

TEXT-BASED EMOTION SENTIMENT ANALYSIS DENGAN PENDEKATAN DEEP LEARNING : AKUISISI TIKTOK TERHADAP TOKOPEDIA

Angga Pornama, Agussalim, Seftin Fitri Ana Wati

Sistem Informasi, UPN "Veteran" Jawa Timur

Jl. Rungkut Madya, Gn. Anyar, Kec. Gn. Anyar, Surabaya, Jawa Timur 60294

anggapornama@gmail.com

ABSTRAK

Akuisisi 75% saham Tokopedia oleh Tiktok memunculkan banyak pro dan kontra. Di samping dampak positif, terdapat kekhawatiran terkait potensi monopoli, dominasi pasar, persaingan tidak sehat, keamanan data pengguna, hingga geopolitik dari China. Pengakuisisian Tokopedia juga disinyalir merupakan langkah TikTok untuk mendapatkan izin dagang *social commerce* meraka. Dengan adanya perdebatan opini publik, dilakukan analisis sentimen untuk mengetahui respon masyarakat terhadap isu terkait. Diadopsi sentimen berdasar emosi guna memahami konteks dari data. Adapun data yang digunakan berasal dari komentar *platform* YouTube yang diambil melalui proses *scraping* di mana didapatkan 3041 baris data setelah proses labeling secara manual. Deep learning dengan model CNN, LSTM, CNN-LSTM, dan LSTM-CNN digunakan sebagai algoritma dalam penelitian dengan melibatkan ekstraksi fitur TF-IDF (*Term Frequency-Inverse Document Frequency*) untuk vektorisasi kata dan augmentasi untuk menyeimbangkan data. Model CNN menunjukkan performa terbaik dengan akurasi tertinggi sebesar 97.03% pada x3 augmentasi, F1-Score 0.97, precision 0.97, dan recall 0.96. Hal ini menunjukkan bahwa model CNN dapat mengidentifikasi kelas emosi pada data lebih baik dibandingkan dengan model lainnya.

Kata kunci : analisis sentimen, deep learning, CNN, LSTM, TikTok, Tokopedia

1. PENDAHULUAN

Pada era digitalisasi yang semakin maju, *e-commerce* dan media sosial telah mengalami perkembangan signifikan dan memainkan peran penting dalam transformasi bisnis dan interaksi masyarakat. Pengguna *e-commerce* di Indonesia pada tahun 2023 mencapai angka 196,47 juta, di mana angka tersebut meningkat sebesar 9,78% dari tahun 2022 [1].

Adapun Tokopedia menduduki posisi teratas dengan 1,25 miliar total pengunjung. Sementara itu, pengguna media sosial secara keseluruhan di Indonesia mencapai angka 167 juta pada 2023, di mana jumlah tersebut setara dengan 60,4% populasi domestik [2].

TikTok merupakan sosial yang paling pesat mengalami pertumbuhan pengguna, per Juli 2024 saja, TikTok memiliki sebanyak 157,6 juta pengguna, naik sebesar 48,4% dibandingkan pada Oktober 2023 dengan 106,2 juta pengguna.

Dalam upayanya melakukan ekspansi bisnis, TikTok mengakuisi 75% saham Tokopedia melalui GoTo (Gojek Tokopedia) dengan nilai 1,5 miliar USD atau sekitar 23,5 triliun rupiah. Dengan akuisisi tersebut, TikTok memiliki kontrol atas Tokopedia dalam kemitraan strategis ini. TikTok menyebutkan bahwa akuisisi ini merupakan upaya mereka dalam mendukung operasional Tokopedia untuk jangka yang panjang. Namun, ada kekhawatiran terkait monopoli, dominasi pasar, persaingan tidak sehat, keamanan data pengguna, hingga geopolitik dari China [3], mengingat sudah banyak negara yang sudah melakukan blokir pada *platform* ini seperti Amerika Serikat, Kanada, hingga India [4].

Meskipun terdapat potensi keuntungan seperti dukungan bagi UMKM, pemasaran, dan branding, masyarakat menilai akuisisi yang dilakukan TikTok adalah cara curang mereka untuk mendapatkan izin dagang di Indonesia mengingat *social commerce* TikTok yang sempat diblokir oleh pemerintah pada Oktober 2023.

Dengan adanya perdebatan opini oleh publik, analisis sentimen menjadi penting dilakukan untuk mengetahui bagaimana respon masyarakat terhadap isu terkait. Melalui analisis sentimen, kita dapat mengidentifikasi pola sikap dan pandangan yang dapat membantu perusahaan, pembuat kebijakan, dan masyarakat untuk memahami dampak dari akuisisi yang sudah dilakukan. Adapun alih-alih hanya melakukan klasifikasi sentimen positif, netral, negatif, identifikasi terhadap emosi seperti perasaan gembira, sedih, bahagia, benci, dan lainnya dapat memberikan pemahaman yang lebih mendalam pada opini pengguna. Dengan pemahaman emosi di balik teks, konteks data yang diperoleh akan menjadi lebih kaya. Dalam penelitian ini, komentar pengguna dari *platform* YouTube digunakan sebagai data primer untuk diolah kemudian dilakukan analisis sentimen berdasar teks emosi.

Terdapat tiga jenis teknik dalam melakukan analisis sentimen: *lexicon based*, *machine learning based*, dan *deep learning based*. Dalam penelitian ini, pendekatan *deep learning* digunakan. *Deep learning* memiliki keunggulan dalam menangani data kompleks dan *non-linear* seperti teks berbahasa natural. Selain itu, *deep learning* dapat mengekstrak representasi fitur yang mendalam dan kompleks dari teks [5] yang membuat model ini mampu mengenali dan memahami

nuansa emosi pada data dengan tingkat akurasi yang lebih tinggi. Adapun algoritma pemodelan deep learning yang digunakan adalah CNN (*Convolutional Neural Networks*) dan LSTM (*Long-short Term Memory*). CNN adalah model dengan arsitektur deep feed-forward yang mampu melakukan generalisasi dengan baik, model ini bekerja dengan membagi bobot parameter yang dilatih sehingga membantu menghindari *overfitting* data [6].

Sementara LSTM dipilih karena model ini mampu mempertahankan kesalahan yang dapat dipropagasi mundur melalui waktu dan layer yang dimiliki [7], memungkinkan model melakukan learning data secara kontinu selama waktu yang dibutuhkan dengan menjaga error tetap konstan.

Berdasarkan uraian di atas, penelitian ini bertujuan untuk mengetahui sentimen masyarakat terhadap akusisi Tokopedia oleh TikTok dengan melibatkan aspek emosi untuk dianalisis secara komprehensif menggunakan pembelajaran deep learning. Pemodelan CNN dan LSTM diharapkan dapat memberikan pemahaman yang akurat terhadap nuansa emosional dalam komentar pengguna, sehingga perusahaan, masyarakat, dan pembuat kebijakan terkait, dapat memahami dampak dari akusisi yang terjadi.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Banyak penelitian yang sudah dilakukan terkait analisis sentimen berdasarkan teks emosi menggunakan deep learning. Seperti dalam penelitian “*Attention-Based Modeling For Emotion Detection And Classification In Textual Conversations* [8].”

Penelitian ini melibatkan pengembangan model untuk mengidentifikasi emosi dalam percakapan teks. Dataset mereka mencakup enam jenis emosi berdasarkan Paul Ekman. Setelah data dikumpulkan, data ditokenisasi dan dimasukkan ke dalam *encoder*. Selanjutnya, *encoder* mengirimkan informasi ke unit Bi-LSTM yang dilatih dengan *average stochastic gradient descent* (ASGD). Klasifikasi data ke dalam kategori emosi yang sesuai dilakukan menggunakan lapisan *dense* dan aktivasi SoftMax, di mana hasil F1 score mencapai angka 75,82%.

Kemudian penelitian dengan judul “*An Efficient CNN-LSTM Model for Sentiment Detection in #BlackLivesMatter* [9].

Di mana dilakukan penggabungan model dari CNN dengan LSTM untuk melakukan analisis sentimen terhadap *hashtag black lives matters* pada dua provinsi di Amerika. Dengan menggabungkan kelebihan CNN dalam pengenalan fitur spasial dan kelebihan LSTM dalam pemahaman konteks, model CNN-LSTM dapat secara efektif menangkap informasi penting dari teks yang diperlukan dalam analisis. Pada penelitian tersebut, model CNN-LSTM mencapai tingkat akurasi tertinggi sebesar 94% dalam mendeteksi berbagai sentimen pada kasus terkait.

2.2. Deep Learning

Menurut [10], *deep learning* adalah studi terkait model-model yang melibatkan jumlah komposisi fungsi atau konsep yang lebih besar daripada yang dilakukan oleh machine learning tradisional. *Deep learning* bekerja berdasarkan prinsip ekstraksi fitur dari data mentah dengan menggunakan beberapa layer atau lapisan untuk mengidentifikasi aspek-aspek yang relevan dengan data inputan.

2.3. Convolutional Neural Network

CNN adalah jenis pendekatan *deep learning* yang digunakan untuk menyelesaikan masalah-masalah kompleks, terutama dalam tugas-tugas visi komputer [11].

CNN terdiri dari beberapa blok arsitektur seperti *convolution layers*, *pooling layers*, dan *fully connected layers*. CNN dirancang untuk secara otomatis dan adaptif mempelajari hierarki fitur spasial melalui propagasi balik dengan menggunakan beberapa blok arsitektur yang sudah disebutkan.

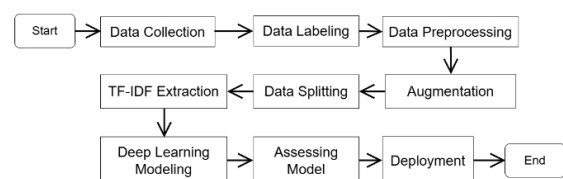
2.4. Long-Short Term Memory

LSTM didefinisikan sebagai arsitektur *recurrent neural networks* (RNN) yang dirancang untuk mengatasi keterbatasan RNN tradisional dalam pembelajaran dan penyimpanan dependensi jangka panjang pada data berurutan [12].

Model LSTM dapat mempelajari cara mengatasi keterlambatan waktu minimal lebih dari 1000 langkah waktu diskrit dengan menggunakan *constant error carousels* (CECs) yang menjaga aliran *constant error* dalam sel-sel tertentu.

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan deep learning dengan beberapa tahapan seperti pada Gambar 1 berikut.



Gambar 1. Metodologi penelitian

Pertama dilakukan pengumpulan data (*data collection*), kemudian pelabelan data (*data labeling*), pra-pemrosesan data (*data preprocessing*) yang meliputi *case folding*, *emoji convert*, *cleansing*, *filtering*, dan *stemming*. Kemudian dilakukan augmentasi data untuk menyeragamkan dan menghindari *overfitting* pada data, selanjutnya pembobotan TF-IDF dan pembagian data uji dan data latih (*data splitting*) dilakukan sebelum pemodelan deep learning menggunakan algoritma CNN dan LSTM. Tahapan setelah proses pemodelan yaitu melakukan evaluasi pada model (*assessing model*), hingga melakukan *deploy* sistem.

3.1. Data Collection

Data yang digunakan dalam penelitian ini merupakan data yang tersedia secara publik pada YouTube. Pengumpulan data dilakukan menggunakan teknik scraping melalui Apps Script pada Google Spreadsheet dengan bantuan YouTube Data API.

3.2. Labeling

Proses pelabelan data dilakukan untuk mengidentifikasi komponen data yang sudah diambil melalui proses *scraping* ke dalam satu kelas sentimen berdasarkan emosi yang telah disiapkan seperti pada Tabel 1.

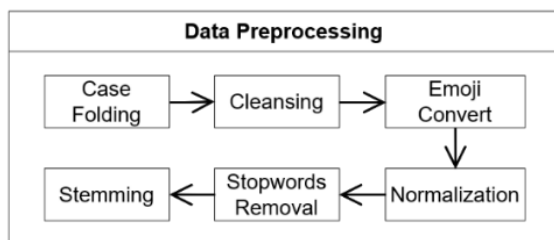
Tabel 1. Kelas emosi label

Label	Emosi
0	Anger
1	Dissapprove
2	Apprehension
3	Acceptance
4	Fear

Dalam penelitian ini, *labeling* dilakukan secara manual dengan 5 aktor yang mengklasifikasikan data teks komentar ke dalam label emosi berdasarkan Plutchik Emotion Model [13] yaitu *anger*, *dissapprove*, *apprehension*, *acceptance*, dan *fear*.

3.3. Preprocessing

Tahapan *preprocessing* merupakan langkah penting yang melibatkan persiapan dan pembersihan data teks guna meningkatkan kinerja algoritma pemodelan. Dalam penelitian ini, terdapat lima tahapan data *preprocessing* yaitu *case folding*, *cleansing*, *normalization*, *stopwords removal*, dan *stemming*.



Gambar 2. Alur data preprocessing

Berdasarkan alur dari Gambar 2, pra pemrosesan data dimulai dengan *case folding* yang merupakan proses mengubah semua karakter dalam sebuah dokumen menjadi huruf yang sama. Hal ini dilakukan untuk standarisasi teks dan untuk mempercepat perbandingan selama tugas pemrosesan teks berlangsung. Kemudian *convert emoji* dilakukan untuk merubah emoji yang terdapat di dalam data teks ke dalam bentuk kata literalnya. *Cleansing* selanjutnya dilakukan untuk mengidentifikasi dan memperbaiki kesalahan serta inkonsistensi data yang terdapat di dalam data teks. Dalam penelitian ini, proses *cleansing* digunakan untuk menghapus karakter angka, simbol, tautan, *emote icon*, tag HTML, tanda baca, dan

karakter kosong atau *white space* yang terdapat pada data. Selanjutnya pada normalisasi, teks diubah menjadi format yang konsisten dan seragam guna mempersiapkan data untuk proses selanjutnya seperti memperbaiki kesalahan ejaan, memperluas singkatan, hingga mengganti slang word dengan kosakata yang bak. Adapun *stopwords removal* dilakukan untuk mengidentifikasi dan menghapus kata-kata umum seperti artikel, konjungsi, preposisi, dan kata ganti dalam data. Sementara *stemming* merupakan tahapan yang bertujuan untuk menghapus imbuhan (*afix*) pada suatu data teks sehingga menyisakan kata dasar yang sesuai.

3.4. Augmentation

Tahapan augmentasi dilakukan untuk mengatasi ketidakseimbangan pada data, teknik ini dilakukan dengan cara *oversampling* pada kelas minoritas (*minority class*) agar sample data dari kelas tersebut menjadi seimbang dengan sample data pada kelas mayoritas (*majority class*).

3.5. Data Splitting

Data splitting atau *train-test split* merupakan proses membagi data menjadi subset untuk pelatihan model dan evaluasi secara terpisah [14].

Dalam penelitian ini, digunakan *K-fold Cross Validation* dengan 3 skenario *train-test split* untuk masing-masing algoritma model yang diuji.

3.6. TF-IDF Extraction

TF-IDF (*Term Frequency-Inverse Document Frequency*) adalah skema pembobotan yang umum digunakan untuk merepresentasikan dokumen teks sebagai vektor [15].

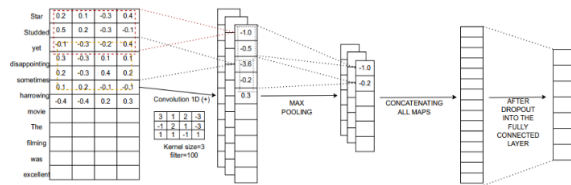
TF-IDF memainkan peran penting dalam menyoroti kata-kata kunci yang paling signifikan muncul dalam suatu korpus, sehingga lebih mudah untuk mengidentifikasi dan memahami konten dari suatu dokumen.

3.7. Deep Learning Modeling

Terdapat dua algoritma model *deep learning* yang digunakan dalam penelitian ini yaitu CNN dan LSTM. Adapun selain berjalan sendiri, kedua algoritma ini juga akan digabung atau *hybrid* oleh peneliti dalam satu model CNN-LSTM dan sebaliknya.

3.7.1. CNN Model

CNN adalah jenis pendekatan *deep learning* yang digunakan untuk menyelesaikan masalah-masalah kompleks, terutama dalam tugas-tugas visi komputer. Arsitektur CNN untuk analisis sentimen berbasis teks terdiri dari komponen-komponen seperti *embedding layer*, *convolutional layer*, *pooling layer*, dan *fully connected layer*.

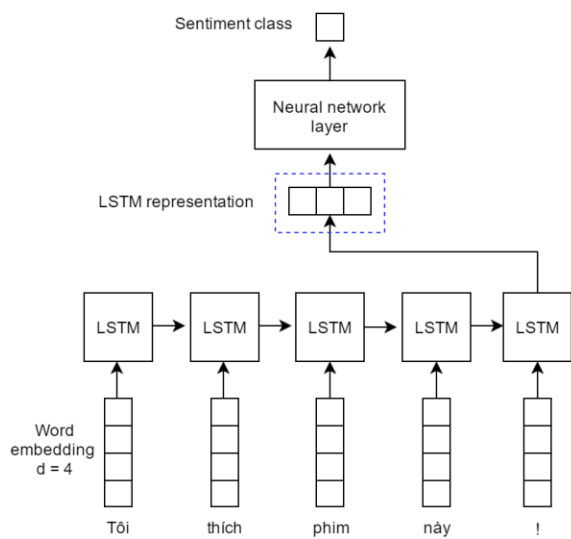


Gambar 3. Arsitektur model CNN

Seperti yang ditampilkan pada Gambar 3, arsitektur dari [16] dimulai dengan input data ke dalam model dengan kalimat yang direpresentasikan sebagai matriks melalui proses *embedding*. Setiap baris matriks adalah vektor yang merepresentasikan sebuah kata. Konvolusi 1D dilakukan pada matriks berukuran kernel 3 bersamaan dengan 4 dan 5. Kemudian *Max-Pooling* dilakukan pada *filter maps*, yang mana kemudian digabungkan dan dikirim ke *fully connected layer* untuk klasifikasi.

3.7.2. LSTM Model

LSTM merupakan arsitektur *recurrent neural networks* (RNN) yang dirancang untuk mengatasi keterbatasan RNN tradisional dalam pembelajaran dan penyimpanan dependensi jangka panjang pada data berurutan.



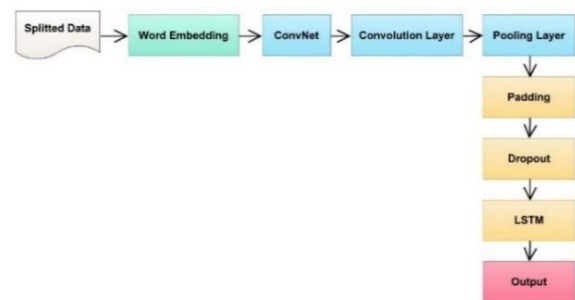
Gambar 4. Arsitektur model LSTM

Pada arsitektur LSTM seperti pada Gambar 4 [17], *memory cell* pada LSTM mampu menyimpan informasi dalam jangka waktu panjang dan memiliki unit pengaturan *non-linear* yang mengatur aliran informasi yang masuk dan keluar dari *cell*. Dengan ini LSTM mampu menangkap hubungan jarak jauh dalam data berurutan secara efektif, memungkinkan LSTM untuk mempertahankan informasi penting.

3.7.3. Hybrid CNN-LSTM Model

Deep Learning Classifier yang menggunakan CNN dan LSTM memberikan urutan input kepada CNN yang digunakan dalam berbagai aplikasi seperti *Computer Vision*. Proses pemodelan *hybrid* algoritma

deep learning CNN-LSTM yang mengadopsi model [9] dapat dilihat pada Gambar 5 berikut.



Gambar 5. Arsitektur model hybrid CNN-LSTM

Alur model CNN-LSTM pada Gambar 5 yang dirancang untuk analisis sentimen dimulai dengan representasi vektor kata melalui proses *embedding word*, diikuti oleh lapisan konvolusi dari CNN untuk mengekstraksi fitur-fitur penting dan *pooling layer* untuk menyederhanakan informasi pada data. Penambahan *padding* dilakukan untuk memastikan integritas informasi selama proses konvolusi, dan lapisan *dropout* digunakan untuk membantu mencegah *overfitting* pada data. Kemudian LSTM digunakan guna memahami urutan kata-kata dan konteks temporal dalam data.

3.8. Evaluate

Proses evaluasi model dilakukan untuk mengukur apakah model algoritma yang digunakan dalam analisis data sudah mampu memahami pola dalam data sehingga memberikan hasil memadai. Dalam penelitian ini, pengukuran kinerja model berupa akurasi, presisi, *recall*, dan nilai *F1-score* menggunakan alat ukur *confusion matrix*.

3.9. Deployment

Tahap *deployment* bertujuan untuk memastikan bahwa aplikasi atau sistem yang dikembangkan dapat beroperasi dengan baik dan dapat digunakan oleh pengguna, dalam penelitian ini, digunakan *flask* yang mana merupakan *framework* web berbasis Python yang ringan dan fleksibel.

4. HASIL DAN PEMBAHASAN

Bagian ini akan membahas hasil dari penelitian di mana mencakup bagaimana dataset dikumpulkan, tahap *preprocessing* data, augmentasi, ekstraksi fitur, pemodelan *deep learning* CNN dan LSTM, evaluasi model, serta *deployment* model ke dalam bentuk web.

4.1. Data Collecting

Proses pengumpulan data dalam penelitian ini menggunakan teknik *scraping* dengan YouTube API pada App Script yang terdapat di dalam *Google Spreadsheet*. Pertama digunakan kata kunci “tiktok tokopedia” untuk mendapatkan list video pada YouTube Stats yang memuat konten terkait akuisisi TikTok terhadap Tokopedia.

Video ID	Comment	Time
myRMT1Xaok	Tokopedia bonekanya tiktok untuk sekarang	2023-12-22T00:28:36Z
myRMT1Xaok	Udah liat zulkihi hasan dimarahi harlison ford? Apa yang mau di	2023-12-22T00:21:08Z
myRMT1Xaok	Yang buat kecewa sebenarnya bukan mergernya Tokopedia den	2023-12-22T00:19:08Z
myRMT1Xaok	@@user-id7zt9im7l owalah thanks	2023-12-28T10:21:03Z
myRMT1Xaok	@@REY-oh4pz iyaa karena tiktok itu ijinya media sosial, bukan	2023-12-28T06:15:01Z
myRMT1Xaok	@@REY-oh4pz gak boleh jualan. Bolehnya cuman media promo	2023-12-27T17:52:21Z
myRMT1Xaok	Maaf bg mau nanya, emang nya sebuah aplikasi kayak tiktok gitu	2023-12-27T14:22:41Z
myRMT1Xaok	Soalnya ga mungkin juga Tiktok shop bisa beroperasi lagi tanpa	2023-12-22T16:52:40Z
myRMT1Xaok	Makasih bang buat penjelasannya...	2023-12-22T00:16:47Z

Gambar 6. Hasil scraping data komentar YouTube

Dari YouTube Stats yang didapatkan, kemudian dilakukan *scraping* guna mendapatkan data komentar dari konten-konten tersebut. Adapun dalam rentang waktu 11 Desember 2023 sampai 20 Februari 2024, dataset yang berhasil dikumpulkan berjumlah 6018 baris data. Hasil dari *scraping* kemudian disimpan ke dalam file .xlsx seperti pada Gambar 6.

4.2. Labeling

Pelabelan emosi dilakukan secara manual dengan 5 pelabel yang masing-masing melakukan *labeling* untuk seluruh row data seperti pada Tabel 2.

Tabel 2. Pelabelan data

Pelabel					Teks
1	2	3	4	5	
4	4	1	4	2	Paling barang import cina diskenariokan menguasai pasar indonesia ..endingnya umkm kita yg hancur😞😞😞

Dari hasil *labeling* secara keseluruhan, peneliti memilih label yang paling banyak muncul pada setiap teks komentar untuk data akhir yang akan digunakan dalam tahapan pemrosesan hingga analisis sentimen seperti pada Tabel 3 yang menampilkan persebaran kelas emosi setelah proses *labeling* dilakukan.

Tabel 3. Sebaran kelas emosi

count	emotion
316	anger
822	dissaprove
1146	apprehension
544	acceptance
213	fear

Berdasarkan tabel di atas, data didominasi oleh emosi yang menyatakan sikap kontra terhadap akuisisi seperti *apprehension*, *dissaprove*, *anger*, dan *fear*.

4.3. Preprocessing

Preprocessing dilakukan menggunakan bantuan dari library Python pada *text editor* di *Google Colaboratory*. Dalam penelitian ini, *preprocessing* memiliki beberapa tahapan yaitu *case folding*, *emoji convert*, *cleansing*, *normalization*, *stopwords removal*, dan *stemming*.

Tabel 4. Hasil *case folding*

Text	Case Folding
Pertanyaannya kemana negara ini saat data masyarakat di pertaruhkan	pertanyaannya kemana negara ini saat data masyarakat di pertaruhkan

Berdasarkan Tabel 4, *case folding* bertujuan untuk menyeragamkan data secara keseluruhan dengan mengubahnya menjadi huruf kecil (*lowercase*).

Tabel 5. Hasil *emoji convert*

Text	Emoji Convert
Bukan jual rugi tapi jual cuan😁	Bukan jual rugi tapi jual cuan (beaming face with smiling eyes)

Proses *emoji convert* melibatkan pembuatan sebuah kamus emoji yang memetakan setiap emoji ke makna deskriptifnya untuk selanjutnya dikonversi ke dalam dataset seperti pada Tabel 5.

Tabel 6. Hasil *cleansing*

Text	Cleansing
Diskon 30% toped bikin saya bbrp x belanja, khusus untuk barang lokal.	diskon toped bikin saya bbrp belanja khusus untuk barang lokal

Berdasarkan hasil pada Tabel 6, *cleansing* dilakukan untuk menghapus karakter yang berpotensi mengganggu proses data lebih lanjut. Dalam penelitian ini, data *cleansing* dilakukan secara bertahap mulai dari menghapus karakter angka, simbol, tautan, tag HTML, tanda baca, dan karakter kosong atau *white space* yang terdapat pada data.

Tabel 7. Hasil *normalization*

Text	Normalization
Pemerintah butuh modal utk kampanye,,, butuh uang SUAP.	Pemerintah butuh modal untuk kampanye butuh uang suap

Normalisasi data bertujuan untuk merubah data teks komentar yang masih belum baku ke dalam bahasa baku. Adapun proses normalisasi dalam penelitian ini dimulai dengan pembuatan kamus berbentuk .csv yang berisikan kata-kata baku bahasa Indonesia. Selanjutnya normalisasi dilakukan dengan bantuan kamus sehingga menghasilkan data seperti pada Tabel 7.

Tabel 8. Hasil *stopwords removal*

Text	Stopwords Removal
Mending blokir tiktok nya sekalian	blokir tiktok

Seperti pada Tabel 8, *stopwords removal* bertujuan menghapus kata henti atau stopwords yang memuat beberapa jenis kata seperti artikel atau kata hubung dalam data.

Tabel 9. Hasil *stemming*

Text	Stemming
Pertanyaannya kemana negara ini saat data masyarakat di pertaruhkan	tanya negara data masyarakat taruh

Ditampilkan hasil dari *stemming* pada Tabel 9, *stemming* sendiri bertujuan untuk menghilangkan afiks atau imbuhan dalam data teks komentar. Digunakan library Sastrawi dalam proses *stemming* pada penelitian ini.

4.4. Augmentation

Tahapan augmentasi dilakukan untuk melakukan *oversampling* pada kelas data emosi yang masuk dalam kategori minor, peneliti melakukan hal ini dikarenakan antara kelas emosi satu dengan yang lain tidak seimbang. Proses menyeimbangkan kelas emosi pada data akan memengaruhi kinerja model dalam pengujian sehingga akurasi yang baik dapat tercapai.

Tabel 10. Skenario augmentasi

Minority class	Augmentation
× 2	6876
× 3	9168

Berdasarkan Tabel 10, digunakan dua skenario augmentasi dalam penelitian yaitu x2 dan x3. Maksud dari perkalian tersebut adalah data yang dikategorikan sebagai kelas minor akan disamaratakan jumlahnya dengan data emosi dari kelas mayor kemudian dikalikan sejumlah skenario yang sudah disebutkan.

4.5. TF-IDF Extraction

Ekstraksi fitur TF-IDF merupakan vektorisasi data dimana dalam tahapan ini data teks akan diubah menjadi representasi numerik dengan basis seberapa penting kata dalam dokumen.

	tiktok	tokopedia	china
0	0.000000	0.000000	0.119972
1	0.289274	0.313286	0.000000
2	0.108401	0.000000	0.000000
3	0.205687	0.000000	0.142949
4	0.000000	0.000000	0.043953

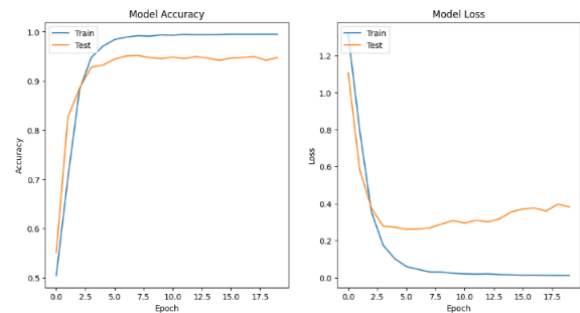
Gambar 7. Ekstraksi TF-IDF

Hasil TF-IDF seperti Gambar 7 menunjukkan dokumen atau setiap kata yang terdapat di dalam data diberikan bobot angka tertentu. Vektor ini akan menjadi representasi numerik teks yang siap digunakan sebagai input pada pemodelan.

4.6. Deep Learning Modeling

4.6.1. CNN Modeling

Pada algoritma model CNN, vektorisasi TF-IDF tidak diterapkan karena representasi kata langsung menggunakan *embedding layer*. CNN bekerja dengan baik pada data yang memiliki struktur spasial, seperti teks yang diubah menjadi vektor *embedding*.

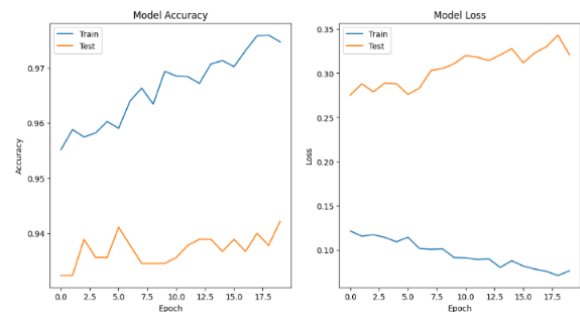


Gambar 8. Plot hasil pemodelan CNN

Pada Gambar 8 menunjukkan grafik perbandingan model ketika dilatih dan diuji dengan komparasi pada *model accuracy* dan *model loss*.

4.6.2. LSTM Modeling

Pada algoritma model LSTM, TF-IDF yang mampu menghasilkan vektor berdimensi tetap diterapkan karena LSTM tidak perlu *embedding* tambahan dalam memproses data.

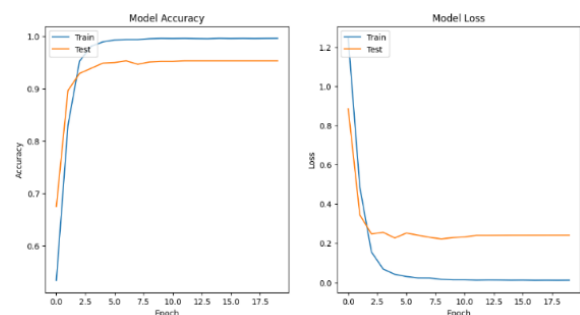


Gambar 9. Plot hasil pemodelan LSTM

Pada Gambar 9 menunjukkan grafik perbandingan model ketika dilatih dan diuji dengan komparasi pada *model accuracy* dan *model loss*.

4.6.3. CNN-LSTM Modeling

Pada algoritma model CNN-LSTM, pertama dilakukan input dari masing-masing model, di mana inputan *layer embedding* digunakan untuk menyimpan proses dari CNN, sementara *layer TF-IDF* adalah inputan yang menyimpan proses dari LSTM. Kedua input ini kemudian digabungkan dalam *concatenate*. Lebih lanjut dropout kembali dibuat di akhir model untuk regularisasi agar model tidak mengalami overfit.

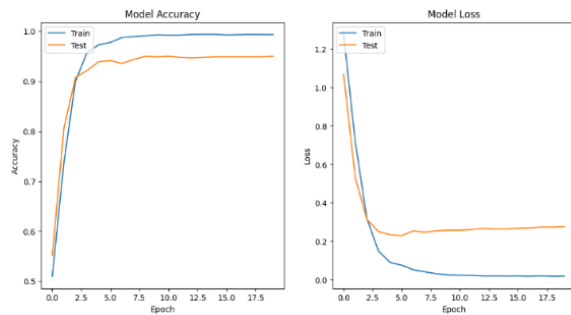


Gambar 10. Plot hasil pemodelan CNN-LSTM

Pada Gambar 10 menunjukkan grafik perbandingan model ketika dilatih dan diuji dengan komparasi pada *model accuracy* dan *model loss*.

4.6.4. LSTM-CNN Modeling

Adapun model *hybrid* pada LSTM-CNN hanya memiliki perbedaan dalam urutan saja. LSTM-CNN memulai dengan menangkap hubungan temporal terlebih dahulu dengan LSTM sebelum kemudian dilakukan ekstraksi fitur lokal menggunakan CNN.



Gambar 11. Plot hasil pemodelan LSTM-CNN

Pada Gambar 11 menunjukkan grafik perbandingan model ketika dilatih dan diuji dengan komparasi pada *model accuracy* dan *model loss*.

4.7. Result Overview

Hasil pemodelan dalam penelitian ini dengan algoritma CNN, LSTM, CNN-LSTM, dan LSTM-CNN pada *K-fold* 0.1 ditampilkan dalam Tabel 11 berikut.

Tabel 11. Hasil pemodelan

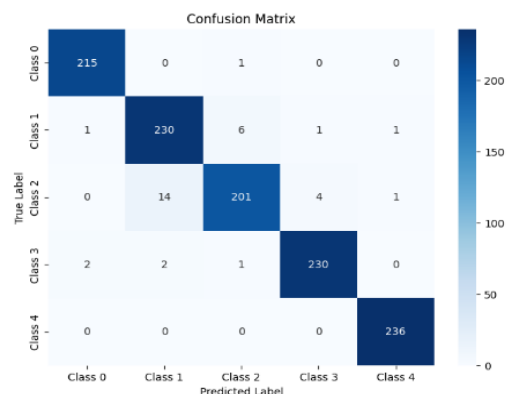
<i>Model</i>	<i>Augmentation</i>	<i>Training Loss</i>	<i>Training Accuracy</i>	<i>Testing Loss</i>	<i>Testing Accuracy</i>
CNN	× 2	0.0148	0.9947	0.4676	0.9055
	× 3	0.0124	0.9956	0.1920	0.9703
CNN + Emoji Convert	× 2	0.0156	0.9938	0.4267	0.9142
	× 3	0.0111	0.9963	0.4478	0.9352
LSTM	× 2	0.1332	0.9510	0.4526	0.9041
	× 3	0.0760	0.9747	0.3210	0.9422
LSTM + Emoji Convert	× 2	0.1399	0.9503	0.4360	0.9041
	× 3	0.1546	0.9434	0.3538	0.9228
CNN-LSTM	× 2	0.0137	0.9952	0.4273	0.9142
	× 3	0.0114	0.9960	0.2406	0.9531
CNN-LSTM + Emoji Convert	× 2	0.0167	0.9958	0.4813	0.9172
	× 3	0.0138	0.9952	0.2956	0.9452
LSTM-CNN	× 2	0.0196	0.9935	0.5283	0.9084
	× 3	0.0160	0.9945	0.3050	0.9487
LSTM-CNN + Emoji Convert	× 2	0.0240	0.9932	0.4956	0.9113
	× 3	0.0243	0.9909	0.3731	0.9390

Berdasarkan hasil pengujian dari empat model yang sudah dipaparkan pada Tabel 11, model LSTM (0.1) dengan augmentasi 9186 menjadi yang terbaik dalam mengatasi *overfitting* data, dimana diantara keempat model, LSTM memiliki persentase gap paling minimum pada hasil akurasi pengujian dan pelatihan yakni 3.25%. Namun, jika dilihat dari akurasi pengujian model tertinggi maka yang terbaik adalah CNN (0.1) augmentasi 9186 dengan nilai akurasi 97.03%, tidak hanya itu, model CNN mempunyai nilai *loss* yang paling kecil yaitu pada angka 0.1920. Oleh karena itu, CNN akan digunakan sebagai model dalam proses deployment.

4.8. Model Evaluate

Setelah proses pelatihan dan pengujian model selesai dilakukan, model kemudian dievaluasi menggunakan *confusion matrix* untuk mengetahui

seberapa baik model memprediksi kelas emosi yang ada. Berikut salah satu contoh *confusion matrix* yang sudah dilakukan.



Gambar 12. *Confusion matrix* CNN

Berdasarkan Gambar 12, evaluasi menggunakan *confusion matrix* pada model CNN menunjukkan dimana secara keseluruhan, model bekerja dengan baik dalam memprediksi kelas yang ada pada data pengujian. Selain menggunakan Confusion Matrix, dalam tahapan evaluasi pemodelan juga diterapkan *classification report* untuk mengetahui nilai *precision*, *recall*, dan *F1-score* dari model yang sudah diuji.

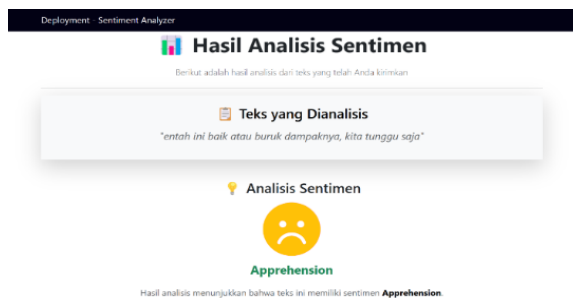
Tabel 12. Hasil *classification report*

Model	Augmentation	Precision	Recall	F1-Score
CNN	× 2	0.92	0.92	0.92
	× 3	0.97	0.96	0.97
LSTM	× 2	0.90	0.91	0.90
	× 3	0.93	0.90	0.91
CNN-LSTM	× 2	0.92	0.92	0.92
	× 3	0.95	0.95	0.94
LSTM-CNN	× 2	0.91	0.91	0.91
	× 3	0.94	0.94	0.93

Berdasarkan Tabel 12, model CNN menunjukkan performa terbaik dengan *F1-Score* tertinggi sebesar 0.97. Model ini juga memiliki nilai *Precision* sebesar 0.97 dan *Recall* sebesar 0.96, menunjukkan kemampuan yang baik dalam mengidentifikasi kelas yang benar.

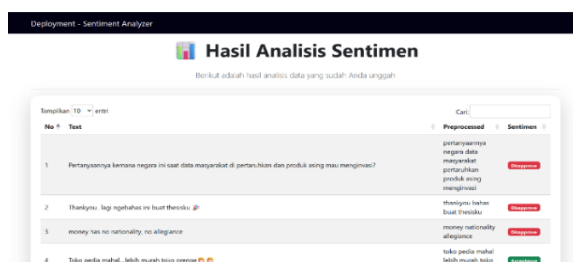
4.9. Deployment

Adapun menu yang ada dalam web *deploy* antara lain menu analisis teks, analisis file, dan *word cloud*.



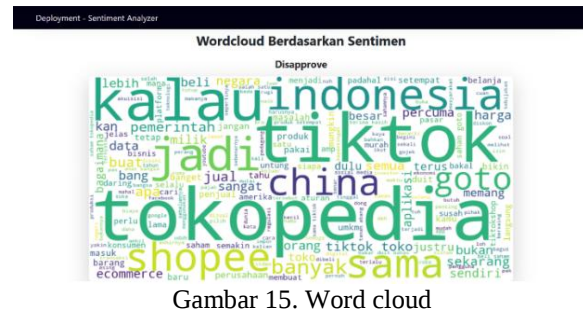
Gambar 13. Menu hasil analisis teks

Pada menu analisis teks, pengguna dapat melakukan input teks untuk kemudian ditampilkan hasil analisisnya yang mana memuat informasi teks yang sudah diinputkan oleh pengguna, hasil *sentiment* sesuai dengan label emosi, serta emoji emosi.



Gambar 14. Menu hasil analisis file

Sementara pada menu analisis file pada Gambar 14, Pengguna dapat melakukan analisis secara masif dengan cara upload data (.csv atau .xlsx) untuk kemudian dianalisis. Hasil analisis dari data yang sudah diunggah adalah teks dari data, hasil preprocessing data, serta hasil sentiment dari data terkait.



Gambar 15. Word cloud

Word cloud yang ditampilkan seperti pada Gambar 15 adalah berdasarkan data file yang sudah diunggah. Dalam *word cloud* ini, pengguna dapat mengetahui kata mana yang paling sering muncul.

5. KESIMPULAN DAN SARAN

Berdasarkan analisis yang telah dilakukan terkait sentimen berdasarkan emosi pada topik akuisisi TikTok terhadap Tokopedia dengan menggunakan metode deep learning, dapat ditarik kesimpulan bahwa data yang berjumlah 3041 didominasi emosi dengan sentimen negatif seperti *apprehension*, *disapprove*, *anger*, dan *fear*. Hal ini menunjukkan adanya banyak sikap kontra atau ketidaksetujuan terhadap tindakan TikTok dalam mengakuisisi Tokopedia. Masyarakat secara mayoritas merasa proses akuisisi yang dilakukan TikTok akan mendatangkan hal-hal negatif terutama bagi ekosistem e-commerce lokal. Adapun Model CNN memberikan performa terbaik untuk klasifikasi data teks dengan akurasi pengujian tertinggi yaitu 97.03% dan *F1-Score* 0.97 saat menggunakan augmentasi x3, hal ini membuktikan peran penting augmentasi dalam membantu mengurangi *overfitting* dan memberikan lebih banyak variasi untuk data pelatihan.

Penulis mengetahui bahwa data yang digunakan saat ini masih sedikit di mana berjumlah 3041. Oleh karena itu, penggunaan *dataset* yang lebih besar dan beragam akan sangat berguna pada penelitian berikutnya. Hal tersebut bukan hanya membantu model dalam generalisasi pola-pola sentimen dengan akurat, tetapi juga memungkinkan model untuk menangkap dinamika sentimen yang lebih luas.

DAFTAR PUSTAKA

- [1] R. Mustajab, "Pengguna E-Commerce RI Diproyeksi Capai 196,47 Juta pada 2023," DataIndonesia.id. Accessed: Jan. 17, 2024. [Online]. Available: <https://dataindonesia.id/ekonomi->

- digital/detail/pengguna-ecommerce-ri-diproeksi-capai-19647-juta-pada-2023
- [2] S. Widi, "Pengguna Media Sosial di Indonesia Sebanyak 167 Juta pada 2023," DataIndonesia.id.
- [3] D. Setyowati, "TikTok Akuisisi Tokopedia, Pasar E-commerce RI Makin dikuasai Asing?," Katadata. Accessed: Jan. 17, 2024. [Online]. Available: https://katadata.co.id/desysetyowati/digital/6576cbc6c6cf0/tiktok-akuisisi-tokopedia-pasar-e-commerce-ri-makin-dikuasai-asing#google_vignette
- [4] M. F. Nurhapy and R. K. Nistanto, "Daftar Negara yang Blokir TikTok Beserta Alasannya," Kompas.com. Accessed: Jan. 17, 2024. [Online]. Available: <https://tekno.kompas.com/read/2023/03/12/10000027/daftar-negara-yang-blokir-tiktok-beserta-alasannya?page=2>
- [5] D. Mahayana and Nuryanti, "Analisis Sentimen Berbasis Aspek dengan Deep Learning Ditinjau dari Sudut Pandang Filsafat Ilmu," *JUMANJI*, vol. 4, no. 2, pp. 70–85, Feb. 2021.
- [6] S. Indolia, A. K. Goswami, S. P. Mishra, and P. Asopa, "Conceptual Understanding of Convolutional Neural Network- A Deep Learning Approach," in *Procedia Computer Science*, Elsevier B.V., 2018, pp. 679–688. doi: 10.1016/j.procs.2018.05.069.
- [7] W. Astriningsih and D. Hatta Fudholi, "Identifikasi Multi Aspek Dan Sentimen Analisis Pada Review Hotel Menggunakan Deep Learning," *Jurnal Teknik Informatika dan Sistem Informasi*, vol. 10, no. 3, pp. 433–443, Sep. 2023, [Online]. Available: <http://jurnal.mdp.ac.id>
- [8] W. Ragheb, J. Azé, S. Bringay, and M. Servajean, "Attention-based Modeling for Emotion Detection and Classification in Textual Conversations," Jun. 2019, [Online]. Available: <http://arxiv.org/abs/1906.07020>
- [9] Ankita, S. Rani, A. K. Bashir, A. Alhudhaif, D. Koundal, and E. S. Gunduz, "An efficient CNN-LSTM model for sentiment detection in #BlackLivesMatter," *Expert Syst Appl*, vol. 193, May 2022, doi: 10.1016/j.eswa.2021.116256.
- [10] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, Massachusetts: MIT Press, 2016.
- [11] L. Alzubaidi *et al.*, "Review of deep learning: concepts, CNN architectures, challenges, applications, future directions," *J Big Data*, vol. 8, no. 1, p. 53, Mar. 2021, doi: 10.1186/s40537-021-00444-8.
- [12] S. Hochreiter and J. Schmidhuber, "Long Short-Term Memory," *Neural Comput*, vol. 9, no. 8, pp. 1735–1780, Nov. 1997, doi: 10.1162/neco.1997.9.8.1735.
- [13] A. S. Imran, S. M. Daudpota, Z. Kastrati, and R. Batra, "Cross-cultural polarity and emotion detection using sentiment analysis and deep learning on covid-19 related tweets," *IEEE Access*, vol. 8, pp. 181074–181090, 2020, doi: 10.1109/ACCESS.2020.3027350.
- [14] I. O. Muraina, "IDEAL DATASET SPLITTING RATIOS IN MACHINE LEARNING ALGORITHMS: GENERAL CONCERNS FOR DATA SCIENTISTS AND DATA ANALYSTS," in *7th INTERNATIONAL MARDIN ARTUKLU SCIENTIFIC RESEARCHES CONFERENCE*, Mardin, Turkey: www.artuklukongresi.org, 2022, pp. 496–504.
- [15] C. Sammut, "TF-IDF," in *Encyclopedia of Machine Learning*, Boston, MA: Springer US, 2011, pp. 986–987. doi: 10.1007/978-0-387-30164-8_832.
- [16] H. Kim and Y. S. Jeong, "Sentiment classification using Convolutional Neural Networks," *Applied Sciences (Switzerland)*, vol. 9, no. 11, Jun. 2019, doi: 10.3390/app9112347.
- [17] Q.-H. Vo, H.-T. Nguyen, B. Le, and M.-L. Nguyen, "Multi-channel LSTM-CNN model for Vietnamese sentiment analysis," in *2017 9th International Conference on Knowledge and Systems Engineering (KSE)*, IEEE, Oct. 2017, pp. 24–29. doi: 10.1109/KSE.2017.8119429.