontribusi utama. Pertama, penelitian ini menyediakan implementasi dan evaluasi sistematis terhadap arsitektur hibrida CNN-LSTM yang dioptimalkan secara khusus untuk deteksi perundungan siber berbahasa Indonesia di media sosial X, dengan desain eksperimental komprehensif yang mencakup beragam skenario pengujian dan studi ablasi terkontrol. Kedua, penelitian ini menyajikan validasi empiris terhadap pentingnya kritis mekanisme Early Stopping dalam mencegah overfitting, dengan mengkuantifikasi dampaknya melalui eksperimen komparatif terkontrol yang mengungkapkan degradasi akurasi hingga 32% tanpa teknik regularisasi ini. Ketiga, penelitian ini melakukan studi komparatif komprehensif lintas paradigma arsitektural yang berbeda, termasuk arsitektur



hibrida (CNN-LSTM), model deep learning tunggal (CNN dan LSTM standalone),

metode machine learning

(SVM, Naive Bayes), dan model berbasis Transformer mutakhir (IndoBERTweet), untuk menyediakan tolok ukur kinerja yang jelas bagi para peneliti. Keempat, penelitian ini mengembangkan model dasar yang efisien dan praktis yang mencapai keseimbangan yang menguntungkan antara kinerja deteksi (akurasi 88,38%) dan efisiensi komputasi (4,97 juta parameter), sehingga menjadikannya sangat cocok untuk deployment di lingkungan dengan sumber daya terbatas.



II. STUDI PUSTAKA

Pemrosesan Bahasa Alami (Natural Language Processing/NLP) telah menjadi salah satu aplikasi utama dari model deep learning karena kemampuan generalisasi dan pembelajaran hierarkisnya. Dalam beberapa tahun terakhir, penelitian tentang deteksi perundungan siber di media sosial telah mengalami kemajuan signifikan, didorong oleh kompleksitas masalah dan urgensi untuk menciptakan lingkungan digital yang aman. Berbagai pendekatan berbasis machine learning dan deep learning telah terus diuji dan dikembangkan di berbagai platform media sosial, dengan tujuan meningkatkan akurasi dan presisi dalam mengklasifikasikan ujaran yang mengandung kekerasan verbal atau emosional.

Pada tahap awal, kerangka kerja berbasis analisis sentimen yang terintegrasi dengan machine learning seperti Random Forest dirancang untuk mengidentifikasi komentar berisiko. Meskipun model-model ini mampu mendeteksi ujaran negatif, mereka masih menghadapi keterbatasan dalam membedakan secara eksplisit jenis pelaku dan jenis agresi verbal. Penggunaan algoritma Random Forest dengan seleksi fitur Information Gain pada data berbahasa Indonesia dari media sosial X mencatat akurasi 72,1% (Mitra et al., 2021), yang menunjukkan peran krusial dari seleksi fitur. Santoso et al. (2023) menunjukkan peningkatan akurasi hingga 84% menggunakan Random Forest untuk klasifikasi komentar intimidatif di Instagram (Masbadi Hatullah Nurnaryo et al., 2022), yang menggarisbawahi pentingnya penyetelan hiperparameter.



Seiring dengan berkembangnya metodologi, model deep learning mulai mendominasi bidang ini.

Andika et al. (2023) menerapkan pendekatan CNN-LSTM pada lebih dari 26.000 komentar YouTube berbahasa Indonesia, mencapai F1-score 0,84 (Santoso et al., 2023). Harish et al. (2023) membandingkan model Transformer, SVM, dan LSTM, dengan Transformer menunjukkan kinerja terbaik dalam klasifikasi komentar intimidatif (Andika et al., 2023). Pendekatan ensemble juga terbukti efektif, dengan Alqahtani dan Ilyas (2024) mencapai akurasi 90,71% untuk klasifikasi multi-label dari enam jenis perundungan siber di media sosial X (Harish et al., 2023). Studi komparatif oleh Ahmad et al. (2024) mengonfirmasi kinerja SVM yang stabil dengan akurasi 85% untuk deteksi ujaran kebencian (Alqahtani & Ilyas, 2024), sementara Akar (2024) menemukan bahwa Random Forest dan XGBoost unggul di antara sembilan algoritma NLP (Ahmad et al., 2024). Tapaskar dan J (2024) mengusulkan kerangka kerja multi-algoritma yang menggabungkan CNN, Naive Bayes, dan SVM untuk meningkatkan presisi klasifikasi (Akar, 2024).

Model deep learning berbasis sekuensial seperti LSTM telah menunjukkan kemampuan luar biasa dalam mengklasifikasikan komentar berbahasa non-Inggris, seperti Bangla, dengan akurasi mencapai 99,8% (Tapaskar & J, 2024). Demikian pula, Purkayastha et al. (2025) mencapai akurasi 91,36% dengan pendekatan hibrida BERT dan machine learning (Aker et al., 2025). Pengembangan model LSTM yang diperkaya dengan teknik interpretasi seperti SHAP dan LIME memungkinkan analisis tingkat keparahan komentar, dengan akurasi hingga 98% (Purkayastha et al., 2025).

Tren terkini juga menyoroti pentingnya penanganan data multibahasa dan pencampuran kode (code-mixing) yang semakin umum di media sosial. Mochamad Alfian Rosid et al. (2024) mengembangkan model hibrida multi-head attention, CNN, dan Bi-GRU yang memanfaatkan embedding FastText dan GloVe untuk deteksi sarkasme dalam teks campuran Indonesia-Inggris, mencapai akurasi 94,60% dan F1-score 94,38% (Prana et al., 2025). Penelitian ini mengonfirmasi bahwa pendekatan multi-model dengan fitur semantik yang kaya sangat penting untuk klasifikasi yang lebih mendalam dan kontekstual, terutama dalam lingkungan multibahasa. Lebih lanjut, perbandingan antara arsitektur hibrida CNN-LSTM dan CNN-GRU menunjukkan bahwa keduanya efektif untuk analisis sentimen, dengan CNN-GRU sedikit lebih unggul dalam presisi dan akurasi, terutama dengan penerapan teknik SMOTE untuk ketidakseimbangan data (Rosid et al., 2024).

Perundungan siber secara eksplisit didefinisikan sebagai perilaku agresif dan disengaja yang dilakukan melalui media elektronik, berulang kali, dan tanpa batasan waktu terhadap korban yang tidak dapat membela diri (Smote, 2025). Bentuknya bervariasi dari ancaman (harassment) hingga peniruan identitas (impersonation), dan dapat menyebabkan gangguan mental seperti stres dan depresi. Arsitektur hibrida Convolutional Neural Network - Long Short-Term Memory (CNN-LSTM) merupakan salah satu pendekatan deep learning yang kuat dalam klasifikasi teks. CNN berfungsi sebagai ekstraktor fitur yang

efisien, mengidentifikasi pola lokal atau n-gram signifikan dari urutan teks, seperti frasa kunci dalam konteks perundungan siber. Sementara itu, LSTM, sebagai jenis khusus dari Recurrent Neural Network (RNN), dirancang untuk menangani masalah dependensi jangka panjang dalam data sekuensial. Setelah CNN mengekstrak fitur, lapisan LSTM menganalisis urutan fitur-fitur tersebut untuk memahami konteks kalimat secara keseluruhan, termasuk bagaimana urutan kata dapat mengubah makna. Kombinasi CNN untuk fitur spasial dan LSTM untuk urutan temporal telah terbukti efektif dalam deteksi perundungan siber.

Meskipun berbagai metode telah diusulkan dan menunjukkan hasil yang menjanjikan, beberapa kesenjangan penelitian masih tersisa yang perlu diatasi. Penerapan model hibrida pada data berbahasa Indonesia masih terbatas pada jenis platform tertentu seperti YouTube atau ulasan aplikasi dan belum secara ekstensif mencakup platform media sosial seperti X, yang memiliki kompleksitas bahasa yang lebih tinggi dan karakteristik unik. Selain itu, penelitian yang secara khusus menangani data dengan karakteristik pencampuran kode (code-mixing) masih terbatas, padahal fenomena ini sangat umum ditemukan di platform tersebut. Eksplorasi strategi ekstraksi fitur yang lebih spesifik, generalisasi model terhadap gaya bahasa lokal dan campuran, serta sensitivitas konteks untuk membedakan antara perundungan siber dan humor masih belum optimal. Oleh karena itu, penelitian ini berupaya mengembangkan model CNN-LSTM yang dioptimalkan untuk mendeteksi komentar perundungan siber berbahasa Indonesia, termasuk yang mengandung pencampuran kode, secara lebih akurat dan kontekstual di media sosial X, serta menyediakan baseline yang kuat untuk pengembangan lebih lanjut.

III. METODE PENELITIAN

Bagian ini menjelaskan metode dan tahapan penelitian yang digunakan secara detail. Penjelasan mencakup tahapan mulai dari studi literatur, pengumpulan dan persiapan data, perancangan dan optimasi



Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.



model

klasifikasi menggunakan arsitektur hibrida CNN-LSTM, hingga evaluasi kinerja model. Alur penelitian secara keseluruhan diilustrasikan pada Gambar 1.

Gambar 1. Tahapan Penelitian

A. Tahapan Penelitian

Penelitian ini dilakukan melalui serangkaian tahapan yang sistematis dan terstruktur. Dimulai dari studi literatur untuk memahami dasar teoretis dan penelitian terkait, diikuti dengan pengumpulan dan persiapan data mentah dari media sosial X. Data kemudian melalui tahap pelabelan manual dan preprocessing untuk membersihkan dan mengubahnya ke dalam format yang sesuai. Selanjutnya, data dibagi menjadi dataset pelatihan, validasi, dan pengujian. Model dasar dengan arsitektur hibrida CNN-LSTM kemudian dirancang dan dilatih menggunakan data pelatihan dan divalidasi. Sebagai tahap penelitian lanjutan, optimasi hiperparameter sistematis dilakukan untuk menemukan konfigurasi model terbaik. Tahap akhir adalah pengujian dan evaluasi model akhir menggunakan data uji untuk mengukur kinerja secara komprehensif dan menarik kesimpulan.



Studi literatur komprehensif dilakukan pada tahap awal. Fokus meliputi deteksi perundungan siber pada teks berbahasa Indonesia, khususnya dari media sosial X, dan penerapan arsitektur Natural Language Processing (NLP) mulai dari machine learning konvensional hingga deep learning termasuk LSTM, CNN, dan model hibrida. Tinjauan literatur ini membangun fondasi teoretis yang kuat, mengidentifikasi kesenjangan penelitian, dan memvalidasi pemilihan metode. Sebagai hasilnya, arsitektur hibrida CNN-LSTM diadopsi. Pilihan ini didasarkan pada kemampuannya untuk menyeimbangkan ekstraksi fitur melalui CNN dengan pemahaman konteks sekuensial melalui LSTM, serta adanya kesenjangan optimasi untuk data teks informal di Indonesia.

B. Pengumpulan Data

Sumber data untuk penelitian ini adalah komentar teks berbahasa Indonesia yang berasal dari media sosial X. Akuisisi data dilakukan melalui Application Programming Interface (API) yang disediakan oleh platform, dengan memfilter komentar berdasarkan kata kunci dan hashtag yang relevan dengan topik. Untuk memastikan kualitas dan orisinalitas dataset, beberapa kriteria ditetapkan, seperti periode pengumpulan data dan memastikan hanya tweet asli, bukan retweet, yang dimasukkan dalam pengumpulan data. Contoh hasil dataset dapat dilihat pada Tabel 1.



Pelabelan manual dilakukan setelah pengumpulan data. Tahap ini sangat penting dalam supervised learning karena model memerlukan ground truth sebagai kunci jawaban untuk pembelajaran. Setiap unit data (tweet) dianalisis dan diklasifikasikan secara individual sebagai 'bullying' (nilai 1) atau 'non-bullying' (nilai 0). Objektivitas dan konsistensi pelabelan dijaga dengan merujuk pada pedoman anotasi yang telah disusun. Dataset akhir sebanyak 13.677 komentar diperoleh setelah pengumpulan dan pelabelan data. Distribusi label dalam dataset ini menunjukkan komposisi data. Contoh dataset yang telah dilabeli dapat dilihat pada Tabel 2.

C. Preprocessing Data

Data yang telah dilabeli memerlukan preprocessing sebelum pelatihan model deep learning. Tahap ini

bertujuan untuk membersihkan, menstandarkan, dan mengonversi teks ke dalam format numerik terstruktur. Tahapan selanjutnya meliputi cleaning,



case folding, tokenization, dan padding.

Text cleaning

adalah tahap preprocessing awal untuk seluruh dataset. Proses ini menghilangkan noise seperti URL, mention (@username), hashtag, dan karakter khusus yang mengganggu pembelajaran model. Semua huruf dikonversi menjadi huruf kecil melalui case folding. Langkah ini memastikan standarisasi dan konsistensi teks sebelum pemrosesan selanjutnya. Setiap kalimat dalam kolom tweet yang telah dibersihkan dipecah

Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.



menjadi unit kata (token) melalui proses tokenization. Kelas Keras Tokenizer digunakan untuk memban

gun

kosakata hanya dari data pelatihan untuk mencegah kebocoran data, dibatasi hingga 15.000 kata unik yang paling sering muncul. Model CNN-LSTM memerlukan input yang seragam, sehingga padding dilakukan pada komentar sekuensial dengan panjang yang berbeda. Setiap sekuens distandarkan menjadi 100 token di mana sekuens pendek ditambahkan dengan nilai nol (post-padding), sedangkan sekuens panjang dipotong di akhir (post-truncation).



D. Pembagian Data

Dataset akhir yang telah melalui preprocessing dibagi menjadi tiga bagian independen untuk keperluan pemodelan: 8.752 data pelatihan, 2.189 data validasi, dan 2.736 data uji. Pembagian ini setara dengan rasio 80:10:10 yang dirancang agar proporsi data pelatihan cukup besar, tetapi tetap menyisakan data yang memadai untuk validasi dan pengujian objektif. Strategi pembagian ini memastikan bahwa model dilatih pada data yang cukup sambil mempertahankan set validasi dan uji yang independen untuk evaluasi kinerja yang tidak bias.

E. Arsitektur Model Yang Diusulkan

Tahap pelatihan model adalah inti dari proses eksperimental, di mana arsitektur hibrida CNN-LSTM dilatih dari awal (from scratch). Berbeda dengan fine-tuning yang mengadaptasi model yang telah dilatih sebelumnya, pendekatan ini melatih model mulai dengan pengetahuan kosong menggunakan bobot acak untuk mempelajari pola secara eksklusif dari dataset yang telah disiapkan, seperti yang diilustrasikan pada Gambar 2.

Tabel 1. Sampel Pengumpulan Data Pengguna Teks

User1 Goblok ini Hoaxs ..Tolong Kalo Jadi Pendukung Anies terus Jangan Gilak Bikan Hoaxs .. Tolong @DivHumas_Polri ini ditindaklanjuti .

User2 Dialog Lingkungan Hidup Ustadz Abdul Somad dan Rocky Gerung di Tanjung Belit <https://t.co/NLTPgeBuka>

User3 Gibran Datang ke Rumahnya Rocky Gerung: Saya Kasih Kopi Oke You Bicara Anak Muda! <https://t.co/ewhKdZIGGa>

Tabel 2. Data Sampel Yang Diberi Label Pengguna Teks Label

User1 Goblok ini Hoaxs ..Tolong Kalo Jadi Pendukung Anies terus Jangan
Gilak Bikan Hoaxs .. Tolong @DivHumas_Polri ini ditindaklanjuti .

1

User2 Dialog Lingkungan Hidup Ustadz Abdul Somad dan Rocky Gerung di
Tanjung Belit <https://t.co/NLTPgeBuka>

0

User3 Gibran Datang ke Rumahnya Rocky Gerung: Saya Kasih Kopi Oke
You Bicara Anak Muda! <https://t.co/ewhKdZIGGa>

0

Tabel 3. Konfigurasi Model Dasar CNN-LSTM
Pengguna Teks

Embedding Layer Trainable, dimensi vektor kata = 128, dilatih dari nol

Conv1D Layer 128 filter, ukuran kernel = 5

BatchNormalization Digunakan untuk stabilitas pembelajaran

Activation ReLU

11

ejournal.unikama.ac.id | Analisis Tingkat Berfikir Kreatif Peserta Didik Dalam Menyelesaikan Soal Matematika Ditinjau Dari Adversity Quotient (AQ)
<http://ejournal.unikama.ac.id/index.php/jtst/article/download/3225/2111>

Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.

12

archive.umsida.ac.id
<https://archive.umsida.ac.id/index.php/archive/preprint/download/4336/30985/34944>

Gambar

2. CNN-LSTM

Tahap Arsitektur dasar yang diusulkan terdiri dari beberapa lapisan yang diatur secara sekuensial. Lapisan pertama adalah embedding layer dengan ukuran kosakata 15.000 kata, dimensi embedding 128, dan panjang sekuens maksimum 100 token. Lapisan ini mengonversi indeks kata menjadi representasi vektor padat. Setelah embedding layer, lapisan konvolusional satu dimensi (Conv1D) dengan 128 filter dan ukuran kernel 5 diterapkan menggunakan fungsi aktivasi ReLU dan same padding untuk mempertahankan panjang sekuens. Lapisan ini mengekstrak fitur lokal dari sekuens yang telah di-embed. Lapisan max pooling dengan ukuran pool 2 kemudian mengurangi dimensi spasial menjadi setengah, mempertahankan fitur yang paling menonjol. Fitur yang telah di-pool dimasukkan ke dalam lapisan LSTM dengan 64 unit, yang menggabungkan dropout 0,2 dan recurrent dropout 0,2 untuk mencegah overfitting. Lapisan LSTM menangkap dependensi jangka panjang dan pola sekuensial dalam teks. Akhirnya, lapisan dense output dengan satu unit dan fungsi aktivasi sigmoid menghasilkan probabilitas klasifikasi biner. Model dikompilasi menggunakan optimizer Adam dengan learning rate 0,001 dan fungsi loss binary crossentropy.



Untuk mencegah overfitting, callback
Early Stopping diimplementasikan untuk memantau validation loss dengan patience 3 epoch,

secara otomatis

mengembalikan bobot model terbaik ketika kinerja validasi berhenti meningkat. Konfigurasi hiperparameter terperinci yang digunakan disajikan pada Tabel 4.



F. Skenario Eksperimental

Untuk mengevaluasi dan menemukan arsitektur yang paling efektif, penelitian ini merancang

beberapa skenario pengujian sistematis. Model dasar yang diusulkan adalah arsitektur hibrida CNN-LSTM dengan konfigurasi yang ditunjukkan pada Tabel 3. Untuk menguji dan memvalidasi kinerja model dasar ini, serta untuk mengeksplorasi potensi peningkatannya, beberapa skenario pengujian sistematis dirancang sebagai berikut.

Skenario pertama berfokus pada evaluasi model dasar, di mana model utama dilatih dan dievaluasi menggunakan callback EarlyStopping untuk menetapkan kinerja sebagai tolok ukur utama. Skenario kedua menganalisis peran EarlyStopping melalui pendekatan komparatif di mana model dasar dilatih untuk epoch penuh tanpa EarlyStopping untuk membuktikan secara empiris dampak overfitting dan memvalidasi pentingnya penggunaan callback ini. Skenario ketiga melakukan studi komparatif lanjutan dengan serangkaian eksperimen tambahan untuk memperkaya analisis, termasuk implementasi pre-trained FastText embeddings dibandingkan dengan embedding yang dilatih dari awal, serta perbandingan dengan model lain yaitu IndoBERTweet untuk mengevaluasi sejauh mana model transformer khusus bahasa Indonesia yang dilatih pada data media sosial X dapat mengungguli atau melengkapi arsitektur CNN-LSTM dalam tugas deteksi perundungan siber. Selain itu, optimasi hiperparameter menggunakan Keras Tuner dengan strategi BayesianOptimization juga dilakukan untuk menemukan konfigurasi terbaik dan memaksimalkan kinerja model yang diuji.

G. Evaluasi Model

Evaluasi kinerja model adalah tahap akhir penelitian. Evaluasi ini secara objektif menguji kemampuan generalisasi model menggunakan dataset uji. Kinerja model dianalisis baik secara kuantitatif maupun kualitatif. Analisis kuantitatif dimulai dengan membuat Confusion Matrix untuk memvisualisasikan hasil prediksi. Dari matriks ini, metrik evaluasi standar untuk tugas klasifikasi dihitung. Accuracy mengukur persentase prediksi yang benar dari semua data, dihitung sebagai rasio true positive dan true negative terhadap jumlah total sampel. Precision mengukur akurasi prediksi positif,



merepresentasikan proporsi



kasus positif yang diprediksi dengan benar di antara semua kasus yang diprediksi positif. Recall mengukur kemampuan model untuk menemukan semua instans positif, merepresentasikan proporsi kasus positif yang diidentifikasi dengan benar di antara semua kasus positif aktual. F1-Score merepresentasikan rata-rata harmonik dari Precision dan Recall, memberikan ukuran seimbang yang mempertimbangkan baik false positive maupun false negative. Area Under the Curve (AUC) mengukur kemampuan model untuk membedakan antara kelas cyberbullying dan non-cyberbullying di berbagai threshold klasifikasi, dengan nilai yang lebih dekat ke 1,0 menunjukkan kinerja diskriminatif yang lebih baik. Selain itu, analisis kualitatif juga dilakukan dengan meninjau beberapa contoh komentar yang salah diklasifikasikan oleh model melalui analisis kesalahan untuk memahami karakteristik teks di mana model mengalami kesulitan. Hasil dari metrik ini akan menjadi dasar utama untuk menganalisis dan menyimpulkan tingkat keberhasilan dan keandalan model hibrida CNN-LSTM yang telah dikembangkan.

Tabel 4. Konfigurasi Hyperparameter Model CNN-LSTM Pengguna Teks

Ukuran Vocabulary	15.000 kata
Panjang Sequence	100 token
Dimensi Embedding (Baseline)	128
Dimensi Embedding (FastText)	300
Conv1D Filters	128
Conv1D Kernel Size	5

LSTM Units 64

Dropout Rate 0.



5

Batch Size 32

Learning Rate 0.001

Optimizer Adam

Embedding Layer Trainable,

dimensi vektor kata = 128, dilatih dari nol

Conv1D Layer 128 filter, ukuran kernel = 5

BatchNormalization Digunakan untuk stabilitas pembelajaran

Activation ReLU

IV. HASIL DAN PEMBAHASAN

Penelitian ini menyajikan evaluasi komprehensif terhadap kinerja arsitektur hybrid CNN-LSTM yang diusulkan untuk mendeteksi ucapan cyberbullying pada dataset media sosial X berbahasa Indonesia. Serangkaian eksperimen dilakukan dalam lingkungan komputasi berkemampuan tinggi menggunakan Google Colaboratory dengan akselerasi GPU Tesla T4, memanfaatkan framework TensorFlow dan Keras. Dataset yang digunakan berjumlah 13.677 titik data, dibagi secara proporsional dengan rasio 80:10:10 menjadi data pelatihan (8.752 sampel), data validasi (2.189 sampel), dan data uji (2.736 sampel). Evaluasi dilakukan berdasarkan tiga skenario pengujian sistematis yang dirancang dalam metodologi: (1) evaluasi kinerja model dasar, (2) analisis strategi pelatihan dan dampak regularisasi, dan (3) studi arsitektur perbandingan dan pembandingan dengan metode terkini.

A. Kinerja Model Dasar

Skenario pertama bertujuan untuk menetapkan kinerja dasar menggunakan model CNN-LSTM yang dilatih dengan representasi kata yang diinisialisasi secara acak dari awal dengan embedding 128 dimensi. Model ini dilengkapi dengan mekanisme Early Stopping untuk secara otomatis menghentikan pelatihan ketika tidak ada peningkatan kinerja pada data validasi, guna memperoleh bobot model optimal dan mencegah overfitting. Konfigurasi hiperparameter yang digunakan secara rinci disajikan dalam Tabel 4.



Stabilitas proses pelatihan dievaluasi melalui visualisasi kurva pembelajaran. Gambar 3a menampilkan grafik Loss pelatihan dan validasi, sementara Gambar 3b menampilkan grafik Akurasi. Kedua grafik menunjukkan konvergensi model yang efektif, di mana mekanisme Early Stopping berhasil menghentikan pelatihan pada epoch ke-5 tepat sebelum Loss validasi mulai menunjukkan tren kenaikan, sehingga mencegah model dari overfitting.



ejournal.unikama.ac.id | Analisis Tingkat Berfikir Kreatif Peserta Didik Dalam Menyelesaikan Soal Matematika Ditinjau Dari Adversity Quotient (AQ)
<http://ejournal.unikama.ac.id/index.php/jtst/article/download/3225/2111>

Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use,

distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original

publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply

with these terms.



archive.umsida.ac.id
<https://archive.umsida.ac.id/index.php/archive/preprint/download/4336/30985/34944>

Tabel

5. Rincian Hasil Evaluasi Model CNN-LSTM
Metrik Nilai

Akurasi 0.

8838

Presisi 0.88

Recall 0.88

F1 Score 0.88

AUC 0.9346

Berdasarkan hasil uji pada set uji yang disisihkan, model dasar menunjukkan kinerja yang sangat solid. Hasil evaluasi kuantitatif dirangkum dalam Tabel 5. Model berhasil mencapai tingkat akurasi 88,38% dengan skor F1 sebesar 88,00%. Untuk memahami distribusi kesalahan prediksi model secara lebih mendalam, analisis dilakukan menggunakan Confusion Matrix yang ditampilkan pada Gambar 4a. Visualisasi ini menunjukkan bahwa model memiliki kemampuan seimbang dalam mengenali kelas mayoritas dan minoritas. Secara spesifik, tingkat Recall untuk kelas positif (cyberbullying) mencapai 90%. Tingkat Recall yang tinggi ini sangat krusial dalam sistem deteksi konten berbahaya untuk meminimalkan risiko komentar bullying yang terlewat (false negatives).

Selain itu, kemampuan diskriminatif model diverifikasi melalui kurva ROC yang ditampilkan pada Gambar 4b. Nilai Area Under Curve (AUC) yang tinggi sebesar 0,9346 menunjukkan bahwa arsitektur CNN-LSTM memiliki keandalan yang sangat baik dalam memisahkan probabilitas antara kelas cyberbullying dan non-cyberbullying pada berbagai ambang batas, mengonfirmasi posisinya sebagai baseline yang kuat untuk penelitian ini.

Tabel 6. Perbandingan Model Dengan Dan Tanpa Early Stopping
Model Akurasi Presisi Recall F1 Score AUC

CNN-LSTM (Baseline) 0.8838 0.88 0.88 0.88 0.9346

(a) (b)

Gambar 3. Kurva Loss (a) dan Akurasi (b) latih validasi pada Model CNN-LSTM dengan Early Stopping.



(a) (b)

Gambar 4.

Confusion Matrix data uji (c) dan Kurva ROC data latih-validasi (d) pada Model Baseline CNN-LSTM dengan Early Stopping.



CNN-LSTM (Tanpa Early Stopping)

0.5592 0.72 0.56 0.46 0.5266

B. Pengaruh Strategi Pelatihan (Earl

y Stopping)

Untuk memvalidasi secara empiris pentingnya strategi regulasi waktu dalam pelatihan Deep Learning, dilakukan uji perbandingan antara model yang dilatih dengan mekanisme Early Stopping (Skenario 1) dan

tanpa mekanisme tersebut (Skenario 2). Hasil eksperimen menunjukkan dampak yang sangat signifikan dari strategi ini terhadap stabilitas model dan generalisasi.

Ketika pelatihan model diizinkan berjalan penuh selama 20 epoch tanpa mekanisme penghentian otomatis, fenomena overfitting yang parah terjadi. Seperti yang terlihat pada Gambar 5a (Loss) dan Gambar 5b (Accuracy), divergensi ekstrem (pemisahan) antara kinerja pelatihan dan validasi mulai terlihat mulai epoch ke-7. Kurva Loss validasi pada Gambar 5a meningkat tajam menjauh dari Loss pelatihan, menunjukkan bahwa model mulai menghafal data pelatihan (memorisasi) dan kehilangan kemampuan generalisasi pada data baru.



Tabel 7. Hasil Optimasi Hyperparameter
Hyperparameter Nilai Baseline Nilai Optimized

Conv1D Filters 128 256

LSTM Units 64 96

Dropout Rate 0.

5 0.5

Learning Rate 0.001 0.0001

Dampak overfitting ini terlihat jelas pada penurunan kinerja pengujian. Meskipun akurasi pada data pelatihan terus meningkat hingga hampir sempurna (99,38%), kinerja model pada data pengujian

(a) (b)

Gambar 5. Kurva Loss (a) dan Akurasi (b) latih validasi pada Model CNN-LSTM tanpa Early Stopping.



(a) (b)

Gambar 6.

Confusion Matrix data uji (c) dan Kurva ROC data latih-validasi (d) pada Model Baseline CNN-LSTM tanpa Early Stopping.



mengalami penurunan drastis. Akurasi anjlok menjadi 55,92%, dan nilai AUC turun menjadi 0,5266, mendekati kinerja tebakan acak. Kegagalan model dalam melakukan klasifikasi yang ben

ar akibat overfitting ditunjukkan melalui Confusion Matrix pada Gambar 6a dan kurva ROC yang hampir datar pada Gambar 6b.



Sebaliknya, dengan implementasi Early Stopping (Model Dasar), pelatihan dihentikan secara otomatis pada titik optimal (sekitar epoch 5), dan bobot model dipulihkan ke kondisi terbaiknya. Strategi ini terbukti berhasil dalam mempertahankan kemampuan generalisasi model dan menghasilkan model yang tangguh. Perbandingan kinerja kuantitatif antara dua kondisi pelatihan ini dirangkum dalam Tabel 6.

C. Pengaruh Embedding Yang Sudah Dilatih Dan Optimasi

Penelitian lebih lanjut dilakukan dalam Skenario 3 untuk meningkatkan kinerja model dasar melalui penggunaan representasi kata yang telah dilatih sebelumnya (pre-trained FastText embeddings) dan optimasi hiperparameter otomatis.

Secara empiris, penggunaan FastText Embeddings (dimensi 300) terbukti memberikan peningkatan dalam stabilitas konvergensi dan kemampuan diskriminasi model dibandingkan dengan inisialisasi acak. Nilai AUC meningkat dari 0,9346 menjadi 0,9559, dan akurasi sedikit meningkat menjadi 88,49%. Peningkatan ini menunjukkan bahwa representasi vektor kata yang dilatih pada korpus Wikipedia Indonesia yang besar membantu model memahami konteks semantik kata-kata langka atau slang yang sering ditemukan di media sosial, yang mungkin tidak terwakili dengan baik jika hanya mengandalkan data pelatihan yang terbatas.

Di sisi lain, upaya optimasi hiperparameter menggunakan metode Bayesian Optimization (Keras Tuner) menghasilkan konfigurasi arsitektur dengan kapasitas lebih besar, yaitu 256 filter pada lapisan CNN dan 96 unit pada lapisan LSTM, tetapi dengan learning rate yang lebih konservatif (0,0001). Meskipun model yang dioptimalkan ini menghasilkan AUC yang sangat kompetitif (0,9507), peningkatan akurasi dan F1-Score relatif marginal dibandingkan dengan baseline. Hal ini menunjukkan bahwa konfigurasi baseline yang diusulkan sudah cukup robust dan optimal untuk karakteristik dataset ini, dan penambahan kompleksitas model tidak selalu berbanding lurus dengan peningkatan akurasi yang signifikan.



Dinamika pelatihan, yang divisualisasikan melalui kurva akurasi dan loss, menunjukkan bahwa model mulai menunjukkan tanda-tanda overfitting setelah epoch ke-2, ditandai dengan peningkatan loss validasi dan penurunan akurasi validasi. Namun, mekanisme Early Stopping secara efektif intervensi dengan menghentikan pelatihan pada epoch 5 dan mengembalikan bobot model ke kondisi optimalnya pada epoch 2. Nilai AUC validasi terbaik sebesar 0,9289 dicapai pada checkpoint ini, mengonfirmasi bahwa strategi Early Stopping berhasil mempertahankan kemampuan generalisasi model. Hasil optimasi hiperparameter disajikan dalam Tabel 7.

Proses optimasi mengidentifikasi konfigurasi dengan filter Conv1D yang digandakan (256 vs. 128) dan unit LSTM yang ditingkatkan (96 vs. 64), serta learning rate yang dikurangi (0,0001 vs. 0,001). Menariknya, model yang dioptimalkan mencapai kinerja yang sebanding dengan model dasar dalam hal akurasi (88,05% vs. 88,38%) dan F1-Score (0,88 untuk keduanya), dengan penurunan akurasi yang sedikit (-0,33 poin persentase). Namun, serupa dengan eksperimen FastText, model yang dioptimalkan menunjukkan kinerja diskriminatif yang lebih unggul dengan AUC 0,9507, mewakili peningkatan relatif 1,61% dibandingkan dengan model dasar. Hal ini menunjukkan bahwa peningkatan kapasitas model (5,21 juta parameter vs. 4,69 juta) dan learning rate yang lebih konservatif memungkinkan kalibrasi probabilitas yang lebih baik di berbagai ambang klasifikasi, meskipun perkiraan akurasi titik tetap serupa. Skor F1 yang sebanding menunjukkan bahwa konfigurasi baseline dan yang dioptimalkan mencapai trade-off precision-recall yang serupa, menunjukkan bahwa hiperparameter baseline sudah cukup sesuai untuk tugas ini. Peningkatan AUC yang marginal, meskipun secara statistik signifikan, datang dengan biaya kompleksitas komputasi dan waktu pelatihan yang lebih tinggi.



D. Studi Arsitektur Perbandingan (Ablation Study)

Untuk memvalidasi efektivitas desain hybrid yang menggabungkan CNN dan LSTM, dilakukan perbandingan komprehensif (ablation study) terhadap berbagai variasi arsitektur Deep Learning. Studi ini membandingkan model hybrid dengan model tunggal (CNN, LSTM, Bi-LSTM) serta variasi hybrid lainnya (CNN-BiGRU). Semua model dalam studi ini menggunakan FastText Embeddings untuk memastikan perbandingan yang adil.

Berdasarkan rekapitulasi hasil eksperimen, model tunggal Bi-LSTM dengan Attention mencatat kinerja tertinggi di antara semua varian Deep Learning klasik yang diuji, dengan pencapaian F1-Score



88,77% dan AUC 0,9577. Keunggulan model ini menunjukkan bahwa mekanisme attention dan pemaha

man

konteks bidirectional sangat efektif dalam menangkap fitur linguistik cyberbullying yang seringkali bergantung padakonteks kalimat yang lengkap.



Meskipun demikian, model hybrid CNN-BiGRU dan CNN-LSTM masih menunjukkan kinerja yang sangat kompetitif dengan perbedaan kinerja kurang dari 0,5% dibandingkan dengan model terbaik. Lebih jauh lagi, arsitektur hybrid ini menawarkan keunggulan dalam efisiensi komputasi. Lapisan konvolusional (CNN) di awal arsitektur berfungsi untuk mengurangi dimensi fitur secara efektif sebelum diproses oleh lapisan recurrent, sehingga mempercepat waktu pelatihan dan inference dibandingkan dengan model Bi-LSTM murni yang harus memproses urutan kata secara lengkap satu per satu.

Arsitektur CNN saja menunjukkan kinerja yang luar biasa dalam hal precision (90,71%), mengindikasikan kemampuan kuat model ini dalam mengidentifikasi instance positif dengan false positive minimal. Konvergensi cepat yang dicapai oleh CNN (berhenti pada epoch 6) mencerminkan efisiensi komputasinya dalam mengekstraksi fitur n-gram lokal. Namun, nilai recall yang sedikit lebih rendah (85,93%) menunjukkan keterbatasan model ini dalam menangkap konteks global kalimat yang lengkap dibandingkan dengan model berbasis sekuensial. Sebaliknya, arsitektur LSTM murni mencapai kinerja yang seimbang di semua metrik tetapi menuntut biaya komputasi yang signifikan, memerlukan 27 epoch untuk konvergen. Hal ini mengonfirmasi sifat iteratif dari pemrosesan sekuensial di LSTM yang lebih lambat dibandingkan dengan pemrosesan paralel di CNN.



Dalam kategori model hybrid, varian CNN-BiLSTM dan CNN-BiGRU menunjukkan kinerja yang sangat sebanding dengan akurasi sekitar 88%. Temuan menarik adalah kesetaraan kinerja antara BiLSTM dengan akurasi 0,8816 dan BiGRU dengan nilai akurasi 0,8834 dengan jumlah parameter yang identik. Hal ini mengindikasikan bahwa mekanisme gating yang lebih sederhana di GRU dapat menyamai efektivitas arsitektur LSTM yang lebih kompleks untuk tugas ini, tetapi dengan efisiensi operasional yang lebih baik. Model yang diusulkan, CNN-LSTM (Baseline + FastText), menempati posisi keseimbangan optimal antara kinerja (F1-Score 88%) dan efisiensi (konvergensi dalam 6 epoch) dengan ukuran model yang relatif kompak.

Temuan paling signifikan adalah munculnya arsitektur BiLSTM + Attention sebagai model dengan kinerja terbaik secara keseluruhan, mencapai F1-Score 88,77% dan AUC tertinggi sebesar 0,9577. Kemampuan mekanisme attention untuk secara dinamis memfokuskan bobot pada kata-kata penting (salient) terbukti memberikan keunggulan diskriminatif, terutama dalam meminimalkan kesalahan prediksi. Meskipun demikian, keunggulan ini datang dengan peningkatan kompleksitas model dan waktu pelatihan yang lebih lama. Mengingat perbedaan kinerja yang marginal di antara semua arsitektur canggi ini (rentang perbedaan akurasi hanya 0,66%), pemilihan model akhir untuk implementasi praktis harus mempertimbangkan faktor efisiensi komputasi dan kendala deployment, di mana arsitektur hybrid CNN-BiGRU atau CNN-LSTM menawarkan alternatif yang lebih ringan dibandingkan dengan model berbasis Attention yang berat. Hasil komprehensif dari ablation study ini dirangkum dalam Tabel 8.

Tabel 8. Rekapitulasi Kinerja Komprehensif Seluruh Model
Model Akurasi Presisi Recall F1-Score AUC

IndoBERTweet 0.915 0.92 0.92 0.915 -

BiLSTM + Attention 0.8871 0.8945 0.8810 0.8877 0.9577

CNN-LSTM (+ FastText) 0.8849 0.89 0.88 0.88 0.9559

CNN 0.8841 0.9071 0.8593 0.8825 0.9566

CNN-BiGRU 0.8834 0.8834 0.8834 0.8834 0.9566

LSTM 0.88 0.88 0.88 0.88 0.9596

CNN-BiLSTM 0.8816 0.8820 0.8816 0.8815 0.9548

CNN-LSTM (Base Line
Early Stopping)

0.8838 0.88 0.88 0.88 0.9346

CNN-LSTM (Keras
Tuner)

0.8805 0.88 0.88 0.88 0.9507

SVM + TF-IDF 0.7818 0.8266 0.7818 0.7747 -

CNN-LSTM (Tanpa
EarlyStopping)

0.5592 0.72 0.56 0.46 0.5266

E.



Benchmarking dengan Model State-of-the-Art



ejournal.unikama.ac.id | Analisis Tingkat Berfikir Kreatif Peserta Didik Dalam Menyelesaikan Soal Matematika Ditinjau Dari Adversity Quotient (AQ)
<http://ejournal.unikama.ac.id/index.php/jtst/article/download/3225/2111>

Copyright

© 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use,

distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original

publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply

with these terms.



archive.umsida.ac.id
<https://archive.umsida.ac.id/index.php/archive/preprint/download/4621/33078/37498>

Untuk

menempatkan kinerja arsitektur CNN-LSTM yang diusulkan dalam konteks yang lebih luas, dilakukan analisis perbandingan mendalam terhadap berbagai pendekatan, mulai dari metode Machine Learning klasik, variasi arsitektur Deep Learning, hingga model transformer state-of-the-art. Evaluasi ini bertujuan untuk memvalidasi keunggulan relatif dari metode yang diusulkan dibandingkan dengan pendekatan yang telah ada sebelumnya.

Sebagai baseline benchmark untuk metode konvensional, Support Vector Machine (SVM) dengan fitur TF-IDF dipilih untuk mewakili pendekatan Machine Learning klasik. Hasil evaluasi menunjukkan kesenjangan kinerja yang substansial antara pendekatan klasik dan Deep Learning. Model CNN-LSTM yang diusulkan mencatat akurasi sebesar 88,38%, mewakili peningkatan signifikan sebesar 10,2 poin persentase dibandingkan dengan akurasi SVM yang hanya mencapai 78,18%. Lebih kritis lagi, peningkatan F1-Score sebesar 10,53 poin persentase pada model Deep Learning menunjukkan keseimbangan yang jauh lebih baik antara precision dan recall. Analisis mendalam terhadap kinerja SVM mengungkapkan pola asimetris, di mana meskipun model dapat mencapai precision tinggi untuk kelas cyberbullying, nilai recall-nya sangat rendah yaitu 60,32%. Hal ini mengindikasikan bahwa model klasik cenderung gagal dalam mendeteksi hampir setengah dari total komentar cyberbullying yang sebenarnya, sebuah kelemahan fatal dalam sistem deteksi yang memprioritaskan keselamatan pengguna. Sebaliknya, arsitektur CNN-LSTM menunjukkan kinerja yang seimbang pada kedua metrik dengan nilai 0,88, mengonfirmasi superioritas representasi fitur otomatis dalam menangkap nuansa bahasa yang kompleks.

Selanjutnya, model yang diusulkan dibandingkan dengan IndoBERTweet, model transformer berbasis BERT yang secara khusus dilatih pada korpus Twitter Indonesia skala masif, mewakili state-of-the-art saat ini. Hasil pengujian memposisikan IndoBERTweet sebagai model dengan kinerja tertinggi di semua metrik evaluasi, dengan pencapaian akurasi dan F1-Score sebesar 91,59%. Keunggulan ini dapat dikaitkan dengan proses pre-training yang ekstensif yang memberikan model pemahaman mendalam tentang pola linguistik, slang, dan struktur informal yang khas dari media sosial Indonesia. Namun, temuan yang patut dicatat adalah bahwa model Deep Learning terbaik yang dikembangkan dalam penelitian ini, yaitu BiLSTM dengan Attention, mencapai akurasi 88,71%, hanya tertinggal sekitar 2,88 poin persentase dari IndoBERTweet.



Kesenjangan kinerja yang relatif sempit ini menjadi sangat signifikan ketika mempertimbangkan disparitas sumber daya komputasi antara kedua pendekatan. Model IndoBERTweet memiliki kompleksitas parameter yang mencapai sekitar 110 juta, jauh melebihi model BiLSTM dengan Attention yang hanya memiliki sekitar 5,32 juta parameter. Selain itu, IndoBERTweet memerlukan 50 epoch pelatihan untuk mencapai konvergensi optimal, sementara model BiLSTM dan CNN-LSTM hanya memerlukan kurang dari 10 epoch. Hal ini menunjukkan bahwa arsitektur berbasis RNN dan CNN yang diusulkan menawarkan efisiensi komputasi yang jauh lebih tinggi, kecepatan inference yang lebih baik, dan kelayakan deployment pada perangkat dengan sumber daya terbatas, tanpa mengorbankan akurasi secara drastis.

Evaluasi komprehensif terhadap semua varian model yang diuji, seperti yang dirangkum dalam Tabel 10, menunjukkan bahwa arsitektur hybrid CNN-LSTM dan variannya secara konsisten menempati posisi kinerja yang kompetitif antara baseline klasik dan model transformer yang berat. Model CNN-BiGRU, misalnya, mencapai akurasi 88,34% dengan efisiensi tinggi, setara dengan model LSTM murni tetapi dengan struktur yang lebih kompak. Sementara itu, penggunaan FastText embeddings di CNN-LSTM terbukti meningkatkan kemampuan diskriminatif model, tercermin dalam peningkatan nilai AUC menjadi 0,9559. Temuan-temuan ini memvalidasi bahwa arsitektur hybrid fundamental yang menggabungkan ekstraksi fitur lokal dari CNN dengan pemodelan sekuensial dari LSTM adalah pendekatan yang robust dan efektif untuk karakteristik linguistik cyberbullying media sosial Indonesia, menawarkan keseimbangan optimal antara kinerja deteksi dan efisiensi sumber daya.

Tabel 9. Contoh Kesalahan Prediksi Model CNN-LSTM Pada Data Uji
No Teks Komentar (Asli) Label Asli Prediksi

Model	Status
Kesalahan	
1 masuk tan anjing Non-Bullying (0)	Bullying (1) False Positive
2 goodness parent actually watching something apple tv brilliant yet another streaming service	Non-Bullying (0)
3 aduh lak onok rawan disalahgunakan men Non-Bullying (0)	Bullying (1) False Positive
4 jomok bawa barak iya oppa kdm iya anjr panggil sayang mau pacar panggil sayang	Bullying (1) Non- Bullying (0)
	False Negative

first tonton mv kambek jeoang skz ya aju
main vocalnya oek

5 halo non moot ijin ikut ya selamat lolo
seleksi kerja tysm moga halus pity sylus nya

bandel ya

F. Analisis Kesalahan Kualitatif

Sesuai dengan desain evaluasi untuk melengkapi metrik kuantitatif, dilakukan tinjauan mendalam terhadap sampel data uji yang salah diklasifikasikan oleh model yang diusulkan (CNN-LSTM). Analisis ini bertujuan untuk mengidentifikasi karakteristik linguistik spesifik dan pola teks yang menjadi tantangan bagi arsitektur CNN-LSTM dalam mengenali cyberbullying.



Tabel 9 menyajikan contoh representatif dari komentar yang mengalami kesalahan klasifikasi, dikategorikan menjadi False Positive (komentar non-bullying diprediksi sebagai bullying) dan False Negative (komentar bullying gagal terdeteksi).

Berdasarkan analisis di Tabel 9, beberapa keterbatasan spesifik dari arsitektur CNN-LSTM dalam menangani data media sosial teridentifikasi. Pertama, Dominasi Fitur Lokal (Efek CNN) di mana kesalahan False Positive pada contoh no. 1 menunjukkan bahwa lapisan CNN sangat sensitif terhadap kata kunci negatif (seperti "anjing"). Meskipun lapisan LSTM bertugas menangkap konteks, fitur n-gram lokal yang diekstraksi oleh CNN terkadang terlalu dominan, sehingga satu kata kasar dapat menyebabkan kalimat netral diprediksi sebagai bullying. Kedua, Keterbatasan Pemahaman Bahasa Code-Mixing di mana model menunjukkan kesulitan dalam menangani kalimat yang sepenuhnya dalam Bahasa Inggris atau campuran (contoh no. 2). Hal ini mengindikasikan bahwa meskipun menggunakan FastText embeddings, representasi semantik untuk frasa asing yang panjang masih belum optimal dalam arsitektur ini. Ketiga, Kompleksitas Pragmatis dan Slang di mana kesalahan False Negative pada contoh no. 4 menyoroti tantangan dalam mendeteksi bullying yang disamarkan dalam slang atau kalimat sarkastik. Arsitektur CNN-LSTM mungkin kehilangan nuansa implisit ini karena operasi pooling di CNN mengurangi dimensi fitur, berpotensi menghilangkan informasi urutan kata yang halus namun penting untuk mendeteksi sarkasme.

Temuan kualitatif ini mengonfirmasi bahwa meskipun model CNN-LSTM memiliki kinerja statistik yang tinggi (F1 lebih besar dari 88%), pengembangan lebih lanjut perlu berfokus pada mekanisme penanganan ambiguitas kata dan pengayaan data pelatihan dengan variasi slang dan bahasa campuran yang lebih beragam. Analisis kualitatif ini memberikan wawasan berharga bahwa peningkatan kinerja di masa depan tidak hanya bergantung pada kompleksitas arsitektur, tetapi juga pada pengayaan data pelatihan untuk mencakup variasi linguistik yang lebih luas.

V. KESIMPULAN

Penelitian ini mengusulkan dan mengevaluasi arsitektur hibrida CNN-LSTM untuk deteksi perundungan siber di media sosial X berbahasa Indonesia menggunakan dataset sebanyak 13.677 komentar. Penelitian ini memperoleh beberapa temuan kunci melalui eksperimen sistematis di tiga skenario. Pertama, mekanisme Early Stopping bersifat kritis untuk mencegah overfitting. Model yang dilatih tanpa mekanisme ini mengalami degradasi kinerja yang parah dengan akurasi menurun hingga 55,92%, sementara model dengan Early Stopping mencapai akurasi 88,38% dan AUC 0,9346. Kedua, embedding FastText yang telah dilatih sebelumnya meningkatkan kemampuan diskriminatif model, meningkatkan AUC menjadi 0,9559. Ketiga, studi ablasi komparatif mengungkapkan bahwa BiLSTM dengan Attention mencapai kinerja tertinggi di antara model pembelajaran mendalam dengan F1-Score 88,77%, meskipun hibrida CNN-LSTM dan CNN-BiGRU menawarkan efisiensi komputasi yang lebih baik dengan akurasi kompetitif masing-masing 88,38% dan 88,34%.

Benchmarking terhadap metode mutakhir menunjukkan bahwa model yang diusulkan secara signifikan mengungguli SVM klasik dengan margin lebih dari 10 poin persentase. IndoBERTweet mencapai akurasi tertinggi sebesar 91,59%, namun arsitektur CNN-LSTM yang diusulkan mencapai 97% dari kinerja tersebut dengan 95% lebih sedikit parameter (5,32 juta berbanding 110 juta), sehingga menjadikannya sangat cocok untuk skenario deployment dengan sumber daya terbatas. Penelitian ini memiliki keterbatasan yang mencakup pembatasan dataset pada media sosial X, pendekatan klasifikasi biner, dan tantangan dalam menangani pencampuran kode serta perundungan implisit seperti sarkasme. Penelitian selanjutnya perlu



memperluas dataset ke berbagai platform, mengimplementasikan klasifikasi multi-kelas untuk tingkat keparahan, mengeksplorasi pendekatan multimodal yang menggabungkan teks dengan analisis visual.

VI. UCAPAN TERIMA KASIH

Penelitian ini dapat terlaksana berkat dukungan dan fasilitas yang disediakan oleh Universitas Muhammadiyah Sidoarjo. Penulis menyampaikan rasa terima kasih yang sebesar-besarnya kepada dosen pembimbing atas bimbingan yang sangat berharga, masukan yang konstruktif, serta wawasan ilmiah yang diberikan selama proses penelitian. Bimbingan tersebut sangat berperan penting dalam membentuk metodologi dan analisis studi ini.

Kami juga mengapresiasi tim peneliti dan rekan-rekan yang telah berkontribusi dalam proses pengumpulan dan anotasi data. Sumber daya komputasi dalam penelitian ini didukung oleh Google Colaboratory dengan akselerasi GPU Tesla T4, menggunakan kerangka kerja (framework) TensorFlow dan Keras. Karya ini didedikasikan untuk mendukung pengembangan lingkungan digital yang lebih aman di Indonesia..



DAFTAR PUSTAKA

Ahmad, S.
R., Insani, N., & Salim, M. (2024).

23

doi.org | Jurnal Informasi Teknologi dan Pendidikan
<https://doi.org/10.24036/jtip.v17i1.807>

Analysis of Cyberbullying on Social Media Using A Comparison of Naïve Bayes, Random Forest, and SVM Algorithms.

Jurnal Teknologi Informasi Dan Pendidikan, 17(1), 75–86. <https://doi.org/10.24036/jtip.v17i1.807>

24

doi.org | Performance Analysis of NLP-Based Machine Learning Algorithms in Cyberbullying Detection
<https://doi.org/10.18185/erzifbed.1474112>

Akar, F. (2024). Performance Analysis of NLP-Based Machine Learning Algorithms in Cyberbullying Detection. Erzincan

Üniversitesi Fen Bilimleri Enstitüsü Dergisi, 17(2), 445–459. <https://doi.org/10.18185/erzifbed.1474112>

Akter, F., Rabbi, F., & Chowdhury, R. R. (2025). Cyberbullying Detection on Social Media Platforms Utilizing Different Machine Learning Approaches. 186(61), 40–50.

Alqahtani, A. F., & Ilyas, M. (2024). An Ensemble-Based Multi-Classification

25

digilib.uin-suka.ac.id | OPTIMALISASI MODEL TRANSFORMER UNTUK IDENTIFIKASI CYBERBULLYING DI MEDIA SOSIAL PADA KONDISI KETERBATASAN SUMBER DAYA KO...
https://digilib.uin-suka.ac.id/id/eprint/72946/1/23206052010_BAB_IV-atau-V_DAFTAR-PUSTAKA.pdf

Machine Learning Classifiers Approach to Detect Multiple Classes of Cyberbullying. Machine Learning

and Knowledge Extraction, 6(1), 156–170. <https://doi.org/10.3390/make6010009>

Andika, A. J., Kristian, Y., & Setiawan, E. I. (2023).

26

doi.org | Deteksi Komentar Cyberbullying Pada YouTube Dengan Metode Convolutional Neural Network – Long Short-Term Memory Network (CNN-LSTM)
<https://doi.org/10.34148/teknika.v12i3.677>

Deteksi Komentar Cyberbullying Pada YouTube Dengan Metode Convolutional Neural Network

– Long Short-Term Memory Network (CNN-LSTM). Teknika, 12(3), 183–188. <https://doi.org/10.34148/teknika.v12i3.677>

27

doi.org | Enhancing Cyberbullying Detection on Platform 'X' Using IndoBERT and Hybrid CNN-LSTM Model
<https://doi.org/10.52436/1.jutif.2025.6.2.4321>

A Hybrid Deep Learning Approach Leveraging Word2Vec Feature
Expansion for Cyberbullying Detection in Indonesian

Twitter. Ingenierie Des Systemes d'Information,
28(4), 887–895. <https://doi.org/10.18280/isi.280410>

Cahyo, D. D. N., Handayani, R., Lestari, V. B., & Febriani, S. (2025). JITE (

28

garuda.kemdiktisaintek.go.id | Garuda - Garba Rujukan Digital
<https://garuda.kemdiktisaintek.go.id/documents/detail/5302713>

Journal of Informatics and
Telecommunication Engineering) Sentiment Analysis of Public Opinion on Online Gambling Through



9(July), 99–115.

El Koshiry, A. M., Eliwa, E. H. I., El-Hafeez,

T. A., & Khairy, M. (2024).

29

doi.org | Detecting cyberbullying using deep learning techniques: a pre-trained glove and focal loss technique [PeerJ]
<https://doi.org/10.7717/peerj-cs.1961>

Detecting cyberbullying using deep
learning techniques: a pre-trained glove and focal loss technique.

PeerJ Computer
Science, 10, 1–33.

<https://doi.org/10.7717/peerj-cs.1961>

Faritha Banu, J., Rajeshwari, S. B., Kallimani, J. S., Vasanthi, S., Buttar, A. M., Sangeetha,

M., & Bhargava, S.

30

doi.org | IASC | Modeling of Hyperparameter Tuned Hybrid CNN and LSTM for Prediction Model
<https://doi.org/10.32604/iasc.2022.024176>

(2022). Modeling of Hyperparameter Tuned Hybrid CNN and LSTM for Prediction Model. Intelligent
Automation

and Soft Computing, 33(3), 1393–1405. <https://doi.org/10.32604/iasc.2022.024176>

Hafiza, A. A., & Setiawan, E. B. (2025). Enhancing Cyberbullying Detection on Platform “X” Using IndoBERT
and Hybrid CNN-LSTM Model. Jurnal Teknik Informatika (Jutif), 6(2), 655–672.
<https://doi.org/10.52436/1.jutif.2025.6.2.4321>

Harish, D., Manimaran, M., Jayashakthi, V. P., & Alamelu, M. (2023). Automatic Detection of Cyberbullying on
Social Media Using Machine Learning.

31

dx.doi.org | Penerapan Kalman Filter Pada Pembacaan Sensor Loadcell Berbasis PLC Siemens S7-1200
<http://dx.doi.org/10.33795/elkolind.v11i13.6158>

2nd International Conference on Advancements in Electrical,
Electronics, Communication, Computing and Automation, ICAECA 2023,
13(3), 185–189.

<https://doi.org/10.1109/ICAECA56562.2023.10201149>

Hasan, M. T., Hossain, M. A. E., Mukta, M. S. H., Akter, A., Ahmed, M., & Islam, S. (2023). A Review on Deep-
Learning-Based Cyberbullying Detection. Future Internet, 15(5), 1–47.
<https://doi.org/10.3390/fi15050179>

Iryani, L., Junaidi,



J., Paisal, P., Purba, M., Umilizah, N., Bakhtiar, B., & Ani, N. (2025).



doi.org | MODEL OF CYBERBULLYING DETECTION ON SOCIAL MEDIA USING MULTI-LABEL DEEP LEARNING: A COMPARATIVE STUDY
<https://doi.org/10.33480/jitk.v10i4.6004>

Model of Cyberbullying

Detection on Social Media Using Multi-Label Deep Learning: a Comparative

Study. JITK (Jurnal Ilmu

Pengetahuan Dan Teknologi Komputer), 10(4), 742–748. <https://doi.org/10.33480/jitk.v10i4.6004>



ejournal.unikama.ac.id | Analisis Tingkat Berfikir Kreatif Peserta Didik Dalam Menyelesaikan Soal Matematika Ditinjau Dari Adversity Quotient (AQ)
<http://ejournal.unikama.ac.id/index.php/jtst/article/download/3225/2111>

Copyright © 2018 Author [s]. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use,

distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original

publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply

with these
terms.

Masbadi Hatullah Nurnaryo, R., Mulaab, M., Oktavia Suzanti, I., Abdul Fatah, D., Cahyani, A. D., & Ayu Mufarroha, F. (2022).



doi.org | Deteksi Komentar Cyberbullying Pada YouTube Dengan Metode Convolutional Neural Network – Long Short-Term Memory Network (CNN-LSTM)
<https://doi.org/10.34148/teknika.v12i3.677>

Deteksi Cyberbullying Pada Data Tweet Menggunakan Metode Random Forest

Dan Seleksi Fitur Information

Gain. Jurnal Simantec, 11(1), 33–40.

<https://doi.org/10.21107/simantec.v11i1.17256>



doi.org | Klasifikasi Jenis Kejahatan berdasarkan Teks Amar Putusan Pengadilan Hukum Pidana KUHP menggunakan IndoBERT
<https://doi.org/10.29408/edumatic.v9i2.30326>

Maulana, M. D., & Aditya, C. S. K. (2025). Perbandingan IndoBERT dan Bi-LSTM Dalam Mendeteksi Pelanggaran Undang-Undang ITE.



myskripsi.ums.ac.id
https://myskripsi.ums.ac.id/media/skripsi/proposal/2025/10/24/Proposal_Skripsi_Almira_Putri_Wibowo_L200220068.pdf

Science And Information Technology
(SINTECH),

8(1), 52–59.

<https://doi.org/10.31598>

Mitra, S., Tasnim,



Document from another user
Comes from another group

A Framework to Detect

and Prevent

Cyberbullying from Social Media by Exploring Machine Learning

Algorithms. 6th International

Conference on Computer, Communication, Chemical, Materials and Electronic Engineering, IC4ME2 2021, February 2023. <https://doi.org/10.1109/IC4ME253898.2021.9768450>

Nasution, M. A. S., & Setiawan, E. B. (2023).



doi.org | Enhancing Cyberbullying Detection on Platform 'X' Using IndoBERT and Hybrid CNN-LSTM Model
<https://doi.org/10.52436/1.jutif.2025.6.2.4321>

Enhancing Cyberbullying Detection on Indonesian Twitter;

Leveraging FastText for Feature Expansion and Hybrid Approach Applying CNN and

BiLSTM. Revue

d'Intelligence Artificielle, 37(4), 929–936. <https://doi.org/10.18280/ria.370413>



, Amrin, J. F., Anwar, Md. M., & Sarker,

I. H. (2025). AI Enabled User-Specific Cyberbullying Severity Detection with Explainability. 1–25. <http://arxiv.org/abs/2503.10650>

Purkayastha, B. S., Rahman, M., & Talukdar, T. I. (2025). Advancing Cyberbullying Detection :

**39**

doi.org | Advancing Cyberbullying Detection: A Hybrid Machine Learning and Deep Learning Framework for Social Media Analysis
<https://doi.org/10.5220/0013436200003929>

A Hybrid Machine Learning and Deep Learning Framework for Social Media Analysis.

2(Iceis), 348–355.
<https://doi.org/10.5220/0013436200003929>

Ramadhan,



V., Pambudi, A., Informatika, T., & Sukabumi, U. M. (2024). IMPLEMENTASI ALGORITMA CONVOLUTIONAL NEURAL NETWORK UNTUK MENGIDENTIFIKASI BERITA HOAKS BERBAHASA INDONESIA. 8(5), 10945–10952.

Rosid, M. A., Siahaan, D., & Saikhu, A.

(2024). Sarcasm Detection in Indonesian-English Code-Mixed Text Using Multihead Attention-Based Convolutional and Bi-Directional GRU. IEEE Access, June, 137063–137079.
<https://doi.org/10.1109/ACCESS.2024.3436107>

**40**

www.mendeley.com | Deteksi Komentar Cyberbullying pa... preview & related info | Mendeley
<https://www.mendeley.com/catalogue/77ff20d6-3b28-3ab3-ace1-773d7100fa70/>

Santoso, H., Putri, R. A., & Sahbandi, S. (2023). Deteksi Komentar Cyberbullying pada Media Sosial Instagram Menggunakan Algoritma Random Forest. Jurnal Manajemen Informatika (JAMIKA), 13(1), 62–72.
<https://doi.org/10.34010/jamika.v13i1.9303>

Shubham Singh. (2025). Twitter Statistics 2025: How Many People Use X [New Data].
<https://www.demandsage.com/twitter-statistics/>

Smote, C. D. I



. (2025). Analisis komparatif klasifikasi sentimen pengguna aplikasi investasi menggunakan algoritma hybrid cnn-lstm, cnn-gru dengan implementasi smote. 6, 112–118.

Tapaskar, V., & J,

M. C. T. (2024). INTERNATIONAL JOURNAL OF A Multi-Algorithm Approach for Cyberbullying Detection. 7(10). <https://doi.org/10.15680/IJMRSET.2024.0710027>