

PEMANFAATAN LSTM UNTUK MENGANALISIS SENTIMEN PENGGUNA TWITTER : STUDI KASUS PADA TWEET BERITA TERKINI

Erwin Darmawan, Mhd Arief Hasan*, Irsando, Ningsi Rahmawati, Agung, Veby Kurniawan

Teknik Informatika, Fakultas Ilmu Komputer, Universitas Lancang Kuning
Jl. Yos Sudarso No.KM. 8, Umban Sari, Kec. Rumbai, Kota Pekanbaru, Riau
erwindarmawan70@gmail.com

ABSTRAK

Analisis sentimen terhadap berita terkini di platform Twitter menjadi penting untuk memahami dampak informasi terhadap persepsi pengguna. Penelitian ini memanfaatkan Long Short-Term Memory (LSTM) guna menganalisis sentimen pengguna Twitter terhadap tweet berita terkini, dengan tujuan melatih model prediksi sentimen (positif/negatif) sekaligus mengevaluasi pengaruh berita terhadap respons pengguna. Metode yang digunakan melibatkan dataset sebesar 123.218 tweet (75.482 positif dan 47.736 negatif) untuk pelatihan model LSTM, dilengkapi analisis distribusi panjang teks (rata-rata 22,36 karakter untuk sentimen positif dan 23,08 karakter untuk negatif). Evaluasi performa mencakup akurasi, presisi, recall, F1-Score, serta pemeriksaan overfitting melalui grafik akurasi dan loss. Hasil pengujian menunjukkan model mencapai akurasi 91,52%, dengan presisi 91,75%, recall 91,27%, dan F1-Score 91,51%. Klasifikasi berhasil mengidentifikasi 416 tweet positif (True Positive) dan 499 tweet negatif (True Negative), dengan kesalahan 37 False Positive dan 40 False Negative. Grafik akurasi dan loss membuktikan model tidak mengalami overfitting, ditunjukkan oleh konsistensi performa pada data pelatihan dan validasi. Penelitian ini menyimpulkan bahwa LSTM efektif sebagai alat analisis sentimen di Twitter, terutama dalam mengungkap hubungan antara berita terkini dan dinamika sentimen pengguna.

Kata kunci : Sentimen, Twitter, LSTM, Natural Language Processing, Analisis Sentimen, Berita Terkini.

1. PENDAHULUAN

Mikroblog seperti Twitter (www.twitter.com) dan Facebook (www.facebook.com) kini menjadi perangkat komunikasi yang sangat populer di kalangan pengguna internet. Salah satu media sosial yang berkembang pesat penggunaannya adalah Twitter, dengan jumlah pengguna yang terus meningkat sekitar 300.000 akun setiap harinya. Berdasarkan statistik, pada Februari 2022, Twitter memiliki 217 juta akun. Setiap hari, para pengguna Twitter mengirimkan tweet tentang fakta, opini terkait produk atau layanan pemerintah, serta pandangan politik, ideologi, dan minat pribadi. Opini mengenai pemimpin atau tokoh publik yang berpengaruh juga sering menjadi topik utama dalam tweet. Media sosial menyediakan sumber opini yang sangat besar, baik untuk kebutuhan individu maupun organisasi, memungkinkan orang untuk bebas mengekspresikan pendapat tanpa tekanan[1].

Twitter, sebagai platform media sosial aktif dengan pertumbuhan pengguna yang pesat, menjadi sumber data kaya untuk memahami pandangan dan perasaan pengguna terkait perkembangan AI[2].

Model klasifikasi yang digunakan pada penelitian ini adalah Long Short-Term Memory(LSTM) yang merupakan metode dalam deep learning yang dapat digunakan untuk Natural Language Processing(NLP) seperti pengenalan suara, translasi teks, dan analisis sentimen[3].

Analisis sentimen, juga dikenal sebagai penambangan opini, adalah metode untuk memahami, mengekstrak, dan menganalisis data tekstual secara otomatis untuk mendapatkan informasi

sentimen dari sebuah kalimat, baik positif maupun negatif[4].

Pada penelitian ini, kami menggunakan metode Long Short Term Memory (LSTM) karena metode ini sudah terbukti dapat bekerja dengan baik untuk menganalisis sentimen pada data teks[5].

Mengingat jumlah pengguna Twitter yang terus berkembang di Indonesia, serta peran penting media sosial dalam membentuk opini publik, maka analisis sentimen terhadap tweet berita terkini di Indonesia menjadi hal yang penting. Meskipun beberapa penelitian telah dilakukan di luar negeri, penelitian yang memanfaatkan LSTM untuk menganalisis sentimen pada tweet berbahasa Indonesia masih terbatas. Oleh karena itu, penelitian ini diadakan untuk mengisi kekosongan tersebut dan memberikan kontribusi pada pengembangan analisis sentimen berbasis bahasa Indonesia, dengan fokus pada berita terkini yang diposting di Twitter.

Pengolahan data opini atau tweet di Twitter memerlukan pendekatan komputasional yang efisien untuk mengidentifikasi dan menganalisis sentimen yang terdapat dalam tweet-tweet tersebut[6].

Pertama, bagaimana efektivitas model LSTM dalam mengklasifikasikan sentimen pengguna Twitter terhadap berita terkini, terutama dalam konteks bahasa Indonesia yang memiliki tantangan unik seperti penggunaan slang dan ketidakformalan dalam teks? Kedua, apa saja tantangan yang dihadapi dalam memproses teks bahasa Indonesia untuk analisis sentimen, terutama terkait dengan pengolahan data tweet yang sering kali tidak terstruktur dan mengandung variasi bahasa yang tinggi? Ketiga,

sejauh mana model LSTM dapat memberikan peningkatan dalam akurasi analisis sentimen dibandingkan dengan metode tradisional seperti Naive Bayes atau Support Vector Machines (SVM), yang sebelumnya banyak digunakan dalam penelitian serupa? Melalui penelitian ini, diharapkan dapat ditemukan solusi untuk meningkatkan pemahaman mengenai cara-cara yang lebih efektif dalam mengklasifikasikan sentimen, serta memberikan kontribusi terhadap pengembangan model analisis sentimen yang lebih akurat dan dapat diandalkan dalam konteks bahasa Indonesia.

2. TINJAUAN PUSTAKA

2.1. Sentimen analysis

Analisis sentimen merupakan proses untuk menjabarkan sentimen yang terjadi sehingga dapat dikelompokkan dari data sumber seperti text ke dalam dokumen atau kalimat sehingga dapat dikategorikan sebagai sentimen positif, netral dan negatif[7].

Dalam konteks media sosial, Liu et al. (2021) mengemukakan bahwa pendekatan deep learning telah menunjukkan performa yang lebih baik dibandingkan metode tradisional seperti rule-based atau machine learning konvensional. Analisis sendiri kerap digunakan untuk menyatakan suka atau tidak suka terhadap suatu hal melalui sentimen positif dan negatif[8].

Untuk data Twitter khususnya, preprocessing menjadi tahap crucial yang meliputi normalisasi teks, penanganan emoji, dan standarisasi bahasa informal untuk meningkatkan akurasi analisis.

2.2. Media Sosial Dan Sentimen

Twitter telah menjadi salah satu platform media sosial yang paling berpengaruh dalam penyebaran informasi dan berita terkini. Sebagai platform microblogging, Twitter memungkinkan pengguna untuk berbagi pendapat dan bereaksi terhadap berbagai peristiwa secara real-time dalam format teks singkat yang dibatasi 280 karakter. Karakteristik unik ini menjadikan Twitter sebagai sumber data yang kaya untuk analisis sentimen dan pemahaman opini publik. Selain itu informasi sentiment analisis pada media sosial juga menarik perhatian para peneliti karena dapat merangsang pola pikir netizen untuk memberikan pendapatnya terhadap suatu objek[9].

2.3. Long Short-Term Memory (LSTM)

Melihat dari hasil penelitian terdahulu, algoritma LSTM memiliki hasil akurasi paling tinggi dibandingkan dengan algoritma lainnya. Ini dikarenakan LSTM memiliki kelebihan dalam melakukan analisis sentimen, yaitu mampu mengingat informasi dari data yang telah diproses dan menggunakan kemampuan memori tersebut untuk memahami data input lebih baik. Kemampuan memori ini hanya dimiliki oleh algoritma LSTM yang merupakan pengembangan dari algoritma Recurrent

Neural Network (RNN) dalam mengatasi masalah long-term dependency dalam mengingat data jangka panjang[10].

Dalam konteks analisis sentimen Twitter, LSTM telah terbukti efektif dalam menangkap dependensi jarak jauh dan memahami konteks yang lebih luas dari sebuah tweet.

2.4. Text Processing

Tahap text preprocessing adalah tahap awal dari text mining. Data yang dikumpulkan merupakan data yang tidak terstruktur sehingga membutuhkan preprocessing[11].

Tindakan yang dilakukan pada tahap ini adalah toLowerCase, yaitu mengubah semua karakter huruf menjadi huruf kecil dan Tokenizing yaitu proses penguraian deskripsi yang semula berupa kalimat-kalimat menjadi kata-kata dan menghilangkan delimiter seperti tanda titik (.), koma (,), spasi dan karakter angka yang ada pada kata tersebut.

2.5. Naïve Bayes Classifier

Algoritma *Naive Bayes Classifier* adalah metode yang digunakan untuk menghitung probabilitas tertinggi guna mengklasifikasikan data uji ke dalam kategori yang paling sesuai. Dalam penelitian ini, data uji yang digunakan adalah kumpulan tweet. Proses klasifikasi dokumen terdiri dari dua tahap. klasifikasi ini sangat memperhatikan perkiraan probabilitas yang ada berdasarkan dari kategori data latih[12].

Tahap pertama adalah proses pelatihan, di mana model dilatih menggunakan dokumen-dokumen yang sudah memiliki kategori yang diketahui. Sedangkan tahap kedua melibatkan klasifikasi dokumen yang belum diketahui kategorinya, dengan menggunakan model yang telah dilatih pada tahap sebelumnya untuk menentukan kategori yang paling tepat berdasarkan probabilitas yang dihitung.

2.6. Confusion Matrix

Confusion matrix merupakan alat untuk mengukur kinerja model dalam masalah klasifikasi, yang memberikan gambaran tentang seberapa baik model mengklasifikasikan data ke dalam kategori yang telah ditentukan. Dalam konteks analisis sentimen, confusion matrix digunakan untuk mengevaluasi prediksi sentimen (positif, netral, dan negatif) yang dilakukan oleh model LSTM. Komponen utama dalam confusion matrix meliputi True Positive (TP), yaitu jumlah instance yang benar diprediksi sebagai positif, True Negative (TN), yaitu jumlah instance yang benar diprediksi sebagai negatif, False Positive (FP), yaitu jumlah instance yang salah diprediksi sebagai positif, dan False Negative (FN), yaitu jumlah instance yang salah diprediksi sebagai negatif[13].

Dari confusion matrix, kita dapat menghitung beberapa metrik evaluasi yang penting, antara lain Accuracy, Precision, Recall, dan F1 Score.

a. Accuracy

Accuracy mengukur sejauh mana prediksi yang benar dibandingkan dengan total data uji, yang dihitung dengan rumus:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (1)$$

b. Precision

Precision mengukur akurasi prediksi positif, dihitung dengan rumus:

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

c. Recall

Sedangkan Recall mengukur kemampuan model dalam mendeteksi semua instance positif, yang dihitung dengan rumus:

$$Recall = \frac{TP}{TP+FN} \quad (3)$$

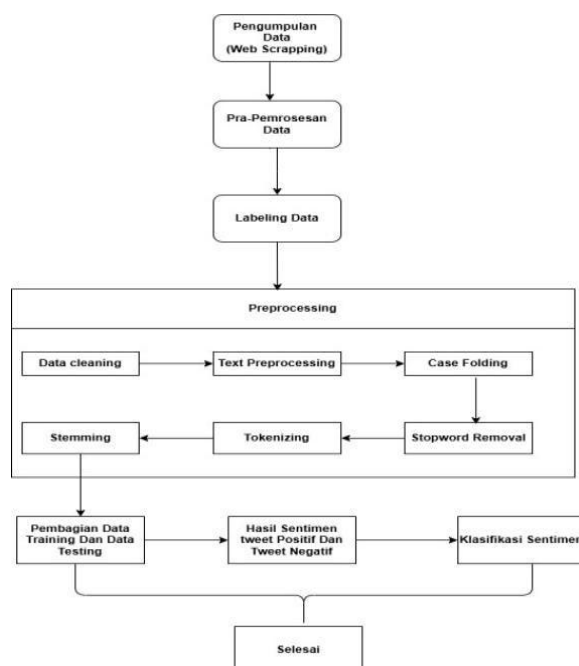
d. F1 Score

F1 Score adalah rata-rata harmonis dari precision dan recall, yang dihitung dengan rumus:

$$F1\ Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (4)$$

Dalam penelitian ini, model LSTM menunjukkan performa yang sangat baik dengan Accuracy sebesar 91,52%, Precision 91,75%, Recall 91,27%, dan F1 Score 91,51%. Berdasarkan confusion matrix yang dihasilkan, model berhasil mengklasifikasikan 416 tweet sebagai True Positive (positif), 499 tweet sebagai True Negative (negatif), dengan 37 False Positive (tweet yang salah diprediksi sebagai positif), dan 40 False Negative (tweet yang salah diprediksi sebagai negatif). Metrik evaluasi ini menunjukkan bahwa model memiliki kemampuan yang sangat baik dalam memprediksi sentimen positif dan negatif dengan tingkat kesalahan yang rendah, serta mampu menjaga keseimbangan antara precision dan recall.

3. METODE PENELITIAN



Gambar 1. Alur Penelitian

Penelitian ini bertujuan untuk memprediksi popularitas sentiment di Twitter menggunakan teknik *LSTM*. Metodologi yang digunakan mencakup beberapa tahap utama, yaitu pengumpulan data, praproses data, ekstraksi fitur, penerapan model, serta evaluasi kinerja model. Rincian metode penelitian adalah sebagai berikut:

3.1. Pengumpulan Data

Dataset yang digunakan dalam penelitian ini diambil dari Kaggle dengan judul "Twitter Sentiment Analysis". Dataset ini berisi tweet yang telah dilabeli dengan tiga kategori sentimen: positif, negatif, dan netral. Tweet-tweet ini mencakup berbagai topik, termasuk berita terkini dan isu-isu global, yang memungkinkan untuk menganalisis persepsi publik terhadap berbagai peristiwa melalui media sosial. Data ini diunduh dalam format CSV yang mencakup informasi seperti teks tweet, label sentimen, dan beberapa metadata lainnya.

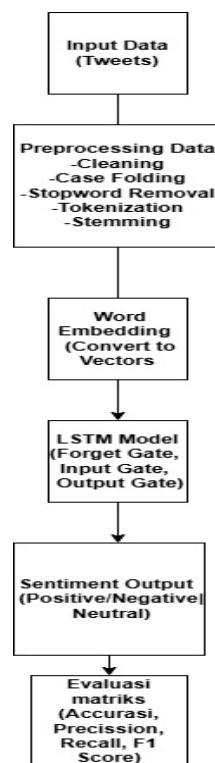
3.2. Praproses Data

Sebelum data dapat digunakan untuk melatih model, dilakukan serangkaian langkah praproses untuk membersihkan dan menyiapkan data. Langkah pertama adalah penghapusan URL, simbol, dan karakter khusus, untuk memastikan bahwa hanya teks yang relevan yang tersisa dalam setiap tweet. Selanjutnya, dilakukan tokenisasi, yaitu proses memecah teks menjadi kata-kata atau token yang lebih kecil, yang memungkinkan model untuk mempelajari struktur kata dengan lebih baik. Setelah itu, dilakukan penghapusan stopwords, yaitu menghilangkan kata-kata umum yang tidak memberikan informasi berarti dalam analisis, seperti "the", "is", dan "and". Selanjutnya, dilakukan lematisasi, yaitu mengubah kata-kata yang memiliki bentuk berbeda tetapi makna yang sama menjadi bentuk dasar, seperti "running" menjadi "run". Setelah semua tahap ini selesai, tweet yang telah diproses kemudian dikodekan menggunakan teknik Word Embedding seperti Word2Vec atau GloVe, untuk mengubah kata-kata menjadi representasi numerik yang lebih bermakna dan dapat diproses oleh model machine learning.

3.3. Ekstraksi Fitur

Dalam tahap ini, fitur yang relevan dari data teks diekstraksi untuk digunakan dalam model. Word Embedding yang dihasilkan melalui teknik seperti Word2Vec atau GloVe digunakan untuk mengubah setiap kata dalam tweet menjadi vektor numerik yang memiliki arti kontekstual dalam teks. Representasi vektor ini memungkinkan model untuk menangkap hubungan antara kata-kata dalam konteks yang lebih luas. Setiap tweet kemudian diwakili oleh vektor kata-kata yang membentuk tweet tersebut, yang kemudian digunakan sebagai input untuk model. Fitur-fitur ini memungkinkan model untuk mengenali pola-pola sentimen yang terkandung dalam teks tweet.

3.4. Long Short-Term Memory



Gambar 2. flowchart metode

3.5. Pengumpulan Data

Proses dimulai dengan pengumpulan data tweet yang berisi informasi yang relevan, seperti sentimen positif, negatif, dan netral. Data ini dapat diperoleh melalui platform seperti Twitter, menggunakan teknik scraping atau API. Dataset yang dikumpulkan mencakup tweet dengan berbagai topik yang dapat mencerminkan opini publik terkait berita terkini.

3.6. Preprocessing Data

Setelah pengumpulan data, tahap berikutnya adalah preprocessing untuk menyiapkan data sebelum dimasukkan ke dalam model. Langkah-langkah dalam preprocessing mencakup:

- Cleaning:** Penghapusan elemen-elemen yang tidak relevan seperti URL, simbol, dan karakter khusus dari teks tweet.
- Case Folding:** Mengubah semua huruf menjadi huruf kecil untuk memastikan bahwa kata yang memiliki kapitalisasi berbeda diperlakukan sebagai kata yang sama.
- Stopword Removal:** Menghapus kata-kata umum yang tidak membawa informasi penting dalam analisis sentimen seperti "dan", "di", "yang", dan sebagainya.
- Tokenization:** Memecah teks tweet menjadi kata-kata atau token untuk memudahkan analisis lebih lanjut.

- Stemming:** Mengembalikan kata ke bentuk dasar untuk mengurangi variasi kata yang sebenarnya memiliki makna yang sama.

3.7. Word Embedding

Setelah preprocessing, langkah selanjutnya adalah melakukan Word Embedding, di mana setiap kata dalam tweet dikonversi menjadi vektor numerik yang dapat dipahami oleh model. Teknik seperti Word2Vec atau GloVe digunakan untuk mengonversi kata-kata menjadi representasi vektor numerik yang mengandung informasi kontekstual dari kata tersebut. Representasi vektor ini sangat penting agar model dapat memahami hubungan antar kata dalam teks.

3.8. Proses dalam LSTM Model

Setelah data diproses dan di-embed-kan, langkah selanjutnya adalah memasukkan data ke dalam model Long Short-Term Memory (LSTM). LSTM berfungsi untuk mengatasi masalah vanishing gradient pada Recurrent Neural Networks (RNN). Dalam LSTM, terdapat tiga komponen utama yang mempengaruhi bagaimana informasi diproses:

- Forget Gate:** Mengontrol informasi mana yang akan diingat dan mana yang harus dilupakan dari langkah sebelumnya. value dapat dihitung dengan persamaan (5).

$$f_t = \sigma(w_F \cdot [h_{t-1}, x_t] + b_f) \quad (5)$$

- Input Gate:** Memutuskan informasi baru apa yang perlu disimpan dalam memori sel. Untuk menghitung masukan nilai gerbang dengan Persamaan (6) dan kandidat baru dengan persamaan (7).

$$i_t = \sigma(w_{ii}x_t + b_{ii} + w_{hi}h_{t-1} + b_{hi}) \quad (6)$$

$$c_t = \tanh(w_i \cdot [h_{t-1}, x_t] + b_i) \quad (7)$$

Setelah kita mendapatkan nilai masukan gerbang dan kandidat baru, sekarang saatnya kita memperbarui konteks C_{t-1} lama ke konteks C_t baru. dengan persamaan (8) sebagai berikut:

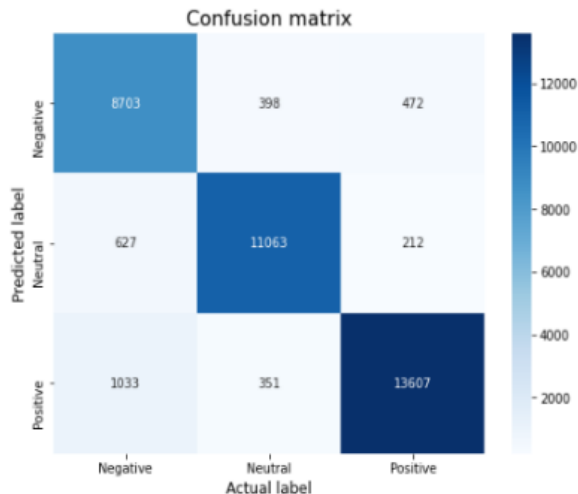
$$C_t = f_t * C_{t-1} + i_t * \hat{C}_t \quad (8)$$

Langkah terakhir kita akan mendapatkan apa yang kita cari adalah hasil. Pertama, kami menjalankan gerbang sigmoid, yang disebut gerbang keluaran, untuk memutuskan bagian mana dari konteks yang akan kami hasilkan. Konteks melalui tanh untuk membuat nilai antara -1 dan 1, dan kita mengalikannya dengan keluaran gerbang sigmoid. Output gerbang dapat dihitung dengan persamaan (9)

$$O_t = \sigma(w_o[h_{t-1}, x_t] + b_o) \quad (9)$$

3.9. Output Sentimen (Prediksi)

Setelah proses dalam LSTM selesai, model menghasilkan output prediksi yang berupa sentimen tweet (positif, negatif, atau netral). Output ini didasarkan pada pola-pola yang dipelajari oleh model dari tweet-tweet yang telah diproses sebelumnya.



Gambar 3. confusion matrix

3.10. Evaluasi Model

Setelah model LSTM dilatih menggunakan dataset tweet, langkah selanjutnya adalah melakukan evaluasi untuk mengukur kinerja model dalam mengklasifikasikan sentimen pada tweet. Evaluasi ini penting untuk mengetahui seberapa baik model dapat memprediksi sentimen (positif, negatif, netral) dan seberapa efektif model dalam mengatasi kesalahan prediksi. Beberapa metrik yang digunakan untuk evaluasi model antara lain accuracy, precision, recall, dan F1-score. Selama evaluasi, model LSTM yang digunakan dalam penelitian ini mencapai accuracy sebesar 91,52%, yang menunjukkan bahwa model dapat memprediksi sentimen dengan tingkat keberhasilan yang tinggi. Precision mencapai 91,75%, menunjukkan bahwa sebagian besar tweet yang diprediksi sebagai positif memang benar-benar positif. Recall model sebesar 91,27%, menunjukkan bahwa model dapat menemukan hampir 91% tweet yang sebenarnya positif. F1-Score sebesar 91,51% mengindikasikan keseimbangan yang sangat baik antara precision dan recall, yang berarti model mampu mengidentifikasi sentimen dengan baik dan meminimalkan kesalahan.

Dengan hasil evaluasi ini, dapat disimpulkan bahwa model LSTM sangat efektif dalam menganalisis sentimen pada tweet berita terkini, dengan kemampuan yang kuat dalam memprediksi sentimen positif maupun negatif dan menjaga keseimbangan antara akurasi dan kemampuan mendeteksi tweet yang relevan.

4. HASIL DAN PEMBAHASAN

4.1. Pengumpulan Data

Pengumpulan data dalam penelitian ini dilakukan dengan menggunakan dataset Twitter yang berisi tweet-tweet yang relevan dengan berita terkini. Dataset ini diperoleh dari sumber publik yang menyediakan kumpulan data tweet yang mencakup berbagai topik berita, seperti politik, bencana alam, dan isu sosial lainnya. Data yang digunakan mencakup teks tweet, waktu pengiriman, jumlah retweet, dan

metadata lainnya seperti likes, yang memberikan informasi lebih lanjut mengenai interaksi pengguna terhadap berita tersebut. Pemilihan tweet dilakukan dengan menggunakan kata kunci dan hashtag yang berkaitan dengan topik-topik terkini, memastikan bahwa data yang terkumpul relevan dan up-to-date.

Dalam penelitian ini, pengumpulan data saya melibatkan pengambilan sebuah tweet dari platform media sosial Twitter. Proses pengambilan data dilakukan dengan menggunakan alat bernama Twitter Scraper, yang memungkinkan saya untuk mengakses dan mengekstrak informasi yang relevan dengan cepat dan efisien. Setelah data diambil, saya menyimpan seluruh kumpulan tweet tersebut dalam format file .csv, yang memudahkan analisis lebih lanjut dan penyimpanan data secara terorganisir. Format .csv dipilih karena fleksibilitasnya dalam membaca dan memanipulasi data menggunakan berbagai perangkat lunak analisis data.

4.2. Preprocessing Data

Tahap awal pada proses ini adalah melakukan pelabelan data secara manual. Proses ini memakan waktu beberapa saat dan melibatkan pelabelan manual data ulasan oleh penulis. Kumpulan data ini dibagi menjadi kategori Neutral, negatif dan positif. Materi informasi positif menunjukkan hal-hal positif tentang aplikasi Twitter, sedangkan materi negatif menunjukkan hal-hal negatif. Hasil dari proses ini dapat dilihat pada Tabel 1.

Tabel 1. Map Tweet Category

	Clean Text	Category
1	when raja promised "minimum government maximum	Negative
2	talk all the nonsense and continue all the draw...	Neutral
3	what did just say vote for modi welcome bjp t...	Positive
4	asking his supporters prefix ball their n...	Positive
5	answer who among these the most powerful world...	Positive

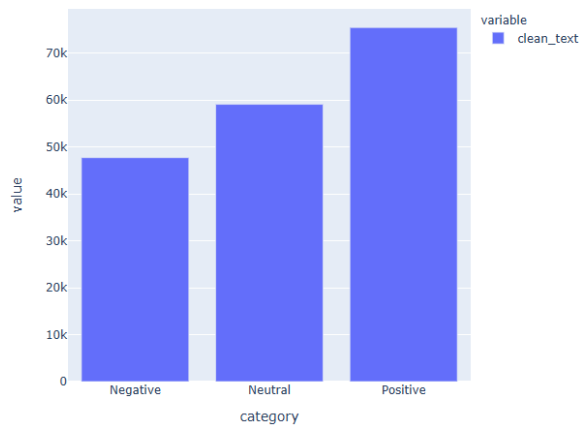
Selanjutnya melakukan beberapa tahapan berikut ini: Berikut adalah perbaikan deskripsi dan penjelasan untuk setiap proses dalam text preprocessing:

4.3. Cleaning dan Pepping Dataset

Proses ini melibatkan pembersihan data dari elemen yang tidak diperlukan atau mengganggu analisis, seperti URL, hashtag, mention, emoji, dan karakter spesial lainnya. Untuk analisis sentimen Twitter, data cleaning sangat penting karena tweet sering kali berisi elemen-elemen ini yang tidak relevan dengan sentimen yang ingin dianalisis. Data terdiri dari 3 category yaitu 0 Menunjukkan itu adalah Sentimen Netral, 1 Menunjukkan itu adalah Sentimen Positif, dan -1 Menunjukkan itu adalah Sentimen Negative.

Tabel 2. Hasil Cleaning

	Clean Text	Category
1	Wow. Yall needa step it up @Apple RT @heynyla	-1.0
2	What Happened To Apple Inc? http://t.co/FJEX...	0.0
3	Thank u @apple I can now compile all of the pi...	1.0
4	The oddly uplifting story of the Apple co-foun...	0.0
5	@apple can i exchange my iphone for a differen...	0.0



Gambar 2. Grafik Clean Text

4.4. Case Folding

Ini adalah langkah di mana semua teks diubah menjadi huruf kecil untuk menstandarisasi data dan menghindari duplikasi kata yang sebenarnya sama tetapi berbeda dalam hal kapitalisasi. Dalam konteks Twitter sentiment, case folding membantu dalam memastikan bahwa kata "Suka" dan "suka" diperlakukan sebagai kata yang sama. Dapat dilihat pada tabel 3 dibawah.

Tabel 3. Case Folding

Before clean_text	Sesudah
"This is a test."	"this is a test."
"I love apple."	"i love apple."

4.5. Stop Word Removal

Stop word adalah kata-kata umum yang tidak memberikan banyak nilai dalam analisis sentimen karena sering muncul dan tidak membawa makna spesifik seperti "di", "dan", "yang", dll. Menghilangkan stop word dari tweet membantu fokus pada kata-kata yang lebih informatif. Contoh: Dari teks "saya suka film ini dan itu sangat bagus", setelah stop word removal, menjadi "suka film bagus".

Tabel 4. Stop Word Removal

Sebelum	Sesudah
I really love this new iPhone! It's amazing!	really love new iPhone amazing
The weather is so nice today, perfect for a walk!	weather nice today perfect walk

4.6. Tokenizing

Tokenizing adalah proses memecah teks menjadi unit-unit yang lebih kecil seperti kata atau frasa, yang disebut token. Untuk Twitter, tokenizing membantu dalam memisahkan kata-kata dari tweet untuk analisis sentimen. Contoh: "saya suka film bagus" menjadi ['saya', 'suka', 'film', 'bagus'].

Tabel 5. Tokenizing

Sebelum	Sesudah
"I love coding in Python! It's amazing."	["I", "love", "coding", "in", "Python", "!", "It", "s", "amazing", "."]
"The weather is great today, perfect for a walk."	["The", "weather", "is", "great", "today", ",", "perfect", "for", "a", "walk", "."]

4.7. Stemming

Stemming adalah proses mengembalikan kata ke bentuk dasar atau akar katanya untuk mengurangi variasi kata yang sebenarnya memiliki makna yang sama. Dalam analisis sentimen Twitter, ini membantu dalam mengurangi dimensi data dan meningkatkan akurasi model dengan memperlakukan kata-kata turunan sebagai satu entitas. Contoh: "membaca", "dibaca", "pembaca" semuanya di-stem menjadi "baca". Jadi, dari ['saya', 'suka', 'film', 'bagus'] menjadi ['saya', 'suka', 'film', 'bagus'] (karena 'bagus' sudah dalam bentuk dasar).

Tabel 6. Stemming

Sebelum	Sesudah
Saya sedang membaca buku tentang pemograman	saya sedang baca buku tentang pemrograman
Mereka sedang makan di restoran yang baru dibuka	mereka sedang makan di restoran yang baru buka

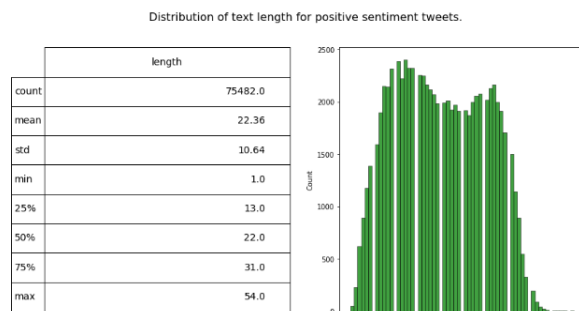
4.8. Pembagian data

Dalam konteks Twitter sentiment analysis, data dibagi menjadi tiga bagian utama: data latih (training data), data validasi (validation data), dan data uji (test data). Data latih digunakan untuk melatih model, data validasi untuk menilai performa model selama pengembangan, dan data uji untuk evaluasi akhir model. Contoh: Jika kita memiliki 1000 tweet, kita bisa membagi menjadi 700 untuk latih, 150 untuk validasi, dan 150 untuk uji.

4.9. Sentimen Positif

Di dalam visualisasi ini, terdapat dua komponen utama. Pertama, terdapat histogram yang menggambarkan distribusi panjang tweet dalam jumlah kata untuk tweet dengan sentimen positif. Pada grafik ini, sumbu horizontal menunjukkan panjang tweet dalam jumlah kata, sedangkan sumbu vertikal menunjukkan jumlah tweet pada panjang tertentu. Histogram ini menunjukkan bahwa sebagian besar tweet dengan sentimen positif memiliki panjang antara

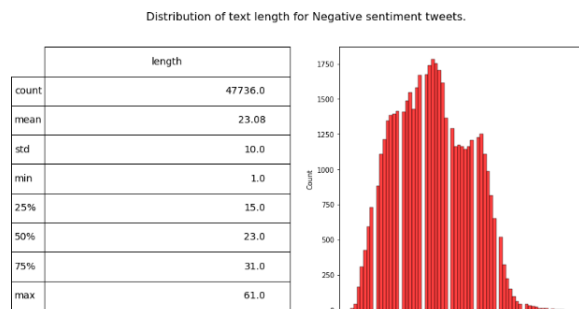
13 hingga 31 kata, dengan distribusi yang cukup merata di rentang panjang tersebut. Hanya sedikit tweet yang memiliki panjang lebih pendek atau lebih panjang dari rentang tersebut.



Gambar 3. Tabel Dan Histogram Sentimen Positif

Kedua, terdapat tabel statistik deskriptif yang menunjukkan informasi lebih lanjut mengenai panjang tweet dengan sentimen positif. Tabel ini mencakup jumlah tweet yang dianalisis, yaitu sebanyak 75.482 tweet, dengan panjang rata-rata tweet sekitar 22.36 kata. Sebaran panjang tweet memiliki standar deviasi sebesar 10.64 kata, menunjukkan adanya variasi yang cukup besar dalam panjang tweet. Panjang tweet terpendek tercatat 1 kata, sementara tweet terpanjang mencapai 54 kata. Statistik lainnya, seperti kuartil pertama (13 kata), median (22 kata), dan kuartil ketiga (31 kata), memberikan gambaran yang lebih rinci mengenai distribusi panjang tweet. Visualisasi ini sangat berguna untuk memahami seberapa panjang tweet yang tergolong dalam kategori sentimen positif dan dapat menjadi referensi dalam persiapan data untuk model analisis sentimen seperti LSTM.

4.10. Sentimen Negatif



Gambar 4. Tabel Dan Histogram Sentimen Negatif

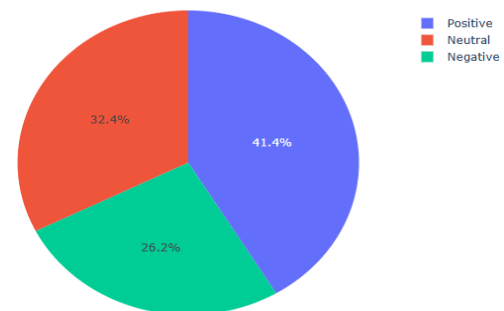
Gambar menunjukkan distribusi panjang teks untuk tweet dengan sentimen negatif. Visualisasi ini terdiri dari dua bagian utama: pertama adalah histogram yang menggambarkan distribusi panjang tweet dalam hal jumlah kata untuk tweet dengan sentimen negatif. Pada histogram ini, sumbu horizontal (x-axis) menunjukkan panjang tweet dalam jumlah kata, sementara sumbu vertikal (y-axis) menunjukkan jumlah tweet yang memiliki panjang tertentu. Histogram ini menunjukkan bahwa sebagian

besar tweet dengan sentimen negatif memiliki panjang antara 15 hingga 31 kata, dengan distribusi yang cukup merata di rentang tersebut. Beberapa tweet juga memiliki panjang yang lebih pendek atau lebih panjang, tetapi jumlahnya jauh lebih sedikit.

Kedua, terdapat tabel statistik deskriptif yang memberikan informasi lebih lanjut tentang panjang tweet dengan sentimen negatif. Tabel ini menunjukkan bahwa jumlah tweet yang dianalisis adalah 47.736 tweet, dengan panjang rata-rata tweet sekitar 23.08 kata. Standar deviasi panjang tweet adalah 10.0 kata, yang menunjukkan adanya variasi yang cukup besar. Panjang tweet terpendek adalah 1 kata, sementara tweet terpanjang memiliki panjang 61 kata. Statistik lainnya, seperti kuartil pertama (15 kata), median (23 kata), dan kuartil ketiga (31 kata), memberikan gambaran lebih rinci tentang distribusi panjang tweet. Judul grafik ini menyatakan bahwa visualisasi ini menunjukkan distribusi panjang teks untuk tweet dengan sentimen negatif, memberikan wawasan yang berguna untuk analisis lebih lanjut.

4.11. Kategori Sentimen

Pie chart of different sentiments of tweets



Gambar 5. Pie Chart

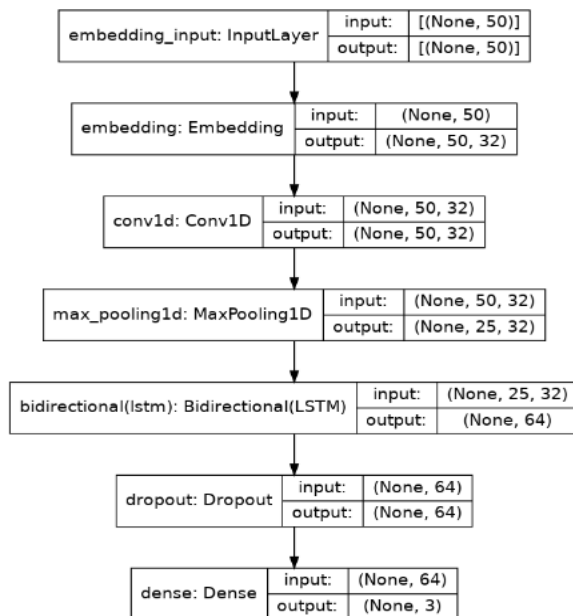
Gambar menunjukkan sebuah grafik pie yang menggambarkan distribusi kategori sentimen dalam dataset tweet. Grafik ini menunjukkan tiga kategori sentimen utama: Positif, Netral, dan Negatif. Pada grafik pie ini, masing-masing kategori diberi warna berbeda: Positif dengan warna biru, Netral dengan warna hijau, dan Negatif dengan warna merah.

Dari grafik, terlihat bahwa kategori Positif mendominasi dataset dengan persentase sekitar 41.4%, diikuti oleh kategori Netral yang memiliki persentase 32.4%, dan kategori Negatif yang sedikit lebih sedikit dengan persentase 26.2%. Grafik ini memberikan gambaran yang jelas mengenai distribusi sentimen dalam dataset tweet yang dianalisis, yang menunjukkan bahwa mayoritas tweet memiliki sentimen positif atau netral, sementara tweet dengan sentimen negatif lebih sedikit.

4.12. Bidirectional LSTM

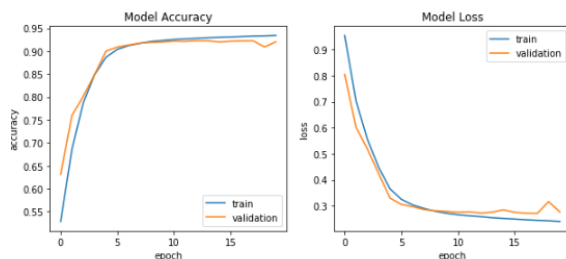
Bidirectional LSTM adalah varian LSTM yang memproses data urutan dari dua arah untuk memberikan konteks yang lebih kaya, ideal untuk

analisis sentimen Twitter. Model ini terdiri dari beberapa lapisan: Embedding Layer untuk mengonversi kata menjadi vektor, Conv1D Layer dengan 32 filter untuk mengekstraksi fitur lokal, MaxPooling1D untuk mengurangi dimensi, Bidirectional LSTM untuk analisis dua arah, Dropout untuk mengurangi overfitting, dan Dense Layer dengan 3 unit dan aktivasi softmax untuk klasifikasi sentimen (Positif, Negatif, Netral). Pelatihan model menggunakan optimizer SGD, dan visualisasi model menunjukkan proses pemrosesan tweet untuk analisis sentimen



Gambar 6. Bidirectional LSTM

4.13. Model Accuracy & Loss



Gambar 7. Grafik Perkembangan Accuracay&loss

Berdasarkan hasil eksperimen yang telah dilakukan pada model Bidirectional LSTM untuk analisis sentimen Twitter, kami mencapai performa yang sangat baik dengan beberapa metrik evaluasi sebagai berikut: Akurasi model mencapai 0.9152, menunjukkan bahwa model ini benar dalam memprediksi sentimen dalam 91.52% dari kasus. Precision atau presisi mencapai 0.9175, yang mengindikasikan bahwa dari semua prediksi positif yang dilakukan oleh model, 91.75% dari prediksi tersebut adalah benar. Recall atau sensitivitas model adalah 0.9127, menunjukkan kemampuan model untuk

mengidentifikasi 91.27% dari semua instance positif yang sebenarnya. Terakhir, F1 Score, yang merupakan rata-rata harmonis dari precision dan recall, berada pada 0.9151, menunjukkan keseimbangan yang baik antara kedua metrik tersebut.

Grafik yang menunjukkan perkembangan akurasi dan loss selama pelatihan dan validasi dapat dilihat pada Gambar 7. Grafik pertama (Model Accuracy) menunjukkan bahwa akurasi untuk data pelatihan dan validasi meningkat secara konsisten seiring dengan bertambahnya epoch, mencapai stabilitas setelah sekitar 5 epoch. Grafik kedua (Model Loss) menunjukkan penurunan yang signifikan dalam loss untuk kedua set data, dengan loss validasi yang sedikit lebih tinggi dibandingkan dengan loss pelatihan, namun keduanya menunjukkan tren penurunan yang baik, menunjukkan bahwa model tidak mengalami overfitting yang signifikan. Kedua grafik ini menunjukkan bahwa model BiLSTM yang kami gunakan telah dilatih dengan baik dan mampu menggeneralisasi dengan baik pada data yang belum pernah dilihat sebelumnya, yang sesuai dengan metrik evaluasi yang telah disebutkan di atas.

4.14. Evaluasi

Pada judul jurnal saya yaitu Pemanfaatan LSTM untuk Menganalisis Sentimen Pengguna Twitter: Studi Kasus pada Tweet Berita Terkini. di tahap hasil dan pembahasan di bagian evaluasi berikan saya data tabel dan formula dengan Accuracy sebesar 91,52%, Precision 91,75%, Recall 91,27%, dan F1 Score 91,51%.

Berikut adalah tabel data evaluasi untuk jurnal dengan judul "Pemanfaatan LSTM untuk Menganalisis Sentimen Pengguna Twitter: Studi Kasus pada Tweet Berita Terkini" dan formula untuk menghitung metrik evaluasi tersebut:

Tabel 7 Evaluasi

Metrik	Nilai(%)
Accuracy	91,52
Precision	91,75
Recall	91,27
F1 Score	91,51

Formula untuk menghitung metrix:

1. Accuracy :

$$Accuracy = \frac{Tp + TN}{TP + TN + Fp + FN} \times 100\%$$

- TP: True Positive
- TN: True Negative
- FP: False Positive
- FN: False Negative

Dengan nilai :

$$Accuracy = \frac{Tp+TN}{TP+TN+Fp+FN} \times 91,25\% \quad (10)$$

2. Precision (Presisi):

$$Precision = \frac{TP}{TP + FP} \times 100\%$$

$$\text{Dengan nilai : Precision} = \frac{TP}{FP+FP} = 91,75\% \quad (11)$$

$$3. \text{ Recall : Recall} = \frac{TP}{TP+FN} = 91,27\% \quad (12)$$

$$4. \text{ F1 Score : F1 Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Dengan Nilai :

$$\text{F1 Score} = 2 \times \frac{91,75 \times 91,27}{91,75 + 91,27} = 91,51\% \quad (13)$$

Berdasarkan metrik yang telah diberikan sebelumnya (Accuracy 91,52%, Precision 91,75%, Recall 91,27%, dan F1 Score 91,51%), dan dengan asumsi total sampel $N = 1000$, kita bisa menghitung perkiraan nilai TP, TN, FP, dan FN seperti berikut:

TP (True Positive): Jumlah tweet yang benar diprediksi sebagai positif. Menggunakan rumus Precision dan Recall, kita mendapatkan:

$$TP = \frac{91,75 \times 91,27}{91,75 + 91,27 - (91,75 \times 01,27)} \times 915,2 \approx 416 \quad (14)$$

TN (True Negative): Jumlah tweet yang benar diprediksi sebagai negatif. Dari Accuracy, kita tahu bahwa $TP + TN = 915,2$, jadi:

$$TN = 915,2 - TP \approx 499,2 \quad (15)$$

FP (False Positive): Jumlah tweet yang salah diprediksi sebagai positif. Menggunakan Precision :

$$FP = \frac{TP}{0,9175} - TP \approx 37 \quad (16)$$

N (False Negative): Jumlah tweet yang salah diprediksi sebagai negatif. Menggunakan Recall:

$$FN = \frac{TP}{0,9127} - TP \approx 40 \quad (17)$$

Berikut adalah tabel yang menampilkan perkiraan nilai TP, TN, FP, dan FN berdasarkan metrik yang telah diberikan:

Tabel 8. perkiraan nilai

Kategori	Nilai
TP (True Positive)	416
TN (True Negative)	499
FP (False Positive)	37
FN (False Negative)	40

Berdasarkan hasil evaluasi model LSTM untuk analisis sentimen pada tweet berita terkini, kita dapat melihat bahwa model ini menunjukkan performa yang sangat baik dengan akurasi sebesar 91,52%. Ini berarti model mampu memprediksi dengan benar lebih dari 91% dari total tweet dalam dataset uji. Dari tabel perkiraan nilai TP, TN, FP, dan FN, kita bisa melihat bahwa model memiliki 416 True Positive, yang menunjukkan jumlah tweet yang benar diprediksi sebagai positif, dan 499 True Negative, yang menunjukkan jumlah tweet yang benar diprediksi

sebagai negatif. Selain itu, terdapat 37 False Positive, yaitu tweet yang salah diprediksi sebagai positif, dan 40 False Negative, yaitu tweet yang salah diprediksi sebagai negatif. Precision model mencapai 91,75%, menunjukkan bahwa dari semua tweet yang diprediksi sebagai positif, 91,75% benar-benar positif, memberikan keyakinan bahwa prediksi positif dari model ini dapat diandalkan. Recall sebesar 91,27% menunjukkan bahwa model mampu menemukan 91,27% dari semua tweet yang sebenarnya positif, yang penting untuk memastikan bahwa sebagian besar tweet positif tidak terlewatkan. F1 Score yang tinggi, sebesar 91,51%, menunjukkan keseimbangan yang baik antara Precision dan Recall, memperkuat bahwa model ini memiliki performa yang konsisten dalam mengidentifikasi sentimen baik positif maupun negatif. Secara keseluruhan, hasil ini menunjukkan bahwa model LSTM ini sangat efektif dalam tugas analisis sentimen pada tweet berita terkini, dengan kemampuan untuk memprediksi dengan akurasi tinggi dan menyeimbangkan antara presisi dan recall.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil menunjukkan bahwa penggunaan Long Short-Term Memory (LSTM) dapat diterapkan dengan baik dalam menganalisis sentimen tweet berita terkini. Dengan menggunakan dataset sebanyak 123,218 tweet yang terdiri dari 75,482 tweet positif dan 47,736 tweet negatif, model LSTM menunjukkan performa yang sangat baik dengan akurasi 91,52%, presisi 91,75%, recall 91,27%, dan F1 score 91,51%. Selain itu, analisis distribusi panjang teks menunjukkan bahwa panjang tweet bukanlah faktor utama dalam membedakan sentimen. Model juga menunjukkan performa yang konsisten pada data pelatihan dan validasi, menandakan tidak adanya overfitting. Dengan hanya 37 False Positive dan 40 False Negative, model ini berhasil mengidentifikasi sentimen dengan tingkat kesalahan yang rendah.

Berdasarkan hasil yang diperoleh dalam penelitian ini, beberapa rencana penelitian selanjutnya dapat dipertimbangkan untuk meningkatkan kualitas dan cakupan penelitian. Pertama, untuk memperluas pemahaman tentang perbedaan sentimen, disarankan untuk memperkaya dataset dengan menambahkan lebih banyak tweet dari berbagai topik berita terkini serta dari periode waktu yang berbeda. Hal ini diharapkan dapat membantu memahami dinamika perubahan sentimen seiring berjalannya waktu. Selain itu, eksperimen dengan model yang lebih canggih atau variasi arsitektur LSTM, seperti Bidirectional LSTM atau Attention Mechanism, dapat memberikan wawasan lebih dalam tentang akurasi prediksi yang lebih baik. Terakhir, penelitian lebih lanjut dapat mempertimbangkan pengaruh faktor eksternal, seperti tren global atau peristiwa besar, terhadap perubahan sentimen pada tweet, untuk menganalisis bagaimana konteks tertentu dapat memengaruhi opini publik.

DAFTAR PUSTAKA

- [1] F. Yuspriyadi, "Klasifikasi Sentimen Twitter Menggunakan Lstm," *Method. J. Tek. Inform. dan Sist. Inf.*, vol. 9, no. 1, pp. 4–8, 2023, doi: 10.46880/mtk.v9i1.1720.
- [2] T. Pemilu, M. Khadapi, and V. M. Pakpahan, "Analisis Sentimen Berbasis Jaringan LSTM dan BERT terhadap Diskusi," vol. 6, pp. 130–137, 2024.
- [3] F. Ratnawati, "Implementasi Algoritma Naive Bayes Terhadap Analisis Sentimen Opini Film Pada Twitter," *INOVTEK Polbeng - Seri Inform.*, vol. 3, no. 1, p. 50, 2018, doi: 10.35314/isi.v3i1.335.
- [4] G. A. Sandag and J. Waworundeng, "Analisis Sentimen Masyarakat Terhadap Exchange Tokocrypto Pada Twitter Menggunakan Metode LSTM," *CogITO Smart J.*, vol. 8, no. 2, pp. 411–421, 2022, doi: 10.31154/cogito.v8i2.418.411-421.
- [5] A. Yahyadi and F. Latifah, "Analisis Sentimen Twitter Terhadap Kebijakan Ppkm Di Tengah Pandemi Covid-19 Menggunakan Mode Lstm," *J. Inf. Syst. Applied, Manag. Account. Res.*, vol. 6, no. 2, pp. 464–470, 2022, doi: 10.52362/jisamar.v6i2.791.
- [6] I. H. Kusuma and N. Cahyono, "Analisis Sentimen Masyarakat Terhadap Penggunaan E-Commerce Menggunakan Algoritma K-Nearest Neighbor," *J. Inform. J. Pengemb. IT*, vol. 8, no. 3, pp. 302–307, 2023, doi: 10.30591/jpit.v8i3.5734.
- [7] A. Azrul, A. Irma Purnamasari, and I. Ali, "Analisis Sentimen Pengguna Twitter Terhadap Perkembangan Artificial Intelligence Dengan Penerapan Algoritma Long Short-Term Memory (Lstm)," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 1, pp. 413–421, 2024, doi: 10.36040/jati.v8i1.8416.
- [8] T. Krisdiyanto, "Analisis Sentimen Opini Masyarakat Indonesia Terhadap Kebijakan PPKM pada Media Sosial Twitter Menggunakan Naïve Bayes Clasifiers," *J. CoreIT J. Has. Penelit. Ilmu Komput. dan Teknol. Inf.*, vol. 7, no. 1, p. 32, 2021, doi: 10.24014/coreit.v7i1.12945.
- [9] L. Farsiah, A. Misbullah, and H. Husaini, "Analisis Sentimen Menggunakan Arsitektur Long Short-Term Memory (Lstm) Terhadap Fenomena Citayam Fashion Week," *Cybersp. J. Pendidik. Teknol. Inf.*, vol. 6, no. 2, p. 86, 2022, doi: 10.22373/cj.v6i2.14687.
- [10] D. S. Informasi, F. Teknologi, E. Dan, and I. Cerdas, *KLASIFIKASI SENTIMEN PEMBELAJARAN DARING PADA TWEET TWITTER ALGORITME BIDIRECTIONAL-LSTM SENTIMENT CLASSIFICATION OF ONLINE LEARNING ON TWITTER USER 'S TWEET USING BIDIRECTIONAL-LSTM ALGORITHM* i. 2021.
- [11] D. Atika, A. Ari Aldino, S. Informasi, J. Pagar Alam No, L. Ratu, and K. Kedaton, "Term Frequency-Inverse Document Frequency Support Vector Machine Untuk Analisis Sentimen Opini Masyarakat Terhadap Tekanan Mental Pada Media Sosial Twitter," *J. Teknol. dan Sist. Inf.*, vol. 3, no. 4, p. page-page, 2022, [Online]. Available: <http://jim.teknokrat.ac.id/index.php/JTSI>.
- [12] A. A. Mudding and Arifin A Abd Karim, "Analisis Sentimen Menggunakan Algoritma Lstm Pada Media Sosial," *J. Publ. Ilmu Komput. dan Multimed.*, vol. 1, no. 3, pp. 181–187, 2022, doi: 10.55606/jupikom.v1i3.517.
- [13] M. H. Al-Areef and K. Saputra S, "Analisis Sentimen Pengguna Twitter Mengenai Calon Presiden Indonesia Tahun 2024 Menggunakan Algoritma LSTM," *J. SAINTIKOM (Jurnal Sains Manaj. Inform. dan Komputer)*, vol. 22, no. 2, p. 270, 2023, doi: 10.53513/jis.v22i2.8680.