

Klasifikasi Sentimen *Tweet* dengan Arsitektur *Hybrid Transformers-CNN* pada *Platform Twitter*

Safrizal Ardana Ardiyansa¹, Abdi Negara Guci², Jemmy Febryan³, Dian Alhusari⁴,
Haidar Ahmad Fajri⁵

safrizal@student.ub.ac.id¹, abdinegara783@gmail.com², danieljemmyfebryan@gmail.com³,

alhusari2@gmail.com⁴, fajrihaidar07@gmail.com⁵

^{1,2,3,4,5}Department of Mathematics, Brawijaya University, Malang, Indonesia

^{1,2,3,4,5}Braincore Indonesia, Jakarta, Indonesia

Informasi Artikel

Diterima : 17 Jan 2025

Direvisi : 25 Apr 2025

Disetujui : 9 Jun 2025

Kata Kunci

Klasifikasi sentimen,
Transformer-CNN,
Twitter

Abstrak

Twitter atau yang sekarang dikenal sebagai X merupakan platform populer yang digunakan untuk mengekspresikan opini terkait tren terbaru, sehingga Twitter atau X menjadi sumber data yang sangat berharga untuk penelitian analisis sentimen. Volume data yang sangat besar membuat analisis manual tidak praktis karena memerlukan waktu yang lama dan sumber daya manusia, sehingga diperlukan proses otomatisasi klasifikasi sentimen melalui *machine learning*. *Machine learning* dapat digunakan untuk mengklasifikasi sentimen dalam skala besar secara cepat dan akurat dengan memanfaatkan pola. Model *machine learning* seperti Transformers-CNN menunjukkan performa paling unggul dengan akurasi mencapai 85,71% pada data uji dan 99,90% pada data latih. Akurasi pada data uji tersebut lebih baik dari arsitektur lain yaitu LSTM, CNN, BERT, Transformers-LSTM, dan LSTM-CNN dengan akurasi sebesar 84,73%; 82,27%; 77,34%; 85,71%; 84,24% berturut-turut. Transformers-CNN juga memiliki waktu pelatihan 30,17 menit yang lebih singkat daripada Transformers-LSTM, namun lebih lama daripada arsitektur lainnya.

Keywords

Sentiment classification,
Transformer-CNN, Twitter

Abstract

Twitter, now known as X, is a popular platform used to express opinions on the latest trends, making it a valuable source of data for sentiment analysis research. The huge volume of data makes manual analysis impractical because it requires a long time and human resources, so it is necessary to automate the sentiment classification process through machine learning. Machine learning can be used to classify sentiment on a large scale quickly and accurately by utilising patterns. Machine learning models such as Transformers-CNN show the most superior performance with accuracy reaching 85.71% on test data and 99.90% on training data. The accuracy on the test data was better than other architectures namely LSTM, CNN, BERT, Transformers-LSTM, and LSTM-CNN with accuracies of 84.73%; 82.27%; 77.34%; 85.71%; 84.24% respectively. Transformers-CNN also has a training time of 30.17 minutes which is shorter than Transformers-LSTM, but longer than the other architectures.

A. Pendahuluan

Platform media sosial seperti Twitter atau yang sekarang dikenal dengan nama X menyediakan ruang bagi untuk berbagi opini atau pemikiran [1], serta pendapat dari setiap pengguna agar dapat terhubung, berkomunikasi, dan berkontribusi pada topik tertentu. Twitter menyediakan cara yang efektif untuk berkomunikasi dan berinteraksi, memungkinkan munculnya pengaruh di antara para pengguna, yang dapat mengubah opini masyarakat [2]. Twitter atau X merupakan salah satu media sosial yang sangat populer di seluruh Indonesia dengan jumlah pengguna mencapai 78 juta di Indonesia [3]. Popularitas Twitter meningkat sebesar 752% pada tahun 2008, mencapai hampir 4,5 juta pengguna di seluruh dunia pada bulan Januari 2009 [4]. Penggunaan Twitter juga telah mengalami peningkatan 10 kali lipat sejak tahun 2009, dengan total lebih dari 500 juta pengguna dengan 316 juta pengguna aktif [5].

Meningkatnya pengguna Twitter atau X dapat dimanfaatkan oleh para pemaku kepentingan. Organisasi nirlaba yang lebih kecil dengan komunitas *online* yang lebih besar juga cenderung menggunakan Twitter sebagai alat komunikasi dengan para pemangku kepentingan [6]. Twitter juga dimanfaatkan oleh para organisasi untuk memberikan dukungan dan kejelasan selama krisis emosional dengan menganalisis tweet dari pemangku kepentingan [7]. Pelaku kepentingan juga memanfaatkan informasi yang tersedia pada platform Twitter atau X untuk mengetahui perubahan psikologi serta perilaku seseorang [1].

Platform Twitter memiliki fitur *short* atau *tweet* yang terdiri dari 140 karakter [2]. *Tweet* telah menjadi subjek dari banyak penelitian analisis sentimen, namun prediksi klasifikasi sentimen Twitter yang canggih berkinerja buruk, dengan akurasi klasifikasi yang dilaporkan di bawah 70% [10], padahal sangat banyak organisasi, perusahaan, dan pemerintah, yang tertarik untuk memahami opini dan perasaan publik tentang produk atau layanan [11]. Proses klasifikasi *tweet* merupakan cara yang mudah untuk mengevaluasi pendapat konsumen tentang suatu produk atau layanan [12]. *Tweet* juga telah memberikan wawasan yang berharga tentang isu-isu yang berkaitan dengan bisnis dan masyarakat [10], [13], [14]. Hal ini dapat membantu pemaku kepentingan dalam pengambilan keputusan strategis seperti perencanaan pemasaran, strategi pembuatan produk, dan peningkatan layanan pelanggan.

Analisis sentimen manual pada *tweet* dilakukan dengan membaca dan menilai setiap *tweet* secara langsung oleh individu untuk menentukan emosi, opini, atau sentimen yang terkandung, yaitu apakah bersifat positif, negatif, atau netral. Proses tersebut membutuhkan pemahaman mendalam karena *tweet* sering menggunakan bahasa yang tidak formal, singkatan, emotikon, atau bahkan sarkasme yang sulit diinterpretasikan secara konsisten oleh manusia [15]. Jumlah data yang sangat besar pada sentimen analisis pada *tweet* memerlukan ribuan hingga jutaan tweet sehingga membuat proses tersebut memerlukan waktu dan tenaga yang banyak serta rentan terhadap bias dari subjektivitas individu.

Pendekatan berbasis teknologi diperlukan untuk mengatasi tantangan tersebut. Teknik komputasi modern dapat mempercepat proses klasifikasi sentimen dan meningkatkan konsistensi hasil analisis [16]. *Machine Learning* (ML) adalah salah satu teknik yang berhasil digunakan untuk melakukan pekerjaan secara otomatis dengan cepat dan waktu yang lebih efisien [15]. ML telah banyak digunakan dalam mengatasi berbagai permasalahan di berbagai bidang, seperti di bidang kesehatan

untuk mendiagnosis serangan jantung [17], di bidang keamanan untuk mendeteksi serangan jaringan internet [18], bahkan bidang politik untuk mendapatkan insight mengenai gagasan para calon wakil presiden [19], [20], sehingga ML juga memiliki potensi yang tinggi dalam proses klasifikasi sentimen.

Natural Language Processing (NLP) merupakan salah satu bagian dari ML yang berfokus dalam memecahkan persoalan *text analytics* [10]. NLP telah berhasil dalam memecahkan berbagai permasalahan seperti *sentiment analysis* [21], [22], [23] dan *text classification* [24], [25]. Klasifikasi sentimen dengan teknik NLP memungkinkan pengambilan keputusan cepat berdasarkan opini masyarakat di media sosial. Hal ini tidak hanya menghemat waktu dan sumber daya, tetapi juga dapat meningkatkan kemampuan untuk memahami pandangan dan opini publik dengan lebih mendalam.

Proses analisis sentimen dengan bahasa Indonesia sudah seringkali dilakukan menggunakan arsitektur model ML konvensional seperti *Naïve Bayes*, *maximum entropy*, *Support Vector Machine* (SVM), dan *decision tree* [26], [27], [28], [29], [30]. Penelitian analisis sentimen dalam bahasa Inggris sudah menerapkan metode deep learning yaitu *Convolutional Neural Network* (CNN) yang berbeda pada penelitian sentimen analisis dalam bahasa Indonesia, dan mendapatkan hasil jauh lebih baik dibandingkan model klasik seperti *Naïve Bayes* dengan peningkatan presisi sebesar 7%, *recall* sebesar 8%, dan *f₁-score* sebesar 9% [30]. Hal ini dapat terjadi karena arsitektur CNN dapat mengekstraksi fitur dari informasi yang tersembunyi dan sulit dideteksi. CNN juga menggunakan konvolusi dari layer sebelumnya sehingga data diekstraksi menjadi fitur dan CNN akan mempertimbangkan hubungan antar fitur tersebut [31]. Penelitian tentang klasifikasi sentimen dengan CNN untuk *tweet* berbahasa Indonesia masih jarang dilakukan [32], sehingga pada penelitian ini akan melibatkan arsitektur metode CNN untuk melakukan sentimen analisis pada *tweet* untuk mendeteksi emosi atau perasaan dari masyarakat berdasarkan *tweet* yang tersedia di Twitter atau X.

CNN dapat digunakan dalam klasifikasi sentimen karena mampu mengekstraksi fitur lokal dari teks, seperti pola kata atau frasa tertentu, dengan cepat dan efisien [33], namun CNN memiliki keterbatasan dalam memahami hubungan antar kata yang lebih kompleks atau konteks global dalam teks. Arsitektur Transformer dalam *deep learning* dapat diimplementasikan untuk mengatasi kekurangan pada CNN [34]. Transformer menggunakan mekanisme *self-attention* yang memungkinkan model menangkap hubungan kata dalam konteks global, baik yang berdekatan maupun berjauhan [35]. Transformer juga memiliki kelemahan, seperti kebutuhan komputasi yang besar dan proses pelatihan yang lambat, terutama untuk *dataset* besar sehingga Transformer juga sering kali terlalu kompleks untuk data sederhana sehingga kurang efisien untuk kasus tertentu [36]. Penggabungan arsitektur CNN dengan Transformer berpotensi menjadi solusi efektif untuk mengatasi kekurangan masing-masing. Arsitektur CNN mampu mengekstraksi fitur lokal dengan efisien, sedangkan Transformer melengkapi analisis dengan pemahaman konteks global. Kombinasi tersebut menghasilkan model yang lebih kuat, efisien, dan akurat dalam klasifikasi sentimen termasuk untuk teks berbahasa Indonesia.

B. Metode Penelitian

Artikel berikut ini merupakan sebuah penelitian eksperimental dengan tujuan untuk mencari arsitektur ML terbaik yang mampu mengklasifikasi *tweet* berbahasa

Indonesia pada Twitter dengan akurasi yang tinggi. Hasil penelitian akan disajikan dalam bentuk visualisasi grafik dari performa pada berbagai jenis arsitektur pada setiap iterasi, serta tabel perbandingan performa dari jenis arsitektur model.

Beberapa arsitektur yang dipilih dalam artikel ini adalah BERT, LSTM, CNN, serta beberapa arsitektur gabungan seperti Transformer-LSTM, LSTM-CNN, dan Transformer-CNN untuk dibandingkan. Sumber *dataset* yang digunakan pada di dalam artikel ini merupakan *tweet* dari platform Twitter atau X. *Dataset* di dalam penelitian ini berjumlah sebanyak 6137 *tweet* dengan 5 jenis label berbeda yaitu *joy*, *sadness*, *fear*, *love*, dan *anger*. Beberapa sampel data yang digunakan di dalam artikel ini dapat dilihat pada Tabel 1.

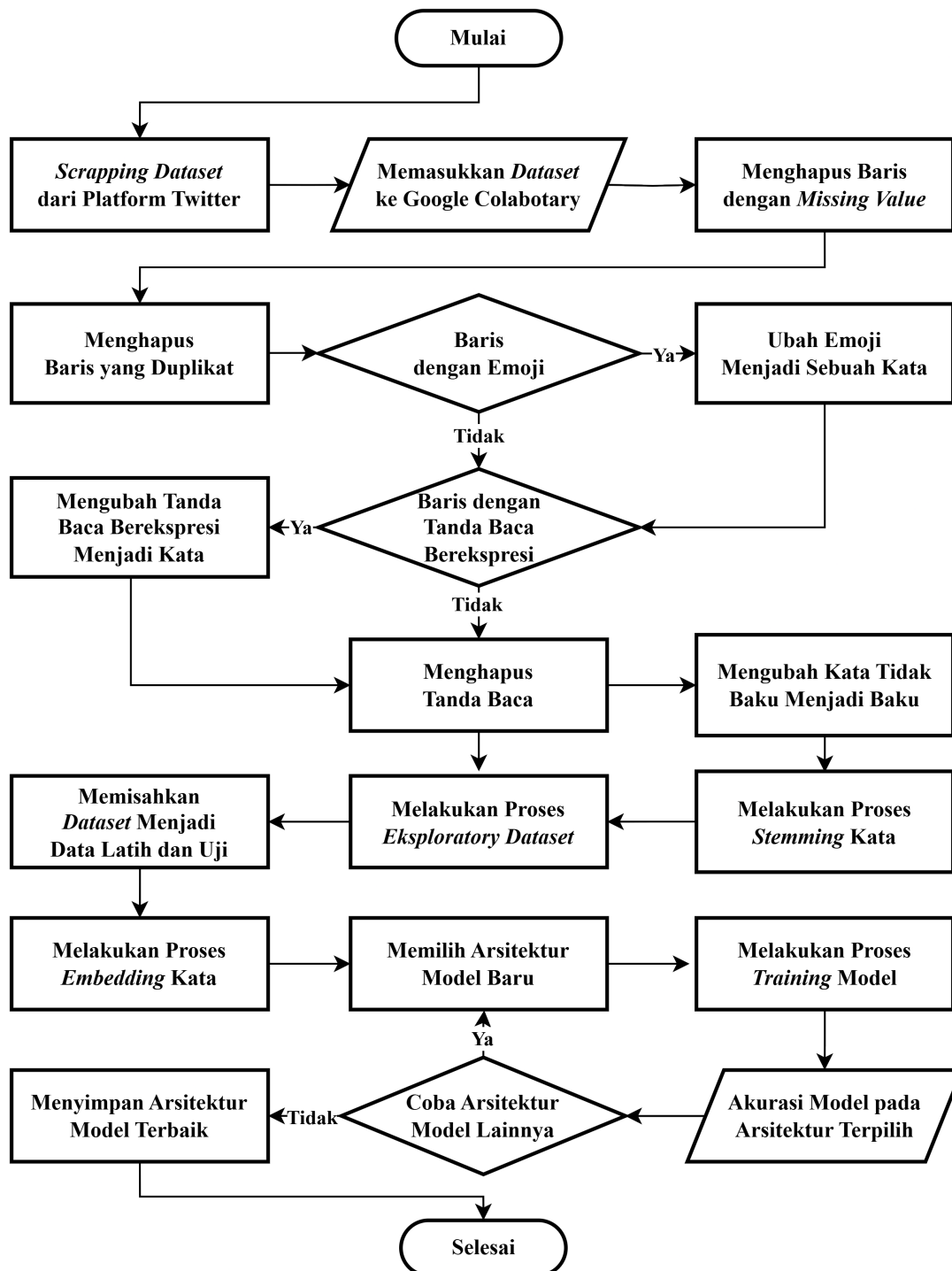
Tabel 1. Sampel *dataset* penelitian

Label	Data
<i>anger</i>	Momen di mana kamu merasa seperti semua yang kamu inginkan adalah meluapkan kemarahanmu tanpa ampun. 🤬🔥 #UnleashTheAnger
<i>anger</i>	Momen di mana kamu merasa seperti amarahmu meresap ke dalam dirimu seperti racun. 🤬🧪 #ToxicAnger
<i>joy</i>	Saat-saat ketika kamu membagikan cerita lucu atau pengalaman lucu dan mengundang tawa dari orang lain. 😂👏 #HumorSharing

Alur di dalam penelitian ini terbagi menjadi beberapa tahapan seperti pada Gambar 1. Hal pertama yang dilakukan adalah melakukan proses *scrapping tweet* melalui platform Twitter untuk mendapatkan *dataset* tersebut. *Dataset* selanjutnya dilakukan proses pelabelan dan dimasukkan ke dalam Google Colabulatory sebelum melakukan proses *preprocessing*. *Preprocessing dataset* dilakukan melalui beberapa tahapan, yaitu penghapusan baris yang mengandung *missing* value dan duplikat. Proses selanjutnya dilakukan pengecekan apakah data tersebut mengandung tanda baca berekspresi senang, sedih, atau marah seperti “:D”, “:(”, “-_-”, dan “:’”. Jika baris mengandung ekspresi-ekspresi berikut, maka akan dilakukan pengubahan menjadi kata senang, sedih, bosan, marah, dan sebagainya. Tanda baca yang tidak mengandung makna selanjutnya akan dihapus. Proses sama juga dilakukan jika ada baris data mengandung emoji seperti “😊”, “😞”, “😡”, dan “😏”. Emoji-emoji yang mengandung makna diubah menjadi kata senang, sedih, marah dan sebagainya.

Preprocessing selanjutnya adalah dengan melakukan proses pengubahan kata yang tidak baku menjadi baku, serta proses *stemming* kata untuk mengubah kata berimbuhan menjadi kata dasar. Proses *eksplorasi dataset* juga dilakukan untuk mengetahui ketidakseimbangan label dan dilakukan proses *embedding* kata untuk mengubah kata-kata baku tersebut menjadi sebuah angka, sehingga menghasilkan vektor pada setiap baris. *Dataset* setelah itu dibagi menjadi dua bagian, yaitu data latih dan data uji dengan proporsi 80% dan 20% secara berturut-turut, kemudian dipilih salah satu arsitektur model. Arsitektur model yang dipilih akan dilakukan proses pelatihan pada data latih, kemudian diuji menggunakan data uji untuk mengetahui akurasi dari arsitektur tersebut.

Proses pemilihan arsitektur terus dilakukan hingga didapatkan performa dari beberapa arsitektur model. Waktu komputasi saat pelatihan juga dicatat untuk mengetahui keefektifan model tersebut. Arsitektur model dengan akurasi tertinggi kemudian disimpan dan dapat digunakan untuk mengklasifikasikan *tweet* ke dalam lima label dengan akurat dan cepat.



Gambar 1. Diagram Alur Penelitian

Konfigurasi spesifikasi *hardware* yang digunakan dalam penelitian ini, serta rincian ukuran data yang diambil dari Twitter atau X ditampilkan pada Tabel 2. Tabel tersebut menginformasikan konfigurasi perangkat yang digunakan di dalam penelitian untuk memastikan ketepatan dan reliabilitas hasil yang telah disajikan.

Tabel 2. Konfigurasi *Hardware*

Nama	Parameter
------	-----------

<i>Memory</i>	8 GB
<i>Processor</i>	Intel Pentium Gold G6400
<i>GPU</i>	-
<i>Language</i>	Python 3
<i>Framework</i>	Jupyter Notebook

C. Hasil dan Pembahasan

Bagian ini akan menjelaskan mengenai hasil dan pembahasan pada penelitian ini yang terdiri dari hasil preprocessing teks, informasi yang didapatkan melalui proses Eksplorasi Data *Analysis* (EDA), dan hasil perbandingan dari arsitektur model.

1. Hasil *Preprocessing* Teks

Beberapa tahapan *preprocessing* yang telah dilakukan di dalam penelitian ini adalah mengubah tanda baca berekspresi menjadi kata, menghapus tanda baca yang tidak digunakan, mengubah emoji berekspresi menjadi kata sehingga didapatkan sebuah *dataset* baru yang sudah berbentuk teks yang tidak mengandung emoji dan tanda baca. Salah satu contoh hasil *preprocessing* yang dilakukan untuk sampel data pada Tabel 1 dapat dilihat pada Tabel 3. Terlihat pada Tabel 3 bahwa sampel data setelah *preprocessing* tersebut sudah tidak mengandung tanda baca, huruf besar, dan emoji.

Tabel 3. Sampel *dataset* setelah *preprocessing*

Label	Data
<i>anger</i>	momen di mana kamu merasa seperti semua yang kamu inginkan adalah meluapkan kemarahanmu tanpa ampun frustrasi marah
<i>anger</i>	momen di mana kamu merasa seperti amarahmu meresap ke dalam dirimu seperti racun marah frustrasi
<i>joy</i>	saat saat ketika kamu membagikan cerita lucu atau pengalaman lucu dan mengundang tawa dari orang lain tertawa megafon

2. Hasil Eksplorasi Data *Analysis*

Proses Eksplorasi Data *Analysis* (EDA) telah dilakukan pada artikel ini untuk mendapatkan informasi banyaknya token atau kata berbeda pada data latih, yaitu sebanyak 11.138 token, dan setelah dilakukan *preprocessing* seperti pengubahan ke huruf kecil, menghilangkan *stopword*, menghilangkan tanda baca, dan pengubahan menjadi kata tidak baku menjadi baku pada *dataset*, maka didapatkan banyaknya token sebesar 5036. Banyaknya presentase kata data uji yang termuat di data latih dapat dilihat pada Tabel 4.

Kata pada data uji pada Tabel 4 yang terkandung juga pada data latih diwarnai dengan *font* putih, sedangkan kata yang tidak terkandung diwarnai dengan warna merah. Kolom pertama pada Tabel 4 merupakan presentase banyaknya kata yang terkandung juga dalam data uji. Informasi presentase ini cukup penting karena jika data uji memiliki presentase yang rendah, maka arsitektur model ML tidak dapat mengklasifikasikan *tweet* tersebut karena model tidak pernah melihat kata tersebut sebelumnya.

Perhatikan bahwa pada Tabel 4 bahwa presentase telah diurutkan dari yang terkecil hingga yang terbesar. Presentase dengan nilai terkecil yaitu sebesar 60%,

sehingga banyaknya token pada data latih sudah cukup digunakan untuk melatih model dalam memprediksi sentimen.

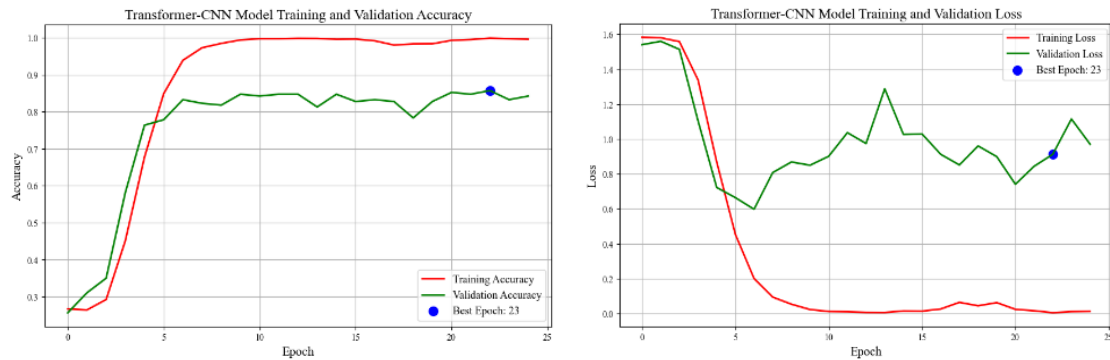
Tabel 4. Presentase kata pada data uji yang termuat pada data latih

Persentase	Data
60%	alhamdulillah happi makin aktif nak 8monthdedek nazhiaputriadhitya sehaterrussayang cintaudek
60%	terima kasih yang sudah datang spare football hari ini unitedindonesia unitedtogeth uici unitedindramayu unitedindonesiachapterindramayu indramayu regenc uniform resourc locat
63%	kurang raku apa kita masih belum datang 1piza ladinadianti pasta piza
66%	happi wed elizasekar moga langgeng sampai kakek nenek wed weddingparti weddingtradision adatjawa
73%	kadang suka kasihan sama tukang parkir bingung begitu kalau markirin jet acu sake gede jadi sekarang prefer naik ufo saja deh lebih minimal tempat parkir keluhanborjui
⋮	⋮
95%	tidak apa apa miss semangat bawa mereka untuk konser lanjut ya dan terima kasih prjuangannya tahun ini semangat miss kita semua tunggu
100%	aku rasa seperti aku ada di ambang ledak karena rasa amarah ini
100%	nama guna nama guna nama guna nama guna waduh cari panggung terlalu jauh bung junjung saja tidak ada yang langsung kayak anda eh by the way anda punya junjung siapa ya
100%	aku sempat tonton berita waktu di thailand cuma karena tidak erti bahasa jadi aku diam saja nde baru sampai bal baru tau yang aku tonton berita tentang itu

3. Perbandingan Arsitektur Model

Model arsitektur yang diusulkan dalam penelitian ini yaitu Transformer-CNN. Arsitektur Transformer-CNN dibandingkan dengan beberapa model lain, seperti Finetuned BERT, LSTM, CNN, Transformer-LSTM, dan LSTM-CNN. Beberapa metrik yang digunakan di dalam penelitian ini adalah *training time* dan *validation accuracy* yang merupakan akurasi pada data uji. Nilai *validation accuracy* tertinggi dengan waktu pelatihan tersingkat merupakan arsitektur model yang terbaik.

Visualisasi dari grafik akurasi dan *loss function* pada saat pelatihan arsitektur Transformer-CNN untuk setiap epoch atau iterasi terlihat pada Gambar 2. Terlihat bahwa Transformer-CNN juga memiliki tingkat konvergensi yang cukup cepat yaitu sudah konvergen pada saat iterasi ke-5. Proses pelatihan terus dilakukan sampai iterasi ke-25 hingga didapatkan akurasi tertinggi pada data uji yaitu sebesar 85.71% ketika iterasi ke-23. Hal tersebut mengakibatkan parameter bobot pada arsitektur Transformer-CNN yang digunakan pada penelitian ini adalah bobot pada saat iterasi ke-23.



Gambar 2. Grafik akurasi dan *loss* pada arsitektur *Hybrid* Transformer-CNN

Berdasarkan pengujian dari keseluruhan model yang telah dilatih pada Tabel 5, terlihat bahwa arsitektur Transformer-CNN memiliki akurasi yang lebih tinggi yaitu sebesar 85.71% pada data uji dan 99.90% pada data latih. Akurasi pada data uji tersebut lebih baik dari arsitektur lain yaitu LSTM, CNN, BERT, Transformers-LSTM, dan LSTM-CNN dengan akurasi sebesar 84,73%; 82,27%; 77,34%; 85,71%; 84,24% berturut-turut. Transformers-CNN juga memiliki waktu pelatihan 30,17 menit yang lebih singkat daripada Transformers-LSTM, namun lebih lama daripada arsitektur lainnya. Meskipun waktu pelatihan membutuhkan waktu sekitar 30.17 menit, hasil yang diperoleh tergolong cukup baik. Hal ini menunjukkan bahwa Transformer-CNN tidak hanya lebih akurat dan juga cukup efisien jika dibandingkan model lain yang diuji. Model ini juga konsisten dalam performanya pada berbagai *dataset*, menunjukkan kemampuannya dalam menangani variasi data. Oleh karena itu, Transformer-CNN menjadi pilihan arsitektur yang paling baik untuk klasifikasi sentimen yang membutuhkan akurasi tinggi.

Tabel 5. Perbandingan performa arsitektur model

Arsitektur Model	Total Parameter	Total Epoch	Training Accuracy	Validation Accuracy	Training Time (minutes)
Fine-tuned BERT	430,735	25	98.31%	77.34%	0.70
LSTM	4,123,781	25	96.02%	84.73%	13.73
CNN	3,307,169	25	98.80%	82.27%	16.05
Transformer-LSTM	5,222,149	25	98.91%	84.73%	43.34
LSTM-CNN	3,851,113	25	94.21%	84.24%	14.48
Transformer-CNN	4,440,553	25	99.90%	85.71%	30.17

D. Kesimpulan

Media sosial seperti Twitter yang kini dikenal sebagai X, menjadi platform populer untuk berbagi pendapat dengan pengguna mencapai 78 juta di Indonesia. Analisis sentimen di Twitter membantu organisasi memahami opini publik, yang berguna dalam pengambilan keputusan strategis, namun klasifikasi sentimen secara manual sulit dilakukan karena volume data besar dan kompleksitas bahasa informal. Machine Learning (ML) dapat mengatasi hal ini dengan mengotomatisasi proses tersebut dengan akurasi yang tinggi. CNN adalah arsitektur dalam ML yang mampu mengekstraksi fitur lokal dari teks, sayangnya CNN memiliki keterbatasan dalam memahami konteks global. Arsitektur Transformer dengan mekanisme self-

attention dapat menangkap hubungan konteks di dalam yang lebih luas, namun memerlukan komputasi besar. Kombinasi CNN-Transformer menawarkan solusi yang efisien dalam klasifikasi sentimen.

Hasil penelitian menunjukkan bahwa arsitektur Transformer-CNN mencapai akurasi pada data uji yang tertinggi sebesar 85.71% dan 99.90% pada data latih. Arsitektur tersebut lebih baik daripada Finetuned BERT (77.34%), LSTM (84.73%), CNN (82.27%), Transformer-LSTM (84.73%), dan LSTM-CNN (84.24%). Hasil ini menunjukkan keunggulan dalam akurasi, walaupun membutuhkan waktu pelatihan yang lebih lama yaitu sebesar 30.17 menit jika dibandingkan dengan LSTM (13.73 menit) dan CNN (16.05 menit). Transformer-CNN juga sudah konvergen pada saat iterasi ke-5. Hal ini menunjukkan bahwa arsitektur tersebut juga memiliki tingkat konvergensi yang cepat, sehingga menjadikannya model terbaik dalam klasifikasi sentimen berdasarkan *dataset* yang digunakan.

E. Ucapan Terima Kasih

Penulis mengucapkan terima kasih sebesar-besarnya kepada Departemen Matematika, Universitas Brawijaya yang memfasilitasi keberlangsungan penelitian ini, sehingga penelitian ini dapat berjalan dengan lancar tanpa adanya suatu kendala apapun.

F. Referensi

- [1] M. M. Mostafa, "More than words: Social networks' text mining for consumer brand sentiments," *Expert Systems with Applications*, vol. 40, no. 10, pp. 4241–4251, Aug. 2013, doi: 10.1016/j.eswa.2013.01.019.
- [2] X. Cheng, C. Chen, W. Zhang, and Y. Yang, "5G-Enabled Cooperative Intelligent Vehicular (5GenCIV) Framework: When Benz meets Marconi," *IEEE Intelligence System.*, vol. 32, no. 3, pp. 53–59, May 2017, doi: 10.1109/MIS.2017.53.
- [3] S. M. Fani, R. Santoso, and S. Suparti, "Penerapan text mining untuk melakukan clustering data tweet akun Bibli pada media sosial Twitter menggunakan K-Means clustering," *Jurnal Gaussian*, vol. 10, no. 4, pp. 583–593, Dec. 2021, doi: 10.14710/j.gauss.v10i4.30409.
- [4] D. Giustini and M. D. Wright, "Twitter: An introduction to microblogging for health librarians," *Journal of the Canadian Health Libraries Association*, vol. 30, no. 1, p. 11, Jul. 2014, doi: 10.5596/c09-009.
- [5] H. Efstathiades, D. Antoniadis, G. Pallis, M. D. Dikaiakos, Z. Szlavik, and R.-J. Sips, "Online social network evolution: Revisiting the Twitter graph," in *2016 IEEE International Conference on Big Data (Big Data)*, Washington DC, USA: IEEE, Dec. 2016, pp. 626–635. doi: 10.1109/BigData.2016.7840655.
- [6] M. D. M. Gálvez-Rodríguez, C. Caba-Pérez, and M. López-Godoy, "Drivers of Twitter as a strategic communication tool for non-profit organizations," *Internet Research*, vol. 26, no. 5, pp. 1052–1071, Oct. 2016, doi: 10.1108/IntR-07-2014-0188.
- [7] J. Brummette and H. Fussell Sisco, "Using Twitter as a means of coping with emotions and uncontrollable crises," *Public Relations Review*, vol. 41, no. 1, pp. 89–96, Mar. 2015, doi: 10.1016/j.pubrev.2014.10.009.

- [8] A. H. Alamoodi *et al.*, "Sentiment Analysis and Its Applications in Fighting COVID-19 and Infectious Diseases: A Systematic Review," *Expert Systems with Applications*, vol. 167, p. 114155, Apr. 2021, doi: 10.1016/j.eswa.2020.114155.
- [9] Y. Qi and Z. Shabrina, "Sentiment analysis using Twitter data: A comparative application of lexicon and Machine-Learning-based approaches," *Social Network Analysis and Mining*, vol. 13, no. 1, p. 31, Feb. 2023, doi: 10.1007/s13278-023-01030-x.
- [10] D. Zimbra, A. Abbasi, D. Zeng, and H. Chen, "The State-of-the-art in Twitter sentiment analysis: A review and benchmark evaluation," *ACM Transactions on Management Information Systems*, vol. 9, no. 2, pp. 1–29, Jun. 2018, doi: 10.1145/3185045.
- [11] M. Al-Smadi, O. Qawasmeh, M. Al-Ayyoub, Y. Jararweh, and B. Gupta, "Deep Recurrent Neural Network vs. support Vector Machine for aspect-based sentiment analysis of Arabic hotels' reviews," *Journal of Computational Science*, vol. 27, pp. 386–393, Jul. 2018, doi: 10.1016/j.jocs.2017.11.006.
- [12] S. S. Jacob and R. Vijayakumar, "Sentimental analysis over twitter data using clustering based Machine Learning algorithm," *Journal of Ambient Intelligence and Humanized Computing*, Jan. 2021, doi: 10.1007/s12652-020-02771-9.
- [13] B. J. Jansen, M. Zhang, K. Sobel, and A. Chowdury, "Twitter power: Tweets as electronic word of mouth," *Journal of the American Society for Information Science and Technology*, vol. 60, no. 11, pp. 2169–2188, Nov. 2009, doi: 10.1002/asi.21149.
- [14] B. Gleason, "Occupy wall street: Exploring informal learning about a social movement on Twitter," *American Behavioral Scientist*, vol. 57, no. 7, pp. 966–982, Jul. 2013, doi: 10.1177/0002764213479372.
- [15] B. Mahesh, "Machine Learning algorithms - a review," *International Journal of Science and Research*, vol. 9, no. 1, pp. 381–386, Jan. 2020, doi: 10.21275/ART20203995.
- [16] V. Umarani, A. Julian, and J. Deepa, "Sentiment analysis using various Machine Learning and Deep Learning techniques," *Journal of the Nigerian Society of Physical Sciences*, pp. 385–394, Nov. 2021, doi: 10.46481/jnsps.2021.308.
- [17] S. A. Ardiyansa, N. C. Maharani, S. Anam, and E. Julianto, "Optimizing heart attack diagnosis using Random Forest with Bat Algorithm and greedy crossover technique," *BAREKENG: Jurnal Ilmu Matematika dan Terapan*, vol. 18, no. 2, pp. 1053–1066, May 2024, doi: 10.30598/barekengvol18iss2pp1053-1066.
- [18] S. A. Ardiyansa, E. Julianto, N. C. Maharani, and H. A. Fajri, "Network attack classification using Neural Network based imputation technique," *The Indonesian Journal of Computer Science*, vol. 13, no. 5, Oct. 2024, doi: 10.33022/ijcs.v13i5.4349.
- [19] E. Julianto, S. A. Ardiyansa, T. S. Tanzi, H. A. Fajri, and N. C. Maharani, "Implementasi deepface untuk analisis ekspresi wajah pada debat calon wakil presiden Republik Indonesia," *Jurnal Sistem dan Teknologi Informasi*, vol. 06, no. 3, pp. 196–206, 2024.
- [20] R. F. Khoiroh, E. Julianto, S. A. Adiyansa, H. A. Fajri, A. A. R. Yasa, and B. Sangapta, "Implementasi speech recognition whisper pada debat calon wakil presiden

- Republik Indonesia," *EXPLORE: Jurnal Informatika dan Komputer*, vol. 14, no. 2, 2024, doi: <https://doi.org/10.35200/ex.v14i2.115>.
- [21] M. Parmar, B. Maturi, J. M. Dutt, and H. Phate, "Sentiment analysis on interview transcripts: An application of NLP for quantitative analysis," in *2018 International Conference on Advances in Computing, Communications and Informatics (ICACCI)*, Bangalore: IEEE, Sep. 2018, pp. 1063–1068. doi: 10.1109/ICACCI.2018.8554498.
- [22] M. A. Shafin, M. M. Hasan, M. R. Alam, M. A. Mithu, A. U. Nur, and M. O. Faruk, "Product review sentiment analysis by using NLP and Machine Learning in Bangla language," in *2020 23rd International Conference on Computer and Information Technology (ICCIT)*, DHAKA, Bangladesh: IEEE, Dec. 2020, pp. 1–5. doi: 10.1109/ICCIT51783.2020.9392733.
- [23] Z. Kastrati, F. Dalipi, A. S. Imran, K. Pireva Nuci, and M. A. Wani, "Sentiment analysis of students' feedback with NLP and Deep Learning: A systematic mapping study," *Applied Sciences*, vol. 11, no. 9, p. 3986, Apr. 2021, doi: 10.3390/app11093986.
- [24] R. Maria, "Machine Learning text classification model with NLP approach," in *Proceedings of the 3d International Conference Computational Linguistics And Intelligent Systems*, Kharkiv, Ukraine, Apr. 2019.
- [25] V. Dogra *et al.*, "A complete process of text classification system using state-of-the-art NLP models," *Computational Intelligence and Neuroscience*, vol. 2022, pp. 1–26, Jun. 2022, doi: 10.1155/2022/1883698.
- [26] S. Anastasia and I. Budi, "Twitter sentiment analysis of online transportation service providers," in *2016 International Conference on Advanced Computer Science and Information Systems (ICACSIS)*, Malang, Indonesia: IEEE, Oct. 2016, pp. 359–365. doi: 10.1109/ICACSIS.2016.7872807.
- [27] I. P. Windasari, F. N. Uzzi, and K. I. Satoto, "Sentiment analysis on Twitter posts: An analysis of positive or negative opinion on GoJek," in *2017 4th International Conference on Information Technology, Computer, and Electrical Engineering (ICITACEE)*, Semarang: IEEE, Oct. 2017, pp. 266–269. doi: 10.1109/ICITACEE.2017.8257715.
- [28] V. A. Fitri, R. Andreswari, and M. A. Hasibuan, "Sentiment analysis of social media Twitter with case of anti-LGBT campaign in Indonesia using Naïve Bayes, decision tree, and Random Forest algorithm," *Procedia Computer Science*, vol. 161, pp. 765–772, 2019, doi: 10.1016/j.procs.2019.11.181.
- [29] A. A. Lutfi, A. E. Permanasari, and S. Fauziati, "Sentiment analysis in the sales review of Indonesian marketplace by utilizing Support Vector Machine," *Journal of Information Systems Engineering and Business Intelligence*, vol. 4, no. 1, p. 57, Apr. 2018, doi: 10.20473/jisebi.4.1.57-64.
- [30] S. Yoo, J. Song, and O. Jeong, "Social media contents-based sentiment analysis and prediction system," *Expert Systems with Applications*, vol. 105, pp. 102–111, Sep. 2018, doi: 10.1016/j.eswa.2018.03.055.
- [31] S. Rani and P. Kumar, "Deep Learning based sentiment analysis using Convolution Neural Network," *Arabian Journal for Science and Engineering*, vol. 44, no. 4, pp. 3305–3314, Apr. 2019, doi: 10.1007/s13369-018-3500-z.
- [32] H. Juwiantho, E. I. Setiawan, J. Santoso, and M. H. Purnomo, "Sentiment analysis Twitter Bahasa Indonesia berbasis Word2Vec menggunakan Deep

- Convolutional Neural Network,” *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 7, no. 1, pp. 181–188, Feb. 2020.
- [33] S. Luo, Y. Gu, X. Yao, and W. Fan, “Research on text sentiment analysis based on Neural Network and ensemble learning,” *Revue D Intelligence Artificielle*, vol. 35, no. 1, pp. 63–70, Feb. 2021, doi: 10.18280/ria.350107.
- [34] D. H. Putra and G. W. I. Noto, “Implementasi generative pre-trained Transformers 2 untuk mengoreksi kesalahan penulisan Bahasa Indonesia pada dokumen jurnal,” presented at the TECHNOPEX-2023, Tangerang Selatan: Institut Teknologi Indonesia, October 25, 2023.
- [35] S. R. Choi and M. Lee, “Transformer architecture and attention mechanisms in genome data analysis: A Comprehensive review,” *Biology*, vol. 12, no. 7, p. 1033, Jul. 2023, doi: 10.3390/biology12071033.
- [36] T. Lin, Y. Wang, X. Liu, and X. Qiu, “A survey of transformers,” *AI Open*, vol. 3, pp. 111–132, 2022, doi: 10.1016/j.aiopen.2022.10.001.